

PD CEN ISO/TR 12489:2016



BSI Standards Publication

Petroleum, petrochemical and natural gas industries — Reliability modelling and calculation of safety systems

bsi.

...making excellence a habit.™

National foreword

This Published Document is the UK implementation of CEN ISO/TR 12489:2016. It is identical to ISO/TR 12489:2013. It supersedes PD ISO/TR 12489:2013 which is withdrawn.

The UK participation in its preparation was entrusted to Technical Committee PSE/17, Materials and equipment for petroleum, petrochemical and natural gas industries.

A list of organizations represented on this committee can be obtained on request to its secretary.

This publication does not purport to include all the necessary provisions of a contract. Users are responsible for its correct application.

© The British Standards Institution 2016.
Published by BSI Standards Limited 2016

ISBN 978 0 580 88682 9

ICS 75.180.01; 75.200

Compliance with a British Standard cannot confer immunity from legal obligations.

This British Standard was published under the authority of the Standards Policy and Strategy Committee on 30 November 2013.

Amendments/corrigenda issued since publication

Date	Text affected
29 February 2016	This corrigendum renumbers PD ISO/TR 12489:2013 as PD CEN ISO/TR 12489:2016.

ICS 75.200; 75.180.01

English Version

**Petroleum, petrochemical and natural gas industries -
Reliability modelling and calculation of safety systems
(ISO/TR 12489:2013)**

Pétrole, pétrochimie et gaz naturel - Modélisation et
calcul fiabilistes des systèmes de sécurité (ISO/TR
12489:2013)

Erdöl-, petrochemische und Erdgasindustrie -
Zuverlässigkeit der Modellierung und Berechnung von
Sicherheitssystemen (ISO/TR 12489:2013)

This Technical Report was approved by CEN on 28 March 2015. It has been drawn up by the Technical Committee CEN/TC 12.

CEN members are the national standards bodies of Austria, Belgium, Bulgaria, Croatia, Cyprus, Czech Republic, Denmark, Estonia, Finland, Former Yugoslav Republic of Macedonia, France, Germany, Greece, Hungary, Iceland, Ireland, Italy, Latvia, Lithuania, Luxembourg, Malta, Netherlands, Norway, Poland, Portugal, Romania, Slovakia, Slovenia, Spain, Sweden, Switzerland, Turkey and United Kingdom.



EUROPEAN COMMITTEE FOR STANDARDIZATION
COMITÉ EUROPÉEN DE NORMALISATION
EUROPÄISCHES KOMITEE FÜR NORMUNG

CEN-CENELEC Management Centre: Avenue Marnix 17, B-1000 Brussels

European foreword

This document (CEN ISO/TR 12489:2016) has been prepared by Technical Committee ISO/TC 67 “Materials, equipment and offshore structures for petroleum, petrochemical and natural gas industries” in collaboration with Technical Committee CEN/TC 12 “Materials, equipment and offshore structures for petroleum, petrochemical and natural gas industries” the secretariat of which is held by NEN.

Attention is drawn to the possibility that some of the elements of this document may be the subject of patent rights. CEN [and/or CENELEC] shall not be held responsible for identifying any or all such patent rights.

Endorsement notice

The text of ISO/TR 12489:2013 has been approved by CEN as CEN ISO/TR 12489:2016 without any modification.

Contents

Page

Foreword	v
Introduction	vi
1 Scope	1
2 Analysis framework	2
2.1 Users of this Technical Report	2
2.2 ISO/TR 12489 with regard to risk and reliability analysis processes	2
2.3 Overview of the reliability modelling and calculation approaches considered in this Technical Report	4
2.4 Safety systems and safety functions	7
3 Terms and definitions	8
3.1 Basic reliability concepts	8
3.2 Failure classification	20
3.3 Safety systems typology	24
3.4 Maintenance issues	25
3.5 Other terms	28
3.6 Equipment-related terms	29
4 Symbols and abbreviated terms	30
5 Overview and challenges	33
5.1 General considerations about modelling and calculation challenges	33
5.2 Deterministic versus probabilistic approaches	35
5.3 Safe failure and design philosophy	35
5.4 Dependent failures	36
5.5 Human factors	37
5.6 Documentation of underlying assumptions	40
6 Introduction to modelling and calculations	41
6.1 Generalities about safety systems operating in “on demand” or “continuous” modes	41
6.2 Analytical approaches	44
7 Analytical formulae approach (low demand mode)	47
7.1 Introduction	47
7.2 Underlying hypothesis and main assumptions	47
7.3 Single failure analysis	48
7.4 Double failure analysis	50
7.5 Triple failure analysis	55
7.6 Common cause failures	56
7.7 Example of implementation of analytical formulae: the PDS method	57
7.8 Conclusion about analytical formulae approach	57
8 Boolean and sequential approaches	58
8.1 Introduction	58
8.2 Reliability block diagrams (RBD)	58
8.3 Fault Tree Analysis (FTA)	59
8.4 Sequence modelling: cause consequence diagrams, event tree analysis, LOPA	61
8.5 Calculations with Boolean models	61
8.6 Conclusion about the Boolean approach	64
9 Markovian approach	65
9.1 Introduction and principles	65
9.2 Multiphase Markov models	68
9.3 Conclusion about the Markovian approach	69
10 Petri net approach	69
10.1 Basic principle	69
10.2 RBD driven Petri net modelling	71

10.3	Conclusion about Petri net approach	74
11	Monte Carlo simulation approach	74
12	Numerical reliability data uncertainty handling	74
13	Reliability data considerations	75
13.1	Introduction	75
13.2	Reliability data sources	76
13.3	Required reliability data	78
13.4	Reliability data collection	80
14	Typical applications	80
14.1	Introduction	80
14.2	Typical application TA1: single channel	82
14.3	Typical application TA2: dual channel	97
14.4	Typical application TA3: popular redundant architecture	110
14.5	Typical application TA4: multiple safety system	119
14.6	Typical application TA5: emergency depressurization system (EDP)	124
14.7	Conclusion about typical applications	135
Annex A	(informative) Systems with safety functions	136
Annex B	(informative) State analysis and failure classification	146
Annex C	(informative) Relationship between failure rate, conditional and unconditional failure intensities and failure frequency	152
Annex D	(informative) Broad models for demand mode (reactive) safety systems	160
Annex E	(informative) Continuous mode (preventive) safety systems	167
Annex F	(informative) Multi-layers safety systems/multiple safety systems	170
Annex G	(informative) Common cause failures	173
Annex H	(informative) The human factor	180
Annex I	(informative) Analytical formulae	186
Annex J	(informative) Sequential modelling	207
Annex K	(informative) Overview of calculations with Boolean models	213
Annex L	(informative) Markovian approach	221
Annex M	(informative) Petri net modelling	239
Annex N	(informative) Monte Carlo simulation approach	248
Annex O	(informative) Numerical uncertainties handling	252
Bibliography		255

Foreword

ISO (the International Organization for Standardization) is a worldwide federation of national standards bodies (ISO member bodies). The work of preparing International Standards is normally carried out through ISO technical committees. Each member body interested in a subject for which a technical committee has been established has the right to be represented on that committee. International organizations, governmental and non-governmental, in liaison with ISO, also take part in the work. ISO collaborates closely with the International Electrotechnical Commission (IEC) on all matters of electrotechnical standardization.

The procedures used to develop this document and those intended for its further maintenance are described in the ISO/IEC Directives, Part 1. In particular the different approval criteria needed for the different types of ISO documents should be noted. This document was drafted in accordance with the editorial rules of the ISO/IEC Directives, Part 2 (see www.iso.org/directives).

Attention is drawn to the possibility that some of the elements of this document may be the subject of patent rights. ISO shall not be held responsible for identifying any or all such patent rights. Details of any patent rights identified during the development of the document will be in the Introduction and/or on the ISO list of patent declarations received (see www.iso.org/patents).

Any trade name used in this document is information given for the convenience of users and does not constitute an endorsement.

For an explanation on the meaning of ISO specific terms and expressions related to conformity assessment, as well as information about ISO's adherence to the WTO principles in the Technical Barriers to Trade (TBT) see the following URL: Foreword - Supplementary information

The committee responsible for this document is ISO/TC 67, *Materials, equipment and offshore structures for petroleum, petrochemical and natural gas industries*.

This first edition of ISO/TR 12489 belongs of the family of reliability related standards developed by ISO/TC 67:

- ISO 14224, *Petroleum, petrochemical and natural gas industries — Collection and exchange of reliability and maintenance data for equipment*
- ISO 20815, *Petroleum, petrochemical and natural gas industries — Production assurance and reliability management*

Introduction

Safety systems have a vital function in petroleum, petrochemical and natural gas industries where safety systems range from simple mechanical safety devices to safety instrumented systems.

They share three important characteristics which make them difficult to handle:

- 1) They should be designed to achieve good balance between safety and production. This implies a high probability of performing the safety action as well as a low frequency of spurious actions.
- 2) Some of their failures are not revealed until relevant periodic tests are performed to detect and repair them.
- 3) A given safety system rarely works alone. It generally belongs to a set of several safety systems (so-called multiple safety systems) working together to prevent accidents.

Therefore improving safety may be detrimental to dependability and vice versa. These two aspects should therefore, ideally, be handled at the same time by the same reliability engineers. However, in reality they are generally considered separately and handled by different persons belonging to different departments. Moreover this is encouraged by the international safety standards, which exclude dependability from their scopes, and the international dependability (see [3.1.1](#)) standard, which excludes safety from theirs. This may lead to dangerous situations (e.g. safety system disconnected because of too many spurious trips) as well as high production losses.

The proof of the conservativeness of probabilistic calculations of safety systems is generally required by safety authorities. Unfortunately, managing the systemic dependencies introduced by the periodic tests to obtain conservative results implies mathematical difficulties which are frequently ignored. The impact is particularly noticeable for redundant safety systems and multiple safety systems. Awareness of these challenges is important for reliability engineers as well as safety managers and decision makers, utilizing reliability analytical support.

Most of the methods and tools presently applied in reliability engineering have been developed since the 1950s before the emergence of personal computers when only pencil and paper were available. At that time the reliability pioneers could only manage simplified models and calculations but this has completely changed because of the tremendous improvement in the computation means achieved over the past 30 years. Nowadays, models and calculations which were once impossible are carried out with a simple laptop computer. Flexible (graphical) models and powerful algorithms based on sound mathematics are now available to handle “industrial size” systems (i.e. many components with complex interactions). This allows the users to focus on the analysis of the systems and assessment of results, rather than on the calculations themselves. All the approaches described in this Technical Report have been introduced in the petroleum, petrochemical and natural gas industries as early as the 1970s where they have proven to be very effective. They constitute the present time state-of-the-art in reliability calculations. Nevertheless some of them have not been widely disseminated in this sector although they can be of great help for reliability engineers to overcome the problems mentioned above. This is particularly true when quantitative reliability or availability requirements need confirmation and/or when the objective of the reliability study lay beyond the scope of the elementary approaches.

The present document is a “technical” report and its content is obviously “technical”. Nevertheless, it only requires a basic knowledge in probabilistic calculation and mathematics and any skilled reliability engineer should have no difficulties in using it.

Petroleum, petrochemical and natural gas industries — Reliability modelling and calculation of safety systems

1 Scope

This Technical Report aims to close the gap between the state-of-the-art and the application of probabilistic calculations for the safety systems of the petroleum, petrochemical and natural gas industries. It provides guidelines for reliability and safety system analysts and the oil and gas industries to:

- understand the correct meaning of the definitions used in the reliability field;
- identify
 - the safety systems which may be concerned,
 - the difficulties encountered when dealing with reliability modelling and calculation of safety systems,
 - the relevant probabilistic parameters to be considered;
- be informed of effective solutions overcoming the encountered difficulties and allowing to undertake the calculations of relevant probabilistic parameters;
- obtain sufficient knowledge of the principles and framework (e.g. the modelling power and limitations) of the well-established approaches currently used in the reliability field:
 - analytical formulae;[\[1\]](#)[\[2\]](#)[\[13\]](#)
 - Boolean:
 - reliability block diagrams;[\[4\]](#)
 - fault trees;[\[5\]](#)
 - sequential: event trees,[\[8\]](#) cause consequence diagrams[\[10\]](#) and LOPA;[\[9\]](#)
 - Markovian;[\[6\]](#)
 - Petri nets;[\[7\]](#)
- obtain sufficient knowledge of the principles of probabilistic evaluations:
 - analytical calculations (e.g. performed on Boolean or Markovian models);[\[1\]](#)[\[2\]](#)[\[3\]](#)
 - and Monte Carlo simulation (e.g. performed on Petri nets[\[7\]](#));
- select an approach suitable with the complexity of the related safety system and the reliability study which is undertaken;
- handle safety and dependability (e.g. for production assurance purpose, see [3.1.1](#)) within the same reliability framework.

The elementary approaches (e.g. PHA, HAZID, HAZOP, FMECA) are out of the scope of this Technical Report. Yet they are of utmost importance and ought to be applied first as their results provide the input information essential to properly undertake the implementation of the approaches described in this Technical Report: analytical formulae, Boolean approaches (reliability block diagrams, fault trees, event trees, etc.), Markov graphs and Petri nets.

This Technical Report is focused on probabilistic calculations of random failures and, therefore, the non-random (i.e. systematic failures as per the international reliability vocabulary IEV 191^[14]) failures are out of the scope even if, to some extent, they are partly included into the reliability data collected from the field.

2 Analysis framework

2.1 Users of this Technical Report

This Technical Report is intended for the following users, in a role defining the scope of work of reliability models (customer or decision-maker), executing reliability analysis or as a risk analyst using these calculations:

- **Installation/Plant/Facility:** operating facility staff, e.g. safety, maintenance and engineering personnel.
- **Owner/Operator/Company:** reliability staff or others analysing or responsible for reliability studies for safety related equipment located in company facilities.
- **Industry:** groups of companies collaborating to enhance reliability of safety systems and safety functions. The use of this Technical Report supports “reliability analytical best practices” for the benefit of societal risk management in accordance with ISO 26000^[54].
- **Manufacturers/Designers:** users having to document the reliability of their safety equipment.
- **Authorities/Regulatory bodies:** enforcers of regulatory requirements which can quote these guidelines to enhance quality and resource utilization.
- **Consultant/Contractor:** experts and contractors/consultants undertaking reliability modelling and probabilistic calculation studies.
- **University bodies:** those having educational roles in society and experts that might improve methods on these matters.
- **Research institutions:** experts that might improve reliability modelling and probabilistic calculation methods.

2.2 ISO/TR 12489 with regard to risk and reliability analysis processes

When a safety system has been designed using good engineering practice (i.e. applying the relevant regulations, standards, rules and technical and safety requirements) it is expected to work properly. After that a reliability analysis is usually undertaken in order to evaluate its probability of failure and, if needed, identify how it can be improved to reach some safety targets.

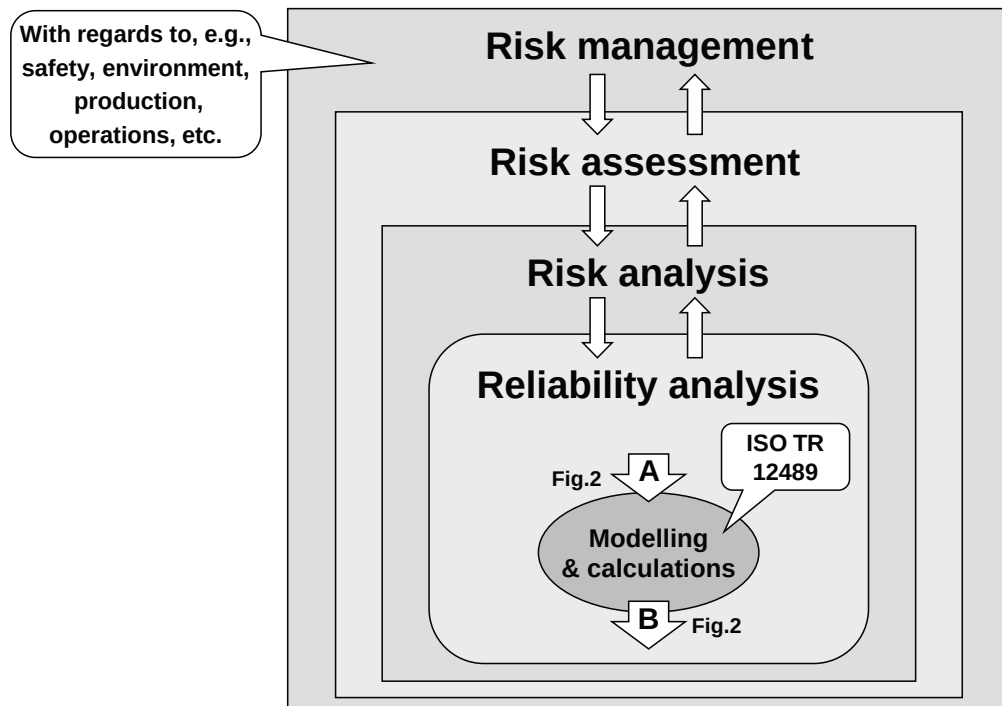


Figure 1 — ISO/TR 12489 within the framework of risk management

Relevant interdisciplinary communication and a good understanding of the safety system life cycle are required to have qualified inputs and correct result interpretations. Applying this Technical Report also requires interaction and compliance with other standards such as ISO 20815^[16] (production assurance), ISO 14224^[15] (reliability data collection) or ISO 17776^[29] and ISO 31000^[28] (risk management). As shown in [Figure 1](#), this Technical Report contributes to the risk management process which encompasses both safety and production (dependability, cf. [3.1.1](#)) aspects and involves different stages such as risk assessment and risk analysis. More precisely, this Technical Report contributes to the probabilistic part (reliability analysis) of the risk analysis stage.

NOTE ISO 20815^[16] gives further information on reliability/availability in a production assurance perspective, while ISO 14224^[15] which is devoted to reliability data collection is another fundamental reference for both safety and production within our industries (within ISO/TC67 business arena). ISO 17776^[29] and ISO 31000^[28] are devoted to risk management.

When such a process is undertaken, the usual steps are the following:

- a) Defining the objective of the study and system boundaries in order to identify the limits of the process and the safety system(s) to be analysed.
- b) Functioning analysis to understand how the safety system works.
- c) Dysfunctioning analysis to understand how the safety system may fail:
 - 1) risk identification and establishment of the safety targets;
 - 2) elementary analyses (e.g. HAZOP, FMEA, etc.);
 - 3) common cause failures identification.
- d) **Modelling and calculations:**
 - 1) **Modelling:**
 - i) **functioning and dysfunctioning modelling**

ii) **common cause/ Common mode failures modelling**

2) **Qualitative analysis:**

3) **Quantitative analysis** (if qualitative analysis is not sufficient).

e) Discussion with field specialists and redesign if improvements are needed.

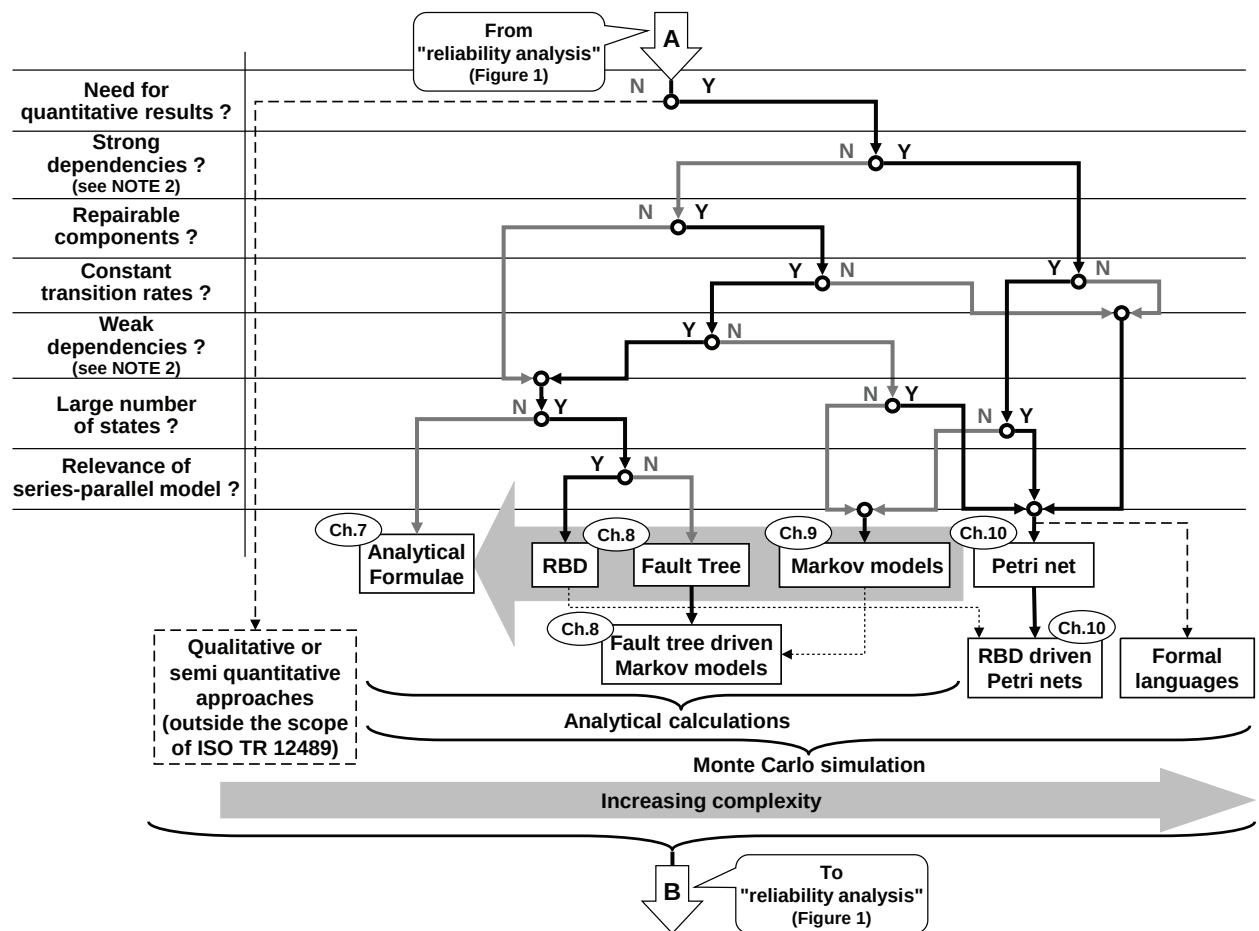
f) **Final results** (weak points, failure contributors, failure probabilities, interpretation, specifications, etc.).

The present Technical Report is focused on the steps written in bold and underlined characters: modelling and calculations [step d)] and final results of interest [step f)]. Nevertheless, step d) and consequently f) can be achieved only if the steps a), b) and c) and consequently e) have been properly undertaken first. Therefore in this Technical Report it is supposed that the limits of the safety system and the objective of the study have been properly identified [step a)], that the analyst has acquired a sound understanding about the functioning [step b)] and dysfunctioning of the safety system under study, that the relevant risk identification and the safety targets have been properly established [and c)] and that field specialists have been invited to give their advice in due time [step e)] to ensure that the final results are close to real life feedback.

This Technical Report also suggests the safety systems and safety functions typically requiring such reliability analysis support in order to utilize resources effectively. See [Annex A](#).

2.3 Overview of the reliability modelling and calculation approaches considered in this Technical Report

[Figure 2](#) gives an overview of the approaches selected for the purpose of this Technical Report and provides some guidelines to select them when the level in difficulty and complexity increases.



NOTE 1 The questions on the left hand side can be used as guidelines to choose an adequate approach to study a given safety system.

NOTE 2 Systems without dependencies do not really exist in the real world but the dependencies may have a negligible impact (weak dependencies) or a strong impact (strong dependencies) on the probability of failure. An example of weak dependency is the use of a single repair team for a topside periodically tested component (because the repair time is negligible compared to the MFDT (Mean Fault Detection Time, see 3.1.35). An example of strong dependency is when a stand-by component starts when another fails.

NOTE 3 "Series-parallel model" refers to a popular model found in numerous text books which uses only series and parallel structures to model the logic of the systems, for example, reliability block diagrams[4].

NOTE 4 The arrow from "Markov" to "Analytical Formulae" through "fault tree" and "RBD" highlights the fact that the analytical formulae are obtained through models mixing Boolean[4][5] and Markov[6] models.

Figure 2 — Overview of reliability modelling and calculation approaches currently used

Other criteria can be used to classify the reliability modelling and calculation approaches:

- the accuracy of results (approximated or exact);
- conservativeness of the results (pessimistic or optimistic);
- the nature of the calculations (analytical or Monte Carlo simulation);
- the nature of the modelling (static or dynamic);
- the user friendliness (graphical or non graphical);
- the input data which can be made available;

- the possibility to update the model after several years by someone else.

The various approaches currently used in reliability engineering have different characteristics (strengths and limitations). It is important for the selection and use of these approaches to be aware of their limitations and conservativeness:

- Analytical formulae:**^{[1][2][13]} analytical methods which provide approximated suitable results when used skilfully. They are useful for quick calculations but the underlying limits and approximations often limit their application to systems of limited complexity. This also limits their application to systems where sequence-dependent failures or other time-dependent failures, such as desynchronized testing (see [3.4.10](#)), are not important contributors to the overall performance. Analytical formulae are generally obtained from underlying Boolean and/or Markovian models.
- Boolean models:** static and graphical models supporting analytical calculations. “Reliability block diagrams” (RBD)^[4] and fault trees (FT)^[5] belong to Boolean models. To some extent, the sequential approaches event trees (ET)^[8], LOPA^[9] or cause consequence diagrams^[10] can also be associated with Boolean models. These approaches provide clear and understandable models for large or complex systems. Boolean models are limited to “two-state” systems (working, failed) and handling of time evolution requires a high level of understanding in probabilistic calculations.
- Markovian models**^[6]: dynamic and graphical models supporting analytical calculations and modelling of sequence-dependent or time-dependent failures. A Markovian model is a “state-transition” model limited to exponentially distributed events. The combinatory explosion of the number of system states limits this approach to small (simple or complex) systems with few states. The impact of approximations performed to deal with larger systems is often difficult to evaluate. Boolean and Markovian approaches can be mixed to model large systems when weak dependencies between the components are involved. This can be achieved by implementing the fault tree driven Markov models (see [Figure 2](#)).
- Petri nets**^[7]: dynamic and graphical models supporting Monte Carlo simulation to provide statistical results associated with their confidence intervals. A Petri net is a “state-transition” model handling any kind of probabilistic distributions. Time-, state- or sequence-dependent failures can be modelled explicitly. The size of the model is linear with regard to the number of components. This makes possible the modelling of very large complex systems. The Monte Carlo simulation computation time increases when low probability events are calculated but probabilities of failure as low as 10^{-5} over one year can be handled with modern personal computers. For large safety systems, the Petri net may become difficult to handle. The use of the RBD driven PN overcomes this difficulty (see [Figure 2](#)).
- Formal languages**^{[11][12]}: dynamic models used to generate analytical models (e.g. Markovian models or fault trees, when possible) or used directly for Monte Carlo simulation. The other characteristics are same as Petri nets except that computations may be slower. They are just mentioned but they are outside the scope of this Technical Report.

Except for bullet e), more details can be found in [Clauses 7](#) to [10](#). All these models can be mathematically described in terms of “finite states automata” (i.e. a mathematical *state machine* with a finite number of discrete states). The system behaviour can be modelled more and more rigorously when going from a) to e) but, of course, every approach can be used to model simple safety systems.

[Figure 2](#) gives advice to the analyst to select the relevant approach in order to optimize the design of a safety system and meet some reliability targets. This choice depends on the safety function, purpose and complexity the analyst has to face. When several approaches are relevant, the analyst may choose his favourite.

A warning may be raised here: using a software package as a black box or a formula as a magic recipe is likely to lead to inaccurate, often non-conservative, results. In all cases the reliability engineers should be aware of the limitations of the tools that they are using and they should have a minimum understanding of the mathematics behind the calculations and a good knowledge of the nature of the results that they obtain (unreliability, point unavailability, average unavailability, frequency, etc.), of the

conservativeness and of the associated uncertainties. Without adequate understanding of the software tool, erroneous results can be obtained through its misuse.

Table 1 — Road map of ISO/TR 12489

Topic	Reference to main report (sub)clause	Reference to annexes
I- General issues		
a) Terms and definitions	3, 4	-
b) General analytical overview	5, 6	B, C, D, E, F
c) Human factors	5.5	H
d) Common cause	5.4.2	G
e) Monte Carlo simulation	11	N
f) Uncertainty	12	O
g) Reliability data	13	-
h) Systems with safety functions	2.4	A
II- Approaches		
a) Analytical formulae	7	I
b) Boolean	8	K
- Reliability Block Diagram	8.2	
- Fault Tree	8.3	
- Sequence modelling	8.4	J
c) Markovian	9	L
d) Petri net	10	M
III- Examples	14	-
IV- Bibliography	End of ISO/TR 12489	-

It is important that the reliability methods and application of those, including the available input data are adapted to the life cycle phase. Uncertainty handling is further addressed in [Clause 12](#).

The human factor is addressed in [5.5](#) and [Annex H](#) in terms of the quantification of the reliability of human performed tasks. This inclusion is intended to support assessment of the pros and cons of including human tasks with the potential for failure in safety systems.

[Table 1](#) gives a road map for these issues and the supporting annexes and supplement [Figure 2](#).

2.4 Safety systems and safety functions

Numerous safety systems are implemented in the petroleum, petrochemical and natural gas industries. They range from very simple to very complex systems, used on-demand or in continuous mode of operation.

[Table A.1](#) gives a non-exhaustive list of safety systems and safety functions which may require reliability modelling in the petroleum, petrochemical and natural gas industries. It has been built in relationship with the taxonomy developed in the ISO 14224[15] standard and covers either safety systems (taxonomy level 5) or other systems with safety function(s). A summary is given below:

- A. Emergency/process shutdown (split in A.1 and A.2)
- B. Fire and gas detection
- C. Fire water

- D. Fire-fighting
- E. Process control
- F. Public alarm
- G. Emergency preparedness systems
- H. Marine equipment
- I. Electrical and Telecommunication
- J. Other utilities
- K. Drilling and Wells
- L. Subsea

NOTE A to G are covered as safety and control systems in Table A.3 of ISO 14224[15]. The list has been extended from H to L to give a broader coverage.

This Technical Report provides a number of reliability modelling and calculation approaches large enough to cope with any kind of safety system like those identified in [Table A.1](#). They can be used when the objectives of the reliability studies lay beyond the scope of the elementary approaches (e.g. PHA, HAZID, HAZOP, FMECA ...) and selected according to [Figure 2](#).

3 Terms and definitions

For the purposes of this document, the following terms and definitions apply.

NOTE 1 Since their introduction more than 50 years ago, the core concepts of the reliability engineering field have been used and adapted for various purposes. Over time this has caused “semantic” drifts and most of the terms have various meanings. They have become so polysemic now that it is necessary to define them accurately to avoid confusion, even when they seem well known.

NOTE 2 The terms are divided into:

- [3.1](#) Basic reliability concepts
- [3.2](#) Failure classification
- [3.3](#) Safety systems typology
- [3.4](#) Maintenance issues
- [3.5](#) Other terms
- [3.6](#) Equipment related terms

Textual definitions are provided as well as, when this is possible, the corresponding mathematical formulae which leave less place to interpretation. Notes are added when clarifications are useful.

3.1 Basic reliability concepts

3.1.1

dependability

ability to perform as and when required

Note 1 to entry: Dependability is mainly business oriented.

Note 2 to entry: IEC/TC 56 which is the international “dependability” technical committee deals with reliability, availability, maintainability and maintenance support. More than 80 dependability standards have been published by the IEC/TC56. In particular, it is in charge of the international vocabulary related to those topics (IEV 191[14]) and also of the methods used in the reliability field (e.g. FMEA, HAZOP, reliability block diagrams, fault trees, Markovian approach, event tree, Petri nets).

Note 3 to entry: The *production availability* is an extension, for production systems, of the classical dependability measures. This term is defined in the ISO 20815[16] standard which deals with *production assurance* and relates to systems and operations associated with drilling, processing and transport of petroleum, petrochemical and natural gas. The relationship between production-assurance terms can be found in [Figure G.1](#) of ISO 20815[16].

[SOURCE: IEC 60050 –191]

3.1.2

safety integrity

ability of a safety instrumented system to perform the required safety instrumented functions as and when required

Note 1 to entry: This definition is equivalent to the dependability of the SIS (Safety Instrumented System) with regard to the required safety instrumented function. Dependability, being often understood as an economical rather a safety concept, has not been used to avoid confusion.

Note 2 to entry: The term “integrity” is used to point out that a SIS aims to protect the integrity of the operators as well as of the process and its related equipment from hazardous events.

3.1.3

SIL

Safety Integrity Level

discrete level (one out of four) for specifying the safety integrity requirements of the safety instrumented functions to be allocated to the safety instrumented systems

Note 1 to entry: Safety integrity level 4 is related to the highest level of safety integrity; safety integrity level 1 has the lowest.

Note 2 to entry: The safety integrity level is a requirement about a safety instrumented function. The higher the safety integrity level, the higher the probability that the required safety instrumented function (SIF) will be carried out upon a real demand.

Note 3 to entry: This term differs from the definition in IEC 61508–4[2] to reflect differences in process sector terminology.

3.1.4

safe state

state of the process when safety is achieved

Note 1 to entry: Some states are safer than others (see [Figures B.1, B.2 and B.3](#)) and in going from a potentially hazardous condition to the final safe state, or in going from the nominal safe condition to a potentially hazardous condition, the process may have to go through a number of intermediate safe-states.

Note 2 to entry: For some situations, a safe state exists only so long as the process is continuously controlled. Such continuous control may be for a short or an indefinite period of time.

Note 3 to entry: A state which is safe with regard to a given safety function may increase the probability of hazardous event with regard to another given safety function. In this case, the maximum allowable spurious trip rate (see [10.3](#)) for the first function should consider the potential increased risk associated with the other function.

3.1.5

dangerous state

state of the process when safety is not achieved

Note 1 to entry: A dangerous state is the result of the occurrence of a critical dangerous failure ([3.2.4, Figure B.1](#)).

3.1.6

safety function

function which is intended to achieve or maintain a safe state, in respect of a specific hazardous event

Note 1 to entry: This term deviates from the definition in IEC 61508-4 to reflect differences in process sector terminology.

3.1.7

safety system

system which realizes one or more safety functions

3.1.8

reliability

$R(t)$

⟨measure⟩ probability for an item to perform a required function under given conditions over a given time interval $z = \frac{S-m}{\sigma}$

Note 1 to entry: This is a time-dependent parameter.

Note 2 to entry: This parameter is related on a continuous functioning from 0 to t .

Note 3 to entry: For non-repairable items, Reliability and Availability are identical.

Note 4 to entry: In IEC 60500-191^[14], the reliability is defined both as ability and as measure.

3.1.9

unreliability

$F(t)$

⟨measure⟩ probability for an item to fail to perform a required function under given conditions over a given time interval $[0, t]$

Note 1 to entry: $F(t)$ is also the probability that the time of the first failure t_f is lower than t : $F(t) = P(t_f \leq t)$. This is in relationship with the occurrence of the first failure.

Note 2 to entry: $F(t)$ is the cdf (cumulative distribution function) of the time to the first failure t_f of the item. It ranges from 0 to 1 when t goes from 0 to infinity.

Note 3 to entry: The unreliability is the complementary of the reliability: $F(t) = 1 - R(t)$

Note 4 to entry: When dealing with safety, $F(t)$ is generally small compared to 1 and this property is used to develop approximated formulae (see [Clause 7](#)).

Note 5 to entry: Unreliability is better to communicate than MTTF.

3.1.10

failure probability density

$f(t)$

⟨measure⟩ probability for an item to fail between t and $t+dt$

Note 1 to entry: $f(t)$ is the classical pdf (probability density function) of the time to the first failure t_f of the item: $f(t) = P(t < t_f \leq t + dt)$

Note 2 to entry: $f(t)$ is the derivative of $F(t)$: $f(t) = \frac{dF(t)}{dt}$

Note 3 to entry: The failure density is linked to the failure rate by the following relation:

Note 4 to entry: $f(t) = \lambda(t)R(t) = \lambda(t)[1 - F(t)]$ (see [Annex C](#) for more details).

3.1.11

instantaneous availability

point availability

$A(t)$

probability for an item to be in a state to perform as required at a given instant

Note 1 to entry: In this Technical Report the word “availability” used alone stands for “instantaneous availability”.

Note 2 to entry: This is a time-dependent parameter.

Note 3 to entry: No matter if the item has failed before the given instant if it has been repaired before.

Note 4 to entry: For non-repairable items, Availability and Reliability are identical.

Note 5 to entry: When dealing with safety, $A(t)$ is generally close to 1 and this property is used to develop approximated formulae (see [Clause 7](#)).

3.1.12

unavailability

instantaneous unavailability

point unavailability

$U(t)$

(measure) probability for an item not to be in a state to perform as required at a given instant

Note 1 to entry: The unavailability is the complementary of the availability: $U(t) = 1 - A(t)$

Note 2 to entry: The unavailability is called “Probability of Failure on Demand” (PFD) by the standards related to the functional safety of safety related/instrumented systems (e.g. IEC 61508[2]).

Note 3 to entry: Note 3 to entry: When dealing with safety $U(t)$ is generally small compared to 1 and this property is used to develop approximated formulae (see [Clause 7](#)).

3.1.13

$\bar{A}(t_1, t_2)$

average availability

(measure) average value of the availability $A(t)$ over a given interval $[t_1, t_2]$

Note 1 to entry: The mathematic definition is the following: $\bar{A}(t_1, t_2) = \frac{1}{t_2 - t_1} \int_{t_1}^{t_2} A(\tau) d\tau$

Note 2 to entry: When $t_1 = 0$ and $t_2 = T$ the average availability becomes $\bar{A}(T) = \frac{1}{T} \int_0^T A(\tau) d\tau$

Note 3 to entry: Mathematically speaking, the average availability is the mathematical expectation of the availability. It does not have the mathematical property of a normal probability and cannot be handled as such.

3.1.14

average unavailability

$\bar{U}(t_1, t_2)$

(measure) average value of the unavailability $U(t)$ over a given interval

Note 1 to entry: The average unavailability is the complementary of the average availability:

Note 2 to entry: $\bar{U}(t_1, t_2) = 1 - \bar{A}(t_1, t_2)$ or $\bar{U}(T) = 1 - \bar{A}(T)$

Note 3 to entry: The average unavailability is called “average Probability of Failure on Demand” (PFD_{avg}) by the standards related to functional safety of safety related/instrumented systems (e.g. IEC 61508[2]): $PFD_{avg} = \bar{U}(T)$ where T is the overall life duration of the system.

Note 4 to entry: Mathematically speaking, the average unavailability is the mathematical expectation of the unavailability. It does not have the mathematical property of a normal probability and cannot be handled as such.

Note 5 to entry: When dealing with safety $\bar{U}(T)$ is generally small compared to 1 and this property is used to develop approximated formulae (see [Clause 7](#)).

3.1.15 probability of failure on demand PFD

unavailability as per [3.1.12](#) in the functional safety standard terminology (e.g. IEC 61508[2])

Note 1 to entry: “Failure on demand” means here “failure likely to be observed when a demand occurs”. This encompasses both the failure occurred before the demand and the failure occurring due to the demand itself. Then this term needs not to be mixed up with the probability of a failure due to a demand (see [3.2.13](#)).

3.1.16 average probability of failure on demand PFD_{avg}

average unavailability as per [3.1.12](#) in the functional safety standard terminology (e.g. IEC 61508[2])

Note 1 to entry: “Failure on demand” means here “failure likely to be observed when a demand occurs”. PFD_{avg} encompasses both the failure occurred before the demand and the failure occurring due to the demand itself. Then this term needs not to be mixed up with the probability of a failure due to a demand (see [3.2.13](#)).

3.1.17 steady state availability asymptotic availability $\bar{A}$ or A_{as}

(measure) limit, when it exists, of the availability $A(t)$ when t goes to infinity

Note 1 to entry: The mathematical definition is the following: $\bar{A} = \lim_{t \rightarrow \infty} A(t)$

Note 2 to entry: Mathematically speaking, the steady-state availability is a probability and can be handled as such.

Note 3 to entry: When it exists, the steady-state availability is also the average availability over the interval $[0, \infty[$:

Note 4 to entry: $\bar{A} \equiv A(\infty) \equiv \bar{A}(T \rightarrow \infty) \equiv \bar{A}(t_1, t_2 \rightarrow \infty)$

Note 5 to entry: Average and steady-state availability should not be confused. The average availability exists in any cases but components with immediately revealed and quickly repaired failures reach quickly a steady-state; periodically tested components have no steady-state.

3.1.18 failure rate $\lambda(t)$

conditional probability per unit of time that the item fails between t and $t + dt$, provided that it works over $[0, t]$

Note 1 to entry: $\lambda(t)$ is the hazard rate of a reliability function: $R(t) = \exp(-\int_0^t \lambda(\tau) d\tau)$ and $\lambda(t) = -\frac{dR(t)}{R(t)dt} = \frac{f(t)}{1 - F(t)}$

Note 2 to entry: $\lambda(t)dt$ is the probability that the item fails between t and $t + dt$, provided that it has working over $[0, t]$. Therefore the failure rate is in relationship with the first failure of the related item.

Note 3 to entry: $\lambda(t)$ is a time-dependent parameter. It is generally accepted that it evolves accordingly a “bathtub” curve: decreasing during the early failures period, becoming (almost) constant during the useful life and increasing during the wear out period.

Note 4 to entry: When the failure rate is constant the following relationship holds: $\lambda(t) \equiv \lambda = \frac{1}{MTTF}$

Note 5 to entry: An individual component with a constant failure rate remains permanently “as good as new” until it fails suddenly, completely and without warnings. This is the characteristic of the so-called catalectic failures (see definition 3.2.9).

Note 6 to entry: The failure rate is linked to the failure density by the following relation:

Note 7 to entry: $\lambda(t) = f(t)/R(t) = w(t)/[1 - F(t)]$ (see [Annex C](#) for more details).

3.1.19

average failure rate

$\bar{\lambda}(t_1, t_2)$, $\bar{\lambda}(T)$

average value of the time-dependent failure rate over a given time interval

Note 1 to entry: The mathematical definition is the following: $\bar{\lambda}(t_1, t_2) = \frac{1}{t_2 - t_1} \int_{t_1}^{t_2} \lambda(\tau) d\tau$ or $\bar{\lambda}(T) = \frac{1}{T} \int_0^T \lambda(\tau) d\tau$

Note 2 to entry: Mathematically speaking, the average failure rate is the mathematical expectation of the failure rate. It does not have the mathematical property of a failure rate as per 3.1.18 and cannot be handled as such.

3.1.20

asymptotic failure rate

λ^{as}

limit, when it exists, of the failure rate $\lambda(t)$ when t goes to infinity

Note 1 to entry: The mathematic definition is the following: $\lambda^{\text{as}} = \lim_{t \rightarrow \infty} \lambda(t)$

Note 2 to entry: Mathematically speaking, the asymptotic failure rate is a failure rate and it can be handled as such.

Note 3 to entry: When it exists, the asymptotic failure rate is also the average failure rate over the interval $[0, \infty[$:

Note 4 to entry: $\lambda^{\text{as}} \equiv \lambda(\infty) \equiv \bar{\lambda}(T \rightarrow \infty) \equiv \bar{\lambda}(t_1, t_2 \rightarrow \infty)$

Note 5 to entry: Average and asymptotic failure rate should not be confused. The average failure rate exists in any case, but:

— components with immediately revealed and quickly repaired failures reach quickly a steady-state corresponding to a constant asymptotic failure rate which is both, on the long term, an average failure rate as per [3.1.19](#) and a failure rate as per [3.1.18](#).

— periodically tested components have no steady-state and therefore the asymptotic failure rate just does not exist. The average of the failure rate can still be evaluated but this average has not the mathematical properties of a failure rate as per 3.1.18 (see also [Figure 32](#), [Figure 33](#) and [Figure 34](#)).

3.1.21

Vesely failure rate

conditional failure intensity

$\lambda_V(t)$

conditional probability per unit of time that the item fails between t and $t+dt$, provided that it was working at time 0 and at time t

Note 1 to entry: The Vesely failure rate is linked to the failure frequency by the following relation:

Note 2 to entry: $\lambda_V(t) = w(t)/A(t) = w(t)/[1 - U(t)]$ (see [Annex C](#) for more details)

Note 3 to entry: In the general case, the Vesely failure rate is not a failure rate as per 3.1.18 and cannot be used as such.

Note 4 to entry: In special cases (e.g. system with immediately revealed and quickly repaired failures), the Vesely failure rate reaches an asymptotic value which is also a good approximation of the asymptotic failure rate

$\lambda^{\text{as}} \approx \lambda_V^{\text{as}}$. In this case it can be used as a failure rate as per 3.1.18.

3.1.22

failure frequency

unconditional failure intensity

$w(t)$

conditional probability per unit of time that the item fails between t and $t+dt$, provided that it was working at time 0

Note 1 to entry: The failure frequency is linked to the Vesely failure rate by the following relation: $w(t) = \lambda_V(t)A(t) = \lambda_V(t)[1 - U(t)]$ (see [Annex C](#) for more details).

Note 2 to entry: $\lambda(t)$, $\lambda_V(t)$ and $w(t)$ should not be confused even if in particular cases they have the same numerical values.

Note 3 to entry: For highly available systems, $A(t) \approx 1$, the following relationship holds: $w(t) \approx \lambda_V(t)$

Note 4 to entry: When the Vesely failure rate reaches an asymptotic value the following relationship holds: $w^{as} \approx \lambda_V^{as} \approx \lambda^{as}$

3.1.23

average failure frequency

$\bar{w}(t_1, t_2)$, $\bar{w}(T)$, $\bar{w}$

average value of the time-dependent failure frequency over a given time interval

Note 1 to entry: The mathematical definition is the following: $\bar{w}(t_1, t_2) = \frac{1}{t_2 - t_1} \int_{t_1}^{t_2} w(\tau) d\tau$ or $\bar{w}(T) = \frac{1}{T} \int_0^T w(\tau) d\tau$

Note 2 to entry: The average failure frequency is also called "Probability of Failure per Hour" (PFH) by the standards related to functional safety of safety related/instrumented systems (e.g. IEC 61508^[2]): $PFH = \bar{w}(T)$ where T is the overall life duration of the system.

Note 3 to entry: Mathematically speaking, this is the mathematical expectation of the failure frequency. It does not have the mathematical property of a failure frequency as per 3.1.22 and cannot be handled as such.

3.1.24

PFH

DEPRECATED: probability of failure per hour

average failure frequency as [3.1.23](#) in the functional safety standard terminology (e.g. IEC 61508^[2] or IEC 61511^[3])

Note 1 to entry: The old meaning "Probability of Failure per Hour" is obsolete and replaced by "average failure frequency". Nevertheless PFH is still in use to keep the consistency with the previous versions of functional safety standards.

3.1.25

hazardous event frequency

accident frequency

$\Phi(t)$

failure frequency as [3.1.23](#) related to the hazardous event (or to the accident)

3.1.26

average hazardous event frequency

average accident frequency

$\bar{\Phi}(t_1, t_2)$, $\bar{\Phi}(T)$, $\bar{\Phi}$

average frequency as [3.1.23](#) related to of the hazardous event (or to the accident)

3.1.27

mean up time

MUT

expectation of the up time

Note 1 to entry: See [Figure 3](#) and also ISO 14224^[15] or IEC 60050-191^[14] for definitions of up time and down time.

[SOURCE: IEC 60050-191]

3.1.28

mean down time

MDT

expectation of the down time

Note 1 to entry: See [Figure 3](#) and also ISO 14224[15] or IEC 60050-191[14] for definitions of up time and down time.

[SOURCE: IEC 60050 -191]

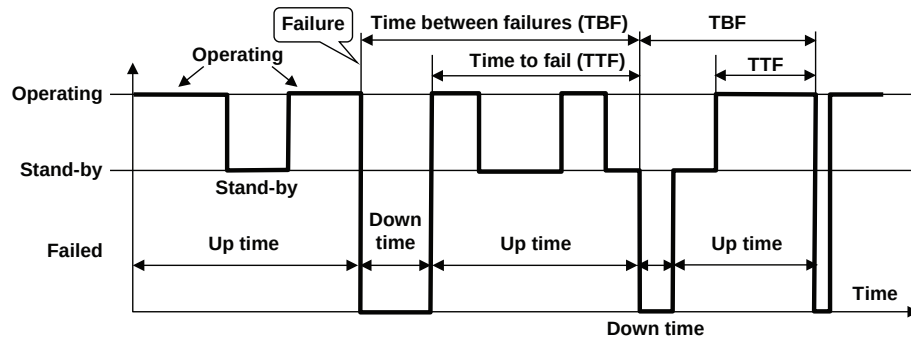


Figure 3 — Illustration of the general behaviour of an item over its lifetime

3.1.29

mean time to failure

MTTF

expected time before the item fails

Note 1 to entry: The MTTF is classically used to describe the time to failure for a non-repairable item or to the first failure for a repairable item. When the item is as good as new after a repair, it is also valid for the further failures (see [Figure 4](#)).

Note 2 to entry: In the cases illustrated by [Figure 4](#), the MTTF may be calculated by the following formulae:

$$\text{MTTF} = \int_0^{\infty} f(t) \cdot t \cdot dt = \int_0^{\infty} R(t) \cdot dt$$

Note 3 to entry: The following relationship holds when the failure rate is constant: $\text{MTTF} = 1/\lambda$.

Note 4 to entry: In the case illustrated by [Figure 3](#) where operating and stand-by failures are mixed, the formulae described in Note 2 to entry are no longer valid.

Note 5 to entry: The MTTF should not be mixed up with the design life time of the item.

Note 6 to entry: Sometimes it may be more understandable to express lifetime in probability of failure (i.e. unreliability, see [3.1.9](#)) during a certain lifespan.

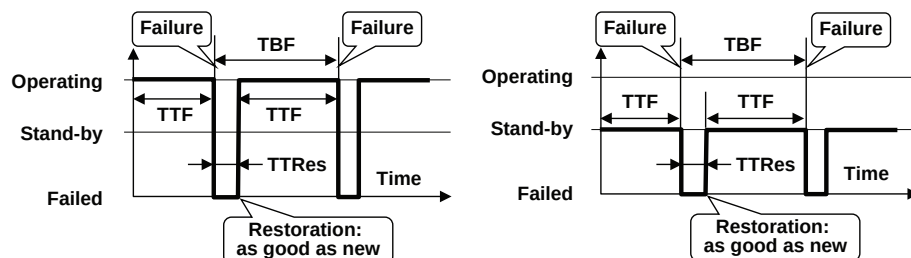


Figure 4 — Particular cases of the behaviour of an item over its lifetime

3.1.30

mean time between failures

MTBF

expected time between successive failures of a repairable item

Note 1 to entry: In the cases illustrated in [Figure 4](#), the MTBF is linked with MTTF and MTTRes by the following relationship: $MTBF = MTTF + MTTRes$. More generally It is also linked to the MUT and MDT by $MTBF = MUT + MDT$.

Note 2 to entry: The acronym MTBF is sometimes defined as the mean *operating* time between failures (e.g. in IEC 191[14]). This is not at all the same and, in this case, the formula described in Note 1 to entry is no longer valid. This is very confusing, therefore the traditional definition of the MTBF is retained in this Technical Report.

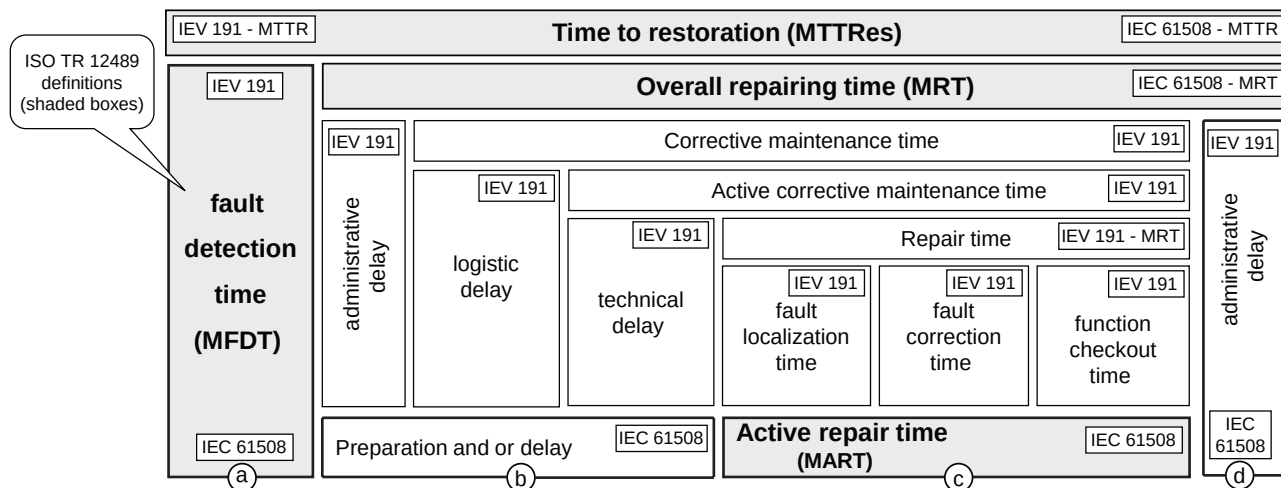


Figure 5 — Repair time taxonomies as per IEC 191[14], IEC 61508[2] and ISO/TR 12489

3.1.31

mean time to repair

MTTR

expected time to achieve the repair of a failed item

Note 1 to entry: This term MTTR is used in ISO 14224[15] and ISO 20815[16] where the fault detection time is not really considered: ISO 14224[15] deals with detected faults (in fact, the actual time spent to detect the fault is never known and cannot be collected); ISO 20815[16] deals mainly with immediately revealed failure where the time spent to detect the faults is close to 0 (i.e. negligible). As the fault detection time is very important for the purpose of this Technical Report there is a need to clearly distinguish between the two following times (cf. [Figure 5](#)):

- 1) the time elapsing from the actual occurrence of the failure of an item to its detection (cf. [3.1.35](#), MFDT);
- 2) the time elapsing from the detection of the failure of an item to the restoration of its function (cf. [3.1.33](#), MRT).

Note 2 to entry: The acronym MTTR is defined as the mean time to restore in the IEC 60500-191[14] or in the IEC 61508[2]. This is not the same as in ISO 14224[15] or ISO 20815[16]. Therefore, in order to avoid any mixed-up, the acronym MTTRes is used in this Technical Report instead of MTTR (cf. [3.1.32](#)).

3.1.32

mean time to restoration

MTTRes

expected time to achieve the following actions: (see [Figure 5](#), [Figure 6](#) and [Figure 7](#)):

- the time to detect the failure a; and,
- the time spent before starting the repair b; and,
- the effective time to repair c; and,
- the time before the component is made available to be put back into operation d

Note 1 to entry: Figure 5 illustrates how the times a, b, c and d defined in the IEC 61508[2] standard are linked to the delays defined in the IEC 60050-191[14] standard. Time b starts at the end of a; time c starts at the end of b; time d starts at the end of c.

Note 2 to entry: Figure 5, Figure 6 and Figure 7 can be used to understand the differences between the definitions of MTTRes, MRT and MART used in this Technical Report.

Note 3 to entry: The MTTRes is linked to the MRT and the MFDT by the following formula: $MTTRes = MFDT + MRT$.

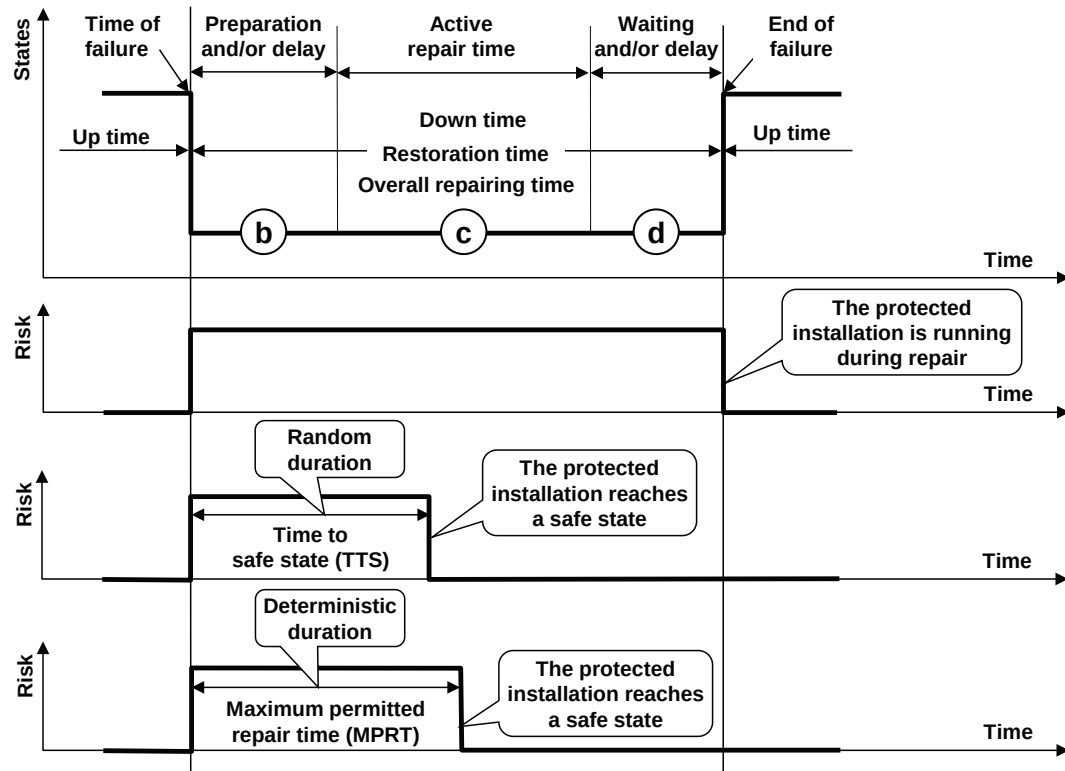


Figure 6 — Illustration of the restoration time and of the risk exposure for immediately revealed failures of an item

3.1.33

MRT

mean overall repairing time

expected time to achieve the following actions:

- the time spent before starting the repair b; and,
- the effective time to repair c; and,
- the time before the component is made available to be is put back into operation d

Note 1 to entry: See Figure 5, Figure 6 and Figure 7.

Note 2 to entry: The terms “repair”, “repairable”, “repaired” used in this Technical Report, unless otherwise specified, are related to the overall repairing time (see Figure 5).

Note 3 to entry: When a safety system operating in demand mode is faulty, the risk disappears as soon as the protected installation is placed in a safe state (e.g. stopped). In this case (see Figure 6 and Figure 7) the MTTS replaces the MRT (see 3.1.36) with regard to the probabilistic calculations.

Note 4 to entry: This definition is in line with IEC 61508[2] but not with IEC 60050-191.[14]

3.1.34

mean active repair time

MART

expected active repair time

Note 1 to entry: The MART is the expected effective time to repair c, (see [Figure 5](#), [Figure 6](#) and [Figure 7](#)).

3.1.35

mean fault detection time

MFDT

expected time needed to detect a fault

Note 1 to entry: The MFDT is the time a) in [Figure 5](#), [Figure 6](#) and [Figure 7](#).

Note 2 to entry: The MFDT is equal to zero for immediately revealed failure; (see [Figure 6](#)) generally negligible for quickly detected failures; (see [Figure 6](#)) depending of the test policy for the hidden failures. In this case it may be the main part of the item down time (see [Figure 7](#)).

Note 3 to entry: The MFDT used in this Technical Report should not be mixed-up with the mean fractional dead time which has the same acronym.

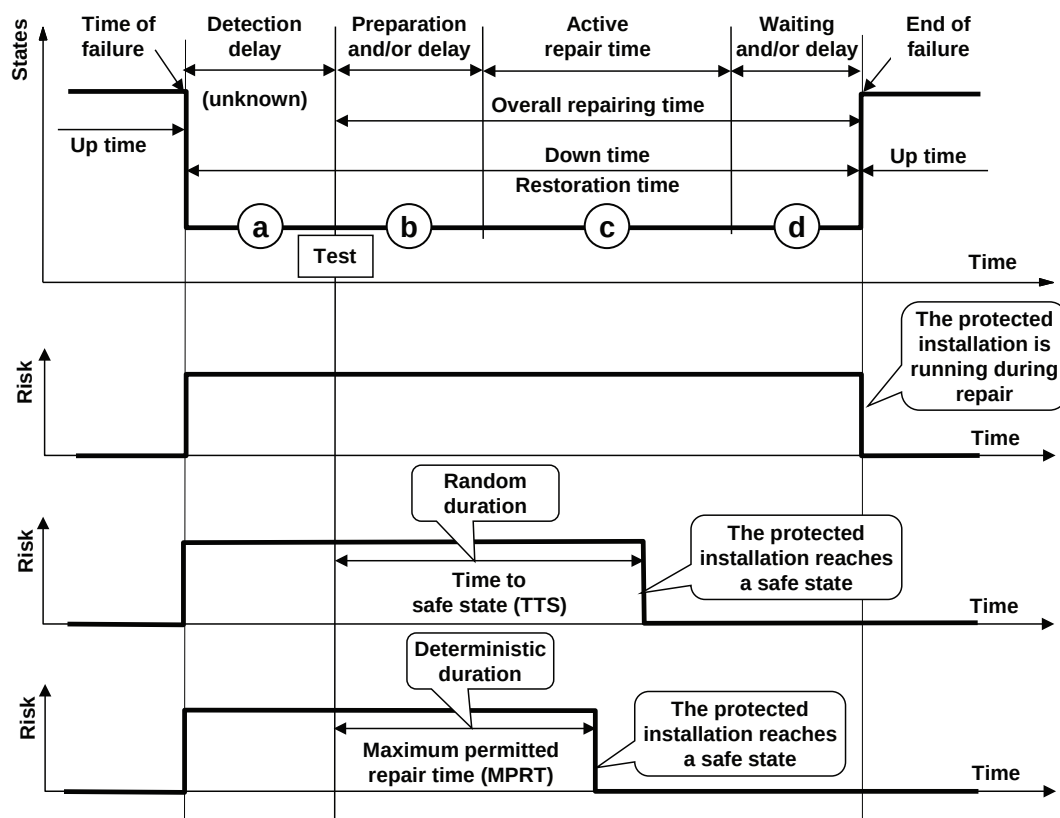


Figure 7 — Illustration of the restoration time and of the risk exposure for the hidden failures of an item

3.1.36

MTTS

mean time to safe state

expected time needed for the protected installation to reach a safe state after a dangerous failure of a safety system has been detected

EXAMPLE When a dangerous fault is revealed for a safety system operating in demand mode, it may be decided to reach a safe state rather than undertake the repair of the fault and this may take some time: a MTTS of 8 h means, for example, that, on average, 8 h are needed to shut down the process. After the shut down, a safe state is reached, the fault is no longer dangerous and it is not necessary to take into account the remaining time spent to complete the repair. This is illustrated in [Figure 6](#), [Figure 7](#) and [Figure B.1](#).

Note 1 to entry: When the MTTS is defined as a maintenance procedure it is necessary to take it into consideration for the probabilistic calculations of hazardous events. In this case the MTTS replaces the MRT (see [3.1.33](#)) with regard to the probabilistic calculations. Reciprocally it is necessary to verify that this MTTS is respected during the actual repair actions in order to keep the probabilistic calculations valid.

Note 2 to entry: The role of the MTTS is close to the role of the MPRT. The difference is that the MPRT is a maximum duration allowed to reach a safe state when the MTTS is the average of the random duration of the TTS needed to reach the safe state when a dangerous fault is revealed (see [Figure 6](#) and [Figure 7](#)). The methods developed in this Technical Report have been focused on average random values (MTTRes, MRT, MTTS) rather than on deterministic values (MPRT), but the MPRT can be easily handled by using Petri nets and Monte Carlo simulations.

3.1.37

maximum permitted repair time

MPRT

maximum time allowed to repair a fault before undertaking an action to make the risk disappearing

EXAMPLE When a dangerous fault is revealed for a safety system operating in demand mode, it may be decided to reach a safe state when a maximum duration has elapsed: a MPRT of 8 h means, for example, that if the repair is not completed after 8 h, the process is shut down. Then a safe state is reached, the fault is no longer dangerous, and it is not necessary to take into account the remaining time spent to complete the repair. This is illustrated in [Figure 6](#), [Figure 7](#) and [Figure B.1](#). When the fault may result of several failure modes, the MPRT allows to repair those within short MRT without shutdown of the process.

Note 1 to entry: When a MPRT is defined as a maintenance procedure it is necessary to take it into consideration for the probabilistic calculations of hazardous events. Reciprocally it is necessary that this MPRT be respected during the actual repair actions in order to keep the probabilistic calculations valid.

Note 2 to entry: The role of the MPRT is close to the role of the MTTS (see 0). The difference is that the MPRT is a maximum duration allowed to reach a safe state and the MTTS is the average duration needed to reach the safe state when a dangerous fault is revealed (see [Figure 6](#) and [Figure 7](#)). The methods developed in this Technical Report have been focused on random repair values (MTTRes, MRT, MTTS) rather than on deterministic values (MPRT), but the MPRT can be easily handled by using Petri nets and Monte Carlo simulations.

3.1.38

mean time to demand

MTTD

expected time before the demand on the safety system occurs

3.1.39

restoration rate

μ

conditional probability per unit of time that the restoration of a failed item ends between t and $t+dt$, provided that it was not finished over $[0, t]$

Note 1 to entry: The following relationship holds when the restoration rate is constant: $MTTRes = 1/\mu$.

Note 2 to entry: The “restoration” rate is in relationship with the restoration time. Similarly the “repairing” rate can be defined in relationship with the “overall repairing” time and the “active repair” rate in relationship with the “active repair” time.

Note 3 to entry: The restoration rate has the same mathematical properties for the restoration as the failure rate for the failures.

3.2 Failure classification

Explanations about the various states and the various failures of a safety system are developed in [Annex B](#).

3.2.1

failure

loss of ability to perform as required

Note 1 to entry: A failure of an item is an event, as distinct from a fault of an item, which is a state (see [Figure 8](#)).

[SOURCE: IEC 60050 –191]

3.2.2

fault

inability to perform as required

Note 1 to entry: A fault of an item is a state, as distinct from a failure of an item which is an event (see [Figure 8](#)).

Note 2 to entry: A fault of an item may result from a failure of the item or from a deficiency in an earlier stage of the life cycle, such as specification, design, manufacture or maintenance.

Note 3 to entry: Qualifying terms may be used to indicate the cause of a fault, such as specification, design, manufacture, maintenance or misuse.

Note 4 to entry: Inability to perform due to preventive maintenance, other planned actions, or lack of external resources does not constitute a fault.

Note 5 to entry: [Figure 8](#) illustrate the relationship between the concepts of failure and fault:

- The *Failure x* occurs at stage 1 and leads to the state *Fault x* which is not detected.
- from stage 2 point of view *Fault x* is a pre-existing fault.
- The *Failure y* occurs at stage 2 and lead to the state *Faults x,y* which is not detected.
- From stage 3 point of view *Fault x,y* is a pre-existing fault.
- and so on.

[SOURCE: IEC 60050 –191]

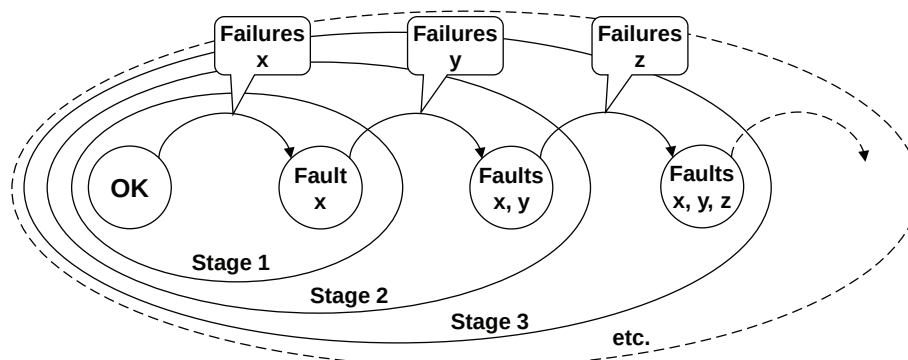


Figure 8 — Relationship between failure and fault concepts

3.2.3

dangerous failure

unsafe failure

failure of a safety system which tends to impede a given safety action

Note 1 to entry: This is a systemic failure in relationship with a given safety action performed by the safety system. Therefore this concept is irrelevant for an individual item on the shelves.

Note 2 to entry: See [Figure B.1](#).

3.2.4

critical dangerous failure

dangerous failure leading to the complete inhibition of the safety action (i.e. leading to a dangerous situation for the protected system)

Note 1 to entry: This is a systemic failure in relationship with a given safety action performed by the safety system. Therefore this concept is irrelevant for an individual item on the shelves.

Note 2 to entry: The same failure of a component belonging to a safety system with internal redundancy may be dangerous or critical dangerous depending on the system state from which it occurs.

Note 3 to entry: The critical dangerous failures that are undetected (e.g. those revealed by periodic tests) are sometimes called *safety critical failures* (cf. ISO 14224[15]). The equipment subject to such potential failures can be identified within a plant and monitored, and the ratio between the number of safety critical failures detected by periodic tests and the corresponding number of tests performed (commonly called "*failure fraction*") is being used for that purpose. This indicator of the average unavailability (PFDavg) due to dangerous undetected failures is established by using test reports. It is important not to mix such failure fraction with other reliability terms.

3.2.5

safe failure

failure of a safety system which tends to favour a given safety action

Note 1 to entry: The concept of safe failure is illustrated in [Figure B.1](#).

Note 2 to entry: A failure is safe only with regard to a given safety function. This is a systemic failure in relationship with a given safety action performed by the safety system. This concept is irrelevant for an individual item on the shelves.

Note 3 to entry: The non-critical safe failures basically increase the probability of success of the safety function. The critical safe failures initiate the related safety actions when this is not needed (see spurious failures).

3.2.6

spurious failure

failure triggering an action in an untimely manner

Note 1 to entry: Critical safe failures (see [Figure B.1](#)) are the typical spurious failures related to safety systems.

Note 2 to entry: A spurious failure does not necessarily imply a spurious trip ([3.4.14](#)) but a spurious trip is always the result of a spurious failure.

3.2.7

critical safe failure

spurious failure of a safety system, due to safe failure(s) of its component(s), triggering the safety action and leading to a spurious safety action

Note 1 to entry: The concept of critical safe failure is illustrated in [Figure B.1](#).

Note 2 to entry: This is a systemic failure in relationship with a given safety action performed by the safety system. This concept is irrelevant for an individual item on the shelves.

Note 3 to entry: The same failure of a component belonging to a safety system may be safe or spurious (critical safe) depending of the system state from which it occurs (e.g. the safe failure of a sensor belonging to 2oo3 is only safe when it occurs in 1st position. It is critical when it occurs in 2nd position).

3.2.8 systemic failure

holistic failure

failure at system level which cannot be simply described from the individual component failures of the system

Note 1 to entry: Systemic/holistic principles have been concisely summarized by Aristotle by "*The whole is more than the sum of its parts*".

Note 2 to entry: Components have only failure modes. Those failure modes become dangerous, safe or spurious only when the components are implemented into a safety "system". This is why dangerous, safe or spurious failures are typical systemic failures. For example the failure "fail to close" of a valve is dangerous only if it belongs to a safety system closing this valve on demand. Otherwise this failure mode does not matter.

Note 3 to entry: "Systematic" failures (i.e. occurring in a deterministic way when given conditions are encountered, see [3.2.17](#)) and "systemic" failures should not be confused.

3.2.9 catalectic failure

sudden and complete failure

Note 1 to entry: This term has been introduced by analogy with the catalectic verses (i.e. a verse with seven feet instead of eight) which stop abruptly. Then, a catalectic failure occurs without warning and is more or less impossible to forecast by examining the item. It is the contrary of failures occurring progressively and incompletely.

Note 2 to entry: Catalectic failures characterize simple components with constant failure rates (exponential law): they remain permanently "as good as new" until they fail suddenly, completely and without warning. Most of the probabilistic models used in reliability engineering are based on catalectic failures of the individual component of the system under study (e.g. Markovian approach).

3.2.10 immediately revealed failure

overt failure

detected failure

evident failure

failure which is immediately evident to operations and maintenance personnel as soon as it occurs

Note 1 to entry: The immediately revealed failures show themselves immediately, but the hidden failures which are quickly detected by specific diagnostic tests are generally considered as immediately revealed failures.

Note 2 to entry: The repair of immediately revealed failures may begin immediately after they have occurred.

Note 3 to entry: The failures which are detected by periodic tests are not considered as immediately revealed failures.

3.2.11 hidden failure

covert failure

dormant failure

unrevealed failure

undetected failure

failure which is not immediately evident to operations and maintenance personnel

Note 1 to entry: Hidden failures do not show themselves when they occur. The occurrence of a hidden failure gives a latent fault which may be revealed by specific tests (e.g. periodic tests) or by the failure of the item to perform its function when required.

Note 2 to entry: The repair of hidden failures cannot begin as long as they have not been detected. The unknown times spent between the failures and their detections belong to the MTTRes.

3.2.12

time-dependent failure

failure occurring with a probability depending of the time

Note 1 to entry: The unreliability $F(t)$ is a typical probability function describing time-dependent failures

3.2.13

failure due to demand

failure occurring on demand

γ, ψ

failure of one item due to a change of its state triggered by an external event (the so-called “demand”)

EXAMPLE 1 Obtaining 2 when launching a dice is an event occurring on demand. The probability of this event is 1/6. It does not depend on the elapsing time but only of the demand itself (i.e. the fact that the dice is launched).

EXAMPLE 2 The failure of an electromechanical relay (e.g. rupture of the spring) when it changes state depends on the number of operations (cycles) rather on the operating time (see IEC 61810-2[49]) and this is the same for the failure of an electronic device due to over voltage when it is switched or the blocking of a diesel engine when it is started, etc.: these are typical examples of failures due to demands (or cycles).

Note 1 to entry: In this Technical Report two kinds of demand are considered: the periodic tests and the demand for an actual safety action. The probability of a failure due to periodic test is a constant number noted γ and the probability of a failure due to one actual demand of the safety action is a constant number noted ψ . Over a given time interval, those probabilities of failure do not depend on the duration but on the number of demands or tests occurring within this interval. The use of γ and ψ is explained in 7.3.

Note 2 to entry: This should not be confused with the “failure on demand” appearing in the term “probability of failure on demand” (see 3.1.14, Note 2 to entry) used in functional safety standards[2] for low demand mode safety systems. In those standards this means “failures *likely to be observed when* a demand occurs”.

3.2.14

common cause failures

CCF

failures of different items, resulting from a single event, where these failures are not consequences of each other

Note 1 to entry: It is generally accepted that the failures occur simultaneously or within a short time of each other.

Note 2 to entry: Common cause failures can lead to common mode failures.

Note 3 to entry: Common cause failures reduce the effect of system redundancy.

Note 4 to entry: Explicit and implicit CCF are defined in 5.4.2.

3.2.15

common mode failures

CMF

failures of different items, occurring in the same way

Note 1 to entry: Common mode failures may have different causes.

Note 2 to entry: Common mode failures can be due to common cause failures.

3.2.16

random failure

failure, occurring in a random way

Note 1 to entry: A random failure may be time or demand dependent. Whether it occurs or not is not predictable with certainty but the corresponding failure rate (see 3.1.18) or probability of a failure due to demand (see 3.2.13) may be predictable and this allows probabilistic calculations.

EXAMPLE Examples of failure mechanisms leading to unpredictable failure occurrence are: hardware random failures resulting from degradation mechanisms; human random failure resulting from error in routine operation, lack of attention, stress, tiredness, etc.

Note 2 to entry: From the probabilistic point of view, the random failures are the contrary of systematic failures (see [3.2.17](#)) which occur in a deterministic way (i.e. with a probability equal to 1) when some conditions are met.

3.2.17

systematic failure

failure that consistently occurs under particular conditions of handling, storage or use

Note 1 to entry: The cause of a systematic failure originates in the specification, design, manufacture, installation, operation or maintenance. Its occurrence is precipitated by particular conditions of handling, storage, use or maintenance (see [Figure G.3](#))

Note 2 to entry: Corrective maintenance without modification will usually not eliminate the failure cause.

Note 3 to entry: A systematic failure can be reproduced by deliberately applying the same conditions, e.g. in verifying the failure cause (from IEC 60050-191 ed3[14]). Systematic failures are non-random failures (see [3.2.16](#)).

Note 4 to entry: In operation, a systematic failure is a manifestation of a systematic fault (i.e. a pre-existing state of the system).

Note 5 to entry: The software systematic failures, called “bugs”, are example of systematic failures: they are due to pre-existing bugs (i.e. faults) and they occur when the input data activate them.

Note 6 to entry: Systematic and systemic (which means “at system level”) failures (see [3.2.8](#)) should not be confused.

[SOURCE: IEC 60050-191]

3.3 Safety systems typology

3.3.1

demand mode of operation safety systems

safety system designed to achieve its safety action only when receiving a specific request from its surrounding environment

Note 1 to entry: Such systems spend most of their time in stand-by position but need nevertheless to be ready to work as soon as a demand occurs.

Note 2 to entry: Such systems are subject to hidden failures. Diagnostic and periodic tests are generally implemented in order to reveal the corresponding latent faults.

Note 3 to entry: When the demand frequency increases, an on-demand mode safety system may be assimilated to a continuous mode of operation systems.

3.3.2

continuous mode of operation safety system

safety system designed to achieve its safety action permanently

Note 1 to entry: With a continuous mode safety system the hazardous event occurs as soon as the safety system fails. This is illustrated in [Figure B.1](#) where the systems states “KO” and “hazardous event” are gathered into a single state.

3.3.3

multiple safety systems

safety system comprising several sub safety systems operating one after the other when the prior ones have failed

Note 1 to entry: Industrial processes often implement multiple safety systems (safety layers). In this case the failure of an intermediate safety layer provokes a demand on the proximate succeeding safety layer and so on. The accident occurs only if the demand is transmitted until the ultimate safety layer and it fails to operate.

3.4 Maintenance issues

3.4.1

maintenance

combination of all technical and administrative actions, including supervisory actions, intended to retain an item in, or restore it to, a state in which it can perform a required function

Note 1 to entry: There are two basic categories of maintenance: corrective maintenance done after a failure has occurred and preventive maintenance (testing, inspection, condition monitoring, periodic) done before a failure has occurred. See also ISO 14224[15], 9.6.

Note 2 to entry: Maintenance activities of either preventive or corrective maintenance category type, is shown in ISO 14224:2006[15], Table B.5.

[SOURCE: ISO 14224]

3.4.2

maintenance concept

definition of the maintenance echelons, indenture levels, maintenance levels, maintenance support, and their interrelationships

Note 1 to entry: The maintenance concept provides the basis for maintenance planning, determining supportability requirements and developing logistic support.

Note 2 to entry: A maintenance echelon is a position in an organization where specified levels of maintenance are to be carried out (e.g. field, repair shop, manufacturer facility).

[SOURCE: IEC 60050-191]

3.4.3

preventive maintenance

PM

maintenance carried out at predetermined intervals or according to prescribed criteria and intended to reduce the probability of failure or the degradation of the functioning of an item

[SOURCE: ISO 14224]

3.4.4

corrective maintenance

maintenance carried out after fault detection to effect restoration

Note 1 to entry: Corrective maintenance of software invariably involves some modification.

Note 2 to entry: Sometimes the *corrective maintenance* is also called *curative maintenance*.

[SOURCE: IEC 60050-191-46-06]

3.4.5

detection method

method or activity by which a failure is discovered

Note 1 to entry: A categorization of detection methods (e.g. periodic testing or continuous condition monitoring) is shown in ISO 14224:2006[15], Table B.4.

3.4.6

maintenance plan

structured and documented set of tasks that include the activities, procedures, resources and the time scale required to carry out maintenance

Note 1 to entry: The maintenance plan should be thoroughly analysed and modelled to produce relevant probabilistic results.

Note 2 to entry: The forecasted probabilistic results established at the design stage are no longer valid if the maintenance plan which has been considered is not thoroughly applied in operation.

Note 3 to entry: The maintenance plan should cover policies for both preventive maintenance (e.g. testing) and corrective maintenance (e.g. minimize downtime, restore lost redundancy).

Note 4 to entry: The maintenance plan is part of an overall Operations and Maintenance plan. It is sometimes called "*maintenance policy*".

[SOURCE: EN 13306]

3.4.7 test policy

set of procedures describing the various test operations (frequencies and procedures) scheduled to reach the safety requirements of a given safety system

Note 1 to entry: The test policy should be thoroughly analysed and modelled to produce relevant probabilistic results.

Note 2 to entry: The forecasted probabilistic results established at the design stage are no longer valid if the test policy which has been considered is not thoroughly applied in operation.

3.4.8 periodic tests

proof tests

planned operation performed at constant time interval in order to detect the potential hidden failures which may have occurred in the meantime

Note 1 to entry: The unsafe hidden failures of a safety system which are not detected by the diagnostic tests may be detected by periodic tests. Such tests are named "proof tests" in the standards dealing with functional safety (e.g. IEC 61508[2]).

[SOURCE: IEC 61508]

3.4.9 diagnostic tests

automatic operations performed at high frequency in order to detect the potential hidden failures as soon as possible when they occur

Note 1 to entry: The unsafe failures of safety system are generally hidden and diagnostic tests may be implemented to detect the larger part of them. As the diagnostic cycle is normally short, the hidden failures detected by diagnostic tests are assimilated to immediately revealed failures.

[SOURCE: IEC 61508]

3.4.10 staggered testing (of redundant items)

test of several items with the same test interval but not at the same time

EXAMPLE [Figure 9](#) shows staggered tests for two item A and B.

Note 1 to entry: When the redundant components of a system are tested at the same time (i.e. when the tests are synchronous) their availabilities are good (just after a test) and bad (just before a test) at the same time. This correlation means that the unavailabilities of the components peak simultaneously. This has a detrimental effect on the system availability which can be cancelled by de-synchronizing the tests. A practical way to do that is staggering the tests (e.g. testing one component in the middle of the test interval of the other); the unavailability peaks are also staggered and this improves the average availability of the system.

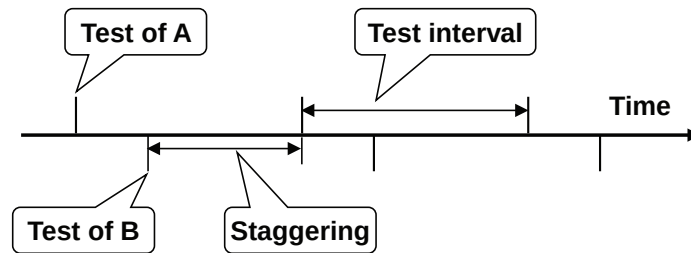


Figure 9 — Example of test staggering for a system made of two redundant items A and B

3.4.11

redundancy

existence of more than one means for performing a required function

Note 1 to entry: The aim of redundancy is to provide backup in case of one or several failures of the means performing a required function.

Note 2 to entry: Redundancy definitions for passive (cold) standby, active (hot) standby and mixed are given in ISO 14224[15], C.1.2.

Note 3 to entry: Redundancy is sometimes (in IEC 61508[2] and IEC 61511[3]) called “fault tolerance”.

[SOURCE: ISO 14224]

3.4.12

spurious activation (of a safety function)

untimely demand of a safety function when this is not needed

Note 1 to entry: The spurious activation of a safety function may be due to the occurrence of one or several safe failures.

3.4.13

spurious action (of a safety system)

result of a spurious activation of a safety function

Note 1 to entry: A spurious safety action is not necessary safe. An example of spurious action is a spurious trip.

3.4.14

trip

shutdown of machinery from normal operating condition to full stop

Note 1 to entry: Two types of shutdown can be identified:

- a) Trip: the shutdown is activated automatically by the control/monitoring system.
 - Real trip: the shutdown is activated as a result of a monitored (or calculated) value in the control system exceeding a pre-set limit.
 - Spurious trip: unexpected shutdown resulting from failure(s) in the control/monitoring system or error(s) imposed by on control/monitoring system originating from the environment or people.
- b) Manual shutdown: the machinery is stopped by an intended action of the operator (locally or from the control room).

Note 2 to entry: Sometimes statements like “equipment trip” or “spurious trip” can be misleading terminology used for failures causing (rotating) equipment shutdown, especially when it is treated as failure mode in reliability data or modelling. A failure mechanism (see Table B.2 of ISO 14224[15]) can be of various types (e.g. mechanical, instrument) and should not be mixed with the term failure modes (of which one is spurious trip). Failure modes are not necessarily instrument-related failures, but could be mechanical failures. See for example failure modes in Table B.6 of ISO 14224[15] for rotating equipment.

[SOURCE: ISO 14224]

3.4.15

reliability data

data for reliability, maintainability and maintenance support performance

Note 1 to entry: Reliability and maintainability (RM) data is the term applied by ISO 14224[15].

[SOURCE: ISO 20815]

3.4.16

generic reliability data

reliability data covering families of similar equipment

Note 1 to entry: See ISO 14224[15] for further details on equipment boundaries and equipment taxonomies that define these families of equipment within the petroleum, petrochemical and natural gas industries.

Note 2 to entry: Plant-specific data on specific equipment could be part of generic reliability databases, but could differ a lot from generic data and should not be mixed with those.

[SOURCE: ISO 14224]

3.5 Other terms

3.5.1

critical state

in a states-transitions model, state belonging to a given class of states and which is distant from the failure class of states by only one transition

Note 1 to entry: This is a mathematical concept in relationship with e.g. Markovian process or Petri nets models.

EXAMPLE The states of a safety system can be sorted out into two classes: class OK when the safety action is available and KO when the safety action is inhibited. In this case a critical state with regards to the safety system failure belongs to the class *OK* and only one failure (i.e. on event) is needed to have a transition to the class *KO*.

3.5.2

dysfunction

impaired or abnormal functioning of an item

Note 1 to entry: This term is built from the Greek prefix “dys” (i.e. “with difficulty”) and the Latin term “functio” (i.e. an activity with a given aim). Primarily used in the medical field, this term is now often used within the technological field as a more generic term than “failure” or “fault”.

3.5.3

dysfunctioning analysis

set of activities aiming to analyse the dysfunctions of an item

Note 1 to entry: This is the counterpart of the functioning analysis which aims to analyse how an item works when the dysfunctioning analysis aims to analyse how an item fails by e.g. identifying, sorting out, characterizing and/or evaluating the probability of occurrence of the dysfunctions.

Note 2 to entry: The term “dysfunctional” analysis is often used as a synonym of “dysfunctioning” analysis.

3.5.4

formal language

set of words, semantics and logical rules with sound mathematical properties

Note 1 to entry: Programming languages are an example of formal language with mathematical properties allowing them to be compiled into computer executable code.

Note 2 to entry: Every reliability model has an underlying formal language behind the graphical elements (e.g. the Binary logic for Boolean models).

Note 3 to entry: Specific formal languages have been developed to model the functioning and the dysfunctioning of industrial systems (e.g. AltaRica^{[11][12]}). According to their powerfulness of modelling and their mathematical properties they can be compiled toward event trees, fault trees, Markov graphs, Petri nets, accident sequences, etc. Some of them can also be directly used for Monte Carlo simulation.

3.5.5

functioning analysis

set of activities aiming to analyse how an item performs as required

Note 1 to entry: This is the counterpart of the dysfunctioning analysis which aims to analyse how an item fails when the functioning analysis aims to analyse how an item works by, e.g. identifying, sorting out and characterizing the various functions related to the item.

Note 2 to entry: The term “functional” analysis is often used as a synonym of “functioning” analysis. However, the term “functional analysis” which has several meanings is not used in this Technical Report to avoid confusion.

3.6 Equipment-related terms

3.6.1

high integrity protection system

HIPS

non-conventional autonomous safety instrumented system with sufficiently high safety integrity (see [3.2.1](#)) to protect equipment against exceeding the design parameters

Note 1 to entry: Deviations from industry standards describing mechanical protection systems (e.g. ISO 23251^[31] = API STD 521^[32], ISO 10418^[33], API RP 14C^[58]) are treated as HIPS. An ultimate protection relying solely on Safety Instrumented Systems (SIS) is qualified as HIPS, irrespective of its required Safety Integrity Level (SIL).

3.6.2

high integrity pressure protection system

HIPPS

HIPS exclusively devoted to protection against overpressure

Note 1 to entry: Alternative terminology: over pressure protection system (OPPS).

Note 2 to entry: A HIPPS can be used as an alternative to, e.g.:

- full pressure rating of downstream equipment, or
- adequately sized mechanical pressure relief devices, or
- design the disposal system for simultaneous reliefs.

3.6.3

HIPPS valve

valve used as a final element in a HIPPS system

3.6.4

subsea isolation valve

SSIV

SIV

valve which closes within a defined time limit derived from the risk assessment in order to reduce consequences of pipeline/riser leak or rupture

Note 1 to entry: The SSIV can be an actuated valve (e.g. remotely controlled subsea valve) or a non-activated valve (subsea check valve). An activated valve is normally designed as fail safe (i.e. closes and remains closed on all failures external to the valve and actuator themselves).

Note 2 to entry: Where the flexible risers are connected directly to the subsea wellhead, the master and wing valve may be considered to represent the SSIV function.

[SOURCE: ISO 14723]

3.6.5

subsurface safety valve

SSSV

device installed in a well below the wellhead with the design function to prevent uncontrolled well flow when actuated

Note 1 to entry: These devices can be installed and retrieved by wireline (Wireline retrievable) and/or pump down methods (TFL-Thru Flow Line) or be integral part of the tubing string (Tubing retrievable).

3.6.6

surface-controlled subsurface safety valve

SCSSV

SSSV controlled from the surface by hydraulic, electrical, mechanical or other means

Note 1 to entry: The SCSSV is sometimes called DHSV (downhole safety valve).

[SOURCE: ISO 14723]

3.6.7

subsurface-controlled subsurface safety valve

SSCSV

SSSV actuated by the characteristics of the well itself

Note 1 to entry: Note 1 to entry: These devices are usually actuated by the differential pressure through the SSCSV (velocity type) or by tubing pressure at the SSCSV (high or low pressure type).

[SOURCE: ISO 14723]

3.6.8

SSSV system equipment

components which include the surface-control system, control line, SSSV, safety valve lock, safety valve landing nipple, flow couplings and other downhole control components

[SOURCE: ISO 10417]

4 Symbols and abbreviated terms

Table 2 — Acronyms used in ISO/TR 12489

Acronyms	Meaning
AC, DC, UPS	Alternative or Direct Current, Uninterruptible Power Supply
APJ	Absolute Probability Judgement
ATHEANA	A Technique for Human Error Analysis
CCF	Common Cause Failure
CMF	Common Mode Failure
CREAM	Cognitive Reliability and Error Analysis Method
FMECA	Failure Modes, Effects and Criticality Analysis
FMEA	Failure Modes and Effects Analysis
FT, FTA	Fault Tree, Fault Tree Analysis
HAZID, HAZOP	HAZard IDentification, HAZard and OPerability study
HIPS	High Integrity Protection System
HIPPS	High Integrity Pressure Protection System
HEART	Human Error Assessment and Reduction Technique
HEP	Human Error Probability

Table 2 (continued)

Acronyms	Meaning
HRA	Human Reliability Analysis
IEC	International Electrotechnical Commission
IEV	International Electrotechnical Vocabulary
ISO	International Organization for Standardization
LOPA	Layer Of Protection Analysis
MDT	Mean Down Time
MART, MRT	Mean Active Repair Time, Mean Repairing Time
MFDT	Mean Fault Detection Time
MPRT	Maximum Permitted Repair Time
MTBF	Mean Time Between Failures
MTTD	Mean Time to Demand
MTTE, MTTR, MTTRes	Mean Time to Fail, Mean Time to Repair, Mean Time To Restore
MTTS	Mean Time To Safe state
MUT	Mean Up Time
OPPS	Over Pressure Protection System
PC	Paired Comparisons
PDS	Reliability of computer-based safety systems (Norwegian)
PFD, PFD _{avg}	Probability of Failure on Demand, Average PFD
PFH	Probability of Failure per Hour (average dangerous failure frequency)
PHA	Preliminary Hazard Analysis
PN	Petri net
RBD	Reliability Block Diagram
SCSSV	Surface-Controlled Subsurface Safety Valve
SDT	Signal Detection Theory
SDV	Shut Down Valve
SIL, SIS	Safety Integrity Level, Safety Instrumented Systems
SLIM-MAUD	Success Likelihood Index Method using Multi-Attribute Utility Decomposition
SSCSV	Subsurface-controlled Subsurface Safety Valve
SSSV	Subsurface Safety Valve
SSIV, SIV	SubSea Isolation Valve
THERP	Technique for Human Error Rate Prediction
TR	Technical Report
TBF, TTF, TTRes	Time Between failures, Time To Failure, Time To Restore
TTS	Time to Safe State

Table 3 — Symbols used in ISO/TR 12489

Symbol	Definition
Time-related symbols	
T	Duration
$t, t_i, [t_1, t_2]$	Instants and interval of time
Event-related symbols	

Table 3 (continued)

Symbol	Definition
C (any letter)	Components C
E (any letter)	Event
$\bar{E}$ or $\neg E$	Complementary event
$e = \text{"true"} \text{ or } \text{"false"}$	Boolean variable indicating the state of an event.
$\{A\}, \{A, B\}, \{A, B, C\}, \dots$ $A \cap B, A \cap B \cap C, \dots$ $A \bullet B, A \bullet B \bullet C$	Combination of events
$??x == Fct(y, z, \dots)$ $??x > Fct(y, z, \dots)$	"Predicates" used in Petri net (see M.1): general form (x: any kind of variable, $Fct(.)$: any kind of function of other variable)
$??E == \text{"true"} \text{ or } ?E,$ $??E == \text{"false"} \text{ or } ?\neg E,$	"Predicates" used in Petri nets for Boolean variable (see M.1)
$!!x = Fct(y, z, \dots)$	"Assertions" used in Petri net (see M.1): general form (x: any kind of variable, $Fct(.)$: any kind of function of other variable)
$!!E = \text{"true"} \text{ or } !E,$ $!!E = \text{"false"} \text{ or } !\neg E,$	"Assertions" used in Petri nets for Boolean variable (see M.1)
Probabilistic-related symbols	
$P(.), P(A), P(t)$	Probability function, Probability of the event A, Probability at time t
$A(t), U(t), R(t), F(t)$	Availability, unavailability, reliability and unreliability at time t
$f(t)$	Failure probability density
$w, w(t), \bar{w}(T)$	Unconditional failure intensity (failure frequency)
$\Phi, \Phi_a, \Phi(t), \Phi_a(t), \Phi_{st}(t)$	Event frequency (accident, spurious safety actions)
$\bar{\Phi}, \bar{\Phi}_a, \bar{\Phi}(T), \bar{\Phi}_a(T), \bar{\Phi}_{st}(T)$	Average event frequency (accident, spurious safety actions)
$N_a, N_a(t), N_f, N_f(t)$	Number of accident or failures
$\Theta_z(T)$	Mean sojourn times in state Z over [0, T]
General notations and subscripts	
$\bar{X}(t_1, t_2), \bar{X}(T), \bar{X}$	Average values over $[t_1, t_2]$, over $[0, T]$ and over $[0, \infty[$
X_{as}	Asymptotic value of X
$X_{du}, X_{dd}, X_{sf}, X_{st},$	X related to <i>dangerous undetected, dangerous detected, safe failures, spurious safety actions (spurious trip)</i> .
$X_a, X_{a,du}, X_{a,dd}, \text{ etc.}$	X related to an hazardous event (accident)
$X_S, X_A, X_{AB}, \text{ etc.}$	X related to a safety system, a component A, a group of component A,B
X_{fs}, X_{ps}	X related to full or partial stroking (valve modelling)

Table 3 (continued)

Symbol		Definition
X_i, X_i^j		i, j indexes used when X is split in several parts.
$X^k, (X_i)^k, (X_i^j)^k$		Power k when k is an integer
$X_{Z,A,du}^3, X_{S,du}^{as}, \text{etc.}$		Indexes, (i, j) , (as, dd, du, sf, st), (S, A, AB) can be mixed
Parameters		
$\lambda, \lambda(t)$		Failure rates (component level)
$\Lambda, \Lambda(t), \Lambda_{eq}, \Lambda_{eq}(t)$		Equivalent failure rates (system level)
$\lambda_v, \lambda_v(t)$		Conditional failure intensity (Vesely failure rate)
λ_d		Demand rate
μ		Repair / restoration rates
γ, ψ		probability of failure due the test itself, Probability of failure due to a demand of the safety action
τ, θ		test interval, first test interval
π		test duration
ω		probability of human error
States		
KO		Class of states where a system is available
OK		Class of states where a system is unavailable
CS		Class of critical states

5 Overview and challenges

5.1 General considerations about modelling and calculation challenges

Reliability modelling and calculations constitute a delicate area because of the human perception, the technical difficulties, the potential “conflicting” interests between safety and production and the cross-functional organizational aspects of this activity within the overall hierarchy of the organization.

Challenges related to the human perception:

- Reluctance to use probabilistic calculations: they are more difficult to comprehend than deterministic rules and many people are circumspect in using them.
- Probabilistic reasoning: it is rather difficult and what seems obvious with a superficial analysis may appear to be inadequate when a sound analysis is performed.
- Working with time constraints leading to (over)simplifications which:
 - promotes superficial quick analyses of individual parts rather than detailed analysis of the system as a whole;
 - encourages practitioners to think that reliability modelling and calculations can be properly made just by applying analytical formulae or using black boxes without really understanding the underlying assumptions and limitations;

- convinces that additions and multiplications is a sufficient mathematical background to undertake probabilistic calculations.
- Poor dissemination of the advances in reliability modelling and calculations and belief that new tools are more difficult to use than formulae - when, actually, the contrary may be the case. A denial of usefulness can even be observed.

Difficulties related to technical issues:

- The basic concepts of reliability engineering seem simple when they conceal difficulties that non reliability experts may not be aware of.
- A system cannot properly be analysed by the simple juxtaposition of its parts: it is important to consider it as a whole. This implies “systemic” approaches and the size of such models may need the use of computers.
- The periodic tests introduce the following key challenges to be addressed:
 - approximations commonly used may not be conservative;

NOTE When periodic tests are implemented, the current approximations become worse when the probability decreases (i.e. when the redundancy increases). This behaviour is contrary to what happens in the usual situation where the approximations are generally better as the probability decreases.

- average probabilities of periodically tested components cannot be directly used to calculate the average probabilities at system level because the average of a product of quantities is not the product of the averages of these quantities.
- systemic dependencies are induced by periodic tests between safety systems working in sequence (see the Note 1 to entry of [3.4.10](#)). This implies that the requirement of independent safety layers is rarely met.
- awareness of the limitations inherent in software tools on the above matters.
- The handling of uncertainties should be done properly. These issues are further addressed in [Clauses 12](#) and [13](#), and also in ISO 20815[16] (D.3.7 and E.2). This relates to:
 - the importance of qualified input data (technical and operational assumptions and reliability data),
 - the quality of the modelling itself (the major focus in this Technical Report) and applicability of software tools used to facilitate calculation,
 - the final calculated result parameters and the need to reveal how robust and sensitive these results are in the above-mentioned framework, requiring possible use of statistical methods.

The traditional split between safety and production in the industrial organizations as well as in international standardization makes it more difficult to handle properly the tight links between safety and production: improving the first one generally damages the other one and reciprocally. Improving safety without production availability considerations is likely to lead to architectures subject to spurious safety actions. Then the safety is achieved to the detriment of the production availability of the installation. This is not the right way to design safety systems because even if the safety aspects constitute the dominant factor, the spurious actions have also to be considered to determine the best compromises. Moreover in the oil and gas industries, a spurious safety action is not necessarily safe: restarting after a spurious trip is often a perilous operation and a safety system provoking too many spurious trips is likely to be promptly disconnected. This may lead to dangerous situations especially when reconnecting the safety system is forgotten. Then combined systemic approaches where both safety and production are thoroughly considered are recommended by the present Technical Report which analyse both the inhibitions and spurious actions of the safety systems.

The probability of failure of the safety systems themselves is obviously a very important topic but it often tends to hide the primary problem which is to bring the probability of accident (or of the hazardous event frequency) below an acceptable level. This cannot be achieved properly just by calculating some

probabilistic measures on individual safety systems. Of course, the implementation of reliable safety systems (e.g. with high SIL requirements) helps to achieve the accident target but only a systemic analysis where all the safety layers are included is able to verify that the accident target is met.

The software tools necessary to accomplish the calculations related to the methods described in this Technical Report should in appropriate manner take into account the issues addressed in this report.

5.2 Deterministic versus probabilistic approaches

Since the beginning of the industrial era in the 19th century, the engineering philosophy has been to design the future systems by using the knowledge accumulated step by step during the past. This know-how built through a “trial and error approach” is now compiled into various state-of-the-art documents: good practices, rules, regulations, standards, etc. It is permanently improved according to the adverse events (failures, incident or accident) occurring in industry.

Waiting for negative experiences to improve the state-of-the-art is not really adequate for high risk industries where the consequences of the accidents are so large that they should be absolutely avoided. This is why alternative complementary methods have been investigated in aeronautics which has used “statistics” to compare the types of planes since World War I and “probabilities” to predict their number of expected accidents since the 1930s. The aeronautics field has been followed by rocket technology during World War II and later by the space industry and telecommunications field, the US army and the nuclear field. More recently it was the turn of the chemical, petrochemical and transportation fields to join this approach which is now in use in the main industrial areas.

The “deterministic” and “probabilistic” approaches are therefore available to design new systems. They are, indeed, very much complementary: it would not be possible to design a new system from scratch just by using the probabilistic approach, but oppositely the deterministic approach is not sufficient to prevent accidents which have not been observed yet. The result is that the deterministic rules constitute the basis of industrial systems design whereas the probabilistic approach provides effective means for in depth analyses, verification and improvement: the aim of the probabilistic approach is not to make the system run properly but to analyse how a system which is reputed to run properly can fail. Therefore the basic assumption of this Technical Report is that the safety systems have been designed according to the current state-of-the-art and run properly, before the probabilistic approaches are undertaken.

5.3 Safe failure and design philosophy

There are two philosophies to design safety systems:

- 1) a loss of energy provokes the safety action (“de-energize to trip” safety systems);
- 2) an emission of energy provokes the safety action (“energize to trip” safety systems).

With regard to the safety action, the first case is more reliable than the second one because the safety action occurs any time the energy is lost somewhere in the safety system. All failures related to energy or signal losses are safe and don’t need to be considered when evaluating the probability of dangerous failures. Therefore, this simplifies very much the analysis to be performed as well as the safety systems modelling.

This philosophy seems perfect and this is the solution which is the most widely implemented in the oil and gas sector. Nevertheless the counterpoint is that it increases the probability of spurious action.

This may be critical when there is no state where the risk completely disappears or when spurious actions have detrimental effects on the installation (e.g. water hammer). In this case, the spurious actions should be limited and the second philosophy could be used for this purpose.

When dealing with de-energize to trip safety system, only the components undergoing dangerous failures are considered within the safety analysis and the other components are generally ignored. Nevertheless they may participate in spurious actions and it is important to consider them when performing the safe or spurious failures analysis. This encompasses, but is not limited to:

- a) loss of power supply, e.g.:

- electric power supply (AC or DC);
- hydraulic power supply (i.e. the isolation valve is “fail-safe close”);
- b) loss of auxiliary supply (UPS, compressed air, etc.);
- c) rupture of the links between the sensors and the logic solvers;
- d) rupture of the links between the logic solvers and the final elements;
- e) any dangerous detected failure which is transformed into a safe failure by the self diagnostic tests.
- f) human actions following false alarms

Then the safe or spurious failure studies of “de-energize to trip” safety system need more detailed analysis than the dangerous failure analysis.

5.4 Dependent failures

5.4.1 Introduction

From a technological point of view (i.e. excluding the human factor considerations dealt in 5.5), there are three main ways to decrease the probability of failure of a given system (including safety systems):

- 1) using more reliable components,
- 2) improving the maintenance strategy (e.g. increasing the test frequency of hidden failures),
- 3) introducing more redundancy to increase “fault tolerance” (see definition 3.4.9).

The first one is practicable only if more reliable components are available on the market, the second one is usable to a small extent (e.g. because tests are costly and may have detrimental effects when their frequency is too high) and therefore the third one is generally implemented. It is effective to the extent that the redundant parts fail independently from each others. Therefore the identification and analysis of dependent failures (i.e. common cause failures and systematic failures) is an important topic when dealing with redundant (or fault tolerant) safety systems.

5.4.2 Common cause failures

From an academic point of view, if the probability of failure of a single component is p , the probability of failure of two independent redundant similar components (e.g. 1oo2) is p^2 , of three independent redundant similar components (e.g. 1oo3) is p^3 , of four independent redundant similar components (e.g. 1oo4) is p^4 , etc. Therefore it seems, in theory, possible to reduce the probability of failure to any level just by implementing the right level of redundancy. Let us take an example where the probability of failure of one component over one year is $p = 10^{-3}$. Over the same period of time this would lead respectively to 10^{-6} , 10^{-9} and 10^{-12} for 1oo2, 1oo3 and 1oo4 systems. These results are correct from a mathematical point of view but completely unrealistic from a physical point of view: in fact 10^{-9} is more or less the probability that system fails due to the collision with a meteorite over one year and 10^{-12} is about 100 times lower than the probability that the failure occurs because the Universe disappears!

Therefore some factors exist which limit the potential probability reduction expected from redundancy: they are linked to the potential for common cause (CCF), common mode (CMF) and systematic failures to cause multiple components to fail.

CCF can be split into the following categories:

- Explicit CCF: an event or mechanism that can cause more than one failure simultaneously is obviously a CCF. Such CCF can generally be explicitly known, analysed and treated as individual failures (e.g. meteorite collision, loss of power supply, fire, flooding).

- **Implicit CCF:** an event or mechanism that only increases the probability of failure of several components is obviously also a CCF. This kind of CCF does not induce the immediate failures of components but only degrade them so that their joint probability of simultaneous failure is increased (e.g. corrosion, premature wear out, over strength, unexpected systemic effects, design error, installation error). Such CCF cannot generally explicitly be known and remains implicit.

Implicit CCFs are far more difficult to identify, analyse and treat than the explicit CCF. It is very likely that some of them will remain ignored until they manifest themselves. They constitute a part of the epistemic uncertainties. ISO 14224[15] advises how such data can be obtained (see ISO 14224[15], Table B.3).

This Technical Report develops three of the pragmatic CCF approaches developed to establish effective safeguards in order to prevent over optimistic evaluations. They are presented in 5.4.2 and Annex G: β factor, shock model and PDS approaches.

NOTE Redundancy is more efficient to fight against CCFs which do not induce immediate failures because, provided that the relevant procedures are implemented to detect those CCFs, they can be discovered and repaired before all the redundant components are affected.

5.4.3 Systematic failures

It is outside the scope of this Technical Report to deal with the systematic failures (see 3.2.17) and the standards dealing with *functional safety* provide valuable information on this subject. A systematic failure results from a pre-existing fault (see 3.2.2) and occurs every time some given conditions are encountered. Then the systematic failures are perfect common cause failures which, like other CCFs do, limit the usefulness of redundancy. Beyond a certain level of redundancy (e.g. two or three channels) the benefit is illusory, especially if the software implemented in the redundant part has been developed on the same specification bases.

Further details are provided in Annex G.

5.5 Human factors

5.5.1 Introduction

The operator is sometimes defined as an unreliable agent of reliability. That means that, on one hand he/she is prone to commit simple action slips as well as more serious (cognitive) errors due to inadequate understanding of the system design and operation when on the other hand it is also acknowledged that he/she is sometimes the only element able, with sufficient understanding, to make good decisions when unscheduled situations occur i.e. improvising when suitable procedures have not been thought out beforehand.

Another interesting issue is that human beings can learn from their errors. Even if the operational situation stays constant the level of experience that the operators have of their interaction with it can affect their performance. They can improve their mental model of the world such that the operational situation is more predictable. Alternatively poor information and experience can result in applying incorrect mental models of the operational situation.

The role of the human operator is predominantly:

- part of the condition monitoring system;
- restoring to the appropriate system state in response to signalled deviations;
- activator of safety systems.

With regard to this Technical Report, there are several situations where the human function is appropriate for carrying out such tasks, e.g.:

- periodic testing;
- calibration;

- inhibition management;
- routine inspection and maintenance;
- repair of failed components;
- restoration / reconfiguration after repair;
- restarting of the process (e.g. restoring a system after a shut down/spurious trip);
- alarm detection and actions following alarms;
- recovery of a failed condition.

NOTE The human operator can be an initiator of unwanted events that make demands on safety systems. This contributes to the demand frequency of those safety systems and this out of the scope of this Technical Report.

In these roles the human operator is not only performing a function but is also a source of undesirable events. For that reason there are pros and cons for having a human in the system. It is therefore of value to model and quantify the effect of this.

Further information can be found in [Annex H](#) and in references [\[36\]](#), [\[37\]](#), [\[38\]](#), [\[39\]](#), [\[40\]](#), [\[41\]](#), [\[42\]](#), [\[43\]](#), [\[44\]](#) and [\[45\]](#)

5.5.2 Human performance considerations

Human error is generally considered to be of three types:

- memory and attention lapses and action slips;
- mistakes in planning actions;
- rule violations.

In addition, three levels of human performance can be considered:

- skill based;
- rule based;
- knowledge based.

These levels are hierarchical. Performance at the knowledge-based level will call upon rule-based procedures which will call upon skill-based operations to execute the procedure. Human performance is most reliable at the skill-based level and least at the knowledge-based level.

When an operator has to assess a given condition such as in inspection and testing (e.g. periodic testing of a safety system) or in monitoring a system state where there is a certain degree of uncertainty a *signal detection* model can be used with four classes of response to the evidence:

- Hit – the operator correctly detects the signal.
- Miss – the operator missed the signal.
- False alarm – the operator identifies a noise event as a signal.
- Correct reject – the operator identifies a noise event as noise.

In human reliability assessment the primary focus is on eliminating systematic causes of human error at these different levels of performance. This is not discussed further in this Technical Report which is more focused on human random failures.

NOTE Systematic causes exist in equipment, man-machine interface and task design, training, procedures, information, communication, stress and fatigue, workload, and underlying organizational and management factors. Other factors to be taken into account are dependencies, for example when the same operator carries out the same task on otherwise independent systems.

5.5.3 Human error modelling

Human errors can be modelled as:

- omissions (omitting to do something);
- commission errors (doing things wrongly);
- acting too late or too early (a subset of commission error).

They are identified by examining procedures associated with a specified task, the initiation of a blow down for example. Commission errors not associated with a requirement for the operator to act and acts of sabotage are not generally modelled.

With regards to this Technical Report, the only human errors of interest are those which affect the functioning of the safety system in some way according to the intended design and operation. So, the analysis begins with the task performed within that function. Thereafter, understanding the task in terms of the performance requirements of the operator, both psychologically and physically, and the capacity of the operator to respond are important for quantification.

The probability of occurrence of a human error is measured by:

HEP = Numbers of human errors occurred/Number of opportunities for error, where HEP is the Human Error Probability.

The operator's opportunity for error can be identified as the whole task or it can be broken down to identify all the opportunities for human error which could result in a failure of interest. An opportunity for error might be, for example, the requirement to return gas detectors to service after maintenance or to close a valve before starting up a system

An operator may forget to close a valve, improvise by not replacing like with like or miss or dismiss signals to act. Such human errors can be identified from historical data but also can be surmised when there is a requirement for human action which could affect system performance. Some human reliability analysis (HRA) methods have human error identification prompts such as:

Can the signal be perceived as unreliable? → Signal ignored → Action omitted.

These undesirable events can be modelled as human error probabilities per demand for action or per task for example, the probability that the operator will not initiate the appropriate safety function in time on the occurrence of a particular blow out indicator.

It is important that the justification behind the quantification describes the conditions under which the operator undertakes the task as changes in these conditions can violate the assumptions made about the human error probability.

These concepts of human performance and human error modelling are important for linking the human performance into the safety system. It is outside the scope of this Technical Report to deal in detail with the human factor but further details may be found in [Annex H](#).

5.5.4 Human Reliability Analysis Methods

See also IEC 61508[2] [Annex A](#) which gives examples of HRA methods.

Because of a lack of real data for the HEP calculation a number of methods have been developed. First generation methods were developed to help provide HEPs required in risk assessment. The most notable publically available methods include:

- THERP: Technique for Human Error Rate Prediction;
- HEART: Human Error Assessment and Reduction Technique.

and expert judgement methods:

- APJ: Absolute probability judgement;
- PC: Paired comparisons;
- SLIM-MAUD: Success likelihood index method using multi-attribute utility decomposition.

Second generation methods start around 1990 and were developed to better model the context in which the operator is working:

- ATHEANA: A Technique for Human Error Analysis;
- CREAM: Cognitive Reliability and Error Analysis Method.

Direct measurement of human failures is also possible.

- Within the framework of SDT - Signal detection theory;
- Conducting experiments on human performance under controlled conditions;
- Use of simulators.

Analysis of past incidents can also identify where human error contributed but not the number of opportunities for error.

Some HEPs and a worked example of the Paired Comparisons method are given in [Annex H](#).

5.6 Documentation of underlying assumptions

For further use of the reliability studies it is of utmost importance to keep the trace of every document which has been used as well as all prerequisites and underlying assumptions which have been adopted during the studies. This is an important part of the reliability analysis. It is often insufficiently done and should be more thoroughly considered.

Documenting underlying assumptions is important in order to:

- Highlight what technical, operational and analytic assumptions the results are based on and the consequences of these assumptions.
- Highlight measures and activities that are important to follow up and implement in design and operation – often being of more importance than the quantitative results themselves.
- Enhance the possibility of understanding the analysis results and the importance of the assumptions.
- Establish a good foundation to evaluate the validity of the results or for future use of the analysis, e.g. when updating the analysis.

Types of assumptions that are important to highlight are for example:

- The sources where the input reliability and maintenance data are taken from, and to what extent the data has been qualified and adjusted to account for a particular application. See ISO 20815^[16], [E.2](#).
- Any assumptions that have been made for the data concerning diagnostic coverage, critical failure modes, safe state, and so on, which are normally not presented (or presented in the same way) in

many data sources. In particular it is important to document and justify/substantiate failure rates that are significantly better than average generic data.

- Any assumptions that have been made in relation to periodic tests and maintenance, including required intervals and specific methods and/or tools to be applied.
- Any assumptions that have been made regarding operating environment.
- Any assumptions and analyses that have been made for evaluating the common cause failures (e.g. β factors) and for reducing them, e.g. diversity or staggered testing.
- All relevant assumptions related to procedures for inhibits and overrides, degraded operation, maximum repair times (if any), and assumptions concerning set points and required human interaction.
- Any assumptions that need to be validated during operation and how often such validation needs to be performed.

6 Introduction to modelling and calculations

6.1 Generalities about safety systems operating in “on demand” or “continuous” modes

A safety system is a *protection* system which protects another system - the *protected* system - against incident or accident. When the protection and the protected systems are not fully separated, it is essential to analyse them together as a whole. The protected system is sometime called “Equipment Under Control” (EUC) (e.g. in IEC 61508[2])

Two different modes of operation are generally considered with regard to how the safety system protects the protected system:

- **demand mode of operation:** the safety system spends most of its time to wait for a demand which occurs randomly when, due to hazardous events, a given physical parameter of the protected system exceeds a pre-established threshold. Reactive safety systems (also called curative safety system) operate on demand mode of operation.
- **continuous mode of operation:** a hazardous event occurs on the protected system as soon as the safety system fails. Preventive safety systems are operating in continuous mode of operation.

Let us consider an overpressure protection system which protects a pipe downstream by closing a valve when the upstream pressure exceeds a given threshold. During normal operation the upstream pressure fluctuates according to the process operation, hazardous events (e.g. overpressure) may occur and the safety system has to be ready to isolate the pipe (e.g. by closing a Shut Down Valve) as soon as a pre-established pressure threshold is reached. It performs its safety action just before the protected system enters in the dangerous zone: this is a typical reactive safety system working on demand mode of operation.

After the overpressure protection system has reacted to a demand by closing a SDV, the dynamic pressure drop decreases and the pressure increases upstream the valve closed between the pressure source and the pipe. After a while this pressure may become higher than the pressure strength of the downstream pipes and a spurious opening of this valve may lead to a dangerous situation. Because of the speed of the pipe pressurization in case of spurious opening, a reactive system may not be fast enough to close another valve before the overpressure occurs. Then another safety system is needed to prevent the spurious reopening of the valve. Its safety action is to continuously maintain the valve in the closed position until the pressure has dropped below a safe value: this is a typical preventive safety system working on continuous mode of operation.

Reactive safety systems are commonly used whereas preventive safety systems have just begun to be implemented in the oil and gas industries.

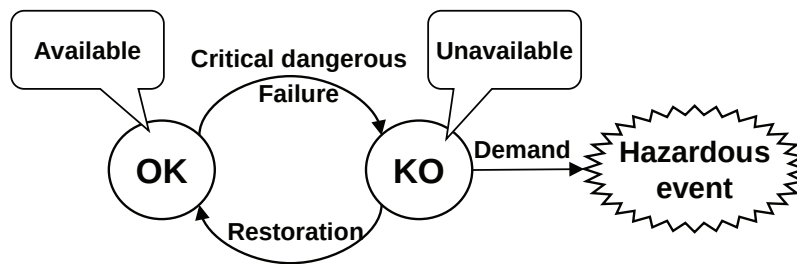


Figure 10 — Broad model for ultimate-on-demand mode safety system

As shown in Figure 10 a safety system working on demand mode may have failed before the occurrence of the demand and has not been repaired in the meanwhile, for a hazardous event occurring. What is important to prevent the hazardous event is that it is ready to operate at the time that the demand occurs. This does not matter if it has previously failed provided it has been repaired in the meanwhile: the “availability” (see definition 3.1.11, 3.1.12, 3.1.13) is the adequate probabilistic parameter to characterize such a safety system. Therefore, if the safety system is unavailable when the demand occurs and if no other succeeding safety system is able to respond, then the hazardous event occurs.

When the average unavailability $\bar{U} = \text{PFD}_{\text{avg}}$ of the safety system is small and the demand rate λ_d is low, then the hazardous event frequency can be calculated as: $\Phi_a(T) \approx \bar{U} \cdot \lambda_d = \text{PFD}_{\text{avg}} \cdot \lambda_d$. This is a widely used formula which is valid with some limitations analysed in Annex D.

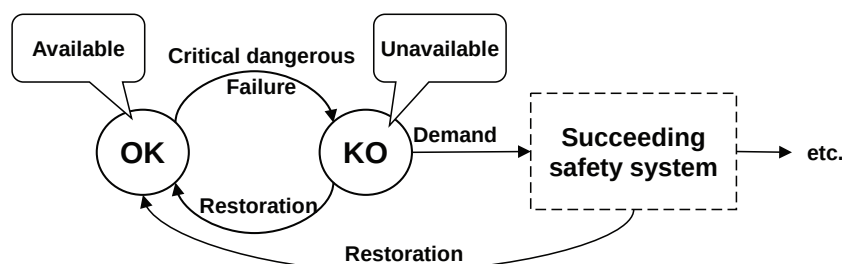


Figure 11 — Broad model for non-ultimate-on-demand mode safety system

In case of multiple safety-systems (see 3.3.3) another succeeding safety system may be able to respond, then, in this case, a new demand is generated to this safety system and the first one may be restored as shown in Figure 11. In this case the above formula gives the demand frequency on the succeeding safety system. The multiple safety systems are analysed in more detail in Annex F.

It is not allowed for a safety system working in continuous mode to fail as long as it is in use (e.g. as long as the pressure has not been lowered behind a safe threshold): the “reliability” (see definitions 3.1.8 and 3.1.9) is the adequate probabilistic parameter to characterize such a safety system.

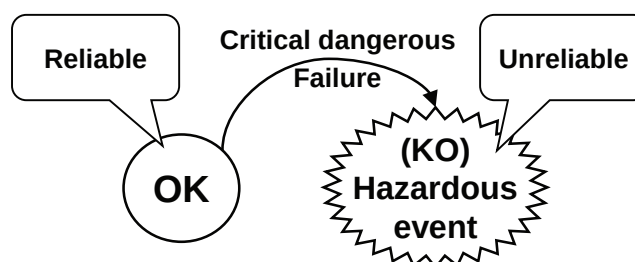


Figure 12 — Broad model for ultimate-continuous mode safety system

As shown in Figure 12 the hazardous event occurs as soon as the safety system fails if no other succeeding safety system is able to respond. Therefore, the critical dangerous failure rate of the safety system is

also the hazardous event rate and the hazardous event frequency $\Phi(t)$ is equal to the failure frequency of the safety system: $\Phi(T) = \bar{w}(T)$. This is analysed in [Annex E](#).

If another succeeding safety system is able to respond, then a demand is generated to this system and the first one may be restored as shown in [Figure 13](#). In this case the above formula gives the demand frequency on the succeeding safety system which, therefore, operates on demand mode. The multiple safety systems are analysed in more detail in [Annex F](#).

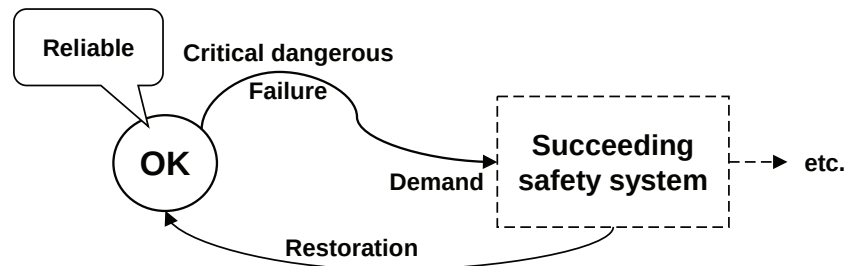


Figure 13 — Broad model for non-ultimate-continuous mode safety system

Even if the aim is the same in any case (i.e. determining the hazardous event frequency), [Figure 11](#) to [Figure 13](#) show that the reliability modelling and calculations depend on the type of the related safety system: average unavailability for safety systems operating in demand mode and unreliability or failure frequency for safety systems operating in continuous mode.

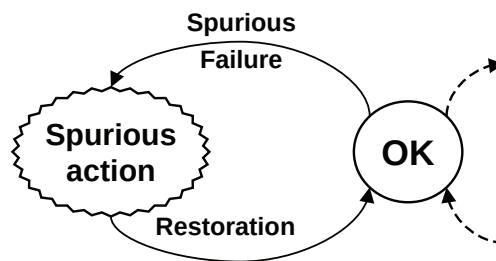


Figure 14 — Broad model for spurious failures

Nevertheless, when the demand frequency is high, the demand mode of operation system can be considered and modelled as a continuous mode of operation safety system. This approximation is conservative.

Both on demand and continuous mode of operation safety systems may trigger the safety action spuriously (e.g. spurious shut down) and the [Figure 14](#) provides a broad model for the spurious failures of any type of safety systems. The similarities between the [Figure 12](#) and [Figure 14](#) suggest that the probabilistic calculations undertaken for evaluating the spurious failure frequency are similar to the calculations undertaken to evaluate the hazardous event frequency in the case on non-ultimate continuous mode safety systems.

Further developments from these broad models are provided in [Annex D](#), [Annex E](#) and [Annex F](#).

NOTE In de-energize to trip design, all the failures able to lead to a spurious action should be considered (see [5.3](#)). This may imply that other components than those considered within the dangerous failure analysis need to be considered.

6.2 Analytical approaches

6.2.1 Basic models (reliability block diagrams)

Basically an individual safety system performs three functions:

- 1) detection that a physical parameter crosses a given threshold;
- 2) decision to trigger the safety action;
- 3) performance of the safety action.

These three functions can be gathered in a single component (e.g. a relief valve) or split between several components (e.g. a safety instrumented system). [Figure 15](#) represents the archetype of the architecture of a safety instrumented system comprising sensors (S), logic solvers (LS) and final elements (V) represented as separate blocks in a so-called block diagram.

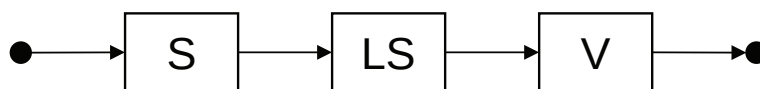


Figure 15 — Basic safety instrumented system/function

Such a block diagram is actually a “reliability block diagram” (RBD)^[4] where each block has only two states (e.g. “OK” and “not OK”, or “KO” as it is denoted here). RBDs are simple and widely used (but should not be mixed up with multistate flow networks) and they are useful to introduce the basics of modelling. Formally they belong to the “Boolean” models which are analysed below (see [Clause 8](#)). Even if a RBD may be close to the actual physical architecture of the safety system, it is an abstract model aiming to represent as well as possible the functioning and the dysfunctioning of the safety systems under consideration.

In [Figure 15](#) the three blocks are in series and the failure of any block implies the failure of the safety function. This basic architecture in three blocks is often preserved for more sophisticated safety systems (e.g. [Figure 16](#) and [Figure 17](#)). In functional safety standard (e.g. IEC 61508^[2]) those blocks are so-called “subsystems”.

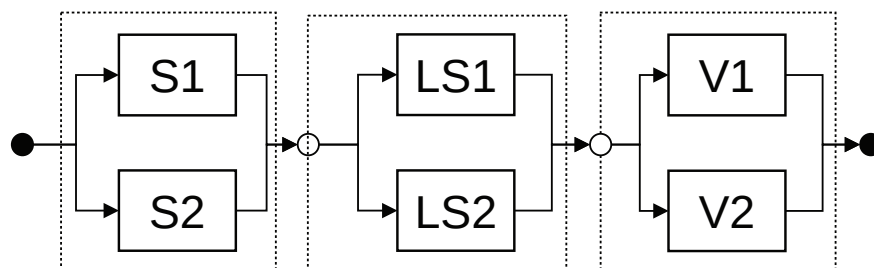


Figure 16 — Model of a safety instrumented system with redundant components - Example n°1

[Figure 16](#) represents a safety system where all the components are redundant in order to decrease its probability of dangerous failure. Redundant components are represented in parallel on the RBD.

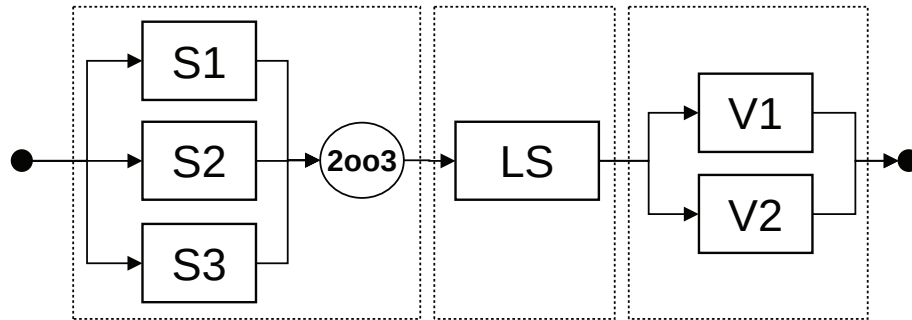


Figure 17 — Model of a safety instrumented system with redundant components - Example n°2

[Figure 17](#) represents a safety system where three redundant sensors are used through a 2oo3 voting logic performed within the logic solver in order to decrease both the probability of dangerous failure and the probability of spurious failures due to the sensors.

Nevertheless all the safety systems are not necessarily split into three parts (e.g. relief valves) and other architectures can be used according to specific needs. [Figure 18](#), for example, proposes an architecture where it is not possible to identify three fully separate parts in series because the two logic solvers get input from only two of the three sensors.

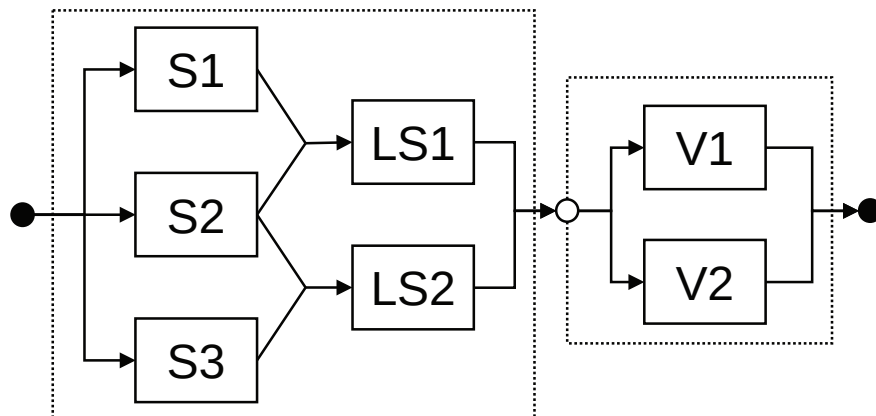


Figure 18 — Model of a safety instrumented system with redundant components - Example n°3

6.2.2 Qualitative analysis (Minimal cut sets)

Beyond the fact that RBDs are flexible enough to represent the most complex architectures, they allow, in addition, the identification of the failure scenarios of the modelled safety system. For this purpose it is “sufficient” to consider the RBD as an electrical circuit and to analyse where it has to be “cut” in order for the whole circuit to be cut. Each failure scenario is then a “cut set” which is said to be “minimal” when all the failures that it comprises are necessary and sufficient to lead to a system failure.

For example, the minimal cut sets of [Figure 15](#) are simply $\{S\}$, $\{LS\}$, $\{V\}$ where S , LS , V represent the events “failures” of the related blocks (components). Similarly:

- Minimal cut sets of [Figure 16](#) are $\{S_1, S_2\}$, $\{LS_1, LS_2\}$, $\{V_1, V_2\}$.
- Minimal cut sets of [Figure 17](#) are $\{S_1, S_2\}$, $\{S_1, S_3\}$, $\{S_2, S_3\}$, $\{LS\}$, $\{V_1, V_2\}$.
- Minimal cut sets of [Figure 18](#) are $\{S_1, S_2, S_3\}$, $\{S_1, S_2, LS_2\}$, $\{S_2, S_3, LS_1\}$, $\{LS_1, LS_2\}$, $\{V_1, V_2\}$.

The list of minimal cut sets is in close relationship with the architecture: this is a systemic representation.

On this example $\{S\}$, $\{LS\}$, $\{V\}$ are single failures when $\{S_1, S_2\}$, $\{S_1, LS_2\}$... are double failures and, $\{S_1, S_2, S_3\}$, $\{S_1, S_2, LS_2\}$... are triple failures. Other multiple failures may be identified in more complex system (e.g. multiple safety systems). As single failures are generally more probable than double failures which are generally more probable than triple failures ..., the analysis of the minimal cut set provides a very effective way to identify the weak points of the system under study from a qualitative point of view: the single failure $\{LS\}$ in Figure 17 is, for example, a weak point candidate. In order not to waste time and money it is wise to analyse $\{LS\}$ and possibly improve it first before looking at $\{V_1, V_2\}$ or $\{S_1, S_2, S_3\}$.

The decision of improvement of the weak points can be based on:

- qualitative architectural constraints (e.g. single failure criteria, minimum redundancy or hardware fault tolerance);
- quantitative probabilistic constraints (e.g. probability or frequency of hazardous events).

6.2.3 Basic considerations about approximated probability calculations

When the minimal cut sets have been identified, then the safety system can be modelled by the RBD formed of its minimal cut sets placed in series. For example, the RBD in Figure 18 can be replaced by the equivalent RBD in Figure 19. The difference is that a given block (e.g. S1) may be represented several times when it belongs to several minimal cut sets.

The sum of the probabilities of the scenarios, i.e. the sum of the probability of the minimal cut sets (in dotted boxes in Figure 19) calculated separately, provides an upper bound of the overall probability of failure of the safety system. When this probability is low (what is normally the case for a safety system) this upper bound is very close to the exact result which, therefore, provides a good conservative approximation (see Sylvester-Poincaré formula in K.1).

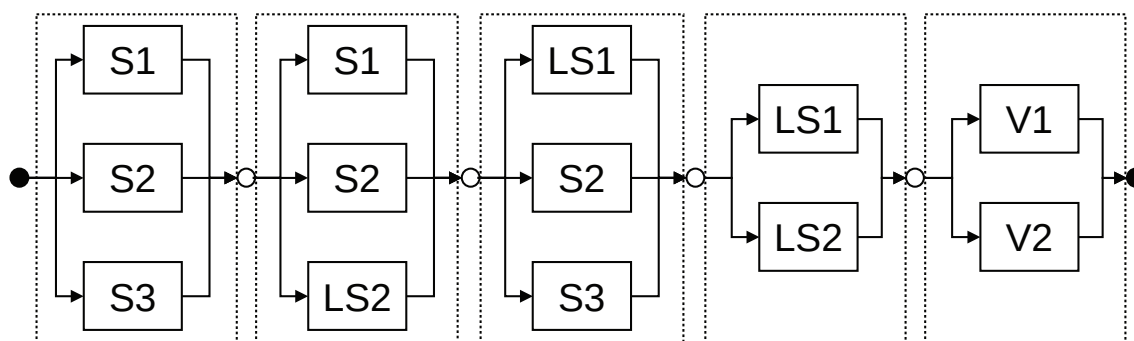


Figure 19 — Minimal cut set representation of example n°3

The following approach can be implemented to develop simplified calculations:

- 1) model the safety system (e.g. by using a RBD);
- 2) identify the failure scenarios (e.g. the minimal cut sets);
- 3) calculate the relevant probability of each minimal cut set;
- 4) add the probabilities obtained at step 3 (conservative estimation).

In particular, this can be used to calculate the unreliability, the unavailability, the average unavailability (i.e. PFD_{avg}) of the whole safety system or the hazardous event probability or average frequency (i.e. PFH).

This may be implemented with any of the approaches described in this Technical Report. Nevertheless it is particularly interesting when the analytical formulae are used as this reduces the problem to the establishment of separate formulae for single failures, double failures, triple failures, etc.

7 Analytical formulae approach (low demand mode)

7.1 Introduction

The use of analytical formulae for probabilistic calculations is wide-spread[2][13][46][47] and the aim of [Clause 7](#) is to illustrate how they can be developed for simple safety systems. These formulae are not intended to be used without a good understanding of what they mean. Therefore some mathematical explanations are given to:

- clarify how they are established;
- allow the user to develop his/hers own formulae if needed.

The underlying model behind the analytical formulae is the minimal cut set theory ([Clause 6](#)) and the Markov theory ([Clause 9](#)).

In order not increase the size of the document only safety systems operating in low demand mode are analysed in this Technical Report, but the principles are the same for safety systems operating in high demand or continuous mode of operation.

In order to make the reading easier, the detailed mathematical developments have been moved in [Annex I](#) and only the main results are presented in the main part of this Technical Report.

7.2 Underlying hypothesis and main assumptions

Let us consider the event $C = A \cup B$. This event C occurs when the event A or the event B occurs and its probability $P(C)$ can be calculated by using the following basic probabilistic formula:

$$P(C) = P(A \cup B) = P(A) + P(B) - P(A \cap B)$$

This formula becomes $P(C) = P(A \cup B) = P(A) + P(B) - P(A).P(B)$ when the events A and B are independent.

This formula is interesting from the probabilistic calculation point of view because:

- When $P(A)$ and $P(B)$ are close to 1, $P(A).P(B)$ is also close to 1. In this case it *cannot* be neglected otherwise probabilities higher than 1 could be obtained.
- When $P(A)$ and $P(B)$ are small compared to 1, $P(A).P(B)$ becomes negligible compared to $P(A)$ and $P(B)$. In this case it *can* be neglected and the following approximation used:

$$P(C) = P(A \cup B) \approx P(A) + P(B)$$

When dealing with safety, the probabilities of good functioning (reliability, availability) of the safety systems and their components are normally high and this is in the case a). Reciprocally, the probabilities of failures (unreliability, unavailability) are small compared to 1 and this is the case b).

Then all along [Clause 7](#) the main underlying hypothesis is that the safety systems and their components are reliable and available enough to apply approximation b):

$$F(t) \ll 1; U(t) \ll 1; \bar{U}(T) \ll 1 \text{ and } R(t) \oplus 1; A(t) \oplus 1; \bar{A}(T) \oplus 1$$

This allows deriving analytical formulae from approximations by keeping the first terms of the Taylor's expansion of the exact mathematical expressions of unreliability or unavailability. This can be done only when relevant assumptions are fulfilled. This Technical Report's assumptions made in developing the analytical formulae are the following:

At component level:

- The test interval τ is constant.

- The test coverage is of 100 % (i.e. all potential failures are detected by the periodic tests).

NOTE The failures not covered by the periodic tests can be modelled by another component with the appropriate characteristics and placed in series with the periodically tested component.

- The failure rate λ is constant (i.e. the failure is catalectic) and $1/\lambda$ is the mean time to fail (MTTF) of the component.
- τ is such that $\lambda\tau \ll 1$. (i.e. the probability of failure within a test interval is low).
- The repair rate μ is constant and $1/\mu$ is the mean repair time (MRT).
- τ is such that $1/\mu \ll \tau$ (i.e. MRT is negligible with regards to the test interval duration).
- The test duration π is constant but it may be equal to 0 in the simplest models.
- τ is such that $\pi \ll \tau$ (i.e. the test duration is negligible with regards to the test interval duration).
- The protected installation is not stopped during tests and repairs of components, otherwise MRT should be replaced by the time elapsing between the detection of the fault and the stopping of the installation (MTTS), and the test duration π would be taken equal to 0.

At overall safety system level:

- The demand rate λ_d is constant and $1/\lambda_d$ is the mean time to demand (MTTD) of the safety system.
- T is such that $\lambda_d T \ll 1$ where T is the period under interest (i.e. the probability of a demand within the period under interest is low).
- $MTTD > 10 \times MTTR_{es}$ (see D.1 and Figure D.5) (i.e. the mean time to restore a failure is 10 times lower than the mean time between demands).

7.3 Single failure analysis

7.3.1 System description

This is the simplest case. It is represented on Figure 20 by a RBD comprising a single block. It models a minimal cut set which is a single failure.

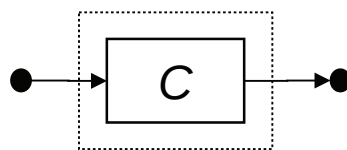


Figure 20 — Single component modelling

The basic reliability parameters are the following:

- λ_{du} : dangerous undetected failure rate (i.e. hidden failures detected by periodic tests);
- λ_{dd} : dangerous detected failure rate (i.e. immediately revealed failures or failures detected by diagnostic tests);
- λ_{sf} : safe failure rate (i.e. able to produce spurious safety actions);

NOTE In de-energize to trip design, all the failures able to lead to a spurious action should be considered (see 5.3). This may imply other components than those considered within the dangerous failure analysis. λ_{st} should be evaluated accordingly.

- μ_{du} : repair rate of dangerous undetected failures ($1/\mu_{du} = MRT_{du}$ after detection);

- μ_{dd} : restoration rate of immediately revealed dangerous failures ($1/\mu_{dd} = \text{MTTRes}_{dd}$);
- μ_{sf} : restoration rate of safe failures ($1/\mu_{sf} = \text{MTTRes}_{sf}$);
- τ : test interval;
- γ : probability of failure due the test itself;
- ψ : probability of failure due to the demand of the safety action itself;
- π : test duration;
- ω : probability of human error.

The usual notation “du” has been used for the dangerous failures remaining hidden until they are detected by a periodic test. The repair rate of these failures does not include the detection delay (see 3.1.33 and Figure 5) which is modelled separately. The usual notation “dd” has been used for immediately revealed or quickly revealed failures. The restoration rate of these failures does include the time needed for the detection (see 3.1.32 and Figure 5).

NOTE γ and ψ are related to actual failures occurring on demand (see 3.2.13). They are not time dependent but demand dependant. Such failures cannot be prevented by tests. This is discussed in Annex I.

7.3.2 Dangerous undetected failures

The mathematical developments are detailed in Annex I and the main results are the following:

- Average unavailability:

$$\bar{U}_{du}(\tau) \approx \frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi$$

- Number of accidents over the time interval $[0, \tau]$:

$$N_{a,du}(\tau) \approx \left(\frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi \right) \lambda_d \cdot \tau$$

- Hazardous event probability over $[0, \tau]$:

$$P_{a,du}(\tau) \approx N_{a,du}(\tau) \approx \left(\frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi \right) \lambda_d \cdot \tau$$

- Average hazardous event frequency:

$$\bar{\Phi}_{a,du}(\tau) = \frac{N_{a,du}(\tau)}{\tau} \approx \left(\frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi \right) \lambda_d$$

NOTE 1 If the protected installation is stopped during tests and repairs (or if MRT_{du} and π are negligible) and if the probability of human errors and of failures occurring on demand are negligible, then the average unavailability is reduced to its first term $\bar{U}_{du}(\tau) \approx \frac{\lambda_{du} \cdot \tau}{2}$.

NOTE 2 For dangerous undetected failure not covered by any test the average unavailability becomes $\bar{U}_{du}(T) \approx \frac{\lambda_{du} \cdot T}{2}$ where T is the duration of the whole period of interest $[0, T]$.

7.3.3 Dangerous immediately revealed/detected failures

The mathematical developments are detailed in Annex I and the main results are the following:

- Average (asymptotic) unavailability:

$$\bar{U}_{dd}(t) = \frac{\lambda_{dd}}{\lambda_{dd} + \mu_{dd}}$$

- Hazardous event probability:

$$P_{a,dd}(T) \approx N_{a,dd}(T) \approx \bar{U}_{dd} \cdot \lambda_d \cdot T = \frac{\lambda_{dd}}{\lambda_{dd} + \mu_{dd}} \lambda_d \cdot T$$

- Average hazardous event frequency:

$$\bar{\Phi}_{a,dd}(T) = \frac{N_{a,dd}(T)}{T} \approx \bar{\Phi}_{a,dd} = \frac{\lambda_{dd} \cdot \lambda_d}{\lambda_{dd} + \mu_{dd}}$$

7.3.4 Spurious failures

The mathematical developments are detailed in [Annex I](#) and the main results are the following:

NOTE 1 With a single component, a safe failure is also a spurious failure which leads to a spurious safety action.

- Average spurious failure frequency:

$$\bar{\Phi}_{st}(T) = \bar{\Phi}_{sf}(T) \approx \frac{1}{MTBF_{sf}} = \frac{1}{MTTF_{sf} + MTTR_{esf}}$$

- When

$$F_{sf}(T) \ll 1, \bar{\Phi}_{st}(T) = \bar{\Phi}_{sf}(T) \leq \frac{1}{MTTF_{sf}} = \frac{\lambda_{sf} T}{T} = \frac{F_{sf}(T)}{T}$$

NOTE 2 Depending of the architecture of the safety systems, a spurious safety action may also provoke a spurious trip of the protected installation.

NOTE 3 When implemented in de-energize to trip architecture, it is very important that the failures of the links with other components or systems are included into the λ_{sf} .

7.4 Double failure analysis

This case is presented in [Figure 21](#) by a RBD comprising two redundant blocks. It models a minimal cut set which is a double failure.

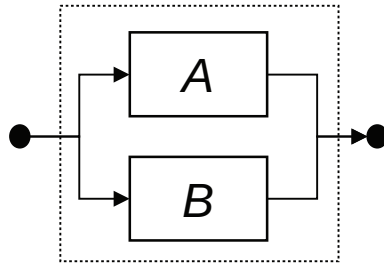


Figure 21 — Double failure modelling

The basic reliability parameters are the following:

- $\lambda_{A,du}, \lambda_{B,du}$: dangerous undetected failure rates of components A and B (i.e. hidden failures detected by periodic tests);
- $\lambda_{A,dd}, \lambda_{B,dd}$: dangerous detected failure rates of components A and B (i.e. immediately revealed failures or failures detected by diagnostic tests);
- $\lambda_{A,sf}, \lambda_{B,sf}$: safe failure rates of components A and B (i.e. failures able to trigger safe or spurious actions);

NOTE In de-energize to trip design, all the failures able to lead to a spurious action should be considered (see [5.3](#)). This may imply other components than those considered within the dangerous failure analysis. $\lambda_{A,sf}, \lambda_{B,sf}$ should be evaluated accordingly.

- $\mu_{A,du}, \mu_{B,du}$: repair rates of dangerous undetected failure of components A and B;
- $\mu_{A,sf}, \mu_{B,sf}$: restoration rates of safe failure of components A and B;
- μ_{dd} : restoration rate of immediately revealed dangerous failure ($1/\mu_{dd} = \text{MTTRes}_{dd}$);
- τ_A, τ_B : test interval of components A and B.

Even for a rather simple system like this one, the number of combinations drastically increases (especially if the extra parameters like $\pi, \omega \dots$ introduced in 7.3.2, are considered). In addition, the tests intervals are not necessarily identical and even if they are identical they can be not performed at the same time. Therefore it is almost impossible to treat all the possible cases. The aim of this clause is to give some advice and to encourage the analysts to understand and develop their own formulae rather than to provide a catalogue of ready-to-use formulae. Further details may be found in Annex I.

7.4.1 System unavailability analysis

The system presented on Figure 21 is failed when both A and B are failed, i.e. when A and B are unavailable at the same time:

$$\text{Safety system unavailability: } U_{AB}(t) = U_A(t) \cdot U_B(t) = [U_{A,du}(t) + U_{A,dd}(t)] [U_{B,du}(t) + U_{B,dd}(t)]$$

This formula comprises three types of terms:

- | | | |
|----|--|--|
| 1) | one term related to dangerous detected failures: | $U_{A,dd}(t) \cdot U_{B,dd}(t)$ |
| 2) | one term related to dangerous undetected failures: | $U_{a,du}(t) \cdot U_{b,du}(t)$ |
| 3) | two terms related to both type of failures: | $U_{a,dd}(t) \cdot U_{b,du}(t), U_{a,du}(t) \cdot U_{b,dd}(t)$ |

$U_{A,dd}(t)$ and $U_{B,dd}(t)$ converge quickly toward asymptotic values u_a, u_b which are also averages values (cf. Annex I) and the above terms can be expressed as:

- | | | |
|----|--|---|
| a) | term related to dangerous detected failures: | $U_1(t) = u_A \cdot u_B$ |
| b) | term related to dangerous undetected failures: | $U_2(t) = U_{A,du}(t) \cdot U_{B,du}(t)$ |
| c) | terms related to both type of failures: | $U_{3,A}(t) = u_A \cdot U_{B,du}(t),$
$U_{3,B}(t) = U_{A,du}(t) \cdot u_B$ |

The term $U_1(t)$ and the terms $U_{3,A}(t)$ and $U_{3,B}(t)$ can be easily evaluated from the results obtained in the single failure analysis. Then only the term $U_2(t) = U_{A,du}(t) \cdot U_{B,du}(t)$ is new. It is related to double dangerous undetected failures and is analysed below.

Therefore, over a given duration T this lead to:

- hazardous event probability: $P_a(\tau) \approx N_a(T) \approx [\bar{U}_1(T) + \bar{U}_2(T) + \bar{U}_{3,A}(T) + \bar{U}_{3,B}(T)] \lambda_d \cdot \tau = \bar{U}_{AB} \cdot \lambda_d \cdot \tau$
- average hazardous event frequency: $\bar{\Phi}_a(T) \approx \bar{U}_{AB} \cdot \lambda_d$.

7.4.2 Double dangerous undetected failures

7.4.2.1 Simultaneous tests

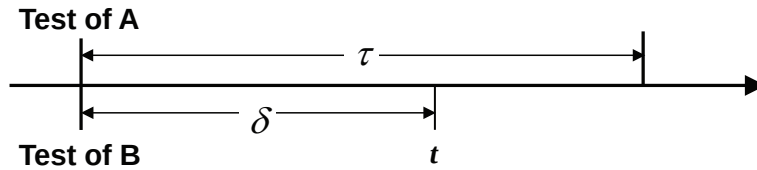


Figure 22 — Two components tested at the same time

The simplest case is presented on [Figure 22](#) where the two components are tested at the same time with the same test interval.

The average unavailability of this system is obtained as the sum of 4 terms:

$$\bar{U}_2(\tau) = \bar{U}_2^1(\tau) + \bar{U}_2^2(\tau) + \bar{U}_2^3(\tau) + \bar{U}_2^4(\tau)$$

The first term is related to the average unavailability due to dangerous undetected failures over the test interval: $\bar{U}_2^1(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du} \tau^2}{3}$

Therefore $\bar{U}_2^1 \oplus \frac{4}{3} \bar{U}_{A,du} \cdot \bar{U}_{B,du}$

INFO BOX: The simple multiplication of the individual average unavailabilities of A and B does not provide a conservative estimation of the average unavailability of the system AB. Nevertheless these kinds of calculations are often observed.

The coefficient 4/3 which is greater than 1, only appears when establishing the formula for the whole system. This is the manifestation of the “systemic dependency” introduced between the components A and B because the periodic tests are simultaneous: therefore the availabilities of A and B are correlated, i.e. good and bad at the same time.

The three other terms are related to the unavailability of the system during the repairs of A, B or A and B:

- 1) A and B under repair after the test: $\bar{U}_2^2(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\mu_{A,du} + \mu_{B,du}} \tau$
- 2) A under repair and B fails before the repair of A: $\bar{U}_2^3(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\lambda_{B,du} + \mu_{A,du}} \tau$ when considering $e^{-\mu_{A,du}\tau} \approx 1$ and $1/\mu_{A,du} \ll \tau$.
- 3) B under repair and A fails before the repair of B: $\bar{U}_2^4(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\lambda_{A,du} + \mu_{B,du}} \tau$ when considering $e^{-\mu_{B,du}\tau} \approx 1$ and $1/\mu_{B,du} \ll \tau$.

NOTE In $\bar{U}_2^2$, the implicit hypothesis is that A and B has their own repair teams. This is an optimistic hypothesis because the failures of A and B are discovered at the same time and that means that the repairs are done at the same time. Therefore if there is only a single repair team the repair policy should be modelled (e.g. in multiplying the MTTRes by 2).

7.4.2.2 Non-simultaneous tests

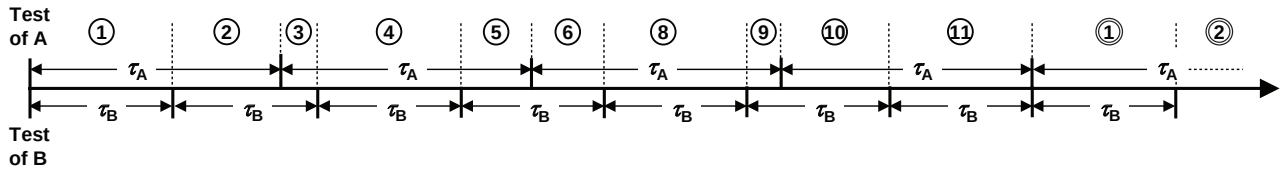


Figure 23 — Example of components not simultaneously tested

Figure 23 shows an example where the components have different tests intervals. Therefore they are not tested at the same time. On this example simultaneous tests occurs every 4 tests of A and 7 test of B. Then the renewal period (i.e. the duration after which the same pattern of intervals comes back) is $T_r = 5\tau_A = 7\tau_B$. Some time intervals start or finish with the test of only one component (intervals 2 to 10), some others start after the tests of the two components (intervals 1 and 4) and the interval 11 ends with the tests of the two components. This implies that the unavailability formulae be established for the 11 different intervals comprised within T_r . This can be done just by using formulae similar to those established above (cf. also [Annex I](#)). A difficulty arises when there is no renewal period or when it is larger than the period of interest. In these cases, it is obligatory to identify and consider all the intervals over the period of interest for the calculations.

The special case where the test intervals have the same duration (τ) is presented in [Figure 24](#):

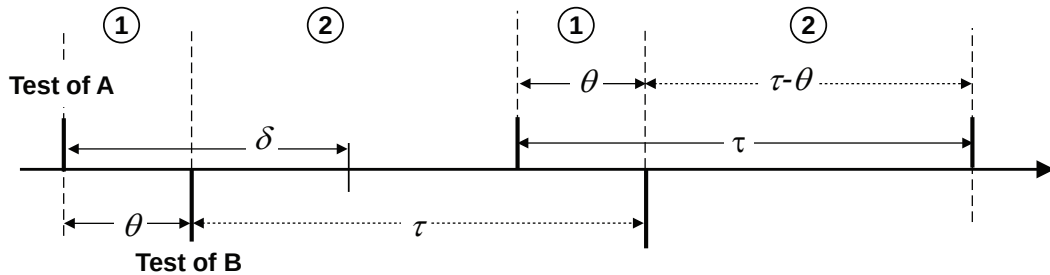


Figure 24 — Example of similar test intervals

All the time intervals begin after the test of one component and end when the other is tested. There is no renewal period but the test pattern comprises only the two types of intervals identified as 1 and 2. On the long range (i.e. after 2 or 3 test intervals) such a system converges toward asymptotic average unavailabilities in each of the intervals. Note that these asymptotic average unavailabilities should not be confused with steady-state unavailabilities because no steady-state exists in this example:

- average unavailability in interval 1: $\bar{U}_{AB,du}^{(1)} \approx \frac{\lambda_{A,du} \lambda_{B,du}}{\tau} \left(\frac{\theta^3}{3} + \frac{\theta^2(\tau - \theta)}{2} \right).$
- average unavailability in interval 2: $\bar{U}_{AB,du}^{(2)} \approx \frac{\lambda_{A,du} \lambda_{B,du}}{\tau} \left(\frac{\tau^3}{3} - \frac{\tau^2\theta}{2} - \frac{\theta^3}{3} + \frac{\theta^3}{2} \right).$
- average unavailability over interval 1+2: $\bar{U}_{AB,du} = \bar{U}_{AB,du}^{(1)} + \bar{U}_{AB,du}^{(2)} \approx \frac{\lambda_{A,du} \lambda_{B,du}}{6} (2\tau^2 - 3\tau\theta + 3\theta^2).$

Therefore the average unavailability is:

- minimum when $\theta = \tau/2$. This gives $\bar{U}_{AB,du} = \lambda_{A,du} \lambda_{B,du} \frac{5}{24} \tau^2.$
- maximum when $\theta = 0$. This gives $\bar{U}_{AB,du} = \lambda_{A,du} \lambda_{B,du} \frac{1}{3} \tau^2.$

INFO BOX: Therefore the average unavailability decreases with 37,5 % when going from simultaneous tests to perfectly staggered tests. Whether this “systemic” improvement is significant or not will depend on the application.

This analysis should be completed with the unavailabilities due to repairs but this similar to what has been already described here above.

7.4.3 Double dangerous detected failures

This is developed in [Annex I](#). The main results are the following:

- Equivalent failure rate: $\Lambda_{AB,dd} \approx \lambda_{A,dd} \frac{\lambda_{B,dd}}{\lambda_{B,dd} + \mu_{A,dd}} + \lambda_{B,dd} \frac{\lambda_{A,dd}}{\lambda_{A,dd} + \mu_{B,dd}}$.
- Average unavailability: $\bar{U}_{AB,dd} \approx \frac{\Lambda_{AB,dd}}{\Lambda_{AB,dd} + \mu_{A,dd} + \mu_{B,dd}}$.
- Hazardous event probability: $P_{a,dd}(T) \approx N_{A,dd}(T) \approx \bar{U}_{AB}(T) \cdot \lambda_d \cdot T$.
- Hazardous event frequency: $\bar{\Phi}_{a,du}(T) \approx \bar{U}_{AB}(T) \cdot \lambda_d$.

7.4.4 Spurious failures

The number of safe failures needed to lead to a spurious action depends on the design of the overall safety system from which the minimal cut set comes from:

- 1oo2 architecture: a single safe failure from of A or B is able to trigger a spurious failure then the spurious failure rate is simply: $\Lambda_{AB,st}^{1F} \approx \lambda_{A,sf} + \lambda_{B,sf}$ where “1F” means “one failure”.
- 2oo3, 3oo4, etc. architecture: a double safe failure (e.g. both A and B) is needed to trigger a spurious safety action then the spurious failure rate becomes: $\Lambda_{AB,st}^{2F} \approx \lambda_{A,sf} \frac{\lambda_{B,sf}}{\lambda_{B,sf} + \mu_{A,sf}} + \lambda_{B,sf} \frac{\lambda_{A,sf}}{\lambda_{A,sf} + \mu_{B,sf}}$ where “2F” means “two failures”.

When the λ_{sf} are low the probability of spurious failure over $[0, T]$ is: $P_{AB,st}(T) \approx \Lambda_{AB,st} T$

The single safe failures have already been analysed in [7.3.4](#). For the double safe failures the results are the following:

- Expected number of failure: $N_{st}^{2F}(T) \approx \frac{T}{MTBF_{sf}^{2F}} = \frac{T}{MTTF_{sf}^{2F} + MTTR_{sf}^{2F}}$.
- Average spurious failure frequency: $\bar{\Phi}_{st}^{2F} = \frac{1}{MTBF_{st}^{2F}}$.

Where $MTTR_{st}^{2F} \approx \frac{1}{\mu_{A,sf} + \mu_{B,sf}}$ and $MTTF_{st}^{2F} = \frac{1}{\Lambda_{AB,st}^{2F}}$.

When the repairs are fast then $MTBF_{st}^{2F} \approx MTTF_{st}^{2F} = \frac{1}{\Lambda_{AB,st}^{2F}}$ and $N_{st}^{2F}(T) \approx \Lambda_{AB,st}^{2F} T$ and $\bar{\Phi}_{st}^{2F} \approx \Lambda_{AB,st}^{2F}$.

7.4.5 Complete formulae for the system made of two components

The analysis leads to a great number of cases when the two components are not tested at the same time or when other parameters than the basic parameters described in 7.3 are used. It is therefore not possible to establish a complete catalogue of detailed formulae for all these cases.

The general philosophy developed in [7.4.1](#) can be used to develop the relevant formulae corresponding to any cases. No further development will be made in this Technical Report.

7.5 Triple failure analysis

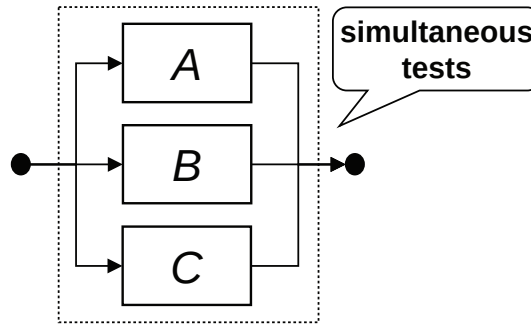


Figure 25 — Triple failure modelling

The analysis becomes very complicated for triple failures because the number of situations to handle increases more or less exponentially with the order of the minimal cut sets under consideration. Nevertheless this analysis can be performed by skilled reliability engineers applying the principles developed above for single and double failures.

Developing in this Technical Report the complete general analysis of triple failures would lead to a great number of pages covered with complicated formulae. This would not be very effective as their meanings would likely be not really understood by users not skilled in probabilistic calculations.

Therefore only the more simplistic situation is analysed hereafter: three components tested at the same time and with only the basic parameters defined in 7.3. None of the extra parameters (γ , ω , π , test staggering) will be considered hereafter.

The basic reliability parameters are the following:

- $\lambda_{A,du}$, $\lambda_{B,du}$, $\lambda_{C,du}$: dangerous undetected failure rates of components A, B and C (i.e. hidden failures detected by periodic tests);
- $\lambda_{A,dd}$, $\lambda_{B,dd}$, $\lambda_{C,dd}$: dangerous detected failure rates of components A, B and C (i.e. immediately revealed failures or failures detected by diagnostic tests);
- $\mu_{A,du}$, $\mu_{B,du}$, $\mu_{C,du}$: repair rates of dangerous undetected failure of components A, B and C;
- μ_{dd} : restoration rate of immediately revealed dangerous failure ($1/\mu_{dd} = \text{MTTRes}_{dd}$);
- τ : test interval of components A, B and C tested at the same time.

As said above, this clause deals only with the simplest case of minimal cut sets of order three and the reader should be aware that this constitutes a simplification which may be non conservative when the assumptions are not fulfilled. The aim of this clause is only to provide some guidance and to encourage the analysts to understand and develop their own formulae rather than to provide a catalogue of ready-to-use formulae. Further details can be found in [Annex I](#).

7.5.1 Dangerous undetected failures

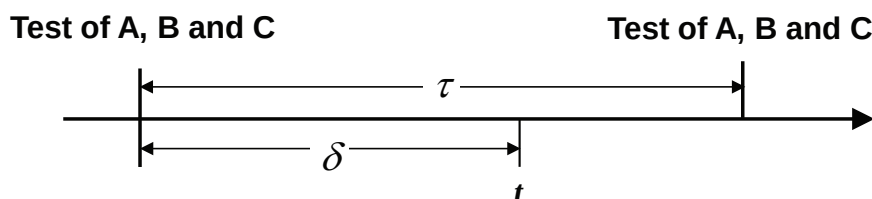


Figure 26 — Three components tested at the same time

In this simplest case, the [Annex I](#) establishes the following results for the triple dangerous undetected failures: $\bar{U}_{ABC,du}(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du} \cdot \lambda_{C,du} \tau^3}{4}$.

It can be observed that $\bar{U}_{ABC,du} \approx \frac{8}{4} \bar{U}_{A,du} \cdot \bar{U}_{B,du} \cdot \bar{U}_{C,du} = 2 \bar{U}_{A,du} \cdot \bar{U}_{B,du} \cdot \bar{U}_{C,du}$.

INFO BOX: Again this means that the simple multiplication of the individual average unavailabilities of A, B and C does not provide a conservative estimation of the average unavailability of the system ABC. Nevertheless these kinds of calculations are often observed.

The coefficient $8/4 = 2$ which is greater than 1, only appears when establishing the formula for the whole system. It is the manifestation of the “systemic dependency” introduced between the components A, B and C because the periodic tests are simultaneous: therefore the availabilities of A, B and C are correlated, i.e. good and bad at the same time.

The system can also be unavailable during the repairs. Seven cases should be considered: A, B and C under repair, A and B under repair and C failing, A and C under repair and B failing, B and C under repair and A failing, A under repair and B and C failing, B under repair and A and C failing, C under repair and A and B failing. The first of these cases is analysed in [Annex I](#).

7.5.2 Dangerous detected failures

The principle is similar as for double failures in [7.4.3](#) but there are seven OK states $\bar{A} \propto \bar{B} \propto \bar{C}, \bar{A} \propto \bar{B} \propto C, \bar{A} \propto B \propto \bar{C}, A \propto \bar{B} \propto \bar{C}, \bar{A} \propto B \propto C, A \propto \bar{B} \propto C, A \propto B \propto \bar{C}$, instead of 3. Therefore the number of formulae is bigger and the formulae are more complicated. They can be developed with the principles presented above and no further guidance is provided in this Technical Report.

NOTE A means “A failed” and $\bar{A}$ means “A not failed” and similarly for B and C.

7.5.3 Spurious failures

According to the logic of system from which comes the minimal cut set under study 1, 2 or 3 individual safe failures may be needed to trigger a spurious failure.

The cases of single and double safe failure has already been analysed above. For a spurious failure due to a triple safe failure then the same principle as those developed for double safe failure can be implemented.

NOTE In de-energize to trip design, all the failures able to lead to a spurious action should be considered (see [5.3](#)). This may imply other components than those considered within the dangerous failure analysis. The analysis should be performed accordingly.

7.6 Common cause failures

[Figure 27](#) shows how common cause failures can be modelled as a single failure. Therefore what has been described in [7.3](#) for single failures applies.

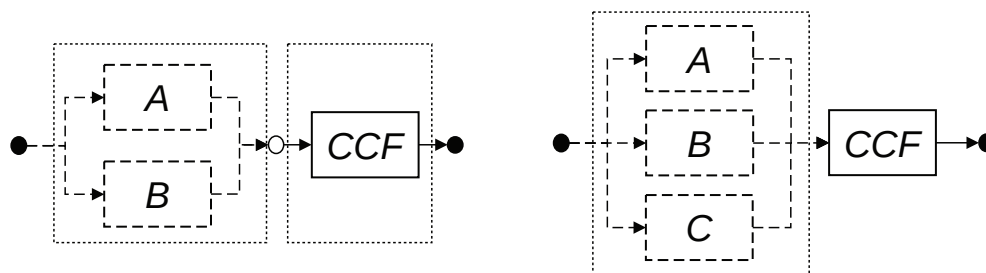


Figure 27 — Common cause failure modelling

A warning should be raised here because [Figure 27](#) models a CCF which is repaired independently of the components A and B. That is to say that the components are only made unavailable when the CCF occurs but not really failed (e.g. a loss of power).

If the occurrence of the CCF implies a separate repair of A and B (e.g. an over tension which break both A and B) this introduces a dependency between CCF, A and B and this cannot be handled by the simple approach described here above. Nevertheless, considering that the mean time to repair the CCF is equal to the sum of the MRTs of the components to be repaired provides a conservative approximation. For example, this can be implemented in order to use, the Boolean approach ([Clause 8](#) and [Annex K](#)) with the CCF models developed in [Annex G](#).

7.7 Example of implementation of analytical formulae: the PDS method

A lot of efforts have been made across industry to implement the analytical formulae approach according to IEC 61508[2]: the PDS method¹⁾ is one of them. It aims to reflect *actual* reliability performance by including the most important parameters that are assumed to influence SIS by implementing relatively simple formulae. The PDS method also recognize the importance of the safe failure (spurious trip) consideration which is within the scope of this Technical Report, and the method therefore provides an interesting example of implementation of a subset of the formulae developed in [Clause 7](#).

The method is described in the *method handbook*[13] primarily used to quantify the safety unavailability of SIS and the loss of production due to safe failure (spurious trip).

The PDS method is used in the oil and gas industries, especially in Norway, but could be applicable to other business sectors. It has been developed during the past 25 years in close co-operation with oil and gas companies as well as vendors and researchers of control and safety systems.

More specifically, the PDS method, in the framework of the assumptions presented in [Clause 7](#), includes simple calculation formulae addressing:

- safety unavailability formulae for SIS, both for low and high demand mode of operation;
- common cause failures (CCF) for different types of voted configurations;
- treatment of some non-hardware or non-random failure;
- downtime due to repair and periodic tests, in light of strategies for degraded operation.

For a more detailed description of the PDS method reference is made to [I.5](#) and to [13]. Regarding the treatment of CCFs in the PDS method, more details are found in [Annex G](#).

7.8 Conclusion about analytical formulae approach

Approximated formulae may be established and applied by skilled reliability engineers but they become more and more difficult to comprehend as the number of parameters of interest or the order of the minimal cut sets increases. The staggering of tests also introduces extra difficulties. In addition, the approximate formulae are generally used to provide average probabilistic values which cannot always be used straightforwardly as input for further conservative calculations.

The main conclusions drawn from this clause are the following:

- Analytical formulae are useful for simple safety system calculations where the degree of complexity and redundancy is limited and the users know what they are doing,
- All assumptions and limitations underlying the formulae are important to state.
- Simply multiplying average unavailability figures together may be non-conservative due to the underlying “systemic” dependencies.

1) PDS is the Norwegian acronym for “reliability of computer-based safety systems”.

- For more complex safety systems alternative calculation methods should be applied and the systemic approaches presented hereafter are more rigorous.

It is paramount that the users understand the formulae and their underlying assumptions and limitations. Since it is optimistic to think that all the users of analytical formulae are going to establish these formulae themselves, they should at least carefully study this Technical Report in order to get sound understanding about their actual possibilities and shortcomings.

[Clause 7](#) has been focused on low demand mode of operation safety systems. The same approach can be applied to develop analytical formulae for high demand or continuous mode of operations but no further guidance is provided in this Technical Report.

8 Boolean and sequential approaches

8.1 Introduction

A safety system has only two states: either it is available to fulfil its reliability function or it is unavailable to do that. Therefore the Boolean approach seems a good candidate to be used as it rightly deals with two-state systems. Several widely used models - reliability block diagrams (RBD)[\[4\]](#) or event tree analysis[\[8\]](#) (ETA) already mentioned in this guide but, with some limitations, also the cause consequence diagrams[\[10\]](#), LOPA[\[9\]](#) or fault tree analysis[\[5\]](#) (FTA) which will be discussed below - belongs to the Boolean approach.

All these models aim at representing the links between the states of the system and the states of its components through the use of an underlying logical function. This is very powerful to model the system architectures and large size systems can be modelled in this way. As the Boolean models are time independent (i.e. static models), difficulties arise when it is necessary to consider time dependencies (e.g. single repair team) or to model the system behaviours as a function of the time (e.g. periodic tests, operational activity, etc.). Nevertheless it is possible in many cases to perform time-dependent calculations provided some strong hypotheses are met.

The main interest of Boolean models is that they provide graphical supports which are helpful for the analysts building the models. To some extent, modelling and calculations can be done separately: the analyst focus on the development of the model (e.g. the utmost importance part of the safety study) and leave an appropriate software package to perform the calculations. Powerful algorithms (e.g. Binary Decision Diagrams) are now available to process very big logical function.

Nevertheless the periodic tests (i.e. safety systems operating on demand mode) are not easy to handle properly because they introduce time issues. This problem can be solved by using the approach discussed below which mixes Boolean and Markovian models in order to model large Markov processes (fault tree driven Markov processes). Some software packages implement such a mixed approach but a minimum knowledge of the underlying mathematics is needed to perform relevant calculations and this is often forgotten: using the best software packages like a black box is likely to provide questionable results.

All calculations are based on the hypothesis that the individual components are independent. This is never completely true, and then this approach is limited to safety systems comprising reasonably independent components.

8.2 Reliability block diagrams (RBD)

The RBD approach is described in the IEC 61078[\[4\]](#) international standard and has been already used above in this Technical Report (see [Figure 15](#) to [Figure 19](#)).

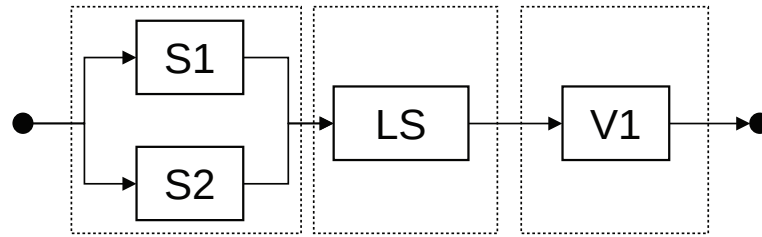


Figure 28 — Example of simple RBD

This is a good tool to model the architecture of safety systems of which it generally remains close. The meaning of a RBD is clear and easily understandable by the engineers from various specialities working for the same project. For example it is clear that [Figure 28](#) represents a safety system made of two redundant sensors, one logic solver and one valve as terminal element. This facilitates the communication between engineers and this is why RBDs are widely known and used. Nevertheless it is more a tool to represent the system architecture than a mean to analyse them.

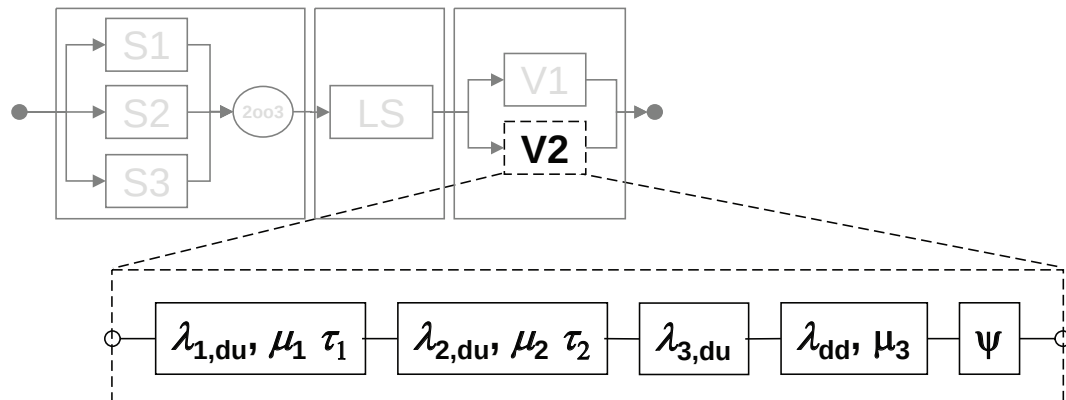


Figure 29 — Example of decomposition at the failure level with a RBD

Reliability block diagrams can be used to model several failure modes for a single component. For example, the [Figure 29](#) shows a component which has been split (from left to right) between:

- two dangerous undetected failure modes tested with different tests intervals;
- one dangerous failure which is not covered by the tests;
- one dangerous detected failure;
- one failure due to the demand of the safety action.

Similar decomposition between DD and DU failures can be observed in the fault tree in [Figure 30](#).

The main concept introduced by the RBDs is this of “minimal cut sets” which has been already introduced in [6.2.2](#). From a mathematical point of view RBDs and fault trees have exactly the same properties but the FTA is a real a method of analysis. Therefore the calculations investigated in the following subclause are devoted to fault tree analysis (FTA).

8.3 Fault Tree Analysis (FTA)

FTA is described into IEC 61025[5]. This is a very important tool because this is practically the only “deductive” (i.e. effect to cause reasoning) method which is used in reliability engineering. It consists in building step by step the logical links between the individual failures and the system failure. A top-down analysis is performed from the top event (unwanted, undesirable or root event) to the individual components failures.

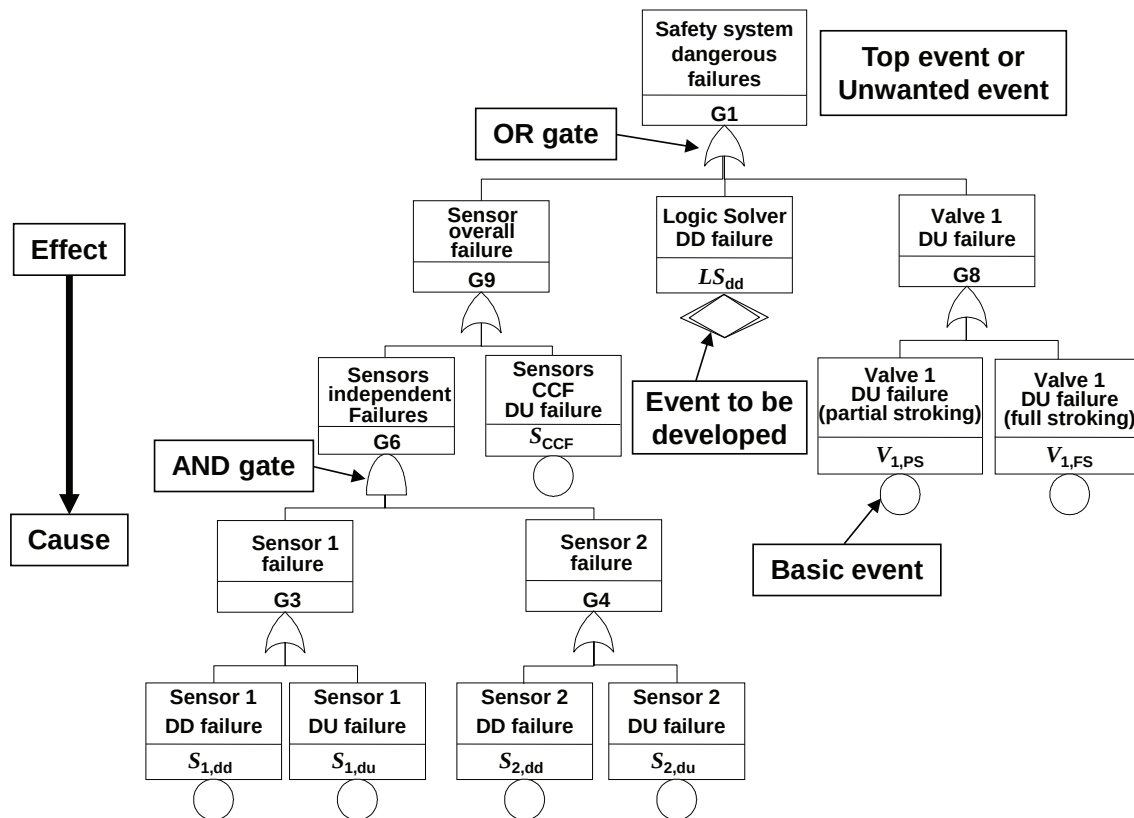


Figure 30 — Fault tree equivalent to the RBD on Figure 28

The fault tree in Figure 30 represents the same system as this modelled by the RBD in Figure 28 but it does not look like the same. It has been built in the following step:

- 1) Selection of the “unwanted” event (i.e. the dangerous failures of the safety system).
- 2) Analysis of the direct “causes” of this “effect” (e.g. Sensor overall failure, Valve DU failure).
- 3) Each identified “causes” is then considered as an “effect” to be analysed.
- 4) Analysis of the direct “causes” of the “effects” identified at step 3.
- 5) Repeat steps 3 and 4 until basics (or elementary) events have been found (e.g. when failure modes as defined in ISO 14224[15], Annex B have been identified).

For example, the failure of the group of sensors is due to the “independent failures of the individual sensors OR a common cause failure. This is modelled by using a logical OR gate. The independent failure of the sensors is caused by the failure of both sensors. This is modelled by using a logical AND gate. The independent failure of a given sensor is due to a dangerous detected (DD) failure or a dangerous undetected (DU) failure. This is modelled by using a logical OR gate.

The process is stopped when reliability data or event data are easily available: these “basic” events are represented by circles). When more investigation is needed this is indicated by using a double diamond. Other standardized graphical symbols are available but are not needed for the purpose of this Technical Report.

The fault tree presented in Figure 30 is a systemic model. It is more detailed than Figure 28 because some failures have been split between dangerous undetected failures and dangerous detected failures and that one common cause failure has been identified. In the same way, the two kind of periodic tests performed on the valve have been identified: the partial stroking which detects if the valve is stuck, whereas the full stroking detects if the valve is stuck or does not close tightly.

This fault tree comprises more detailed minimal cut sets than the RBD in [Figure 28](#): $\{S_1, S_2\}$ has been split between 4 minimal cut sets and $\{V_1\}$ has been split between 2 minimal cut sets:

$$\{S_1, S_2\} \equiv \{S_{1,dd}, S_{2,dd}\}, \{S_{1,dd}, S_{2,du}\}, \{S_{1,du}, S_{2,dd}\}, \{S_{1,du}, S_{2,du}\}$$

$$\{V_1\} \equiv \{V_{1,ps}\}, \{V_{1,fs}\}$$

This level of detail should be difficult to represent with a RBD when it is very easy with a fault tree.

The above notation is not simple to manipulate then a more practical notation of minimal cut sets is used in the remaining parts of this Technical Report:

- Conjunction (logical “AND”): $\{A, B\} \equiv A \cap B \equiv A \bullet B$
- Disjunction (logical “OR”): $\{\{C\}, \{D\}\} \equiv C \cup D \equiv C + D$

8.4 Sequence modelling: cause consequence diagrams, event tree analysis, LOPA

The reliability block diagrams[10] or the fault trees[5] described above are focused on individual safety systems but in reality, single safety systems rarely work alone and, when a demand occurs (the so-called “initiating event”) several safety layers run in sequence to prevent the occurrence of the related hazardous events. Those multiple safety systems are further discussed in [Annex F](#).

The first attempt to deal with such sequences of operations has been published in 1971 by the Danish Atomic Energy Commission[10]: this is the so-called “Cause-Consequence Diagram” model. Few years after, the well known “event tree”[8] approach appears in the famous “Rasmussen report” (Wash 1400) where it is implemented to evaluate the probability of accident of nuclear power plants. At the present time the last version of this kind of approach is the very fashionable “Layer of Protection Analysis” (LOPA)[9].

All these models may be represented by fault trees and this is why they are associated with the Boolean approach. Nevertheless, and contrarily to the Markovian or Petri net approaches, the mathematics behind the Boolean approach fundamentally deal with static behaviours and are not really adapted for sequence calculations. Therefore it is essential to use the available approximations based on the Boolean mathematics cleverly and cautiously with a full knowledge of the limitations and of the side effects on the results.

Those models are described in [Annex J](#) as well as their links with the fault tree modelling.

8.5 Calculations with Boolean models

8.5.1 Unavailability calculations

INFO BOX: FT calculations based on average unavailabilities of periodically tested components (e.g. with PFD_{avg} of the form $\lambda\tau/2$) leads to non conservative results. The non conservativeness increases when the redundancy increases. That means, for example, that the safety system is likely to be qualified for a higher SIL than it is really. Some available software packages perform these kinds of calculations. This is a part of the duty of the analyst to be cautious and to verify that software he/she uses implement sound algorithms.

Basically, a logical formula is a static model which does not contain any concept of evolution along the time. Nevertheless the temporal dimension can be introduced when the events evolve independently from each others: the probability of the logical formula at time t_i can be obtained from the probability of each event at time t_i . It is fundamental that the events are independent all the time (i.e. always independent when the time is elapsing) and therefore only unavailability calculations can be handled in this way.

When the components unavailabilities reach asymptotic values, the unavailability of the top event reaches also an asymptotic value which is also the long term average unavailability. This is illustrated in [Figure 31](#).

This is the typical behaviour of systems with only failures which are completely and quickly repairable, i.e. system of which the failures are quickly detected and repaired (like immediately revealed failures). When these hypotheses are met, the asymptotic values are reached rather quickly after a duration equal 2 of 3 times time the larger MTTRes of the components. In this case only one calculation of $U_S(t)$ is needed provided that t is large enough to ensure that all the asymptotic values have been reached. This approximation is conservative.

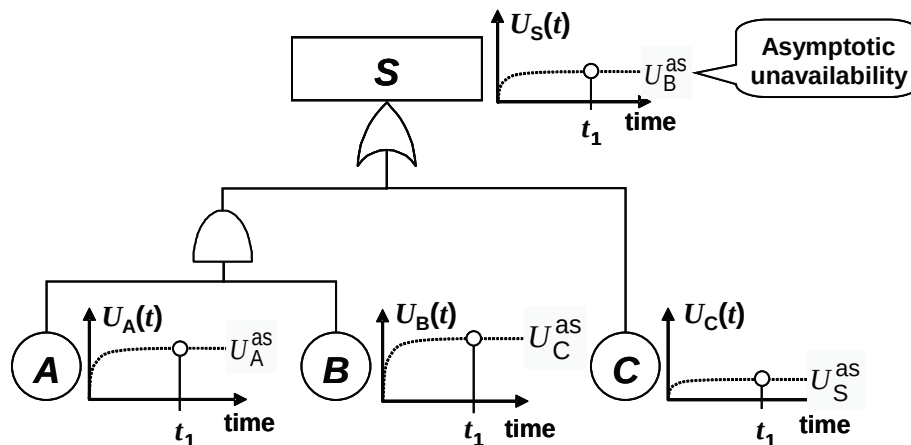


Figure 31 — Average unavailability as asymptotic value of the instantaneous unavailability

INFO BOX: This is, in principle, the case of the dangerous detected failures as well as immediately revealed safe failures. This is of the duty of the analyst to verify that the hypotheses related to this particular case are met and that the calculations can be performed in this way.

In the general case, calculating the average unavailability $\bar{U}_S(T)$ over the period $[0, T]$ cannot be done directly because combination of the average unavailabilities of the components through the logical function does not give the average probability of the top event of the fault tree. This implies that the instantaneous unavailability $U_S(t)$ be established first. This can be done by using the principle illustrated on [Figure 32](#): for a given instant t_1 the unavailabilities $U_A(t_1)$, $U_B(t_1)$ and $U_C(t_1)$ of components A, B and C are calculated independently from each others. Then they are combined through the fault tree according to underlying logical function to obtain $U_S(t_1)$. Doing that for a relevant number of time values over a period of interest gives the whole curve $U_S(t)$ over this period of time.

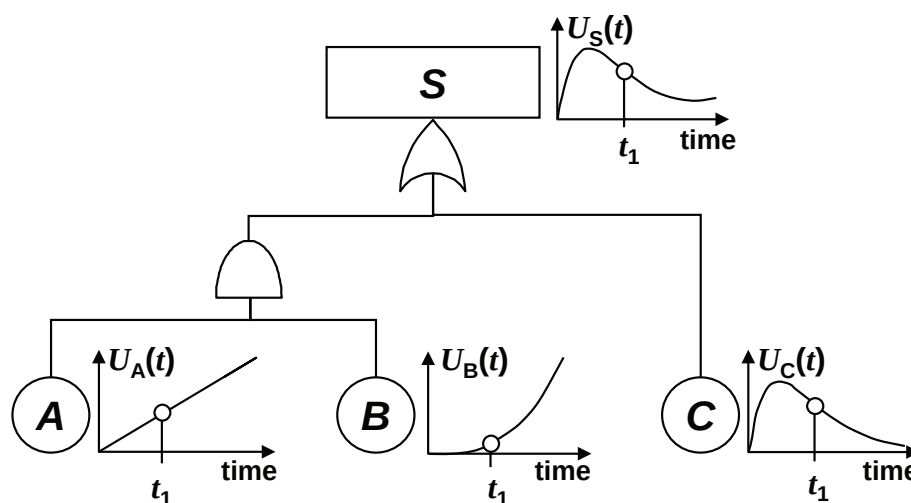


Figure 32 — Unavailability calculation

Then $\bar{U}_S(T)$ can be obtained by calculating the average of $U_S(t)$ as shown in [Figure 33](#). This is the relevant method to use when the periodic tests performed to detect and repair the dangerous undetected failures are modelled.

Such an example is presented in [Figure 34](#) where a fault tree driven Markov model is implemented. This is a “mixed” approach where the logic is described by the fault tree and the component behaviours are modelled by Markovian multiphase models (see [Clause 9](#) and [Annex L](#)). This is valid when the various input events in the fault tree are independent or subject to weak dependencies (e.g. single repair team).

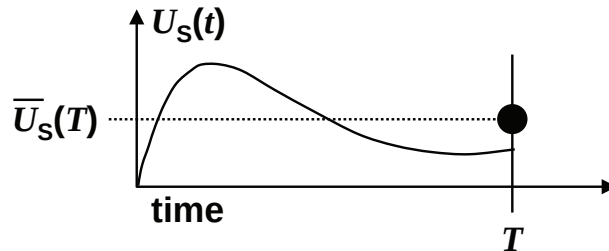


Figure 33 — Average unavailability as average of the instantaneous unavailability curve

In this example, the unavailabilities of A and B are typical saw-tooth curve related to periodically tested failures when the unavailability of C continuously increases because it is not repairable ([Figure 34](#)).

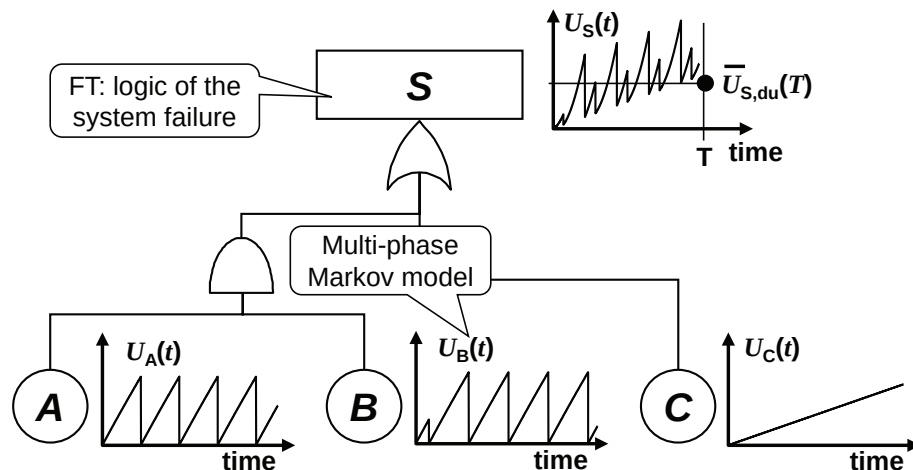


Figure 34 — FT driven Markov process modelling a safety system with periodically tested components

Even with this simple example, the result is a quite complex saw-tooth curve. Its average is represented by the black bullet on the top right hand side of the figure. Achieved in this way the fault tree calculations provide relevant results where the systemic dependencies due to the correlations introduced by the periodic tests (cf. [7.4.2](#) and [7.4.3](#)) are taken into account properly.

It should be noted that for this example the average unavailability (i.e. the PFD_{avg} of the modelled safety system) is not a good indicator as it is continuously increasing.

Further developments are provided in [Annex K](#).

8.5.2 Unreliability and frequency calculations

Evaluating the reliability and the failure frequency with Boolean models conceals a subtle difficulty that analysts are often not aware about. It is illustrated in [Figure 35](#) where A and B are repairable components. The whole system fails when both A and B are failed at the same time: that is to say, when B fails when A is under restoration (case illustrated in [Figure 35](#)) or when A fails when B is under restoration.

The reliability $R_S(t)$ of the system S is the probability that it does not fail before t . Therefore, from a mathematical point of view, the restoration of the system after a failure is excluded from the calculations. This implies that the failures of A and B can only be repaired as long as the whole system has not failed. When this occurs, the repairs of A or B are excluded of the reliability calculation because the repair of one of them would restore the whole system. This is another kind of systemic dependency which should be considered for the calculations.

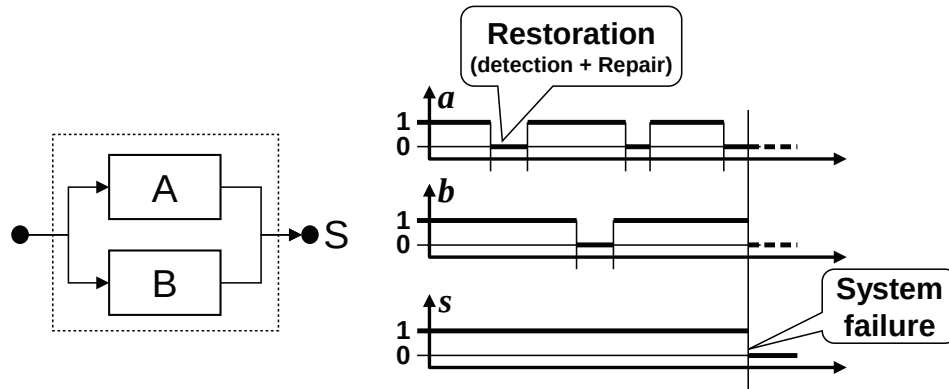


Figure 35 — Links between component and overall system failures

The calculation of the failure frequency of the safety system (the so-called unconditional failure intensity) and the hazardous event frequency are also not very easy to perform. They involve the calculation of Birnbaum importance factors which, themselves, are obtained from conditional probabilities. This is almost impossible to do this by hand but very powerful algorithms are now available to achieve such calculation by using the relevant software packages.

Further developments about these topics are provided in [Annex K](#).

8.6 Conclusion about the Boolean approach

The Boolean approach is a generic analytical approach which encompasses several approaches (reliability block diagrams, fault tree, cause consequence diagrams, event trees and LOPA) which are relatively well known by most reliability engineers. This is the approach to investigate first when the limits of the analytical formulae are reached. When selecting and using software tools, one should understand the mathematical approach and its appropriateness for the application.

The Boolean approach primarily provides the unavailability $U(t)$ of the modelled system from the input unavailabilities $U_i(t)$ of its components. Therefore the adaptation to SIL related calculations for periodically tested systems implies the following developments:

- establish formulae for $U_i(t)$ inputs;

NOTE The $U_i(t)$ formulae can be obtained by using the Markovian approach described in [Clause 9](#) and [Annex L](#). This approach mixing Boolean and Markov models is implemented, for example, in the fault tree driven Markov processes (see [Figure 34](#)).

- calculation of $\bar{U}(T)$ over the period of interest $[0, T]$ in order to obtain the PFDavg (safety systems working in low demand mode of operation);
- calculation of the failure frequency $w(t)$ and of its average $\bar{w}(T)$ over the period of interest $[0, T]$ in order to obtain the PFH (safety systems working in high demand or continuous mode of operation).

The calculations are relevant only when the components are independent but, in fact, it works quite well when the dependencies are weak. It is no longer usable when the dependencies increase and other approaches like those described in the following clauses should be used.

NOTE The dependencies are weak for example for the safety systems installed on the topsides of offshore platforms where each failure is individually and quickly repaired after detection. The dependencies are strong for example for the safety systems installed on a subsea facility where each failure is not quickly repaired after detection nor necessarily individually repaired.

The Boolean approach is rather easy to implement but the underlying mathematics which are not really simple should be understood so that fundamental assumptions behind the method are not violated. Except in very simple models calculations are not possible by hand but powerful algorithms have been developed (e.g. binary decision diagram based algorithms^[51]) which are able to process large models with hundreds of components. They are so much powerful that this makes it possible to handle the uncertainty calculations (see [Clause 12](#)). Thanks to its graphical features the analysts can focus on the building of accurate graphical models and forget the details of the mathematics - to the extent that the basic mathematical principles are deeply understood to avoid misusing and inadequate calculations.

9 Markovian approach

9.1 Introduction and principles

The Markovian approach has been developed at the beginning of the past century and is a popular approach for probabilistic calculations in numerous scientific fields. Its use in reliability engineering is described in IEC 61165^[6]. This approach has already been mentioned many times in this Technical Report as illustration in [Clause 3](#) and [Clause 5](#) and possible fault tree input in [Clause 8](#). It is also the underlying model behind the formulae developed in [Clause 7](#). Further explanation can be found in reference^[24]

It is a state-transition model and its principle is to identify the various states of the system under interest and to analyse how the system jumps from a state to another state. This allows building graphical models - called Markov graphs - which represent the dynamic behaviour of the system which is modelled.

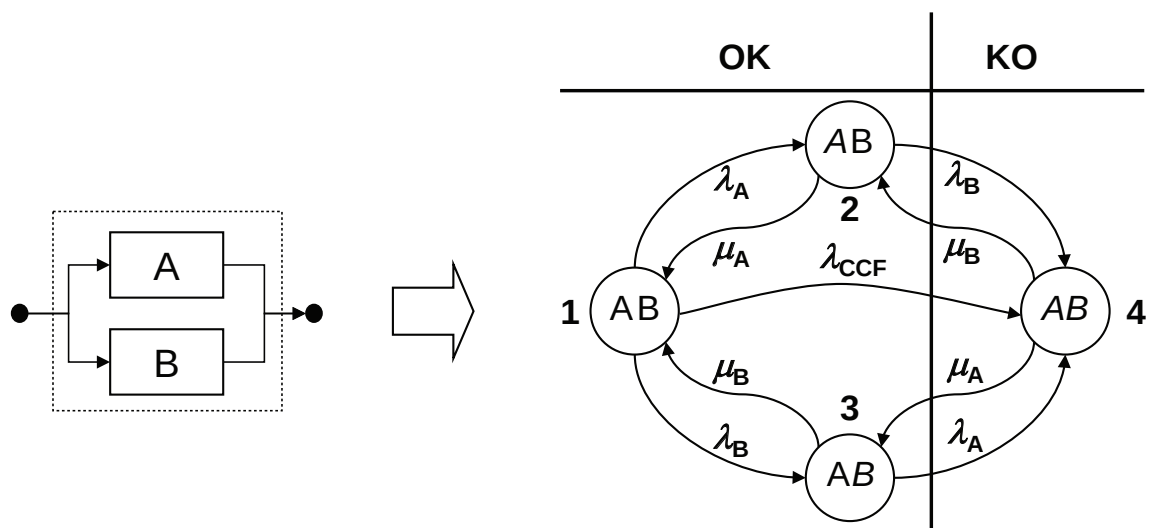


Figure 36 — Example of simple Markov model

[Figure 36](#) gives an example of simple Markov graph. It models the small redundant system on the right hand side made of two components A and B which have repairable failures which are immediately detected. Then, $MFDT \approx 0$ and $MTTRes \approx MRT$.

The states are represented by circles and the jumps (transitions) by arrows. For example the transition from state 1 to state 2 models the failure of component A and the transition from state 3 to state 1 models the restoration (i.e. detection + repair) of component B.

The transition between states 1 and 4 models a common cause failure between A and B. This is the kind of common cause failure which breaks the components A and B which afterwards are individually restored.

When the transitions between states are characterized by transition rates like the constant failure (λ_A, λ_B) or restoration (μ_A, μ_B) rates then the underlying stochastic process is a “homogeneous Markov process” which allows analytical calculations. When the transition rates are time dependent the underlying stochastic process is a “semi-Markov process” which is rather difficult to handle analytically. Fortunately this can be modelled by using Petri nets (see [Clause 10](#)) and calculated by Monte Carlo simulation (see [Clause 11](#)). Therefore the remaining part of the Markov clause deals only with homogeneous Markov processes.

Mathematically speaking a homogeneous Markov process is equivalent to a conventional set of linear differential equations (cf. [Annex L](#)). Powerful algorithms exist to solve such equations in order to calculate the probability $P_i(t)$ of each state at time t and also the cumulated time spend in each state $\theta_i(T)$ over a given period $[0, T]$.

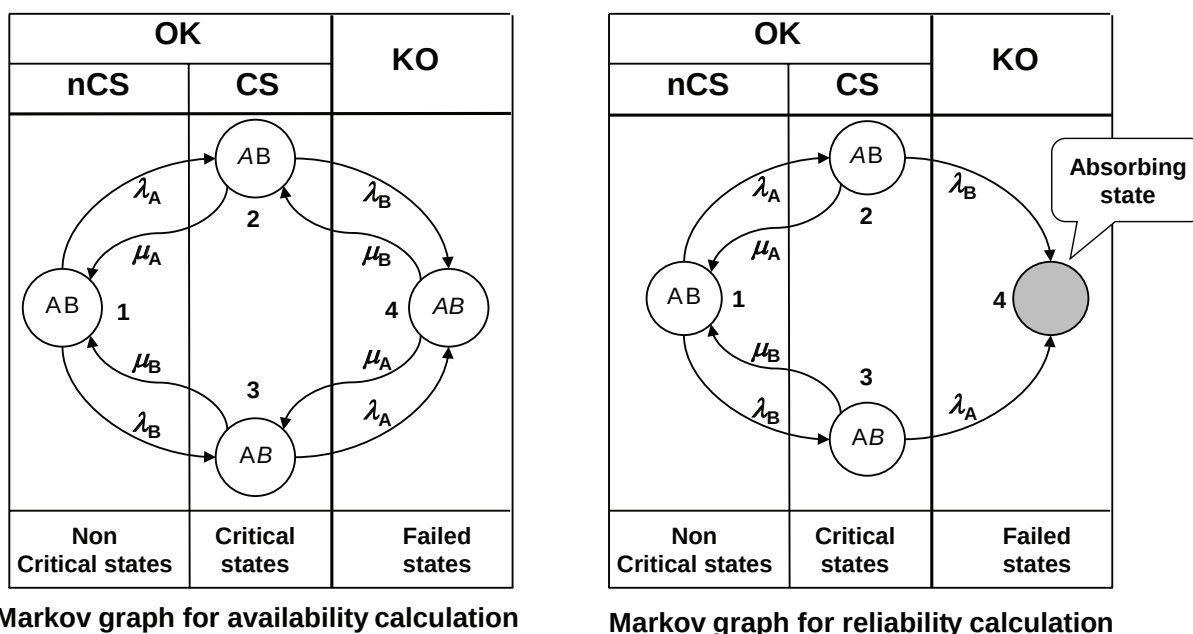


Figure 37 — Type of Markov graphs and the three classes of states

Two classes OK (available) and KO (unavailable) have been identified in [Figure 36](#). In [Figure 37](#) the class OK has been split into two sub-classes CS (critical states) and nCS (non critical states) for the purpose of the following explanations. Note that the CCF transition has been ignored to simplify the explanations.

This [Figure 37](#) illustrates the two types of Markov graphs which are of utmost importance for unavailability and reliability calculations:

- On the left hand side, it is possible to come back to the class OK after being entered in the class KO (i.e. the overall system failures are repaired). This is an availability graph giving:
 - system unavailability at time t : $U_S(t) = P_4(t)$;
 - system availability at time t : $A_S(t) = P_1(t) + P_2(t) + P_3(t)$;

- average unavailability over $[0, T]$: $\bar{U}_S(T) = \Theta_4(T)/T$.
- On the right hand side it is no longer possible to come back to the class OK after being entered in the class KO (i.e. the overall system failures are not repaired). The state 4 becomes “absorbing” and this is typical of a reliability graph. This gives:
 - system unreliability at time t : $F_S(t) = P_4(t)$;
 - system reliability at time t : $R_S(t) = P_1(t) + P_2(t) + P_3(t)$.

The critical class CS gathers the states which are just at the border between OK and KO. They are important because they are in direct relationship with the main reliability parameters of the system under study.

From the reliability graph it is obtained:

- Failure density: $f_S(t) = \lambda_B \cdot P_2(t) + \lambda_A \cdot P_3(t)$.
- Equivalent failure rate: $\Lambda_{eq}(t) = f_S(t)/R_S(t)$.

From the availability graph it is obtained:

- Failure frequency (unconditional failure intensity): $w_S(t) = \lambda_B \cdot P_2(t) + \lambda_A \cdot P_3(t)$.
- Average failure frequency: $\Phi_S(T) = \bar{w}_S(T)$.
- Vesely failure rate (conditional failure intensity): $\Lambda_{V,S}(t) = w_S(t) / A_S(t)$.

All the above formulae come directly from the definitions established in Clause 3. It is importance to notice that the formulae of $R_S(t)$, $F_S(t)$, $f_S(t)$, $\Lambda_{eq}(t)$ are identical to the formulae for $A_S(t)$, $U_S(t)$, $w_S(t)$, $\Lambda_{V,S}(t)$. The differences in their values only come from the properties of the Markov graph (with or without absorbing state).

When the failures are quickly detected and repaired $A_S(t)$, $U_S(t)$, $w_S(t)$ and $\Lambda_{V,S}(t)$ reaches asymptotic values. In this case, the Vesely failure rate provides a good approximation for the equivalent failure rate: $\Lambda_{eq}^{as} \approx \Lambda_V^{as}$.

Other reliability parameters can be obtained from the Markov graphs:

- Mean time to fail: $MTTF_S = MTTF_S(\infty) \approx \frac{\bar{A}_S(\infty)}{\bar{w}_S(\infty)}$.
- Mean time to restore: $MTTRes_S = MTTRes_S(\infty) \approx \frac{\bar{U}_S(\infty)}{\bar{w}_S(\infty)}$.
- Mean time between failures: $MTBF_S = MTTF_S + MTTRes_S$.

The faster the repair and the faster is the convergence. Practically, the asymptotic values are reached after two or three time the longest MTTRes of the components comprised within the system.

The number of states increases exponentially when the number of components of the modelled system increases. This is main limitation of the Markovian approach. Using this approach on industrial size systems imply approximations difficult to handle therefore in this Technical Report only the use of this approach for components (to provide fault tree input) or small safety systems is analysed. For larger safety systems the analysts are encouraged to use the mixed fault tree × Markov approach described previously when this is possible or to jump to the Petri nets and the Monte Carlo simulation described below.

Further developments are given in [Annex L](#).

9.2 Multiphase Markov models

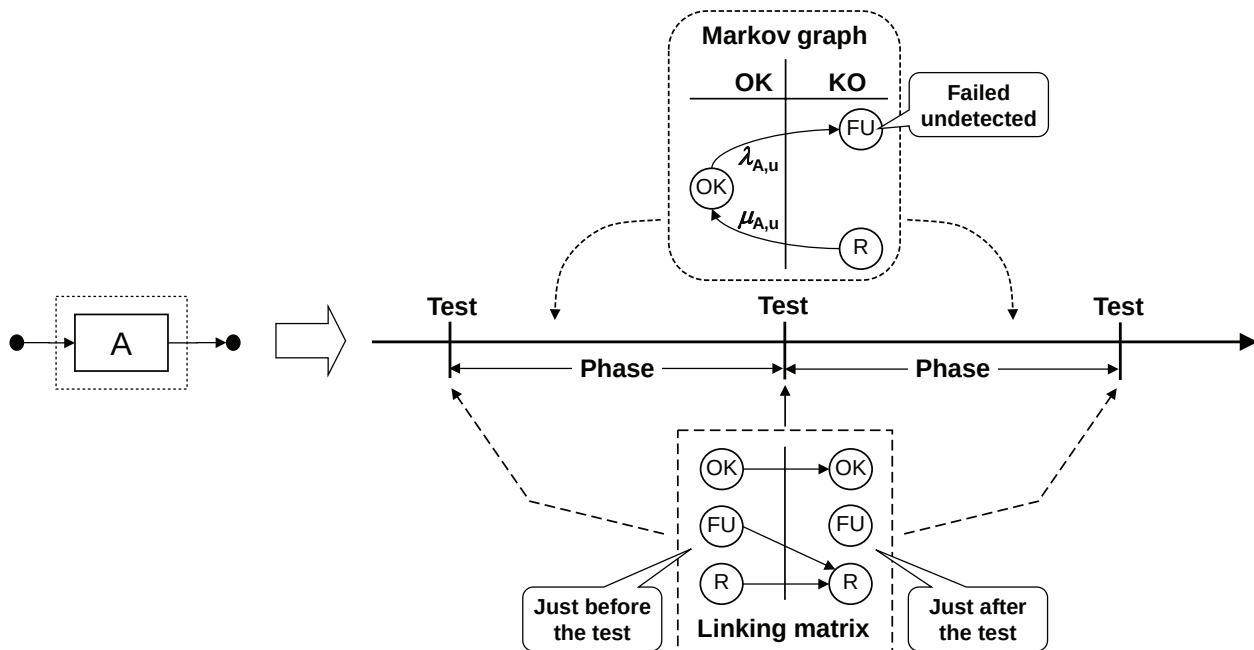


Figure 38 — Example of simple multiphase Markovian model

The Markov graphs developed in 9.1 only deal with random events (failures and restoration of immediately revealed failures). This is not valid to model the deterministic test intervals of periodically tested failures. In this case it is necessary to switch to the “multiphase” Markovian approach[24].

A simple example is provided in Figure 38 to illustrate the principles. The time is divided in successive test intervals. Each of them constitutes one phase which begins just after a test has started and finishes just before the next test is going to start. During the test interval the system behaves as a three states Markovian process:

- OK: the system is perfect;
- R: the system is under repair;
- FU: the system is failed but the failure is not detect yet.

When a test is performed the failures are detected and repairs can start. This modelled by the “linking” matrix which is applied to the states of the system:

- If the state is OK before the test it remains OK after the test;
- If the state is FU before the test it becomes R after the test (repair starts);
- If the state is R before the test it remains R after the test (repair goes on).

This model explicitly splits the *time to restore* between the *fault detection time* and the *repairing time*. In particular $\mu_{A,u} = 1/\text{MRT}_u$.

This models works even if the time to repair is not negligible compared to the test interval duration and if the probability of failure is high during the test interval. It is valid both for undetected dangerous and safe failures. This is why the notation *FU* (failed and undetected), $\lambda_{A,u}$ and $\mu_{A,u}$, is used.

From the graph in Figure 38 it is found that:

- system unavailability: $U_S(t) = P_{FU}(t) + P_R(t)$.

- system average unavailability: $\bar{U}_s(T) = \frac{\Theta_{FU}(T) + \Theta_R(T)}{T}$.

These results are illustrated in [Figure 39](#) where several curves have been plotted. The curve 1 is the classical saw-tooth curve (low failure rate, quick repair after detection). The curves 2, 3 and 4 illustrate what happens when the failure rate increases and the repair rate decreases. Finally, the curve 5 illustrates the case where the test interval decreases. At the limit, if the test interval goes to 0, the curve 5 would become similar to this of an immediately revealed failure. Therefore this model encompasses both dangerous undetected failures and detected dangerous failures.

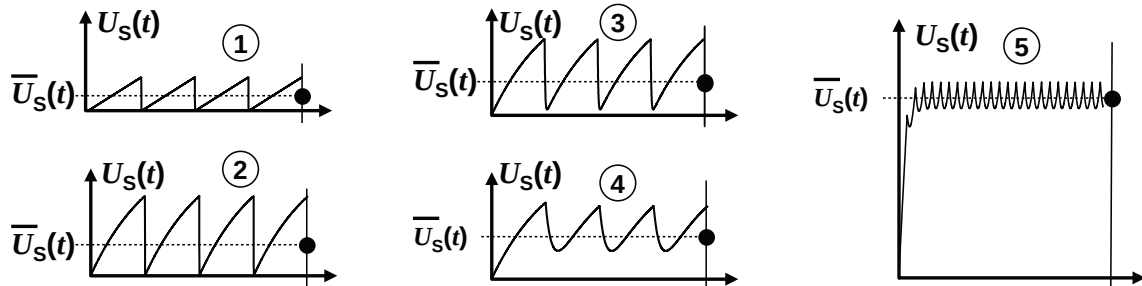


Figure 39 — Various saw-tooth curves derived from multiphase Markov model

Such models may be used as input in fault tree driven Markov processes as mentioned in [Clause 8](#). This is the main use of the Markovian approach in this Technical Report.

Nevertheless further developments are given in [Annex L](#). They are useful to improve the above multiphase Markov models and to evaluate the reliability, failure frequency and spurious failure frequency of simple safety system as well as the hazardous event frequency.

9.3 Conclusion about the Markovian approach

The Markovian approach gives a clear illustration of core concepts defined in [Clause 3](#) (availability, reliability, frequency, equivalent failure rate, etc.) and this provides a good help for the users to acquire a sound understanding of these concepts in order that they are properly used when dealing with the modelling and calculation of safety systems.

The number of states increases dramatically with the number of components and beyond 3 or 4 components this approach becomes unmanageable if the models are built by hand. Therefore the use of Markov graphs is particularly effective for small complex systems as this has been shown in this clause. Nevertheless methods and tools are available to generate automatically the Markov graphs of big systems and to simplify and/or approximate those Markov graphs in order to push the limits away. These approximations need skilled reliability engineers to be properly implemented and are beyond the scope of this Technical Report.

From safety system point of view, this is in conjunction with the fault tree approach that the use of Markov modelling has proven to be very efficient when the components are reasonably independent: multiphase Markov models can be used as fault tree input to model periodically tested components. This mixed approach, named FT driven Markov process, is described in [8.5.1](#).

More details about Markovian approach can be found in [Annex L](#) and References [\[6\]](#) and [\[24\]](#).

10 Petri net approach

10.1 Basic principle

The Petri net approach has been standardized in the IEC 62551[\[7\]](#)

This part is described in more detail in [Annex M](#) and further detail can be found in Reference [\[25\]](#)

The basic graphical elements of Petri nets (PN) according to the IEC 62551^[7] standard are represented in [Figure 40](#).

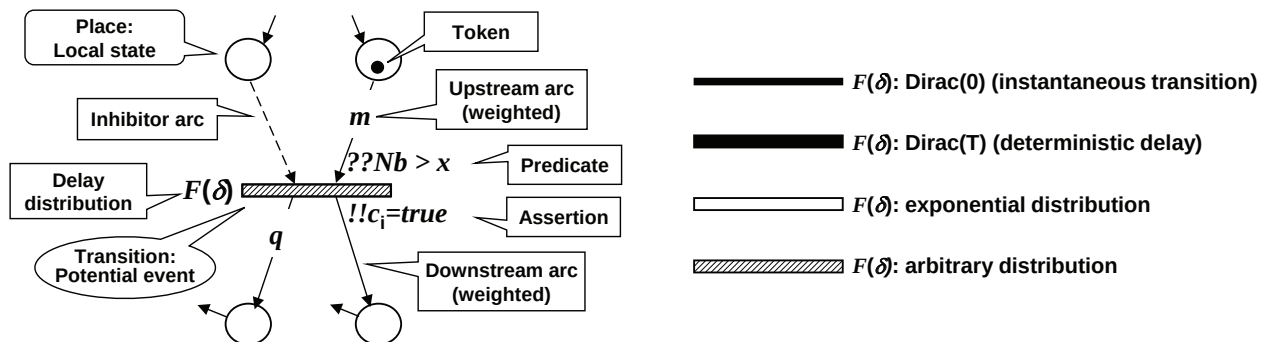


Figure 40 — Basic graphical elements for Petri nets

A Petri net is split within two different kinds of graphical elements:

- static elements (place, transitions, arcs) which are simple drawings;
- dynamic elements (token, predicates and assertions) which define the current state of the modelled system and change accordingly.

The transitions represent the events which can occur. When random delays are introduced as $F(\delta)$ in [Figure 40](#), the Petri net becomes “stochastic”. This allows implementing a Monte Carlo simulation (cf. [Clause 11](#)) to perform the probabilistic calculation.

The principle is the following:

- validate a transition according to the rules described in [Annex M](#) (i.e. determine if the attached event can occur);
- use random numbers to evaluate when it is going to be fired (i.e. determine when the attached event will occur);
- fire the transition when the delay is elapsed according to the rules described in [Annex M](#) (i.e. trigger the attached event).

Proceeding in this way allows modelling the dynamic behaviour of large and complex system. This may be used for production assurance (cf. ISO 20815^[16]) as well as safety systems calculations as this is done in this Technical Report.

10.2 RBD driven Petri net modelling

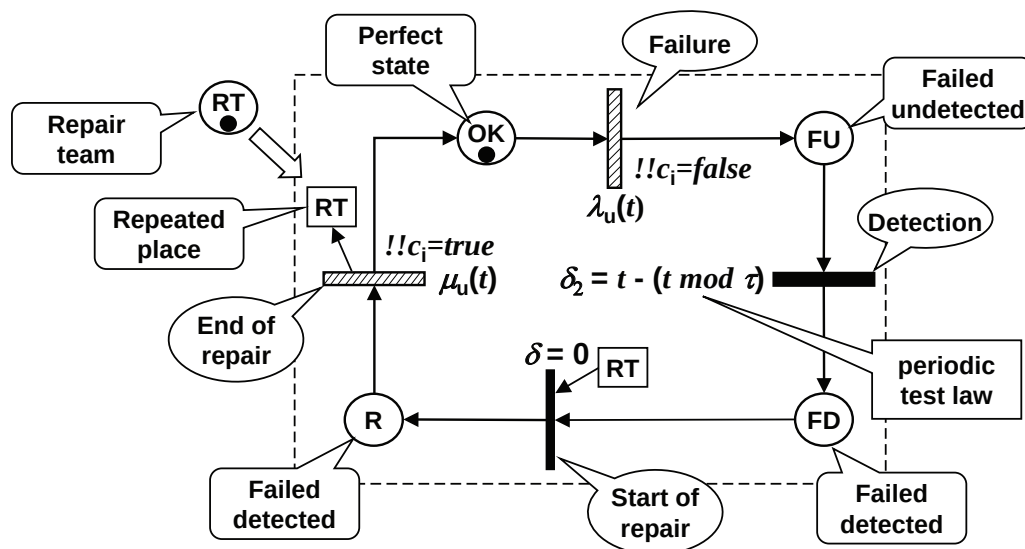


Figure 41 — Basic example of a periodically tested component

As discussed in 6.2.1, the safety systems have basically only two states (OK, KO) and therefore using reliability block diagrams (RBD) seems a natural way of modelling. The approach of using an underlying virtual RBD to drive the building of the PN itself is detailed below. It has proven to be very effective because:

- it helps in building the PN;
- it helps to understand the behaviour of the PN;
- it helps to master large models thanks to the use of libraries of blocks;
- it allows calculations which are far beyond the capacity of the common RBDs;
- it does not really limit the modelling powerfulness of Petri nets.

Figure 41 is related to a sub-PN modelling an individual periodically tested component:

- the component has four states: OK, FU (failed and undetected), FD (Failed and detected) and R (under repair);
- the component has four transition:
 - failure, governed by a failure rate (not necessarily constant);
 - detection, governed by a periodic test law;
 - start of repair, governed by the availability of the repair team (cf. repeated place RT);
 - end of repair, governed by a repair rate (not necessarily constant);
- the “state” variable c_i models the state of the component. It becomes false when the component fails and becomes true again when it is repaired.

NOTE This model explicitly split the restoration time between the detection and repair times.

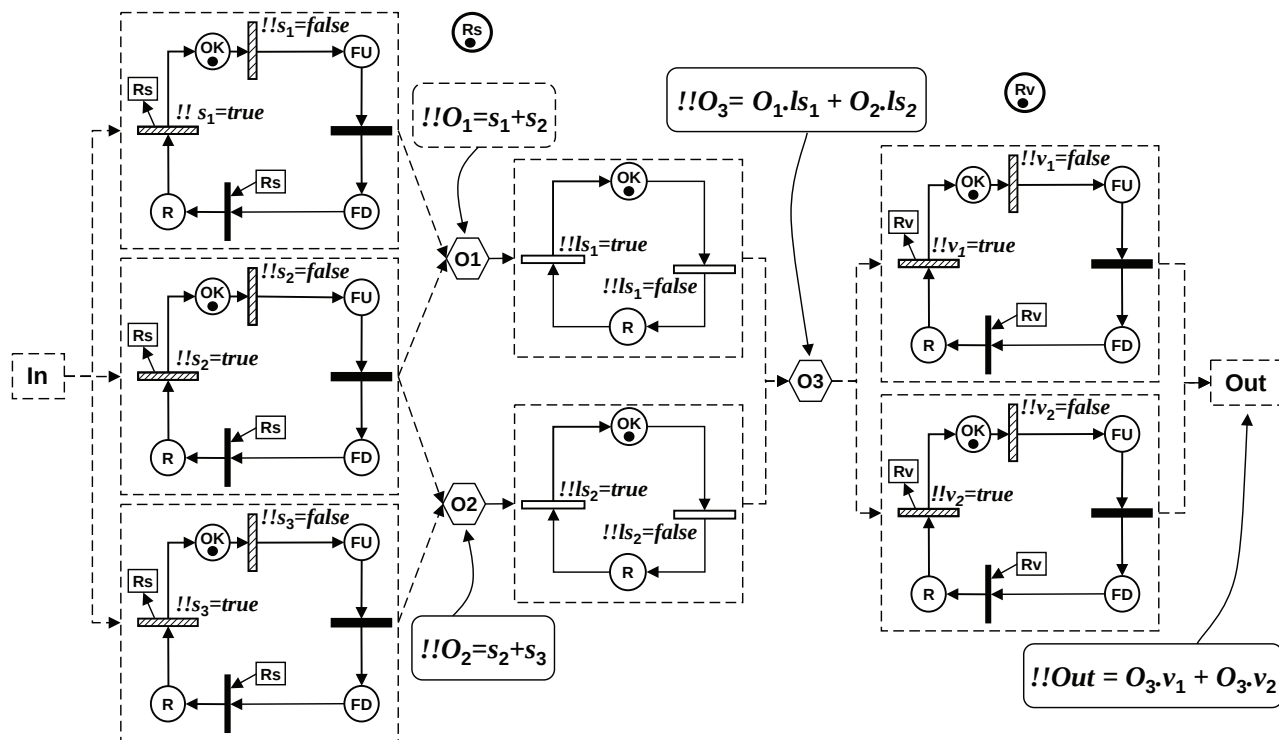


Figure 42 — Example of RBD driven PN

Such small Petri nets can be built for different cases and used as block within a RBD like what has been done in Figure 42. The underlying virtual RBD is represented in dotted lines. It is a straightforward re-transcription of the RBD presented in Figure 18:

- The blocks of the RBD representing the sensors (left hand side) and the final elements (right hand side) are directly derived from the sub-PN described in Figure 41. The logic solver (in the middle) has been modelled with a simpler sub-PN related to immediately revealed failures.
- The logics of the RBD is achieved through four global assertions related to the nodes O1, O2, O3 and Out which use the state variables of each component.
- The three sensors share the repair team Rs when the two valves share the repair team Rv. The repair teams are supposed to be readily available for the logic solvers.

The use of this approach allows developing Petri nets for complex safety systems while keeping the readability of the model and avoiding the mistakes.

The remaining problem is to extract the probabilistic parameters of interest from this model. This role is devoted to the auxiliary Petri net presented in Figure 43.

The link between this auxiliary sub-PV and the model in Figure 42 is done through the variable Out which gives the current state of the overall safety system:

- At the beginning the safety system is OK.
- When the variable Out becomes false then the transition “failure” is fired immediately.
- The token is removed from OK and one token is put in KO indicating that the system is now unavailable.
- The first time that KO is marked, the transition “Unreliability” is immediately fired
- This removes the tokens in KO and M and adds one token in KO and F. Then the marking of KO is unchanged.

- When the variable *Out* becomes true then the transition “restoration” is fired and the token in KO is removed and one token is added in OK. One cycle has been achieved.
- And the process continues until the defined criterion is achieved (e.g. time limit).

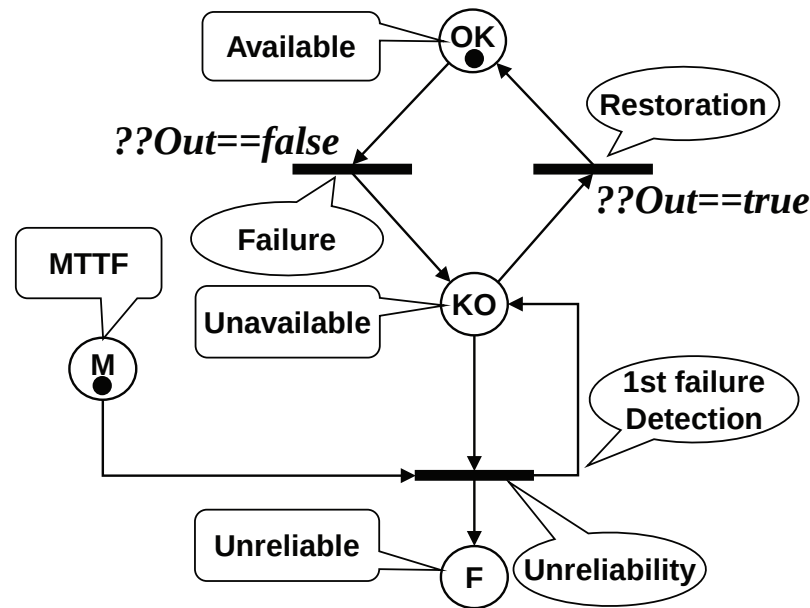


Figure 43 — Auxiliary PN for unavailability/reliability/frequency calculations

It should be noted that the transition “Unreliability” is a “one shot” transition: after the first shot, the token in M definitively disappears and the transition will be never fired again.

When animated by a Monte Carlo simulation this model provides all the relevant probabilistic measures which are needed:

- Availability $A_S(t)$: place OK marked at t .
- Unavailability $U_S(t)$: place KO marked at t .
- Average availability $\bar{A}_S(T)$: mean marking of place OK.
- Average unavailability $\bar{U}_S(T)$: mean marking of place KO.
- Unreliability $F_S(T)$: firing frequency of the transition “Unreliability”.
- $MTTF_S$: Mean marking of place M (provided that T is large enough to have at least one failure for each simulated history).
- Failure frequency (not ultimate safety layer): Number of firings of the transition “Failure” divided by T .
- Failure frequency (ultimate safety layer)
: approximately frequency of the transition “Unreliability” divided by T .
: $\approx 1 / MTTF_S$
- MUT_S : mean marking of OK multiplied by T and divided by the number of firings of the transition “failure”.
- MDT_S (i.e. $MTTRes_S$): mean marking of KO multiplied by T and divided by the number of firings of the transition “Restoration”.
- $MTBF_S = MUT_S + MDT_S$.

Therefore, when using Petri net modelling it is useless to distinguish between, low demand, high demand, and continuous mode of operation or spurious failures. The modelling principle is the same in any cases.

The above consideration summarizes the main features of Petri net modelling. Further guidelines are given in [Annex M](#) for hazardous event frequency calculations and common cause failure modelling.

10.3 Conclusion about Petri net approach

The stochastic Petri nets with predicates and assertions constitute a powerful graphical approach allowing building accurate functioning and dysfunctioning behavioural models of complex industrial systems. They prove to be efficient in modelling safety systems when the limits of the Boolean approach are reached (e.g. strong dependencies between the events) or when the limit of the Markovian approach are reached (e.g. non constant failure or repair rates, large number of states, etc.).

The Petri nets approach is rather easy to understand and the modular approach allows building, step by step, large models in a sound and controlled way (i.e. with a deep understanding of the dynamic behaviour of the model). This allows to remain as close as possible of the behaviour of the actual safety system.

Contrarily to the Markovian approach, no combinatory explosion occurs and the size of the models remains linear with the number of components.

Even if they were developed in the 1960s Petri nets are not well known by reliability engineers but they are more and more used as shown in IEC 62551^[7] which describes the PN approach and its use. They constitute the best compromise between the powerfulness of modelling and the intellectual investment needed to use them. Reliability engineers are encouraged to learn and use this approach which proves to be very powerful in association with Monte Carlo simulation (see [Clause 11](#)) in the framework of probabilistic calculations related to safety systems.

Further developments are given in [Annex M](#) and references^[7] and ^[25]

11 Monte Carlo simulation approach

The Monte Carlo simulation is not a modelling approach but a way of calculation which can be used with any model described in this Technical Report. It consists in building a random game and to play it a great number of times in order to obtain wanted results through statistical calculation. This is a powerful tool which still works when analytical approaches have reached their limits.

Any of the models described in this Technical Report can be used within Monte Carlo simulation but the Petri nets are the most effective to do that.

Contrarily to analytical approaches the Monte Carlo calculation is not linked to the size of the models. It depends mainly on the number of events which occur during the simulations. With the current computation power of personal computers there is no problem to safety system up to SIL4 (e.g. with average unavailabilities in the range of 10^{-4} - 10^{-5} over one year). Several safety barriers running in sequences may be also handled without too many difficulties. It may be difficult to calculate very low hazardous event frequencies. This can be solved by using faster computers or by mixing the Monte Carlo simulation results with analytical calculations (cf. [6.1](#)).

Further developments are given in [Annex N](#) and reference^[25]

12 Numerical reliability data uncertainty handling

There are different ways of describing the relevant uncertainties, which reflects the varying definitions of risk in ISO standards (e.g. ISO 31000^[28] vs. ISO 20815^[16]), and a clarification before starting on the assessments should be made to make sure relevant aspects are captured. Regardless of which one is selected, the presentation of uncertainties should be described beyond the use of probabilities, to also include uncertainties related to e.g. expert judgements and assumptions in the reliability modelling and calculations. The various features of uncertainty handling with respect to assumption, modelling and results were addressed briefly in [5.1](#).

In the context of this Technical Report, the treatment of uncertainties relates to the variability and the (expected) frequency of safety relevant failures. Therefore this Technical Report focuses mainly on the numerical uncertainties of input reliability data. Those uncertainties are analysed in [Clause 13](#) and the present clause explains how they can be propagated throughout any model in order to evaluate their impact on the results. Further developments are proposed in [Annex N](#).

NOTE ISO 20815[16], D.3.7 does also state the importance of how to handle uncertainty as part of an analysis procedure. It is not specifically safety system related, but still useful to recognize for the modelling covered in this Technical Report.

The accuracy of the probabilistic calculations is directly linked to the numerical uncertainties of the reliability parameters which are used as inputs. Therefore it is beneficial and mostly required to take those uncertainties into account in order to evaluate their impact on the results which are found.

Except in very simple case the analytical approaches do not work and the Monte Carlo simulation provides the only way to do that. This is illustrated in [Figure 44](#) on a simple fault tree.

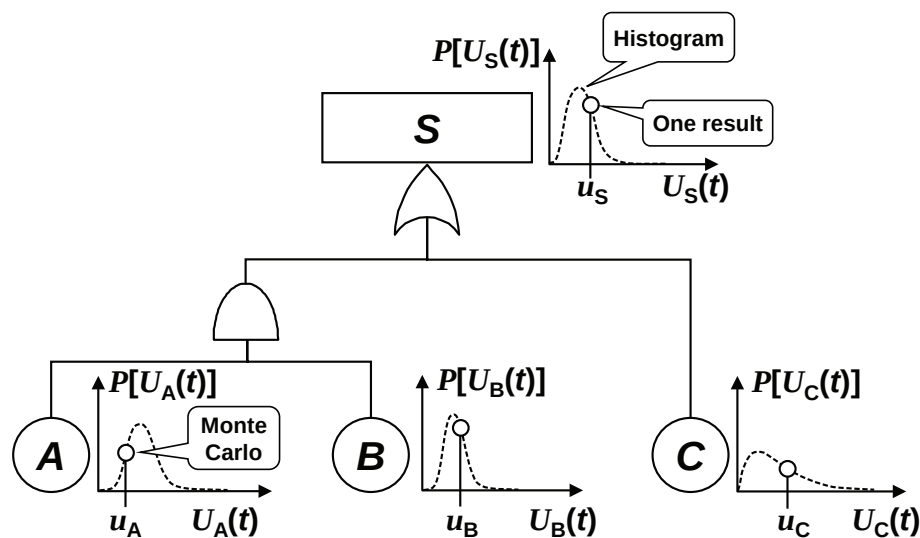


Figure 44 — Uncertainty propagation

When handling the uncertainties, the input reliability parameters are no longer deterministic and they are replaced by random variables. Therefore the probabilities of failure of the leaves (e.g. unavailabilities) become also random variables in turn. The probability density functions (*pdf*) are more or less sharp or flat according to the degree of uncertainty which is involved. The result at the top of the fault tree is no longer deterministic. Thank to the Monte Carlo simulation it is obtained as a histogram which is more or less sharp or flat according to the impact of the input uncertainties.

This principle is very general and may be applied on any kind of models: analytical formulae, Markov processes and even Petri nets. Further developments are provided in [Annex O](#).

13 Reliability data considerations

13.1 Introduction

Establishing relevant reliability data, as input to the reliability model, is an essential part of any quantitative reliability analysis. It is also one of the most challenging tasks and raises a series of questions:

- Are reliability data available for the equipment under consideration?
- Are the available data relevant for the specific application, under the given operating and environmental conditions?

- Should manufacturer/ vendor supplied data be combined with experience/field feedback data for a limited population or should generic data from larger (but not necessarily fully comparable) populations of components be applied?
- Are the assumptions underlying the given data known, such as e.g. system boundaries, what components that are included, type of failure modes covered, energized or de-energized to trip, etc.?
- Is there a maintenance strategy in place to support the assumption about the components being in their useful life period?
- What uncertainties are associated with the data, for example due to the quality of failure recording, the quality and coverage of periodic tests, the number of operating hours, and so on?

NOTE These issues have been addressed in ISO 20815[16] and in ISO 14224:[15]

ISO 14224[15] provides fundamental guidance on the data collection principles and reliability data vocabulary in general (e.g. equipment taxonomies and equipment class attributes like equipment specific failure modes) also to be used when estimating reliability data for analysis usage.

ISO 14224[15] and ISO 20815[16] emphasize the need for undertaking appropriate data qualification adapted to the study in question.

The reliability data are established from available data sources and this may need to combine several reliability data collections. Methods exist e.g. the Bayesian[27] approach which are beyond the scope of this Technical Report provides the framework to do that but no further development are provided.

Good reliability data need qualified reliability data collection of historic equipment performance. ISO 14224[15] has been developed to ensure this with regards to petroleum, petrochemical and natural gas industries. It should be noted that ISO 14224[15] can also be used when estimation of non-recorded data are made, to format the predictions at the same format as one would have done by historic data estimation.

13.2 Reliability data sources

Reliability data sources can be of various types, as classified in Table 4.

Table 4 — Classification of reliability data sources

Source of data	Description
1. Generic data (see 3.4.16)	Reliability data covering families of similar equipment. Such generic data can be: Computerized database of data, typically grouped in data tables with several equipment attributes. The related historic data collection may be done according to published standards (e.g. ISO 14224[15] for petroleum, petrochemical and natural gas industries) Published data handbook[19], sometime simplified versions of computerized databases. The formats may depend on the publisher. Such data handbooks would normally be historic data, i.e. operating field experience. Published data handbook based on expert judgement, but not on historic operating data or underlying database. The data may apply for one specific industry, or it may be collected from several industries. NOTE Some industrial initiatives may provide other data, (e.g. failures detected during test) that might be useful for establishing reliability data input.

Table 4 (continued)

Source of data	Description
2. Operator/ company specific data	Reliability data or reliability indicators based on operating field experience of a single company. Such data can be established by one operator/oil company from: • One or several of its installations, • Its own interpretation of different data sources, • Key Performance Indicators (KPI) NOTE 1 Such operator/company specific data may be part of an industrial cooperation generic reliability database, or purely own company data. NOTE 2 The data may also be part of a company database that might comply with an international standard (e.g. ISO 14224[15]). NOTE 3 ISO 14224[15] Annex E exemplifies KPIs. NOTE 4 The events recorded in a CMMIS system are not reliability data, but could be used to establish some reliability indicators (e.g. KPI).
3. Manufacturer data	Reliability data produced by a particular manufacturer data for a particular product. Such data may be based on: • Operating field experience from • The manufacturer himself. The data can be aligned or not with an international standard (e.g. ISO 14224[15]). • The users (e.g. specific or generic data mentioned above) • Component FMECA/ studies, • Laboratory testing, e.g. accelerated life time testing, reliability testing. This may apply for new technology equipment for which experience data not yet exist. Such pre-operational data should normally be entitled 'pre-operational/test reliability data', as opposed to actual field operating experience. See also IEC 61164[50] for statistical tests and estimation methods for reliability growth.
4. Expert judgement (see references[26] and[27])	Expert judgement would involve • General advice from safety system equipment expert • Use of statistical expert analysis methods (e.g. Delphi, etc.) to utilize a variety of qualified sources as input to the reliability analysis
5. Human error data	Various sources of human error data exist and H.2 gives some advice about human error probabilities.

INFO BOX: It is a weakness in the industry that too little attention is given to the quality of the input data. Therefore, undertaking the qualification of reliability data found in data sources is vital for the credibility of the results in risk decision making. This qualification is out of the scope of this Technical Report, but ISO 20815[16], Annex E gives advice on this matter.

1. Generic data:

Generic data are often (but not necessarily), see Table 4 based on operational experience from a number of installations and a number of comparable equipment types, such as e.g. flame detectors from different vendors. In such case the generic data reflect some kind of average expected field performance for the equipment type under consideration.

At early project stages generic data are often selected due to lack of detailed information as all equipment features decisions have not yet been made. However at later project stages one should preferably apply valid application or equipment specific data – if well documented and considered relevant.

2. Operator/company specific data

Authorities require that the companies keep control of their safety barriers throughout the entire lifecycle of an installation. Consequently, it is often required for the operators to collect installation specific failure data during maintenance and operation. During modification analyses such data are of particular relevance for the purpose of documenting the performance history of given equipment. However, since the statistical confidence in data from only one installation may often be poor (or all potential failure events may not have occurred so far at the installation), reliability analyses are seldom based on such data alone. However, for some equipment where the number of installed units is high, e.g. fire and gas detectors, it may be relevant to apply installation specific data only.

3. Manufacturer data

It is often stated by the analyst that supplied manufacturer data are significantly “better” than comparable generic data (i.e. lower failure rates). This may have several reasons, such as varying equipment quality, failure modes included and the definition of equipment boundaries. Another important aspect, however, is that failures due to environmental stress, due to mal-operation, installation failures, maintenance errors, etc. have frequently been excluded from the manufacturer data. This is understandable since manufacturers are in the business of selling and does not want to include failures that may be attributed to factors external to the equipment itself. Also, if the vendor charges for failure analysis this is a disincentive to return the failed components. Another aspect is the fact that feedback from the operators using the equipment may be poor (especially beyond the warranty period) and in such case it is difficult for the manufacturer to establish a good failure rate estimate. Consequently, using data from manufacturers may involve too low failure rates and as such needs to be carefully considered. It is therefore advisable to deploy the ISO 14224^[15] principles to strengthen the quality of the data and the communication on these matters.

When using manufacturer data the reliability engineer should remember to add failures due to connection blockage which are often included in field experience data but excluded from manufacturer data.

4. Expert judgement data

The use of experts to estimate reliability data requires qualification of that the expert is an equipment expert, and understands the methods being used in reliability data estimation. ISO 14224^[15] provides good guidance on such matters, even though not all type of equipment is covered. If experts are used it would also be beneficial to undertake independent reviews and also be conscious in how reliability data are being communicated, e.g. number of observable events for a given fleet of equipment units at an installation for a certain period of time would be more practical than estimating failure rates in 10^{-6} per hrs. Separate methods exist for expert data analysis.

5. Human error data

Observations or other information can be used to quantify the failed human interactions:

- generic data (look-up tables);
- collected data (calculated human error probability) specific to the task;
- estimation methods (expert judgement);
- combination of the above.

13.3 Required reliability data

NOTE 1 ISO 14224^[15] [Clause 8](#), [Table 8](#) and [Annex B](#) provide valuable information about the topics below.

The required reliability data for reliability modelling of safety systems and safety functions depend on the method selected and reliability parameters to be used. This would require addressing the following data:

- 1) Failure data e.g.:
 - i) failure modes and related failure rates or probability of failure occurring on demand;
 - ii) failure causes;
 - iii) failure mechanisms;
- 2) Operational and Maintenance data e.g.:
 - i) detection method;
 - ii) maintenance activity;
 - iii) downtime and related restoration and repair rates;
 - iv) active maintenance time and active repair rate;
 - v) test intervals and test policy (including test duration, staggering, partial stroking for the valves, etc.).

INFO BOX: When establishing reliability input data, verify that the battery limit of the item in the reliability database is not the same as the battery limit of the item in the reliability model (see the note below). This will strengthen analysis quality on this matter. Therefore it is important to establish that the reliability data to be used are aligned with the taxonomy implemented in the database, or at least some remarks on this matter is given when documenting data qualification in a specific reliability study.

NOTE 2 For example, if a compressor is considered, it is important that the analyst verifies if the motor is included (or not included) in the reliability data. ISO 14224^[15], 8.2 provides valuable information about the taxonomy levels. According to this standard, for example, a solenoid valve inside a HIPPS valve is a component (taxonomy level 8), while the valve itself is an equipment unit (taxonomy level 6). A subsea valve (e.g. SSIV) can be a component (level 8) as part of a subsea isolation system (subunit at taxonomy level 7) and part of an export manifold (taxonomy level 6). The item to be modelled should be on similar taxonomy level.

INFO BOX: Erroneous equipment count combined with inaccurate failure rate per item can give improper results. Then it is important to quantify some key items to make sure proper system understanding exists, prior to establishing reliability data. Inadequate interpretation about the physics of the failure may also lead to improper results.

NOTE 3 ISO 14224^[15], Annex F contains furthermore information on critical/dangerous failures of some safety systems/components that describes also the failure physics to make sure correct interpretation of data sources used.

INFO BOX: The quality of reliability data depends on the relevance of the data collection with regard to the needs of the reliability study and of the quantity of information collected. Therefore there always are some uncertainties on input reliability data, especially when no failures have been observed. These uncertainties can be measured by confidence intervals obtained by processing the reliability data samples by classical statistics calculations.

NOTE 4 ISO 14224^[15], Annex C describes how the Chi square distribution can be used to establish confidence interval as well as to estimate the failure rates when zero failures has been observed.

13.4 Reliability data collection

Even with perfect calculation tools it is not possible to achieve good probabilistic estimations without good reliability data. Therefore the user of this Technical Report should be ready to make a sufficient effort to establish the reliability data needed for the probabilistic calculations.

It is beyond the scope of this Technical Report to detail how to proceed to actually collect reliability data but requirements and guidance are given in ISO 14224[15]. This is also supported by ISO 20815[16], [E.1](#). Those standards are very helpful to:

- Use the available reliable data.
- Identify gaps in the data needed, and thus specify later reliability data collection to be done.
- Produce recommendations to assist the operator to take further steps in data collection.

Finally it is vital that the reliability study reveal uncertainties and quality constraints in a data dossier and identify how this may impact the results. See [Clause 12](#) and [Annex O](#) for further discussion about uncertainty handling and [26][27] about expert judgement handling.

The reliability databases have not been necessarily designed specifically for the modelling and calculation of the safety systems. Therefore some needed reliability data may be difficult to find without specific studies (e.g. the probability of failure due to a demand - see [3.1.15](#)).

14 Typical applications

14.1 Introduction

The aim of this clause is threefold:

- a) Illustrating the various approaches described within this Technical Report.
- b) Proposing a pedagogic approach of the current modelling and calculation problems encountered when dealing with safety system associated with relevant solutions.
- c) Highlighting the modelling possibilities and limits of the various approaches.

The typical applications chosen to do that are *overpressure protection systems* operated in demand mode. However the architectures and hypothesis implemented are general enough to cover most of the safety systems identified in [Annex A](#). Then, the content of this clause is representative of most of the reliability studies of safety systems performed in petroleum, petrochemical and natural gas industries as well as in other industries. These typical applications are summarized in Table 5.

Assumptions are chosen to make the typical applications as simple as possible and as close as possible to actual applications. The complexity is increased step by step to progressively introduce the difficulties and the solutions. This is intended to help the users to gradually learn how to handle the content of this Technical Report. Then it is encouraged to start the reading of this clause from the beginning.

NOTE The reliability data have been chosen to highlight the impact of the assumptions on the results. They are not to be used for actual reliability studies.

The average unavailability is the parameter usually used as a probabilistic measure of the efficiency of the safety system operated in demand mode when the average failure frequency is the parameter usually used for the safety systems operated in continuous mode. However both parameters (i.e. average unavailability and average failure frequency) are evaluated in this clause in order to illustrate how the various approaches can be used for both types of probabilistic calculations. In addition, the average spurious action frequency is also modelled and calculated for some of these typical applications.

Table 5 — Summary of the typical applications analysed in this Technical Report

Typical Application		Simplified description
TA1		Single channel: Simplest safety loop (1 sensor + 1 logic solver + 1 final element)
	1	Periodic tests perfects and at the same time protected installation shut down during tests and repair
	2	Idem TA1-1 + different test interval
	3	Idem TA1-1 + protected installation not shut down during tests and repair + test off line of the sensors
	4	Idem TA1-1 + Periodical test coverage < 100 %
TA2		Dual channel: Two parallel channels similar to TA1
	1	Same assumptions as TA1-1
	2	Idem TA2-1 + different test intervals + partial and full stroking of valves
	3	Idem TA2-1 + protected installation not shut down during tests and repair + test off line of the sensors
	4	Idem TA2-1 + staggering of the tests
TA3		Popular redundant architecture: 2oo3 of sensors, 1 logic solver, 1oo2 of valves
	1	Same assumption as TA1-1
	2 to 4	(For memory)
	5	Idem TA3-1 + Subsea HIPPS assumptions (rig mobilization, production stopped during valve repair, etc.)
	6	Idem TA3-2 + switching from 2oo3 to 1oo2 on detection of a sensor is failure
TA4		Multiple safety system: two dependent safety loops operating sequentially
TA5		Emergency depressurization system of an hydrocracking unit: EDP system

14.2 Typical application TA1: single channel

14.2.1 TA1-1: basic case

14.2.1.1 TA1-1: description and assumptions

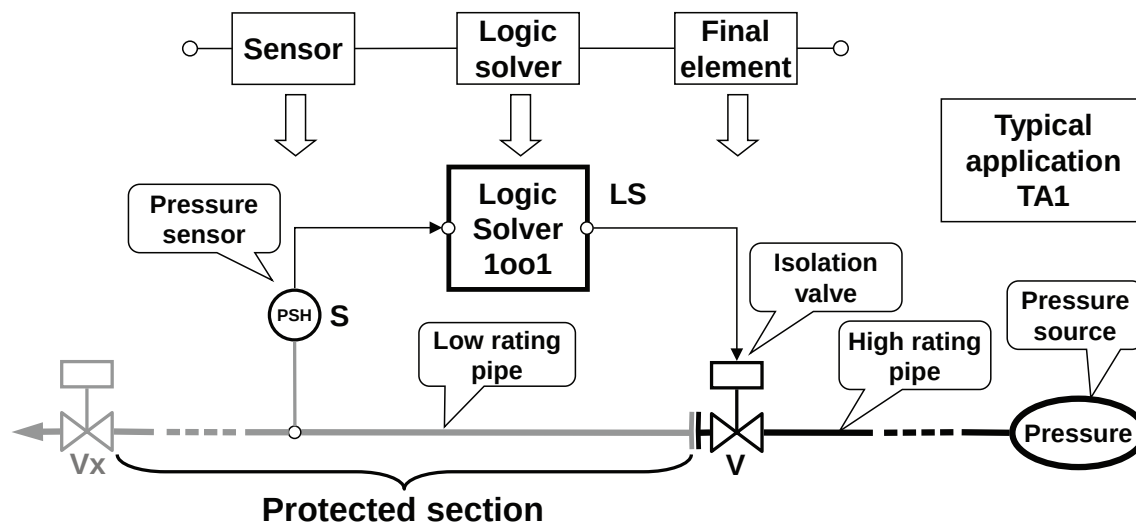


Figure 45 — Simplest basic architecture - Typical application TA1-1

The simplest architecture of a safety instrumented system is illustrated in [Figure 45](#): one single sensor, one single logic solver and one single final element. It is popular in process industry for common safety loops with low to moderate reliability requirements (SIL1 to SIL2).

This architecture is implemented to design the safety system presented in [Figure 45](#):

Type of system: pressure protection system (see [Annex A](#)).

Mode of operation: low demand mode.

Layout: one pressure sensor (S), one logic solver (LS) and one isolation valve (V) organized in series.

Design philosophy: de-energize to trip.

Hazardous event: over pressurization and burst of the low rating section of the installation.

Functioning description: when the pressure exceeds a given threshold, the sensor (S) sends a signal to the logic solver (S) which in turn commands the valve (V) to close.

NOTE In reality the isolation valve is piloted by a solenoid valve which actually receives the signal from the logic solver. This has not been represented here in order to simplify the example.

Assumptions:

- Periodical tests are perfects and performed at the same time.
- Installation stopped during periodical tests and repair.
- Dangerous detected and undetected failures of a given component are independent.
- Constant failure rates.
- Components as good as new after repairs.

Table 6 — TA1 reliability parameters

Parameter	Component			
	Pressure sensor	Logic solver	Isolation valve	De-energize to trip elements
Dangerous undetected failure rate	$\lambda_{du,Psh} = 3,0 \cdot 10^{-7} \text{ h}^{-1}$	NA	$\lambda_{du,V} = 2,9 \cdot 10^{-6} \text{ h}^{-1}$	
Dangerous detected failure rate	$\lambda_{dd,Psh} = 3,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{dd,Ls} = 6,0 \cdot 10^{-7} \text{ h}^{-1}$	NA	
Periodical test interval (hours)	8 760 h	NA	8 760 h	
Safe failure rate	$\lambda_{sf,Psh} = 3,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{sf,Ls} = 6,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{sf,V} = 2,9 \cdot 10^{-4} \text{ h}^{-1}$	$\lambda_{sf,deg} = 4,0 \cdot 10^{-6} \text{ h}^{-1}$
Maintenance time (h) MRT or MTTR ^a $\mu = 1/\text{MRT or } 1/\text{MTTR}$	15 h $\mu_{Psh} = 1/15$	15 h $\mu_{LS} = 1/15$	40 h $\mu_V = 1/40$	24 h $\mu_{deg} = 1/24$
^a In this clause the mean repair time, MRT, is related to the repairs of the dangerous undetected failures after they have been detected by periodical tests. The mean time to restoration, MTTR, is related to dangerous detected failures and includes both the detection and the repair times. Nevertheless as the mean fault detection time, MFDT, of the dangerous detected failures is generally negligible compared to the MRT, it can be considered in this clause that MTTR and MRT have the same numerical values.				

The values presented in the above table are proposed for illustration purposes only.

14.2.1.2 TA1-1: analysis

As this is a “de-energize to trip” safety system, the isolation valve closes each time the energy is lost somewhere. This occurs, e.g. in case of:

- loss of power supply, e.g.:
 - electric power supply (AC or DC);
 - hydraulic power supply (i.e. the isolation valve is “fail-safe close”);
- loss of auxiliary supply (UPS, compressed air, etc.) electric supply;
- open circuit between the sensor and the logic solver;
- open circuit between the logic solver and the isolation valve;
- any dangerous detected failure which is transformed into a safe failure by the self diagnostic tests.

Therefore, with the “de-energize to trip” philosophy, all of the failures related to energy or signal losses are safe and don't need to be considered when evaluating the probability of dangerous failures. This simplifies very much analysis to be performed as well as the modelling. The counterpoint is that this increases the probability of spurious safety actions.

According to the assumptions, the installation is stopped during the repairs. Then the risk disappears during the repairs of the dangerous failures ‘undetected as well as detected’).

According to the assumptions, the installation is stopped during the periodical tests. Then the risk disappears during the periodical tests of dangerous undetected failure.

Then the exposure to the risk will exist only during the fault detection times but they are equal to zero for immediately revealed failures and normally negligible (e.g. few seconds) for detected failures.

Finally with the above assumption, only the fault detection times of dangerous undetected failures are considered in the hazardous event probability calculations.

If T_t is the cumulated time spent to perform periodical tests and repairs over a period $[0, T]$ then $[T - T_t]$ is the overall duration of the exposure to the risk. Therefore, and provided that a periodically tested component remains available for its safety function during the periodical tests, neglecting T_t in the calculations (i.e. considering that the duration of the exposition to the risk is equal to T) is conservative.

The assumption that the dangerous detected and undetected failures of the same component are independent should be addressed because when a dangerous undetected failure occurs on a given component this can inhibit the detected failures on this component and vice versa. Therefore, over a given period $[0, T]$, the actual time during which a dangerous detected failure (or a dangerous undetected failure) can occur may be lower than T . When the probability of dangerous detected and undetected failure is low, the impact is negligible. However, neglecting these aspects (which is equivalent to use T in the calculations) is conservative.

14.2.1.3 TA1-1: probabilistic calculations

As this safety system operates on demand, its average unavailability (also called PFD_{avg} for a safety instrumented system) is the relevant parameter to calculate.

In this very simple case any approach described in this Technical Report may be used.

14.2.1.3.1 Analytical formulae

The hypothesis that all the components are periodically tested at the same time allows to gather all the components within a macro component with the following parameters:

- Dangerous undetected failure rate: $\lambda_{du} = \lambda_{du,Psh} + \lambda_{du,V} = 3,2 \cdot 10^{-6} \text{ h}^{-1}$.
- Periodical test interval: $\tau = 8\,760 \text{ h}$.

Therefore with all the assumptions described above, this is the simplest case which can be encountered. The analytical formulae can be straightforwardly applied:

Average unavailability: $\bar{U}_{du}(\tau) \approx \frac{\lambda_{du} \cdot \tau}{2} = 1,4 \cdot 10^{-2}$.

14.2.1.3.2 Fault tree approach

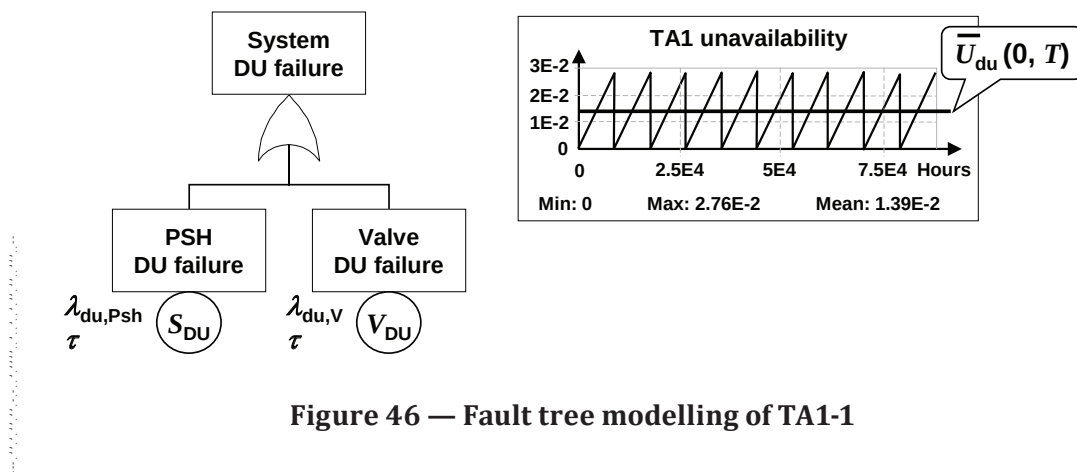


Figure 46 — Fault tree modelling of TA1-1

In the fault tree presented in [Figure 46](#) the instantaneous unavailabilities of the two leaves are similar to the saw tooth curves developed with the multi-phase Markovian approach in [Clause 9](#). This allows calculating the well known saw-tooth curve presented on the right hand side of [Figure 46](#).

Using the same reliability parameters as above gives an average unavailability $\bar{U}_{du}(0,T)$ over 10 periodical test intervals of 8 760 h (i.e. $T = 10$ years), which obtains: $\bar{U}_{du}(0,T) = 1,39 \cdot 10^{-2}$.

This result is very close to the result obtained with the formula (see 14.2.1.3.1) which is slightly conservative.

Another result which can be obtained straightforwardly from the above fault tree is the average dangerous failure frequency of TA1-1. The result is given in Figure 47 which presents the time dependant failure frequency. This is also a saw tooth curve with an average value of $\bar{w}_{du}(0, T) = 3,16 \cdot 10^{-6}$ failure per hour. Numerically speaking, this is very close to the dangerous undetected failure rate Λ_{du} calculated above in 14.2.1.3.1.

Such a result is required when analysing the functional safety of instrumented safety systems operating in continuous or high demand modes. This is so-called PFH of the safety system.

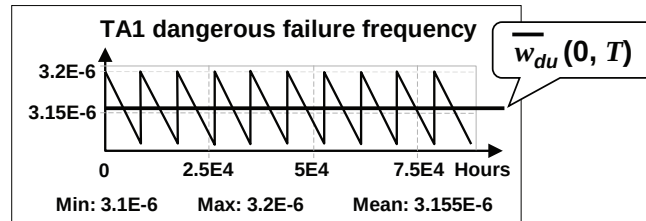


Figure 47 — Average dangerous failure frequency of TA1-1

No repair occurs within a test interval then when the system as failed once no further failure can occur. This explains why the failure frequency decreases (see Figure 47) during the test interval. If the test interval is long enough the failure frequency tends to zero when the probability of failure will tends to one.

Of course, this model is not really interesting for such simple cases but is useful when further developments are expected.

14.2.1.3.3 Multi-phase Markovian approach

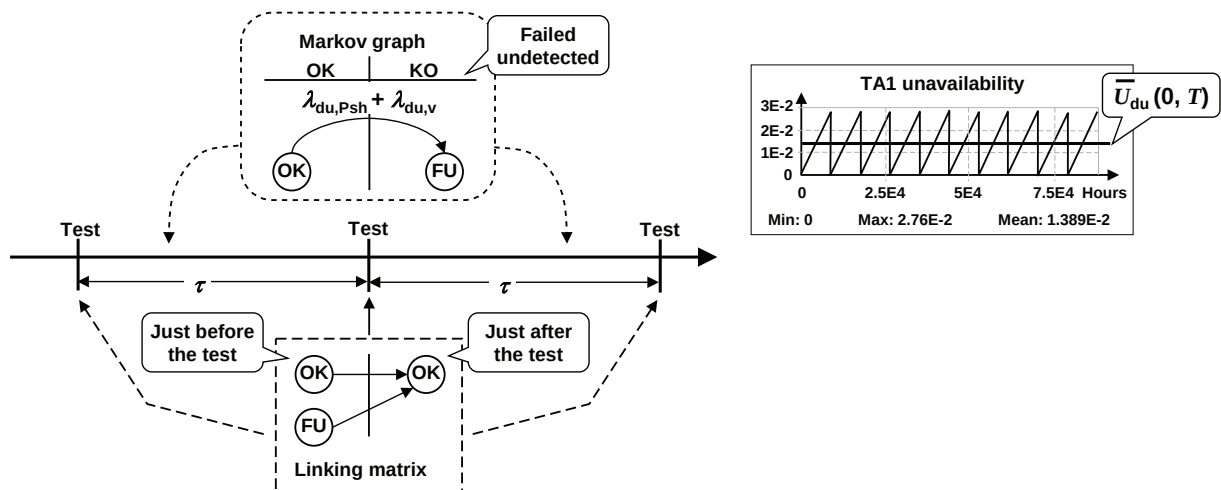


Figure 48 — Multi-phase Markov model of TA1-1

Figure 48 is a simplified version of the multi-phase Markov model developed in 9.2. It gives the same result as the fault tree over the same period of 10 years (i.e. Ten periodical test intervals of 8 760 h). It provides exactly the same saw tooth curve which has been obtained with the fault trees. Again an average unavailability is found: $\bar{U}_{du}(0, T) = 1,39 \cdot 10^{-2}$

This Markovian model gives also the same result as the fault tree for the dangerous failure frequency and like for the fault tree approach, it is not really interesting for such simple cases but it may be useful if further developments are forecasted.

14.2.1.3.4 Petri net and Monte Carlo simulation approach

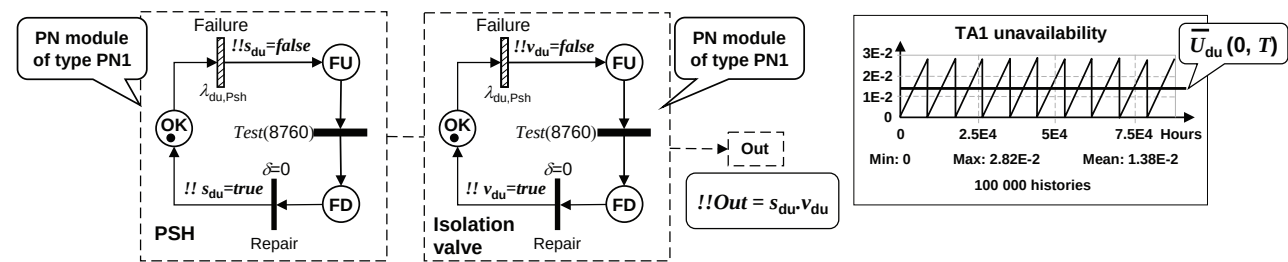


Figure 49 — PN modelling of TA1-1 with PN modules of type PN1

As shown in [Figure 49](#), Petri net modelling and Monte Carlo simulation can also be used for dealing with simple systems. This PN is made of two similar sub-PN (PN1 modules) organized through a virtual reliability block diagram (in dotted lines) as this has been explained in [Clause 10](#).

[Table 7](#) shows the evolution of the results when the number of histories (i.e. the number of Monte Carlo simulations) increases. Even when this number is low (e.g. 100), the result is not far from the results given with the other methods but the confidence is not very accurate. When the number of histories increases the confidence increases and with a big number of histories (e.g. 10⁶) a result which is very close to the previous results is obtained. Therefore the result is the same as with other approaches providing that the number of simulated histories is high enough.

The PN presented in [Figure 49](#) is also able to provide the average dangerous failure frequency of TA1-1. For doing that it is necessary to count how many times the variable Out changes from “true” to “false” and 0,276 changes over 87 600 h with a 90 % confidence interval of [0,275 0 – 0,276 7] is found. This leads to $\bar{w}_{du}(0,T)=3,15 \cdot 10^{-6}$ [3,14 10⁻⁶ – 3,16 10⁻⁶]. This is very close to the previous results

Table 7 — Accuracy of Monte Carlo simulation versus number of histories

Number of histories	Average unavailability $\bar{U}_{du}(0,T)$	90 % confidence interval of $\bar{U}_{du}(0,T)$	Standard deviation of $\bar{U}_{du}(0,T)$
100	1,40 10 ⁻²	[9,41 10 ⁻³ – 1,85 10 ⁻²]	2,77 10 ⁻²
1000	1,37 10 ⁻²	[1,21 10 ⁻² – 1,52 10 ⁻²]	2,96 10 ⁻²
10 000	1,35 10 ⁻²	[1,30 10 ⁻² – 1,40 10 ⁻²]	2,99 10 ⁻²
100 000	1,38 10 ⁻²	[1,36 10 ⁻² – 1,40 10 ⁻²]	3,00 10 ⁻²
1 000 000	1,38 10 ⁻²	[1,38 10 ⁻² – 1,39 10 ⁻²]	3,01 10 ⁻²

Such modelling seems not really interesting for such simple cases but it allows introducing basic modules which can be used for more complex modelling.

Nevertheless it also provides results which cannot be obtained with the other approaches, for example the distribution of the average unavailability, which is actually a random variable: the *N* different histories performed during the Monte Carlo simulation provide *N* different values of the cumulated down time of the safety system over the period of interest. Therefore *N* values of the average unavailability over this period can be derived and a statistical sample of the average unavailability is obtained.

This sample is, conventionally, characterized by its expected value which is the $\bar{U}_{du}(0,T)$ given above and by its scattering around this expected value which is measured by the standard deviation. This has been indicated in the right hand side column of [Table 7](#): it can be seen that the sample of the average unavailability is rather scattered around its expected value $\bar{U}_{du}(0,T)$: this scattering is due to the use of exponential laws.

It should be noted that the standard deviation is almost independent of the number of histories whereas the width of the confidence interval decreases when the number of histories increases. These two parameters should not be mixed-up.

The statistical sample may also be processed in order to build a histogram representing its distribution. This has been done in [Figure 50](#) which presents the various percentiles of the sample.

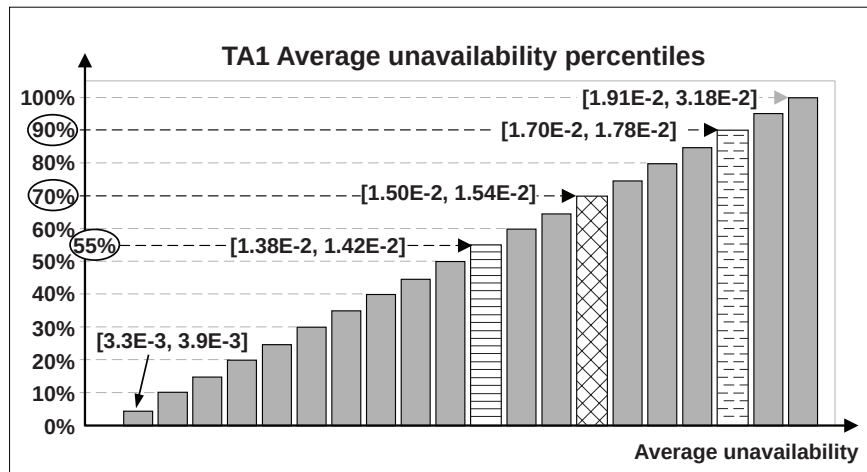


Figure 50 — Average unavailability percentiles of TA1-1

It should be noted that, when the probability of failure is low, most of the histories provide unavailabilities equal to 0 which are not really useful to build a histogram. This is the case of our example calculated for 10 years. Fortunately this can be overcome by considering that, in the case of TA1-1 that $\bar{U}_{du}(0,T)$ converges toward an asymptotic value which has already been reached after 10 years (It is reached after two or three periodic test intervals; i.e. Three years in our example). Then $\bar{U}_{du}(0,T) = 1,39 \cdot 10^{-2}$ established above for 10 years does not change (with regard to the Monte Carlo simulation variations) if the calculation were performed over 100 or 1 000 years.

If a time T long enough to have a high probability of failure is chosen, then an accurate histogram can be obtained because all the classes will be populated. This is what has been done for the histogram in [Figure 50](#) which has been established over a period of interest of 1 000 years.

This histogram shows that the average value, $1,39 \cdot 10^{-2}$, is just a little bit higher than the median value: about 45 % chances to be worse and 55 % chance to be better than this average value.

NOTE The accuracy of the Monte Carlo simulation ([Table 7](#)) should not be mixed up with the percentiles ([Figure 50](#)). This are two different concepts and, for example, from accuracy point of view, ([Table 7](#)) shows that the average unavailabilities has 90 % chances to be comprised within $[1,38 \cdot 10^{-2} - 1,39 \cdot 10^{-2}]$ and from scattering point of view, the histogram in [Figure 50](#) shows that there are 55 % chances that the actual value of the average unavailability be better than this average value.

It also shows that the true value of the average unavailability has about 70 % chances to be better than $1,52 \cdot 10^{-2}$ and about 90 % chances to be better than $1,74 \cdot 10^{-2}$. The above considerations are useful when conservative estimations are wanted.

Another thing which cannot be done with the other approaches is to handle non constant failure rates. This is straightforward with the Monte Carlo simulation.

For example the exponential law can be replaced by Weibull law with the same MTTF:

- $\lambda_{du,Psh} = 3,0 \cdot 10^{-7} \Rightarrow \text{MTTF} = 3,333 \cdot 10^6 \text{ hours}$
- $\lambda_{du,V} = 2,9 \cdot 10^{-6} \Rightarrow \text{MTTF} = 3,448 \cdot 10^5 \text{ hours}$

Using Weibull laws with the above average values and shape parameters equal to 3 the results presented in [Table 8](#) are obtained.

Table 8 — Accuracy of Monte Carlo simulation versus number of histories (Weibull)

Nb of histories	Average unavailability $\bar{U}_{du}(0, T)$	90 % confidence interval of $\bar{U}_{du}(0, T)$	Standard deviation of $\bar{U}_{du}(0, T)$
10 000	5,55 10 ⁻⁴	[4,53 10 ⁻⁴ – 6,57 10 ⁻⁴]	6,22 10 ⁻³
100 000	5,63 10 ⁻⁴	[5,23 10 ⁻⁴ – 5,93 10 ⁻⁴]	6,09 10 ⁻³
1 000 000	5,56 10 ⁻⁴	[5,45 10 ⁻⁴ – 5,66 10 ⁻⁴]	6,00 10 ⁻³

It is evident that these results are very different from the previous ones. They seem stable but a question arises: are they relevant?

The left hand side of [Figure 51](#) provides the answer: very good availability in the first years compensates worse results in the last years of the period. This is due to the shape factors equal to 3 of the Weibull laws which model components with continuously increasing failure rates (e.g. ageing components). The instantaneous failure rate start from 0 and after that, it is continuously increasing. Therefore according to the period of observation the results can be better, similar or worse than with a constant failure rate. In our example, T is short enough to obtain a better result but this is not a general observation.

Anyway, the average unavailability is continuously increasing as shown on the right hand side of the [Figure 51](#). Therefore, even if this result is on line with common practices (e.g. IEC 61508[2]) it is not really meaningful.

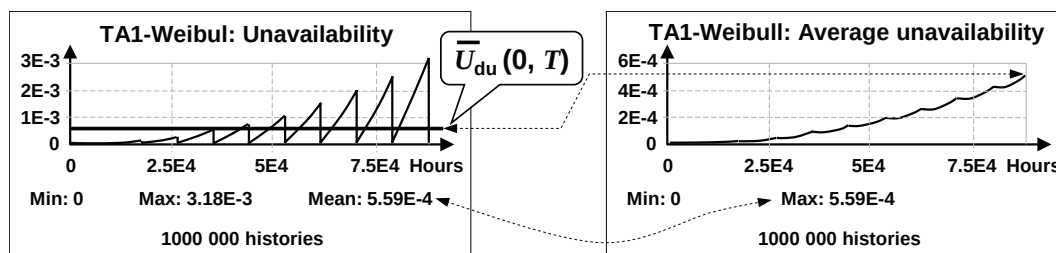


Figure 51 — TA1-1 with Weibull laws

The average dangerous failure frequency can be calculated as this has been shown above: 0,011 6 changes over 87 600 h with a 90 % confidence interval of [0,011 4 – 0,011 8] are found. This leads to $\bar{w}_{du}(0, T) = 1,32 \cdot 10^{-7}$ with a 90 % confidence interval [1,30 10⁻⁷ – 1,35 10⁻⁷]. As this may be anticipated from the unavailability results, this frequency is different from the exponential case. It is lower over [0, 876 00] but would increase if this time interval was increased.

Then, and even in the simplest cases, the reliability analyst should verify that the hypothesis of constant failure rates is realistic when he/she decides to implement formulae, Boolean or Markov calculations.

14.2.1.4 TA1-1: spurious failure analysis

When considering the spurious actions of a safety system it is important to first think about the safe failures of the sensor(s), logic solver(s) and final element(s). But, as already mentioned in [14.2.1.2](#) other potential spurious actions are inherent to the “de-energize to trip” philosophy, e.g. the losses of power (AC or DC electricity, hydraulic power, compressed air, etc.) or of the command-control links.

The failure of power supply and auxiliary systems impact a lot of systems and should be analysed as such. On the contrary, the safe failures of the wires, cables, pipes, etc. linking the various components of the safety system under study which are not necessarily included in the safe failures of these components, actually belong to the safety system. Therefore they should be addressed within the spurious failure

analysis. It is beyond of the scope of this Technical Report to analyse in detail the links between the components. Then, for the purpose of illustrating the use of the various approaches this will be modelled by a global safe failure rate $\lambda_{sf,deg\text{tt}}$ aggregating all the safe failures due to the rupture of the physical links between the components.

For the simple TA1-1 example, each safe failure will lead to a spurious action. Then the spurious action rate is given by: $\Lambda_{st} = \lambda_{sf,Psh} + \lambda_{sf,Ls} + \lambda_{sf,V} + \lambda_{sf,deg\text{tt}} = 3,84 \cdot 10^{-4} \text{ h}^{-1}$

According to 1.2.3 the average spurious failure frequency is given by

$$\bar{\Phi}_{st}(T) = \bar{\Phi}_{sf}(T) = \frac{N_{sf}(T)}{T} \approx \frac{1}{MTBF_{sf}} = \frac{1}{MTTF_{sf} + MTTR_{es_{sf}}}$$

where $MTBF_{sf}$ is then the mean time between spurious actions, $MTTF_{sf}$ the mean time to a spurious action and $MTTR_{es_{sf}}$ mean time to restore from a spurious action.

So:

- $MTTF_{sf} = 1 / 3,84 \cdot 10^{-4} = 2604 \text{ h}$
- $MTTR_{es_{sf}}$ is not known and remains to be evaluated.

As the safe failures trigger a spurious action, the fault detection time is equal to zero and $MTTR_{es_{sf}} \equiv MRT_{sf}$

MRT_{st} is not known but, on the hypothesis that:

- the individual MRT are exponentially distributed (i.e. the corresponding repair rates are constant) and,
- the components are independent (i.e. not repaired by the same repair team)

it can be approximated by: $MRT_{sf} \approx \frac{\sum_i \lambda_{sf,i} MRT_{sf,i}}{\sum_i \lambda_{sf,i}} = 34 \text{ h}.$

Then the average spurious failure frequency is: $\bar{\Phi}_{st}(T) \approx \frac{1}{MTTF_{sf} + MTTR_{es_{sf}}} = 3,8 \cdot 10^{-4}$ spurious failures per hour, e.g. 3,3 spurious actions per year.

The spurious failure frequency can also be calculated by using fault trees as shown in Figure 52. When the system is under repair no spurious failure can occur. Then when the unavailability of the system increases, the ratio of time spent under repair increases and the frequency of failure decreases. This explains the shape of the curve in Figure 52: as the system is in perfect state at $t = 0$ (unavailability = 0), the spurious failure frequency decreases from this point until an asymptotic value is reached. This asymptotic value $\Phi_{st,as} = 3,79 \cdot 10^{-4}$ is very close to the result found above with the formula.

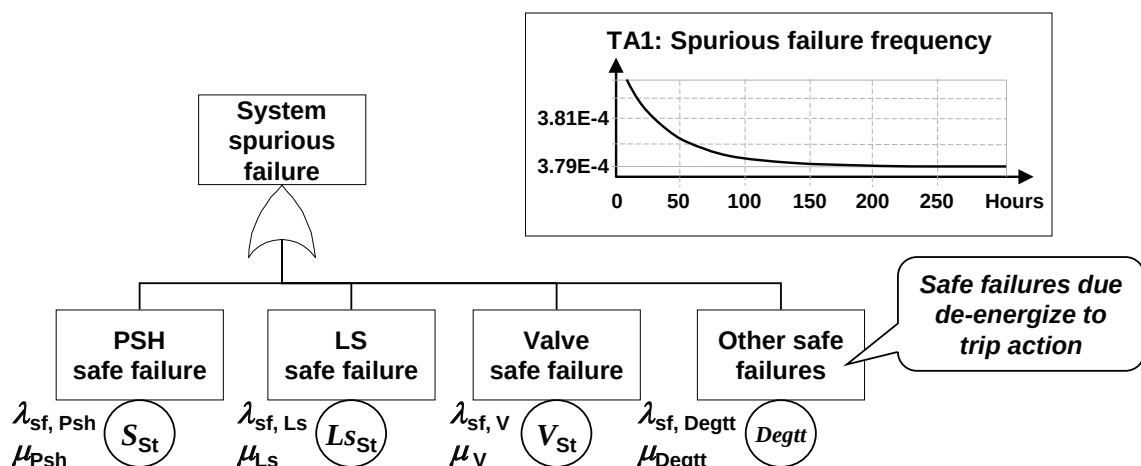


Figure 52 — TA1-1 Spurious failure frequency $\Phi_{st}(t)$ as function of the time

If the components were not independent with regard to the repair, the above calculations would not be relevant, mathematically speaking. Nevertheless as $1/\text{MTTF}_{\text{sf}} = 3,84 \cdot 10^{-4}$ provides a conservative value of the average spurious failure frequency in any case, it is not really useful to develop more sophisticated models here.

14.2.2 TA1-2: different periodical test intervals

14.2.2.1 TA1-2: description and assumptions

The assumptions of this typical example are exactly the same as TA1 except for the periodical tests of the components which are not performed with the same interval. In addition, two kinds of periodical tests are performed on the isolation valve:

- partial stroking tests to detect that the valve is able to move;
- full stroking tests to detect that the valve is tight after closure.

Table 9 — TA1-2 Reliability parameter

Parameter	Component		
	Pressure sensor	Isolation valve: Full stroking	Isolation valve: Partial stroking
Dangerous undetected failure rate	$\lambda_{\text{du, Psh}} = 3,0 \cdot 10^{-7} \text{ h}^{-1}$	$\lambda_{\text{du, Vfs}} = 4,4 \cdot 10^{-7} \text{ h}^{-1}$	$\lambda_{\text{du, Vps}} = 2,46 \cdot 10^{-6} \text{ h}^{-1}$
Periodical test interval	$\tau = 8\,760 \text{ h}$	$\tau_{\text{Vfs}} = 17\,520 \text{ h}$	$\tau_{\text{Vps}} = 4\,380 \text{ h}$

14.2.2.2 TA1-2: probabilistic calculations

14.2.2.2.1 Analytical formulae

According to [Table 9](#), the safety system may be in a dangerous fault state according to 3 failure modes. The average unavailabilities due to each of those failure modes can be calculated by a formula of the form $\lambda \cdot \tau / 2$ where λ is the dangerous undetected failure rate and τ the corresponding periodical test interval. This gives:

- failure due to PSH: $1,31 \cdot 10^{-3}$
- failure tested by full stroking: $3,85 \cdot 10^{-3}$
- failure tested by partial stroking: $5,39 \cdot 10^{-3}$

The components being in series, the above average unavailabilities can be added to evaluate the average unavailability of the whole safety system: $\bar{U}_{\text{du}}(0, T) = 1,06 \cdot 10^{-2}$.

The fault tree in [Figure 53](#) models the three dangerous undetected failure modes identified above. Due to the various periodical tests intervals the saw-tooth cure is more complicated than for TA1 but, again the result for the average unavailability over 10 years is very close to the result found by using the formulae:

Average unavailability: $\bar{U}_{\text{du}}(0, T) = 1,05 \cdot 10^{-2}$

14.2.2.2.2 Fault tree approach

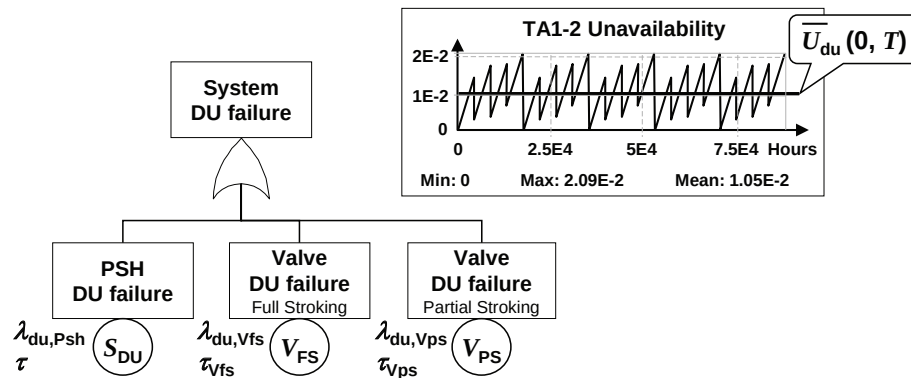


Figure 53 — Fault tree modelling of TA1-2

Using the same fault tree provides the average dangerous failure frequency (see Figure 54): $\overline{w}_{du}(0, T) = 3,17 \cdot 10^{-6}$.

NOTE The failure frequency decreases when the time elapse because a periodically tested component can fail only once during the test interval. Then when its probability to be failed increases its failure frequency decreases. When a test is performed the failures are repaired, then the component can fail again and the failure frequency increases.

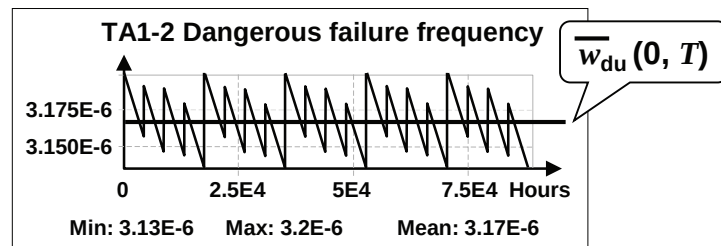


Figure 54 — Average failure frequency of TA1-2

14.2.2.2.3 Markovian approaches

The multi-phase Markov model implies 8 states and 4 different phases. This is not really complicated but not as simple as fault tree. As it is not really useful to detail this model, no further development are provided here but the result would be exactly the same as with fault trees.

14.2.2.2.4 Petri net and Monte Carlo simulation approach

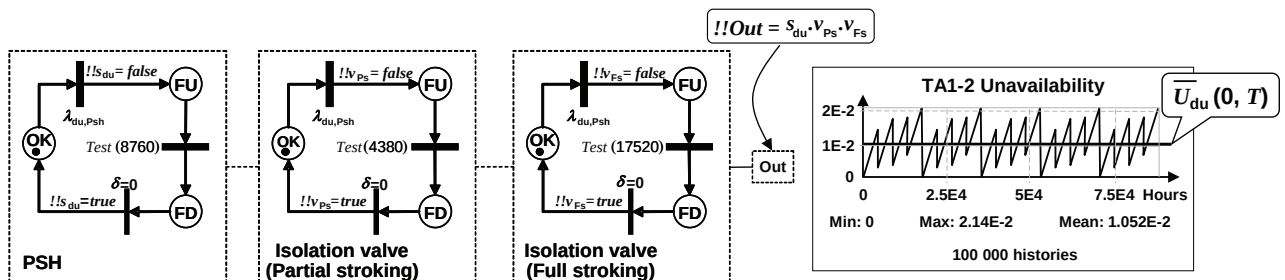


Figure 55 — TA1-2 PN modelling

Figure 55 gives the PN modelling the TA-2 safety system. This model is very easily derived from the previous model developed for TA1-1 as it uses the same basic modules PN1. It provides for the average

unavailability the same result as formulae and fault tree: $\bar{U}_{du}(0,T)=1,05 \cdot 10^{-2}$ with a 90 % confidence interval $[1,04 \cdot 10^{-2} - 1,07 \cdot 10^{-2}]$ and a standard deviation of $2,8 \cdot 10^{-2}$.

The same PN provides the average number of failure over 10 years: 0,278 [0,277 - 0,279]. This leads to $\bar{w}_{du}(0,T)=3,17 \cdot 10^{-6}$ $[3,16 \cdot 10^{-6} - 3,18 \cdot 10^{-6}]$. This is similar to the result found by fault tree.

14.2.3 TA1-3: protected installation running during tests and repairs

14.2.3.1 TA1-3: assumptions

The assumptions of this typical example are exactly the same as TA1 except that:

- The installation is not shut down during the repair of the sensor and of the logic solver.
- The sensor is periodically tested off line and is no longer available for its safety function during the periodical test. The periodical test duration lasts $\pi = 2$ h.

This implies that the risk does not disappear during the repair of the sensor or during the periodical test of the sensor. Therefore:

- the MRT of the sensor is considered for its detected and undetected dangerous failures;
- the MRT of the logic solver is considered for its detected dangerous failures;
- the duration of the periodical test interval of the sensor is considered.

14.2.3.2 TA1-3: probabilistic calculations

14.2.3.2.1 Analytical formulae

Compared to the TA1-1 calculations, the above assumptions add three unavailability terms:

- average unavailability during sensor repair: $(\lambda_{du,Psh} + \lambda_{dd,Psh}) \times MRT_{Psh} = 4,55 \cdot 10^{-4}$;
- average unavailability during logic solver repair: $\lambda_{dd,Ls} \times MRT_{Ls} = 9,00 \cdot 10^{-6}$;
- average unavailability during sensor periodical tests: $\pi / \tau = 2 / 8\,760 = 2,28 \cdot 10^{-4}$.

The sum of the three above terms is $6,92 \cdot 10^{-4}$.

Finally, this obtains: $\bar{U}(\tau) \approx 1,4 \cdot 10^{-2} + 6,92 \cdot 10^{-4} = 1,47 \cdot 10^{-2}$.

Therefore the impact of the new hypotheses is about 10 % of the average unavailability.

14.2.3.2.2 Fault tree approach

The unavailability laws used to calculate this fault tree are those which have been established in the Clause 9 dealing with the Markovian approach. The saw tooth curve shows the peaks of the unavailability due to the periodical tests of the sensors which inhibit the safety function when they are performed. Such peaks can be acceptable provided that compensating measures are implemented and that their duration is short.

Again, the fault tree gives results close to those found with the formulae: $\bar{U}(\tau) = 1,46 \cdot 10^{-2}$.

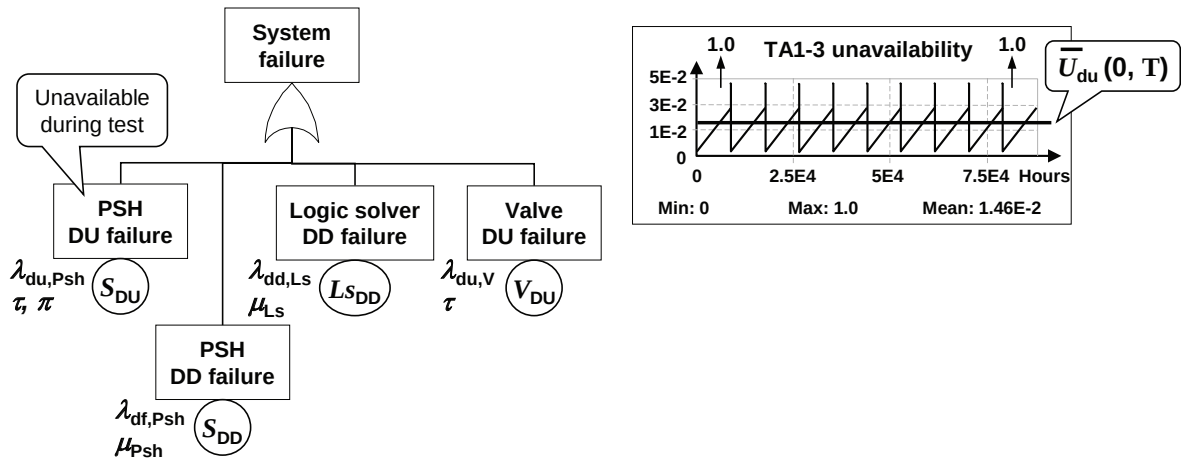


Figure 56 — Fault tree modelling of TA1-3

As shown in Figure 57, the fault tree gives: $\bar{w}(\tau) = 1,39 \cdot 10^{-4} ..$

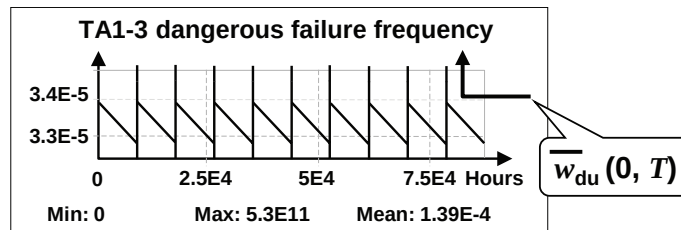


Figure 57 — Dangerous failure frequency of TA1-3

NOTE When a periodical test is performed the PSH jumps from the up to the down state and the number of system failures is instantaneously increased by 1. This is not a random phenomenon and the unavailability jump is equivalent to a Heaviside function whose derivative is a Dirac *distribution*. This introduces peaks which are just ignored when using ordinary integral calculation and a Lebesgue integral has to be used instead to calculate the failure frequency. This leads to tedious mathematical developments which are beyond the scope of this Technical Report. This problem should, nevertheless, not be ignored by the analyst. If the unavailability jumps are not properly handled, the number of failures will be strongly underestimated (by about 9 in the TA1-3 typical application).

14.2.3.2.3 Markovian approach

The Markov model implies about 16 states and two different recurrent phases. It is too big to be presented here.

It should be noted that the difficulty to handle the peaks to calculate the failure frequency arises also with the Markovian approach.

14.2.3.2.4 Petri net and Monte Carlo simulation approach

The PN model of the typical application TA1-3 is presented in Figure 58. This is the opportunity to introduce two new modules:

- Modelling of the periodical test duration (PN2 on the left hand side of Figure 58). Note that the modelling of periodical tests has been split: a dangerous failure has occurred on the right part of PN2 and no dangerous failure has occurred on the left part of PN2.
- Dangerous detected failures (PN3 in the middle of Figure 58).

In order to be consistent with the previous calculations the periodical test duration is modelled by a constant delay of 2 h and the repair time by a repair rate equal to $1/\text{MRT}$.

This model gives: $\bar{U}(0,T) = 1,46 \cdot 10^{-2}$ with a 90 % confidence interval $[1,4 \cdot 10^{-2} - 1,5 \cdot 10^{-2}]$ and a standard deviation of $3,1 \cdot 10^{-2}$. This is similar with the result obtained by fault tree.

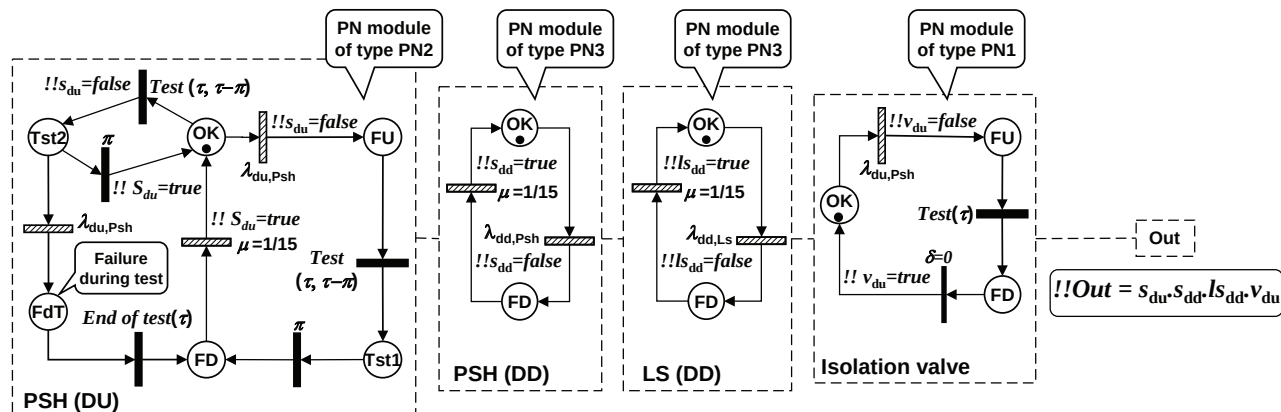


Figure 58 — PN modelling of TA1-3

The same PN provides also the average number of failure over 10 years: 12,63 $[12,62 - 12,64]$ i.e. $\bar{w}(0,T) = 1,442 \cdot 10^{-4}$ $[1,441 \cdot 10^{-4} - 1,443 \cdot 10^{-4}]$. This is similar but slightly greater than the $\bar{w}(\tau) = 1,39 \cdot 10^{-4}$ found with the fault tree approach.

It should be noted that the Monte Carlo evaluates the number of failures straightforwardly (just by a simple *counting*) when the fault tree approach needs rather complex analytical calculations: this explain this slight difference.

14.2.4 TA1-4: imperfect periodical test coverage

14.2.4.1 TA1-4: description and assumptions

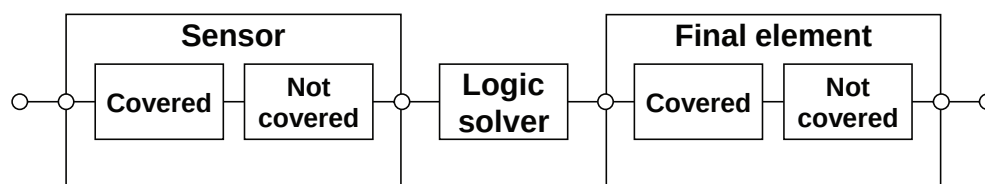


Figure 59 — Reliability block diagram for TA1-4

The assumptions of this typical example are exactly the same as TA1 except that the coverages of the periodical tests are not 100 %. The data in [Table 10](#) are used in the probabilistic calculations.

Table 10 — TA1-4 Reliability parameter

Parameter	Component			
	Pressure sensor (periodical test coverage 90 %)		Isolation valve (periodical test coverage 95 %)	
	Failures covered by periodical tests	Failures not covered by periodical tests	Failures covered by periodical tests	Failures not covered by periodical tests
Dangerous undetected failure rate	$\lambda_{du,Psh} = 2,7 \cdot 10^{-7} \text{ h}^{-1}$	$\lambda_{dnc,Psh} = 3 \cdot 10^{-8} \text{ h}^{-1}$	$\lambda_{du,V} = 2,76 \cdot 10^{-6} \text{ h}^{-1}$	$\lambda_{dnc,V} = 1,45 \cdot 10^{-7} \text{ h}^{-1}$

14.2.4.2 TA1-4: probabilistic calculations

14.2.4.2.1 Analytical formulae

The hypothesis that all the components are periodically tested at the same time allows gathering all the failures covered by the periodical tests within a macro component with the following parameters:

- Dangerous undetected failure rate: $\Lambda_{du} = \lambda_{du,Psh} + \lambda_{du,V} = 3,03 \cdot 10^{-6} \text{ h}^{-1}$
- Periodical test interval: $\tau = 8\,760 \text{ h}$

The failures which are not covered by periodical tests may be also gathered into a macro component. The unavailability due to these failures can be calculated in the same way as the periodically tested failures: if the failure is present at time T , then, in average it has occurred at $T/2$. Finally:

- Not covered dangerous failure rate: $\Lambda_{dnc} = \lambda_{dnc,Psh} + \lambda_{dnc,V} = 1,75 \cdot 10^{-7} \text{ h}^{-1}$.
- Time interval: $T = 87\,600 \text{ h}$ (i.e. 10 years).

Therefore the analytical formulae can be straightforwardly applied:

$$\text{Average unavailability: } \bar{U}_{du}(T) \approx \frac{\Lambda_{du} \cdot \tau}{2} + \frac{\Lambda_{dnc} \cdot T}{2} = 1,32 \cdot 10^{-2} + 7,67 \cdot 10^{-3} = 2,09 \cdot 10^{-2}.$$

This result may be compared with the average unavailability of $1,39 \cdot 10^{-2}$ found for TA-1. The imperfect periodical tests have a negative impact which here is of about 50 %. Contrarily to TA1, the average unavailability is not independent of the period T of utilization of the safety system and the negative impact increases when T increases.

Using the same reliability parameters as above gives an average unavailability $\bar{U}_{du}(0,T) = 2,08 \cdot 10^{-2}$ over 10 periodical test intervals of 8 760 h (i.e. $T = 10$ years). This is pretty close to the result obtained with the formula which, again, is slightly conservative.

Figure 60 shows how the unavailability evolves when the time is increasing. The probability of failure is greater at the end of the period of interest than at the beginning. This is more or less the same observation as with the Weibull law above: the average is not really a good indicator of the risk undergone by the operator at a given instant. This important aspect is completely ignored when dealing only with the analytical formulae.

14.2.4.2.2 Fault tree approach

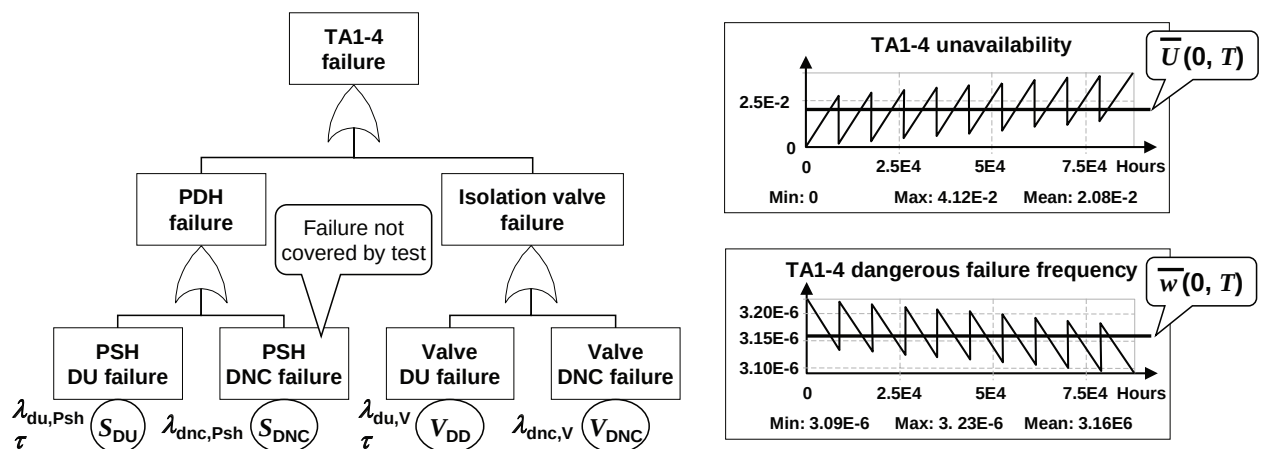


Figure 60 — Fault tree modelling of TA1-4

Figure 60 gives also $\bar{w}_{du}(0,T) = 3,158 \cdot 10^{-6}$ of TA1-4. Surprisingly it is similar to this of $3,156 \cdot 10^{-6}$ found for TA1 but as can be seen in Figure 60, this frequency decreases when T increases. This seems to be

paradoxical but the paradox disappears when realizing that the dangerous failures not covered by periodical tests can occur only once when those which are detected can occur several times over $[0, T]$. With our data, this effect is not perceptible yet.

14.2.4.2.3 Markovian approach

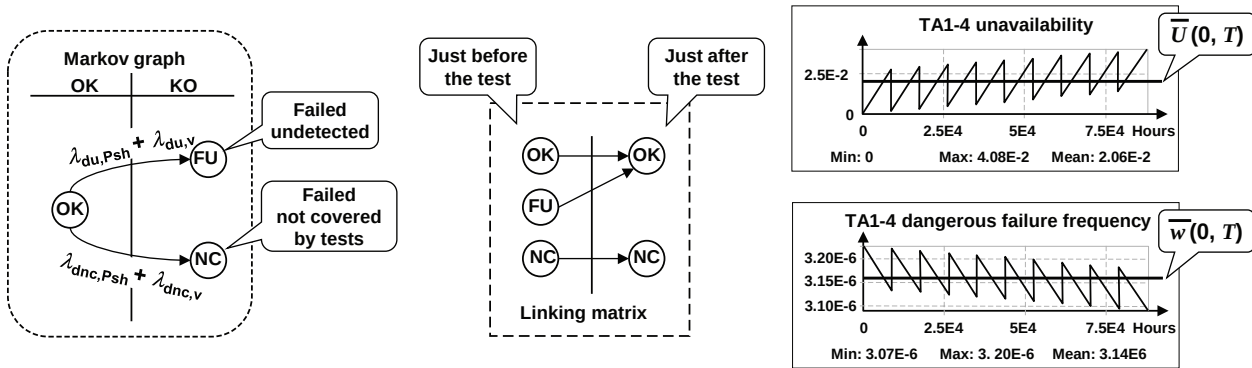


Figure 61 — Multi-phase Markov model of TA1-4

Figure 48 can be very easily modified to cope with the uncovered dangerous failures just by changing the Markov graph and the linking matrix as shown in Figure 61.

The results obtained are similar as those obtained by formulae or fault tree.

14.2.4.2.4 Petri net and Monte Carlo simulation approach

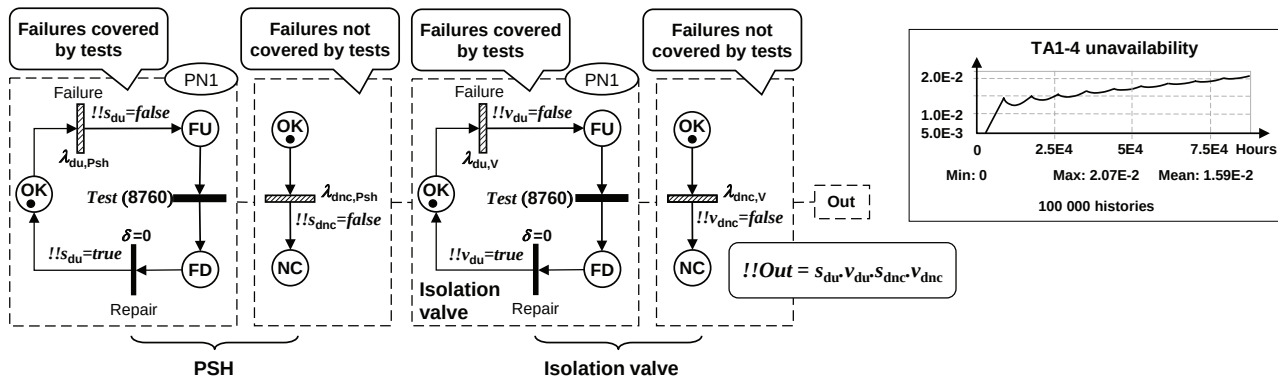


Figure 62 — Petri net modelling of TA1-4

The PN modelling the typical application TA1-4 is shown in Figure 62. This is the PN presented in Figure 49 where the failures not covered by the periodical tests (sub-PN with single transitions) have been added.

On the right hand side of the figure the average unavailability as a function of time is given. This average unavailability does not reach an asymptotic value. It increases permanently when t increases. The average unavailability for 10 years is given by the maximum value of this curve: $\bar{U}_{du}(0, T) = 2,06 \cdot 10^{-2}$.

More accurately the Monte Carlo simulation gives: $2,07 \cdot 10^{-2}$ with a 90 % confidence interval of $[2,03 \cdot 10^{-2} - 2,11 \cdot 10^{-2}]$. This is similar with the results previously found.

This model gives also 0,274 failures over 10 years with $[0,272 - 0,277]$ as 90 % confidence interval. This is equivalent to an average failure frequency $\bar{w}_{du}(0, T) = 3,13 \cdot 10^{-6}$ failure per hour with a 90 % confidence interval of $[3,1 \cdot 10^{-6} - 3,16 \cdot 10^{-6}]$. Again this is similar with the previous results.

14.3 Typical application TA2: dual channel

14.3.1 TA2-1: basic case

The typical application TA1-1 is not very reliable and it may be needed to decrease its probability of failure occurring on demand. This is generally achieved by introducing redundancy and a straightforward solution to do that is to duplicate the TA1-1 safety system in order to obtain the dual channel architecture shown in [Figure 63](#).

14.3.1.1 TA2-1: description and assumptions

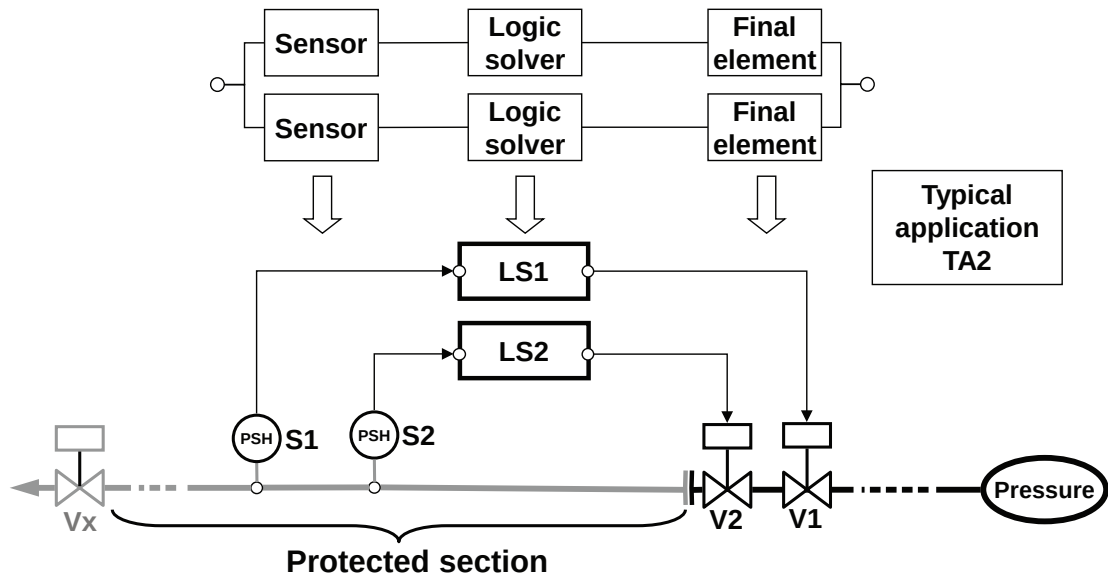


Figure 63 — Dual channel example TA2-1

The top of the [Figure 63](#) illustrates an architecture made of two channels running in parallel. Each of them is made on one single sensor, one single logic solver and one single final element. Such dual channel architecture is commonly used for moderate to high reliability requirements (SIL2 to SIL4).

This architecture is implemented at the bottom of [Figure 63](#):

Type of system: pressure protection system (see [Annex A](#)).

Mode of operation: low demand mode.

Layout: two channels running in parallel (1oo2 configuration). Each channel comprises one pressure sensor (S_i), one logic solver (LS_i) and one isolation valve (V_i) organized in series.

Design philosophy: de-energize to trip.

Hazardous event: over pressurization and burst of the low rating section of the installation.

Functioning description: when the pressure exceeds a given threshold, the sensors ($S1$ and/or $S2$) send a signal to their respective logic solvers ($S1$ or $S2$) which in turn commands their respective isolation valve ($V1$ or $V2$) to close.

NOTE In reality, the isolation valves are piloted by a solenoid valves which actually receives the signal from the logic solvers. This has not been represented here in order to simplify the example.

Assumptions (same as TA1-1):

- Periodical tests perfects and performed at the same time.

- Installation stopped during periodical tests and repair.
- Dangerous detected and undetected failures of a given component are independent.
- Constant dangerous failure rates.
- Components as good as new after repairs.

Table 11 — TA2 Reliability parameters

Parameter	Component			
	Pressure sensor	Logic solver	Isolation valve	De-energize to trip elements
Dangerous undetected failure rate	<i>Independent</i> $\lambda_{du,Psh} = 2,85 \cdot 10^{-7} \text{ h}^{-1}$	<i>NA</i>	<i>Independent</i> $\lambda_{du,V} = 2,755 \cdot 10^{-6} \text{ h}^{-1}$	
	CCF ($\beta = 5 \%$) $\kappa_{du,Psh} = 1,5 \cdot 10^{-8} \text{ h}^{-1}$		CCF ($\beta = 5 \%$) $\kappa_{du,V} = 1,45 \cdot 10^{-7} \text{ h}^{-1}$	
Dangerous detected failure rate	<i>Independent</i> $\lambda_{dd,Psh} = 2,85 \cdot 10^{-5} \text{ h}^{-1}$	<i>Independent</i> $\lambda_{dd,Ls} = 5,7 \cdot 10^{-7} \text{ h}^{-1}$	<i>NA</i>	
	CCF ($\beta = 5 \%$) $\kappa_{dd,Psh} = 1,5 \cdot 10^{-6} \text{ h}^{-1}$	CCF ($\beta = 5 \%$) $\kappa_{dd,Ls} = 3,0 \cdot 10^{-8} \text{ h}^{-1}$		
Periodical test interval	8 760 h	<i>NA</i>	8 760 h	
Safe failure rate	$\lambda_{sf,Psh} = 3,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{sf,Psh} = 3,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{sf,V} = 2,9 \cdot 10^{-4} \text{ h}^{-1}$	$\lambda_{sf,deggt} = 4,0 \cdot 10^{-6} \text{ h}^{-1}$ (one channel)
Maintenance time MRT or MTTRes	15 h	15 h	40 h	24 h

The values presented in [Table 11](#) are proposed only for illustration purposes.

Compared to the reliability data used for the TA1-1 typical application, the failure rates have been split between an independent part and a common cause failure part. A beta factor of 5 % (e.g. rather high) has been used to do that.

14.3.1.2 TA2-1: analysis

The analysis performed in [14.2.1.2](#) for TA1-1 (related to the de-energize to trip design and the disappearance of the risk during periodical tests and repairs) is also valid for TA2-1.

Finally with the above assumption:

- the TA2-1 safety system is unavailable when both channels are unavailable at the same time;
- only the fault detection times of the dangerous undetected failures is considered;
- a spurious failure occurs when one channel or the other produce a spurious closure of its own isolation valve.

14.3.1.3 TA2-1: probabilistic calculations

As this safety system operates on demand, its average unavailability (also called PFD_{avg} for a safety instrumented system) is the relevant parameter to calculate.

With the strong assumptions described above, any approach described in this Technical Report may be used.

14.3.1.3.1 Analytical formulae

The hypothesis that all the components are periodically tested at the same time allows gathering all the components within a macro component with the following parameters:

- Independent dangerous undetected failure rate: $\lambda_{du} = \lambda_{du,Ps} + \lambda_{du,V} = 3,04 \cdot 10^{-6} \text{ h}^{-1}$.
- CCF for dangerous undetected failure rate: $K_{du} = \kappa_{du,Ps} + \kappa_{du,V} = 1,6 \cdot 10^{-7} \text{ h}^{-1}$.
- Periodical test interval: $\tau = 8\,760 \text{ h}$.

Therefore with all the assumptions described above, this is the simplest case which can be encountered with redundant systems. The analytical formulae can be straightforwardly applied:

$$\text{Average unavailability: } \bar{U}_{du}(\tau) \approx \frac{\lambda_{du}^2 \tau^2}{3} + \frac{K_{du} \tau}{2} = 2,36 \cdot 10^{-4} + 7,01 \cdot 10^{-4} = 9,37 \cdot 10^{-4}.$$

14.3.1.3.2 Fault tree approach

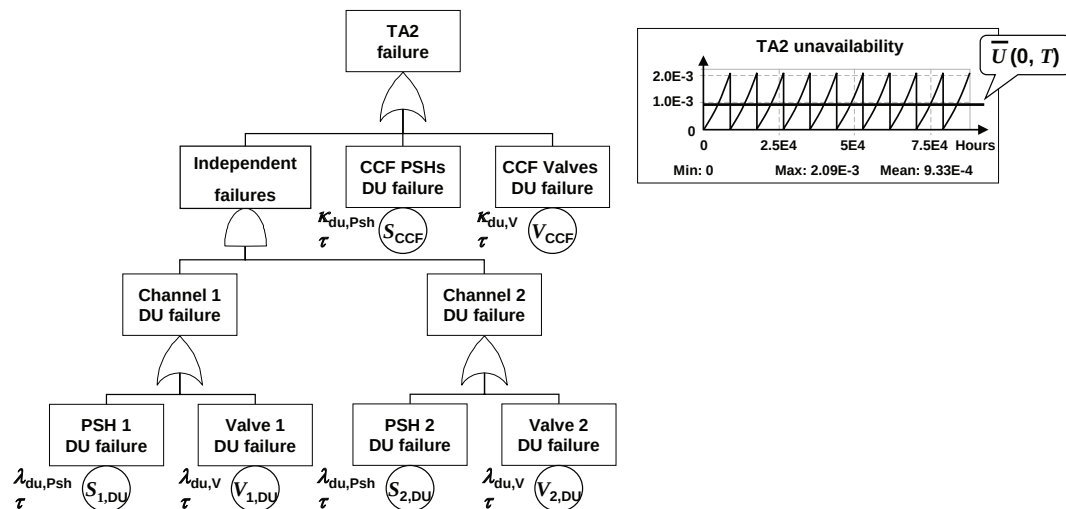


Figure 64 — Fault tree modelling of TA2

In the fault tree presented in Figure 64 two sub fault trees similar to the TA1-1 fault tree linked through a logical AND gate (modelling of the independent failures) and two leaves modelling the common cause failures between the redundant components can be recognized.

Using the same reliability parameters as above for calculating the average unavailability $\bar{U}_{du}(0,T)$ over 10 periodical test intervals of 8 760 h (i.e. $T = 10$ years) $\bar{U}_{du}(0,T) = 9,33 \cdot 10^{-4}$ is obtained. This result is pretty close to the result obtained with the formula.

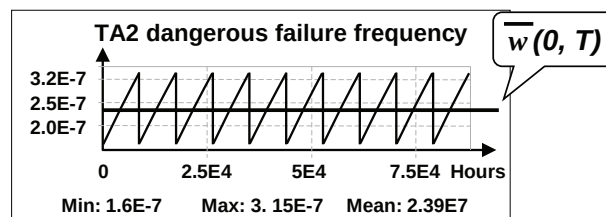


Figure 65 — Dangerous failure frequency of TA2-1

As shown in [Figure 65](#), using the same fault tree provides the average dangerous failure frequency: $\bar{w}_{du}(0,T)=2,39\ 10^{-7}$. It should be remarked that the shape of this curve is the opposite of this of [Figure 47](#). This is a side effect of the PSH redundancy (AND gate of the fault tree).

Of course, this model is not really interesting for such simple cases but its usefulness becomes evident when it is used for further calculations.

14.3.1.3.3 Multi-phase Markovian approach

With the strong assumptions adopted above, the number of potential states drastically decreases and it is even possible to build the multi-phase Markov graph presented in [Figure 66](#). With the same reliability parameters, it gives a result very close to the previous ones: $\bar{U}_{du}(0,T)=9,20\ 10^{-4}$.

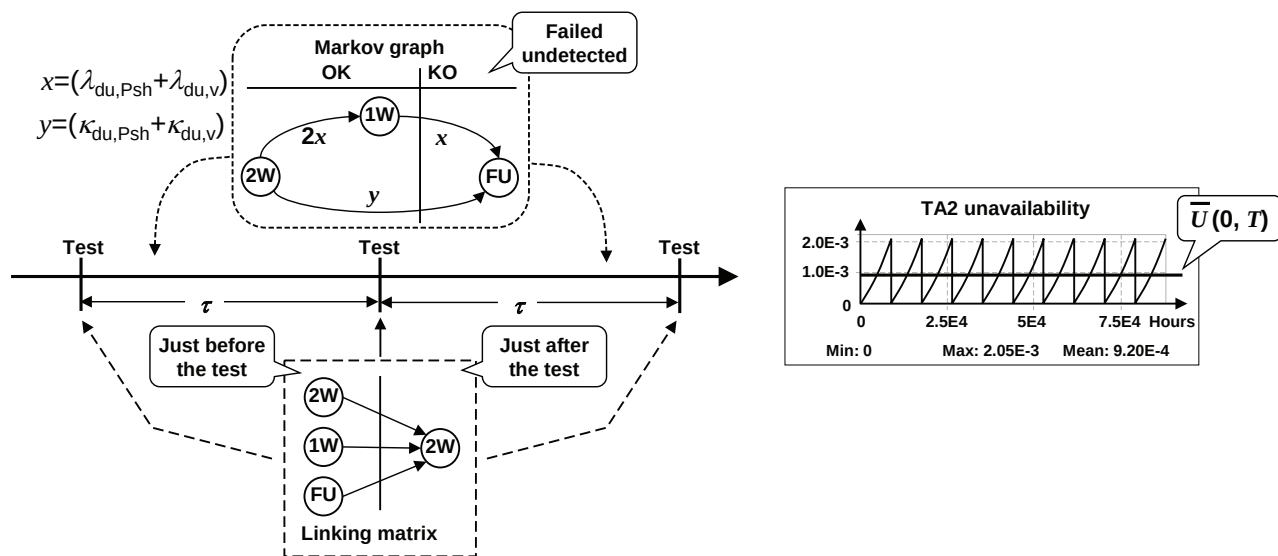


Figure 66 — Multi-phase Markov model of TA2-1

As shown in [Figure 67](#), using the same multi-phase Markov graph provides the average dangerous failure frequency: $\bar{w}_{du}(0,T)=2,34\ 10^{-7}$. The shape of the curve and the result are the same as with the fault tree.

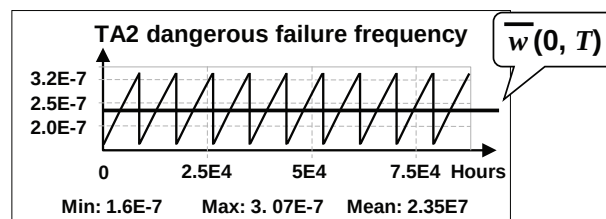


Figure 67 — Dangerous failure frequency of TA2-1

14.3.1.3.4 Petri net and Monte Carlo simulation approach

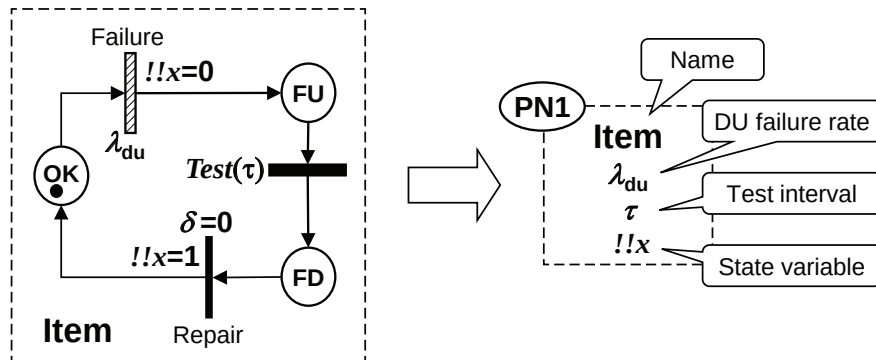


Figure 68 — Simplified presentation of the PN1 module (instantaneous repair)

The implementation of the RBD driven PN models (see 10.2) can be achieved through standardized sub-PN module: Figure 68 represent such a module named PN1 which has been already used in the previous subclauses. It models the behaviour of an item with periodically tested undetected failures and instantaneous repair when the failure has been discovered (i.e. the installation is stopped in order to suppress the risk during repair).

NOTE PN1 is a module usable even when the repair time is not equal to 0. In our case where $\delta = 0$, the two transitions “Test” and “repair” could have been aggregated in a single one.

In order to simplify the PN models presented hereafter, this module will be represented by a box containing its characteristics: name of the modelled item, undetected failure rate, periodical test interval and the state variable of the related item.

As shown in Figure 69 the TA2-1 example can be built just by using the above module. It highlights the modular structure of the PN modelling TA2 and the virtual reliability block diagram used as guideline to build the overall PN model. The systemic aspects of the functioning of TA1 are modelled through logical equations using the values of the state variables of the various components (in Figure 69 the “+” represents the logical “OR” and the “.” the logical “AND”). At the end the value of variable “Out” is calculated as function of the others: it is equal to 1 when the system is available and to 0 when it not available. Then its average value over 10 years gives the average availability $\bar{A}$ of TA1 and the complementary value $\bar{U} = 1 - \bar{A}$ gives the average unavailability.

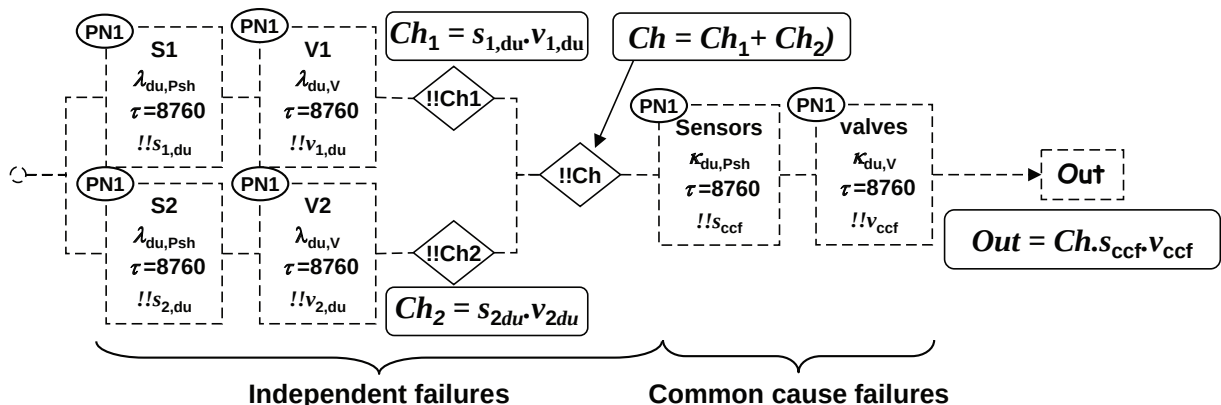


Figure 69 — Modular structure of the PN modelling TA2-1

The simulation of the PN embedded in this modular representation gives:

$\bar{U}_{du}(0, T) = 9,30 \cdot 10^{-4}$ with a 90 % confidence interval of $[8,9 \cdot 10^{-4} - 9,68 \cdot 10^{-4}]$ and a standard deviation of $7,6 \cdot 10^{-3}$.

This model provides also a number of failures of $2,1 \cdot 10^{-2}$ over the 10 years with a 90 % confidence interval of $[2,03 \cdot 10^{-2} - 2,18 \cdot 10^{-2}]$, i.e. an average failure frequency $\bar{w}_{du}(0,T)=2,41 \cdot 10^{-7}$ with a 90 % confidence interval of $[2,3 \cdot 10^{-7} - 2,49 \cdot 10^{-7}]$.

Again these results are similar to the previous ones obtained by the other approaches.

This model illustrates the use of the PN approach and Monte Carlo simulation on a simple application. Such model is useful when the assumptions are made more realistic.

14.3.1.4 TA2-1: spurious failure analysis

A safe failure of one of the channel is sufficient to trig a spurious safety action. Therefore the spurious failure rate is the double of the safe failure rate established for TA1:

Spurious action rate: $\Lambda_{st} = 2 (\lambda_{sf, Psh} + \lambda_{sf, Ls} + \lambda_{sf, V} + \lambda_{sf, degtt}) = 2 \times 3,84 \cdot 10^{-4} \text{ h}^{-1} = 7,68 \cdot 10^{-4}$.

According to 1.2.3 the average spurious failure frequency is given by

$$\bar{\Phi}_{st}(T) = \bar{\Phi}_{sf}(T) = \frac{N_{sf}(T)}{T} \approx \frac{1}{MTBF_{sf}} = \frac{1}{MTTF_{sf} + MTTR_{sf}}$$

where $MTBF_{sf}$ is then the mean time between spurious actions, $MTTF_{sf}$ the mean time to a spurious action and $MTTR_{sf}$ mean time to restore from a spurious action.

This gives:

- $MTTF_{sf} = 1 / 7,68 \cdot 10^{-4} = 1\,302 \text{ h}$
- $MTTR_{sf} = MRT_{sf}$ is the same as for TA1 (i.e. 34 h).

Then the average spurious failure frequency is: $\bar{\Phi}_{st}(T) \approx \frac{1}{MTTF_{sf} + MTTR_{sf}} = 7,48 \cdot 10^{-4}$ spurious failures per hour, i.e. 6,6 per year. It is the double of the spurious failure of TA1-1 and this highlight a well known phenomenon observed in reliability analysis: when the redundancy is increased to improve the probability of failure then the probability of spurious action is increased.

The spurious failure frequency can also be calculated by using fault trees and the result is as shown in Figure 70. This figure shows that the spurious failure frequency converge toward an asymptotic value $\Phi_{st,as} = 7,48 \cdot 10^{-4}$ which is the same as the result found above with the formula.

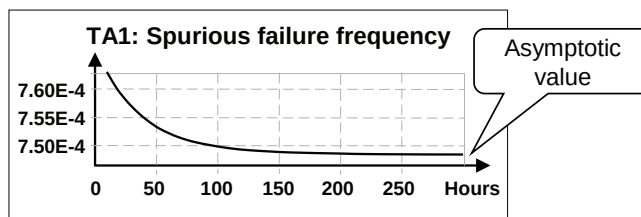


Figure 70 — TA2 Spurious failure frequency $\Phi_{sf}(t)$ as function of the time

14.3.2 TA2-2: different periodical test intervals

14.3.2.1 TA2-2: assumptions

The assumptions of this typical example are exactly the same as TA2-1 except for the periodical tests of the components which are not performed with the same interval. In addition, two kinds of periodical tests are performed on the isolation valves:

- partial stroking tests to detect that the valve is able to move;

- full stroking tests to detect that the valve is tight after closure.

Table 12 — TA2-2 Reliability parameter

Parameter	Component		
	Pressure sensor	Isolation valve: Full stroking	Isolation valve: Partial stroking
Dangerous undetected failure rate	<i>Independent</i> $\lambda_{du,Psh} = 2,85 \cdot 10^{-7} \text{ h}^{-1}$	<i>Independent</i> $\lambda_{du,Vfs} = 4,13 \cdot 10^{-7} \text{ h}^{-1}$	<i>Independent</i> $\lambda_{du,Vps} = 2,34 \cdot 10^{-6} \text{ h}^{-1}$
	CCF ($\beta = 5 \%$) $\kappa_{du,Psh} = 1,5 \cdot 10^{-8} \text{ h}^{-1}$	CCF ($\beta = 5 \%$) $\kappa_{du,Vfs} = 2,18 \cdot 10^{-8} \text{ h}^{-1}$	CCF ($\beta = 5 \%$) $\kappa_{du,Vps} = 1,2310^{-7} \text{ h}^{-1}$
Periodical test interval	8 760 h	17 520 h	4 380 h

14.3.2.2 TA2-2: probabilistic calculations

14.3.2.2.1 Analytical formulae

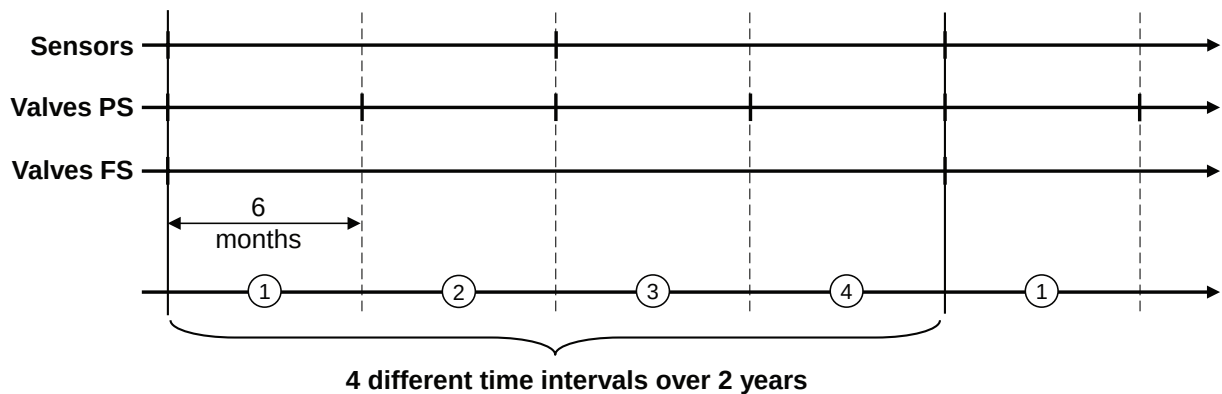


Figure 71 — Periodical test pattern for TA2-2

Figure 71 shows the various periodical tests performed to detect the failures of the sensors and of the valve (FS: full stroking and PS: partial stroking). It identifies 4 different types of periodical tests intervals.

Due to the different periodical tests intervals, it is no longer possible to aggregate the failures modes of a single channel as this has been done in 14.3.1.3.1. As explained in 6.2.3 in this case, it is still possible to use the minimal cut sets of the safety system.

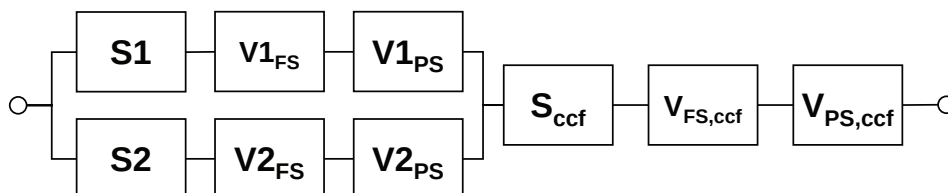


Figure 72 — Reliability block diagram modelling TA2-2

The minimal cut sets can be identified thanks to the reliability block diagram built in Figure 72 which represents the logics of TA2-2. If it is noted:

- S_i : dangerous failures of the sensors,
- $V_{i,FS}$: dangerous failures of the valves periodically tested by full stroking,

- $V_{i,PS}$: dangerous failures of the valves periodically tested by partial stroking and,
- S_{ccf} , $V_{FS,ccf}$ and $V_{PS,ccf}$: corresponding common cause failures between the channels,

then the following minimal cut sets are found which can be sorted into 3 categories:

- 1) Single failures: S_{ccf} , $V_{FS,ccf}$, $V_{PS,ccf}$
- 2) Double failures periodically tested at the same time: $S_1.S_2$, $V_{1,FS}.V_{2,FS}$, $V_{1,PS}.V_{2,PS}$
- 3) Double failures not periodically tested at the same time:
 - $S_1.V_{2,FS}$, $S_2.V_{1,FS}$
 - $S_1.V_{2,PS}$, $S_2.V_{1,PS}$
 - $V_{1,FS}.V_{2,PS}$, $V_{1,PS}.V_{2,FS}$

There is no problem to handle minimal cut sets of categories 1 and 2: this may be done in the similar way that this has been done above. This is more difficult for the minimal cut sets of the category number 3. For the minimal cut set $S_1.V_{2,PS}$, for example, 4 different periods of times (see [Figure 71](#)) can be identified:

- 1) Periodical test of V2 (partial stroking) and periodical test of S1 and
- 2) Periodical test of V2 (partial stroking) and S1 not tested since 3 months
- 3) Periodical test of V2 (partial stroking) and S1 not tested since 6 months
- 4) Periodical test of V2 (partial stroking) and S1 not tested since 9 months

The calculation $S_2.V_{1,PS}$ implies also 4 intervals, the calculation of $S_1.V_{2,FS}$ or $S_2.V_{1,FS}$ implies 2 intervals each and the calculation of $V_{1,FS}.V_{2,PS}$ or $V_{1,PS}.V_{2,FS}$ implies the calculations of 8 intervals each.

Finally, the exact calculation implies to calculate 34 terms ($3 + 3 + 4 + 4 + 2 + 2 + 8 + 8$) to evaluate the average unavailability of TA2-2. Formulae can be established for each of these intervals (see [7.4.2](#)). This is not complex but very tedious.

Therefore more simplified approximated calculations are often used. Unfortunately, as shown hereafter, they lead to non conservative results.

Non conservative approximation 1:

- independent failure of one channel:

$$u = \frac{\lambda_{du,Psh} \cdot 8760}{2} + \frac{\lambda_{du,Vfs} \cdot 17520}{2} + \frac{\lambda_{du,Vps} \cdot 4380}{2} = 9.99710^{-3}$$

- Approximation for independent failures dual channel: $\bar{U}_{ind1} = u^2 = 9.9910^{-5}$
- Common cause failures: $c = \frac{\kappa_{du,Psh} \cdot 8760}{2} + \frac{\kappa_{du,Vfs} \cdot 17520}{2} + \frac{\kappa_{du,Vss} \cdot 4380}{2} = 5.2610^{-4}$

This lead to: $\bar{U}_{approx1} = \bar{U}_{ind1} + c = 6.2610^{-4}$.

Non conservative approximation 2:

- Double failure of sensors: $u_1 = \frac{(\lambda_{du,Psh} \cdot 8760)^2}{3} = 2.0810^{-6}$
- Double failure of valves (full stroking): $u_2 = \frac{(\lambda_{du,Vfs} \cdot 17520)^2}{3} = 1.7510^{-5}$
- Double failure of valves (partial stroking): $u_3 = \frac{(\lambda_{du,Vps} \cdot 4380)^2}{3} = 3.5110^{-5}$
- Approximation for independent failures dual channel: $\bar{U}_{ind2} = u_1 + u_2 + u_3 = 5.4610^{-5}$

- Common cause failures: $c = \frac{\kappa_{du,Psh} \cdot 8760}{2} + \frac{\kappa_{du,Vfs} \cdot 17520}{2} + \frac{\kappa_{du,Vss} \cdot 4380}{2} = 5.26 \cdot 10^{-4}$

This lead to: $\bar{U}_{approx2} = \bar{U}_{ind2} + c = 5.81 \cdot 10^{-4}$

The accuracy of $\bar{U}_{approx1}$ and $\bar{U}_{approx2}$ is analysed in the next subclause where exact calculations are made with a fault tree. The user of this Technical Report interested to make exact calculations by using formulae is invited to implement the formulae developed in [Clause 7](#) and [Annex I](#).

14.3.2.2.2 Fault tree approach

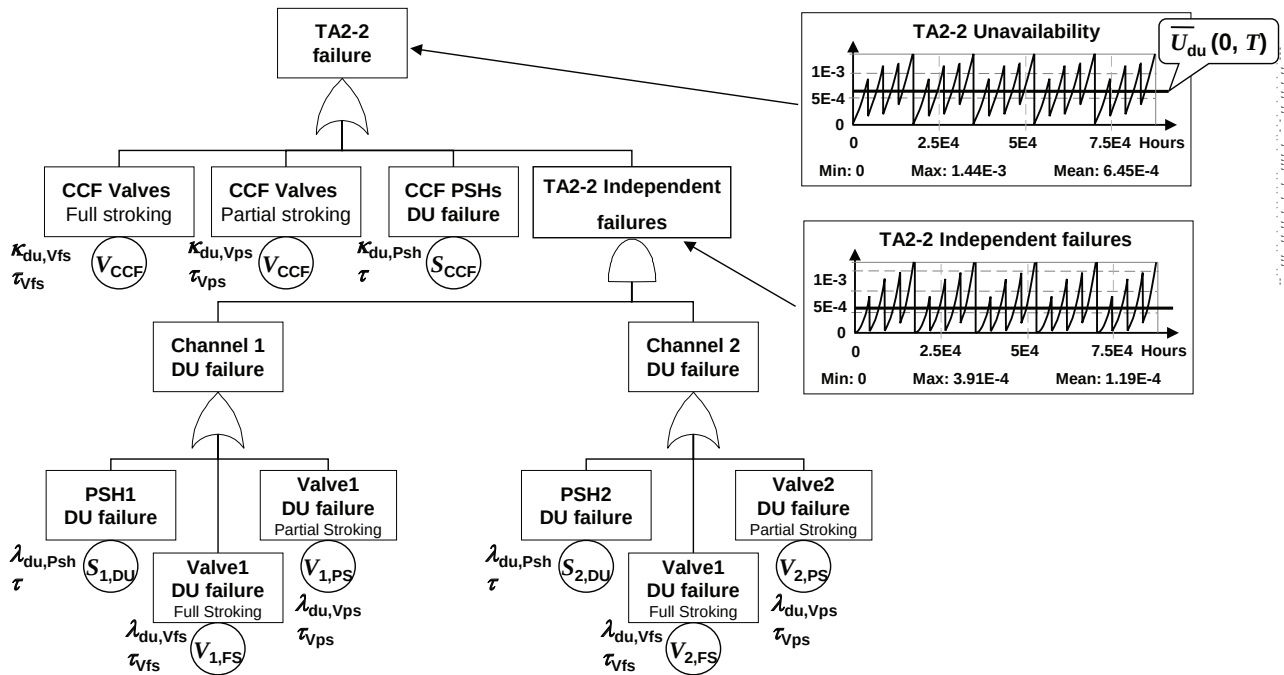


Figure 73 — Fault tree modelling of TA2-2

The fault tree in [Figure 73](#) models the independent failures and the common cause failures of the TA2-2 typical application. The result for the average unavailability over 10 years is the following:

Average unavailability: $\bar{U}_{du}(0,T) = 6.45 \cdot 10^{-4}$

This is 3 % higher than $\bar{U}_{approx1}$ and 10 % of $\bar{U}_{approx2}$. A quick conclusion could be that, even not conservative, the approximations are not so bad. But, looking at the average unavailability due to the independent failure it is found that the exact result ($1.19 \cdot 10^{-4}$) is 16 % higher than $\bar{U}_{ind1}$ and 54 % of $\bar{U}_{ind2}$.

Therefore the approximations seem valid only because the very bad calculation of the independent failure is hidden behind the common cause failures. This problem would be even more important if the beta factor has been taken to 1 % where the exact result would have been 9 % higher than $\bar{U}_{approx1}$ and 30 % of $\bar{U}_{approx2}$.

14.3.2.2.3 Markovian approaches

The multi-phase Markov model implies $2^9 = 512$ states and 8 different phases. Even if some states can be aggregated thanks to the symmetries, it is too big to be done by hand.

14.3.2.2.4 Petri net and Monte Carlo simulation approach

The [Figure 74](#) gives the layout of a PN modelling the TA2-2 safety system. It implements the basic module presented in [Figure 68](#) and it may be easily derived from [Figure 69](#).

It provides for the average unavailability the same result as fault tree: $\bar{U}_{du}(0,T) = 6.46 \cdot 10^{-4}$ with a 90 % confidence interval $[6,35 \cdot 10^{-4} - 6,57 \cdot 10^{-4}]$ and a standard deviation of $6,7 \cdot 10^{-3}$.

It also provides a result the average unavailability due to the independent failures: $1,21 \cdot 10^{-4}$ with a 90 % confidence level $[1,17 \cdot 10^{-4} - 1,24 \cdot 10^{-4}]$ and a standard deviation of $2,31 \cdot 10^{-3}$. This is very close to the result obtained by fault tree.

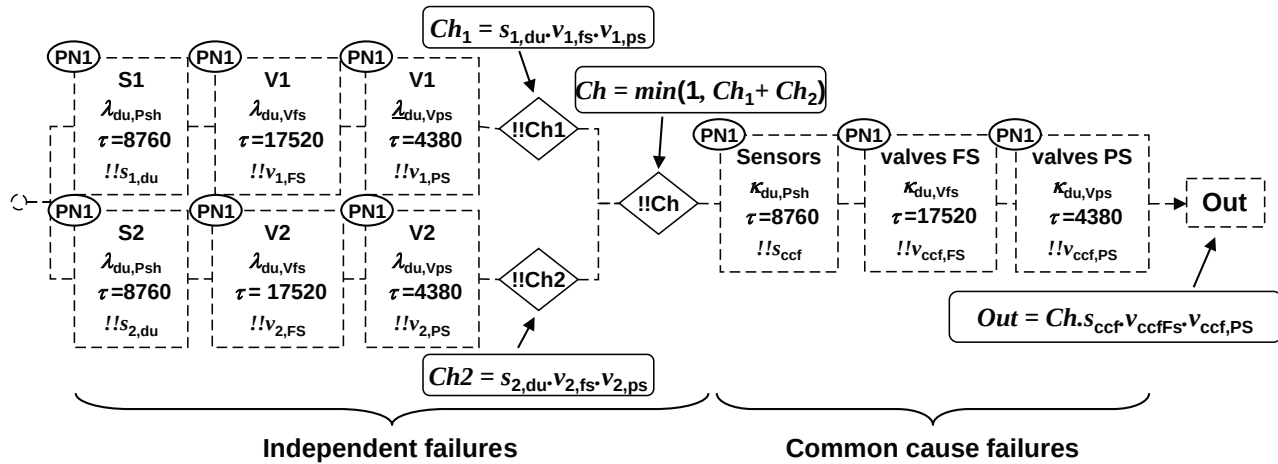


Figure 74 — TA2-2 PN modelling

14.3.3 TA2-3: protected installation running during tests and repairs

14.3.3.1 TA2-3: assumptions

The assumptions of this typical example are exactly the same as TA2-1 except that:

- the installation is not shut down during the repair of the sensors and of the logic solver.
- the sensors are periodically tested off line and are no longer available for their safety function during the periodical test. The periodical test duration lasts 2 h.

This implies that the risk does not disappear during the repairs of sensors and logic solvers or during the periodical tests of the sensor. Therefore:

- the MRTs of the sensors is considered for their detected and undetected dangerous failures,
- the MRTs of the logic solvers is considered for their detected dangerous failures,
- the duration of the periodical test of the sensors is considered.

14.3.3.2 TA2-3: probabilistic calculations

14.3.3.2.1 Analytical formulae

Taking the above assumption into account leads to calculations more complicated than in 14.3.2.2.1. This introduces more terms due to the periodical test duration of the sensors and the repairs of dangerous detected and undetected failures (see [Clause 7](#) and [Annex I](#)).

Therefore there has been no attempt to implement analytical formulae for handling the TA2-3 typical application.

14.3.3.2.2 Fault tree approach

The unavailability laws used to calculate this fault tree are those which have been established in [Clause 9](#) dealing with the Markovian approach.

The saw tooth curve shows the peaks of the unavailability due to the periodical tests of the sensors which inhibit the safety function during 2 h each time that they are performed. The acceptability of such peak depends of the possible compensating measures which may be implemented and that of their duration.

This fault tree gives the average unavailability: $\bar{U}(0,T)=1.19 \cdot 10^{-3}$.

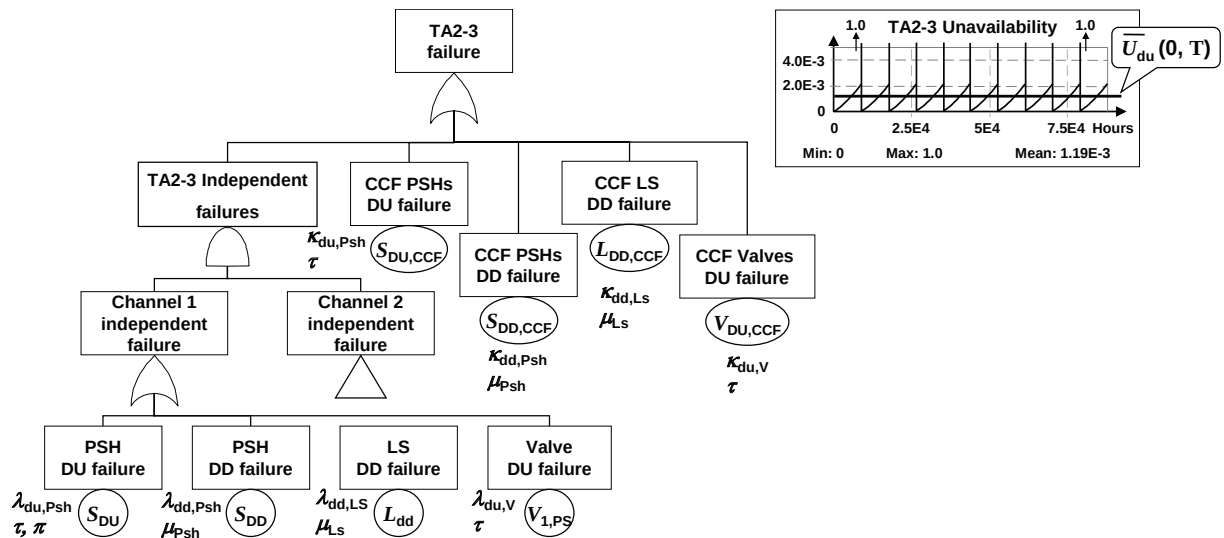


Figure 75 — Fault tree modelling of TA2-3

14.3.3.2.3 Consecutive periodical tests

As shown in [Figure 75](#) the peaks go to 1. This is due to the hypothesis that the periodical tests of the sensors are performed at the same time. In fact this hypothesis is not realistic as in reality they are likely to be done one after the other both from practical performance point of view and to avoid to completely inhibit the safety system during those periodical tests. Therefore the periodical tests of the sensors S1 and S2 are staggered in order to be consecutively performed as shown in [Figure 76](#).

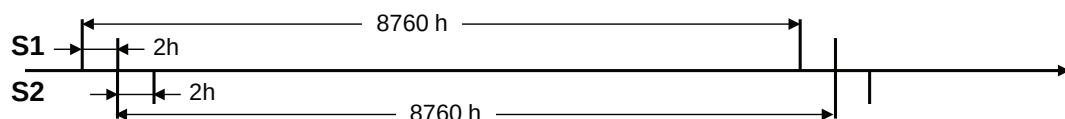


Figure 76 — Consecutive periodical tests of the sensors

This slight modification in the model leads to the results illustrated in [Figure 77](#).

The average unavailability has been improved to $\bar{U}(0,T)=9.68 \cdot 10^{-4}$. In addition, the peaks toward 1 have disappeared. They now reach only $2,96 \cdot 10^{-3}$ and this is far more satisfactory than in the previous case.

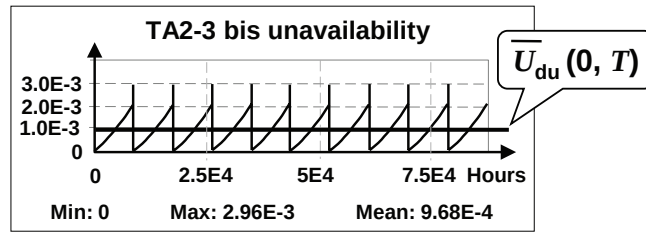


Figure 77 — Effect of consecutive periodical test performance on unavailability peaks

The Figure 78 shows the maximum staggering which is possible when using two redundant items. It slightly improves the average unavailability: $\bar{U}(0, T) = 9.29 \cdot 10^{-4}$.

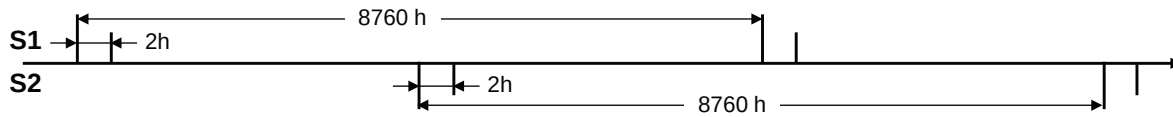


Figure 78 — maximum staggering of the periodical tests of the sensors

14.3.3.2.4 Petri net and Monte Carlo simulation approach

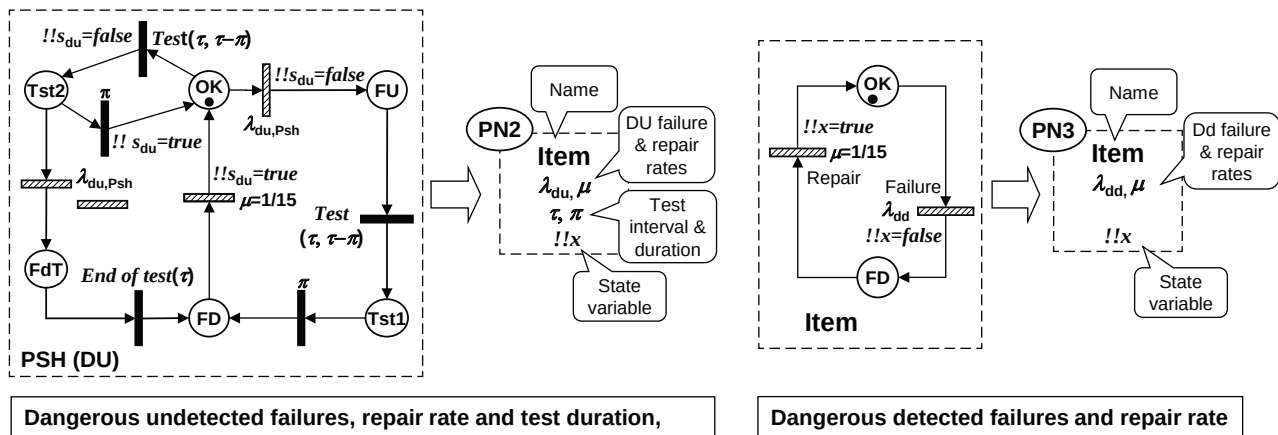


Figure 79 — Sub-PN needed to model TA2-3

Figure 79 represent the two modules which have already been introduced in Figure 56. They can be used for building the PN modelling TA2-3

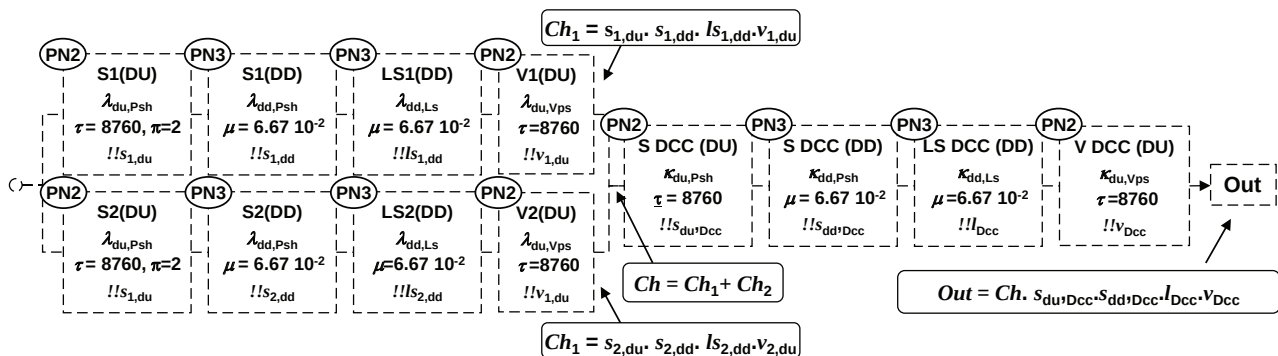


Figure 80 — TA2-3 PN modelling

This model gives: $\bar{U}(0,T)=1.19 \cdot 10^{-3}$ with a 90 % confidence interval $[1,15 \cdot 10^{-3} - 1,23 \cdot 10^{-3}]$ and a standard deviation of $7,8 \cdot 10^{-3}$. This is similar with the result obtained by fault tree in [14.3.3.2.2](#).

The same PN as in [Figure 80](#) may be used to evaluate the average unavailability when the periodical tests of the sensors are performed consecutively instead of simultaneously. This may be done just using a law including a value 8 762 before performing the 1st periodical test of S2 (and keeping a test interval of 8 760 h).

14.3.4 TA2-4: Test staggering

14.3.4.1 TA2-4: assumptions

The assumptions of this typical example are exactly the same as TA2-1 with, in addition:

- The periodical tests of the sensors are staggered (see [Figure 78](#)) and S2 is periodically tested in the middle of the periodical test interval of S1.
- The periodical tests of the valves are staggered and V2 is periodically tested in the middle of the periodical test interval of valve V1.
- Each periodical test of a component is an opportunity to detect the related potential common cause failures.

This last assumption is valid only if the relevant periodical test procedures are implemented.

14.3.4.2 TA2-4: probabilistic calculations

14.3.4.2.1 Analytical formulae and Markovian approach

Staggering the periodical tests increases the number of intervals to be considered and therefore using the formulae or the Markovian approaches is even more difficult than above.

14.3.4.2.2 Fault tree modelling

The fault tree is exactly the same as this presented in [Figure 64](#). Only the laws of the dangerous failures of the sensors and of the valves are slightly modified to take the periodical tests staggering into account.

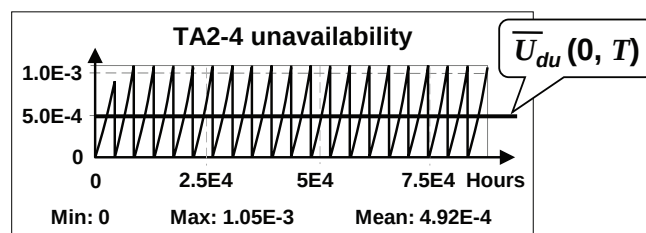


Figure 81 — Effect of periodical test staggering

As shown in [Figure 81](#) the effect of periodical test staggering is very beneficial on the average unavailability. In this example there is a decrease of 47 % of the average unavailability from $\bar{U}(0,T)=9.33 \cdot 10^{-4}$ for TA2 to $\bar{U}(0,T)=4.9 \cdot 10^{-4}$ for TA2-4.

This decrease comes from the decrease of the average unavailability due to the independent failures (see [1.3.2.3](#)) and also of the periodical test frequency of common mode failures which has been multiplied by 2.

The maximum value of the instantaneous unavailability has been divided by 2 from $2,1 \cdot 10^{-3}$ to $1,0 \cdot 10^{-3}$. From SIL analysis point of view, the PFD_{avg} of TA2 as well as TA2-4 are in the range of SIL3 applications but the instantaneous risk for the operators just before performing the periodical tests is lower for TA2-4 than TA2. In addition the instantaneous unavailability of TA-2 remains lower than the SIL3 upper bound all the time what is not the case for TA2: TA2-4 is a so-called *permanent* SIL3 safety system.

14.3.4.2.3 Petri net and Monte Carlo simulation approach

The same PN developed in [Figure 69](#) for TA2 can be used for calculating the unavailability of TA2-4. Only the adjustment of the laws modelling the periodical tests are needed to cope with the new assumptions.

This model gives: $\bar{u}(0,T) = 4,94 \cdot 10^{-4}$ with a 90 % confidence interval $[4,72 \cdot 10^{-4} - 5,15 \cdot 10^{-4}]$ and a standard deviation of $4,0 \cdot 10^{-3}$. This is similar with the result obtained by fault tree.

14.4 Typical application TA3: popular redundant architecture

14.4.1 TA3-1: basic case

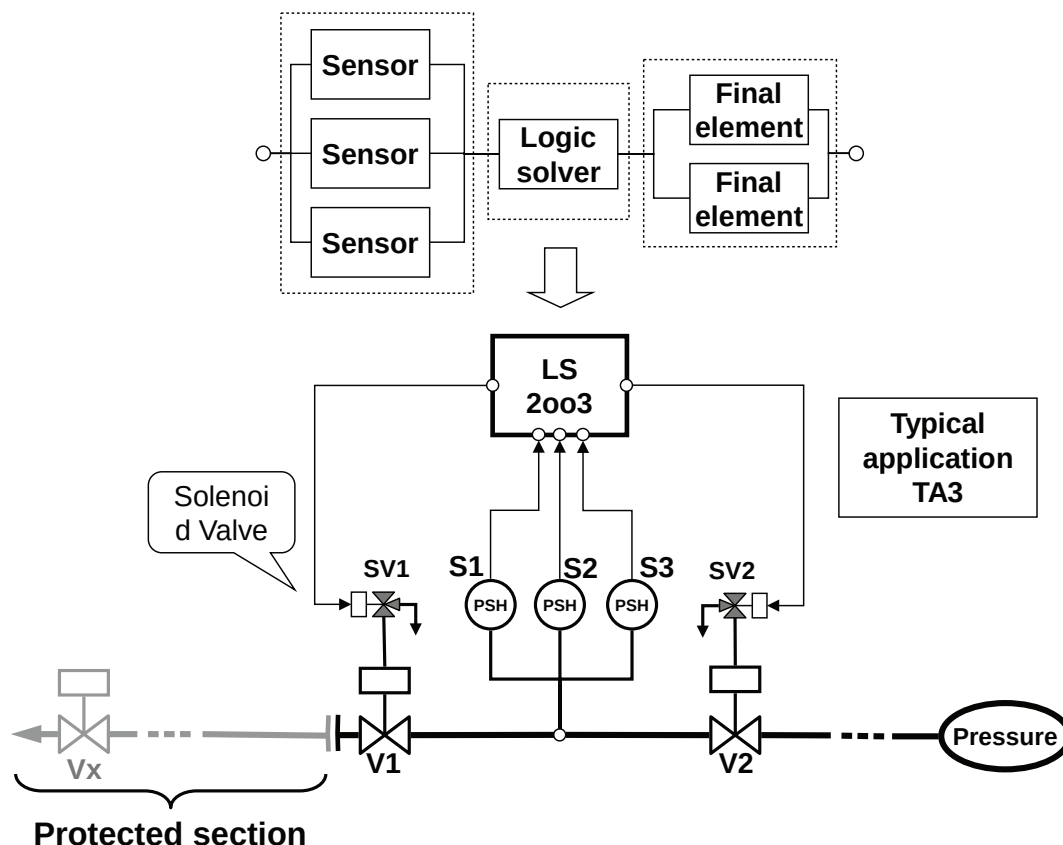


Figure 82 — Typical application TA3-1

The RBD at the top of [Figure 82](#) illustrates a popular architecture of safety instrumented system made of three parts (subsystems) organized in series: three sensor working in a 2oo3 configuration, one single logic solver and two single final element.

Such architecture is used in process industry for common safety loops with moderate to high reliability requirements (SIL3 to SIL4). This may be a HIPPS architecture. A typical safety system implementing this architecture is the pressure protection system represented at the bottom of this [Figure 82](#):

Type of system: pressure protection system (cf. [Table A.1](#)).

Mode of operation: low demand mode.

Layout: three pressure sensors organized in 2oo3 (S1, S2, S3), one logic solver (LS) and two isolation valve (V1, V2) organized in series and piloted by two solenoid valves (SV1, SV2).

Design philosophy: de-energize to trip.

Hazardous event: over pressurization and burst of the low rating section of the installation.

Functioning description: when the pressure exceeds a given threshold, the sensors send signals to the logic solver. When the logic solver receives at least two signals, then it commands the solenoid valves to open. This releases the hydraulic pressure maintaining the isolation valves open and they close under the action of springs.

The role of this system is the same as the typical application TA1: protect a low rating section of pipe on the left hand side of the figure against overpressure coming from the pressure source represented on the right hand side. Contrarily to TA1 and as in reality, the solenoid valves which actuate the isolation valves have been represented.

The sensors impulse nozzles have been located between the isolation valves to avoid drilling holes in the weak section of the pipe. This does not change the functioning of the system with regard to its safety function.

14.4.1.1 TA3-1: assumptions

Assumptions:

- Periodical tests perfects and performed at the same time.
- Installation stopped during periodical tests and repair.
- Dangerous detected and undetected failures of a given component are independent
- Constant dangerous failure rates.
- Components as good as new after repairs.

These assumptions are the same as those adopted for TA1.

Table 13 — Reliability parameters

Parameter	Component			
	Pressure sensor	Logic solver	Isolation valve	Solenoid valve
Dangerous undetected failure rate	$\lambda_{du,Psh} = 3,0 \cdot 10^{-7} \text{ h}^{-1}$	NA	$\lambda_{du,V} = 2,6 \cdot 10^{-6} \text{ h}^{-1}$	$\lambda_{du,SV} = 8 \cdot 10^{-7} \text{ h}^{-1}$
Dangerous detected failure rate	$\lambda_{dd,Psh} = 3,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{dd,LS} = 6,0 \cdot 10^{-7} \text{ h}^{-1}$	NA	NA
Periodical test interval (τ)	8 760 h	NA	8760 h	8760 h
Safe failure rate	$\lambda_{sf,Psh} = 3,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{sf,LS} = 6,0 \cdot 10^{-5} \text{ h}^{-1}$	$\lambda_{sf,V} = 2,6 \cdot 10^{-4} \text{ h}^{-1}$	$\lambda_{sf,SV} = 8 \cdot 10^{-5} \text{ h}^{-1}$
Maintenance time MRT or MTTRes	15 h	15 h	48 h	20 h

The values presented in this table are proposed only for illustration purpose. Compared to the reliability data used for TA-1 the data for the isolation valves has been split between the isolation valve itself and its pilot valve.

14.4.1.2 TA3-1: analysis

The analysis made for TA1-1 and developed in [14.2.1.2](#) is also valid for TA3-1 and is not re-copied here.

14.4.1.3 TA3-1: probabilistic calculations

As this safety system operates on demand, its average unavailability (also called PFD_{avg} for a safety instrumented system) is the relevant parameter to calculate. With the above simplifying assumptions this is a very simple case and any approach described in this Technical Report may be used.

14.4.1.3.1 Analytical formulae

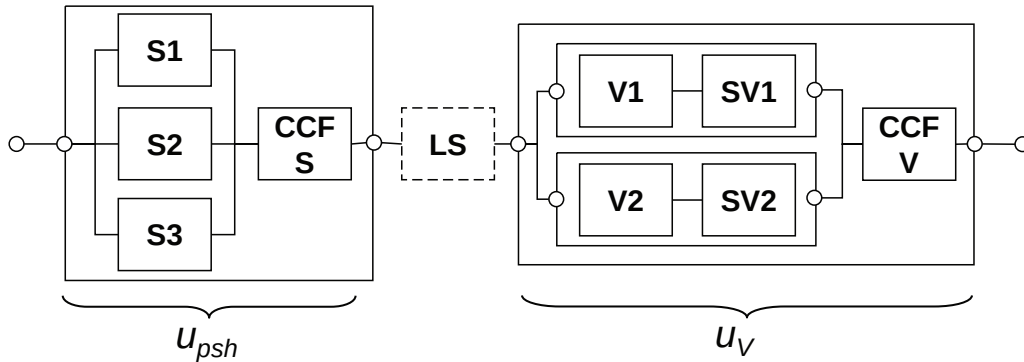


Figure 83 — Equivalent reliability block diagram of TA3-1

Contrarily to the typical application TA2, the typical application TA3-1 can be split between 3 independent parts in series (Figure 83).

With the above assumptions, only the dangerous undetected failures necessitate to be considered. Therefore the average unavailability can be approximated by the sum of the average unavailability of the sensors and of the valves as shown in Figure 83.

According the Figure 83 the average unavailability due to the sensors is made of two terms (independent and common cause failure). Applying the formulae for double and single failures (see Clause 7 and Annex I) with a beta factor of 5 % leads to:

$$u_{psh} = 3 \frac{[\lambda_{du,psh}(1-\beta)\tau]^2}{3} + \frac{\lambda_{du,psh}\beta\tau}{2} = 7.19 \cdot 10^{-5}$$

According to Figure 83 the average unavailability due to the valves is also made of two terms (independent and common cause failure) but the failure rates of the isolation valve and of its related solenoid valve have to be aggregated first. This aggregation is possible as both are periodically tested at the same time. If $\lambda_{du,Vg} = \lambda_{du,V} + \lambda_{du,SV}$:

$$u_V = \frac{[(\lambda_{du,Vg} + (1-\beta)\tau)^2]}{3} + \frac{\lambda_{du,Vg}\beta\tau}{2} = 8.29 \cdot 10^{-4}$$

Finally $\bar{U}_{du}(\tau) = 9.0 \cdot 10^{-4}$ is obtained for TA3-1.

14.4.1.3.2 Fault tree approach

Figure 84 models the typical application TA3 and: $\bar{U}_{du}(\tau) = 9.0 \cdot 10^{-4}$.

Again, the interest of such a fault tree is that it can be used as it is to for other calculations (e.g. periodical test staggering, periodical test durations, etc.) difficult (or impossible) to perform just by using formulae.

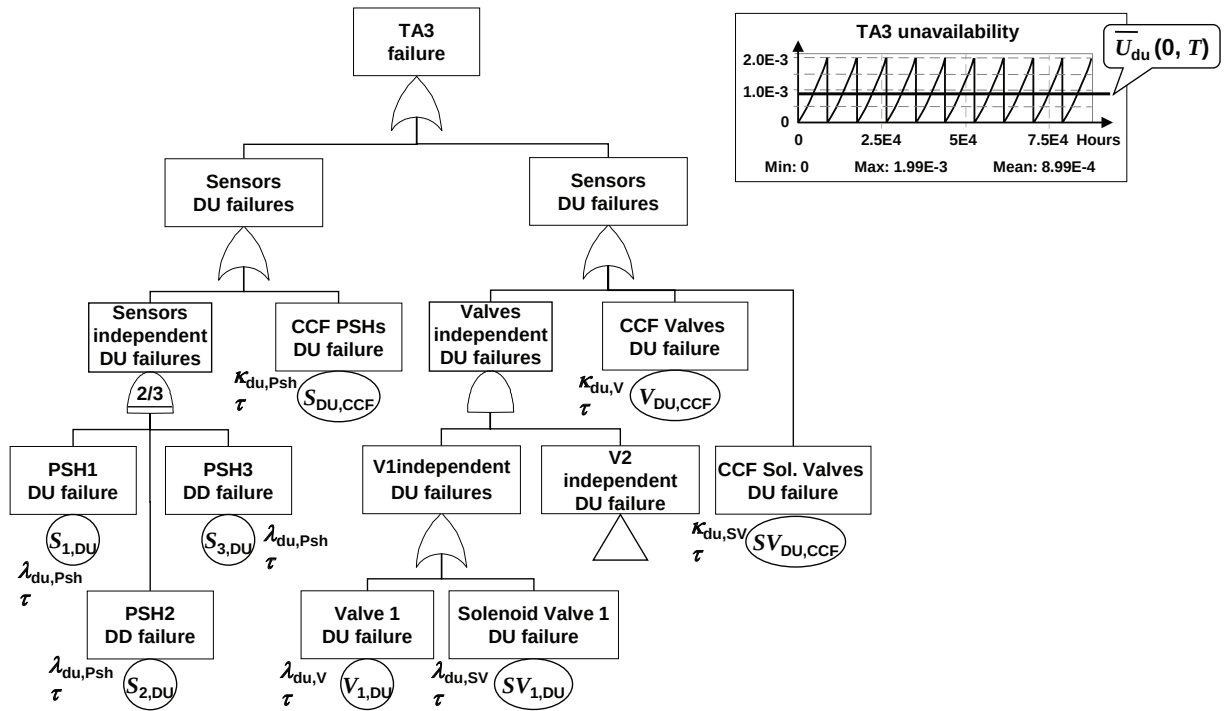


Figure 84 — Fault tree modelling TA3-1

14.4.1.3.3 Petri net and Monte Carlo simulation approach

The reliability block diagram presented in Figure 83 can be used as guideline to build the PN modelling the TA3 example. The blocks of this diagram have just to be fulfilled by modules of the type described in Figure 68.

This modelling is very similar to other PN models previously presented. Therefore it has no been developed here but it will give the same results as fault tree or formulae.

14.4.1.4 TA3-1: spurious failure analysis

As explained in 14.2.1.2, 14.2.1.4 and 14.3.1.4 the analysis of the safe and spurious failures of de-energize to trip safety systems implies:

- the safe failures of the components considered in the dangerous failure analysis,
- and also the identification of all the failures of other components which de-energize some parts of this safety system.

For the TA3-1 example this would imply a detailed analysis which, except for the 2oo3 sensor subsystem, would not add very much value to this subclause with regard with what has already been explained.

For a 2oo3 logic voter, two dangerous failures are needed to inhibit its safety function and two safe failures are needed to provoke its spurious action. Therefore the 2oo3 constitute a good balance between the dangerous and the safe failures.

As a general law, a voting logic which works in "*r out of m*" for the dangerous failure works in "*(m-r+1) out of m*" for the safe failures: *m-r+1* dangerous failures inhibit the "*r out of m*" when *r* safe failures provoke its spurious action. Then the same number of dangerous and safe failures is needed when *m* is odd and *r = (m+1)/2*: then the 1oo1, 2oo3, 3oo5, 4oo7, etc. belongs to the same family as the 2oo3. The 2oo3 is very popular because it is redundant but not too much.

Coming back to TA3-1, a spurious action will occur if two safe inputs from the sensors to the logic solver drop to zero. This implies the safe failures of the sensors themselves but also the rupture of the

links between the sensors and the logic solver, the loss of electricity feeding the sensors (if any), the obstruction of the impulse lines, etc. This depends on the implementation of the sensors and of the policy of actions when a dangerous failure has been detected by a diagnostic test.

14.4.2 TA3-2, TA3-3 and TA3-4 (reminder)

Modelling TA3 with the same assumptions as TA2-2, TA2-3 and TA2-4 can be derived straightforwardly from what has been previously done for TA2-2, TA2-3 and TA2-4. This does not add very much value to this subclause and therefore no further development is done here.

14.4.3 TA3-5: subsea HIPPS

The architecture presented in [Figure 82](#) may be used to design a subsea HIPPS which is not as easy to maintain as a topside system because it is necessary that a subsea intervention vessel be mobilized before the repair can start: according to the location of the system this may take some days to some months.

In addition, when a dangerous detected failure is detected, the production is not necessarily stopped at once.

14.4.3.1 TA3-5: assumptions

Assumptions:

- Periodical tests perfects and performed at the same time.
- Dangerous detected and undetected failures of a given component are independent.
- Constant dangerous failure rates.
- Components as good as new after repairs.
- The production is not shut down during the repair of the sensor, of the logic solver and of the solenoid valves.
- It is necessary that a maintenance rig be mobilized for the repair operations and the production is not shut down when waiting for the maintenance rig.
- The sensors and solenoid valves are periodically tested off line and are no longer available for their safety function during the periodical tests. The periodical test duration lasts 2 h.
- The production is stopped during the maintenance of the isolation valves.

14.4.3.2 TA3-5: probabilistic calculations

The above assumptions introduce dependencies between the events which are not possible to model by using analytical formulae and for which the impact is difficult to “feel” in order to decide if it is negligible or not. In addition, those dependencies cannot be modelled by fault trees and the corresponding Markov graph would be too large to be easily handled in the framework of these simple examples.

Therefore only the PN approach is available and new sub-PN modules are needed to cope with the new needs. They are developed hereafter.

The sub-PN PN4 represented in [Figure 85](#) mixes the modelling of several events:

- dangerous undetected failures;
- common causes of dangerous undetected failures;
- dangerous detected failures;
- common causes of dangerous detected failures;
- periodical tests duration (item unavailable during the periodical tests);

- maintenance mechanisms allowing waiting until the maintenance team (i.e. the maintenance rig) is available to start a repair and to perform only one repair at the same time.

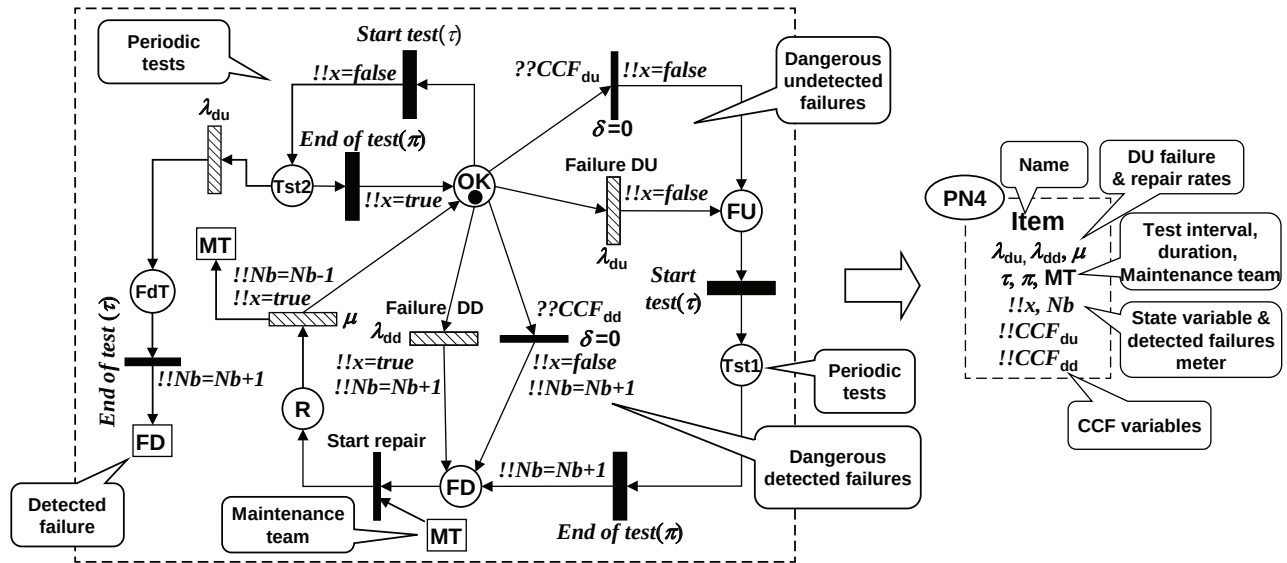


Figure 85 — PN4 module with detected, undetected and common cause failures

PN4 is designed to work in cooperation with sub-PN modelling the common cause failures and the maintenance team behaviour (see Figure 86).

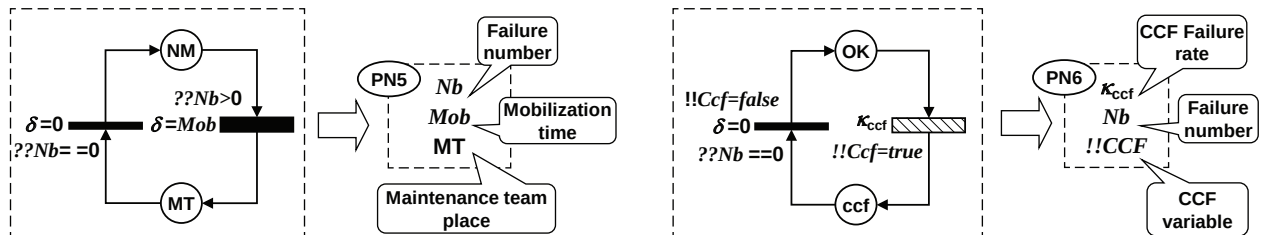


Figure 86 — PN5 and PN6 modules for maintenance team and CCF modelling

The PN5 module on the left hand side of Figure 86 models the mobilization of the maintenance team (i.e. the maintenance rig for sub-sea HIPPS): when the number of failure to be repaired becomes greater than zero, then the maintenance team is mobilized. After a mobilization delay it is available to perform the repairs. It is demobilized when the number of failure has dropped to zero. The number of failure is counted thanks to the variable Nb which is updated by modules of the type PN4 (see Figure 85). When the maintenance team is available one token is placed in MT. This token is used to start one repair in a module of the type PN4. Every other repair has to wait until that the token come back in this place MT after that the undertaken repair is finished. This mechanism insures that only one repair is done at the same time.

The PN6 module, on the right hand side of Figure 86, models a common cause failure. It works in conjunction with sub-PN of the type PN4 presented in Figure 85. When the common cause failure occurs, the Boolean variable $!!Ccf$ becomes “true” and is used to instantaneously fire the corresponding transitions of the various sub-PN of the type PN4 which handle this variable.

According to the needs the sub-PN PN4 can be adapted to items with only dangerous detected failure (PN7), only dangerous undetected failures (PN8) or instantaneous periodical tests and repairs (PN9) just by removing the useless parts as this is shown in Figure 87. This demonstrates the flexibility of the PN modelling.

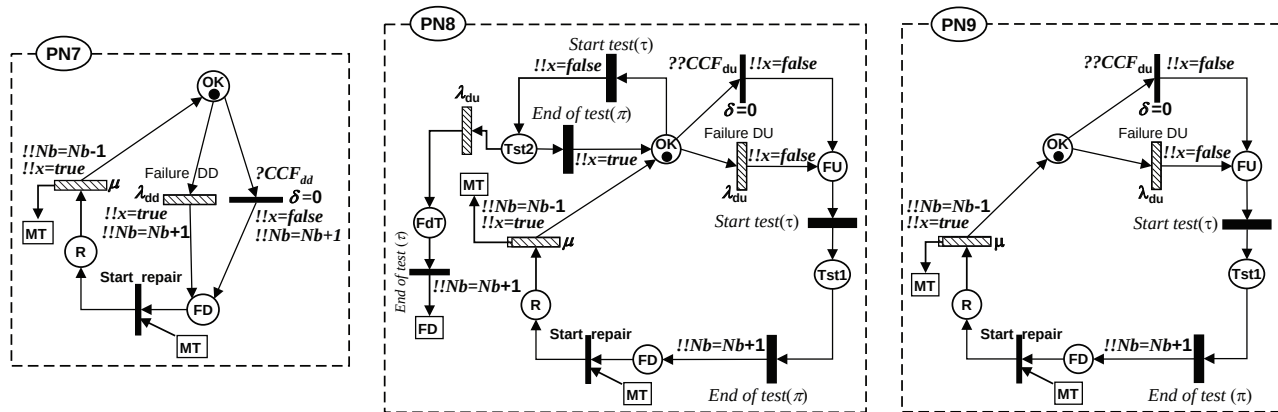


Figure 87 — Modules PN7, PN8 and PN9 derived from module PN4

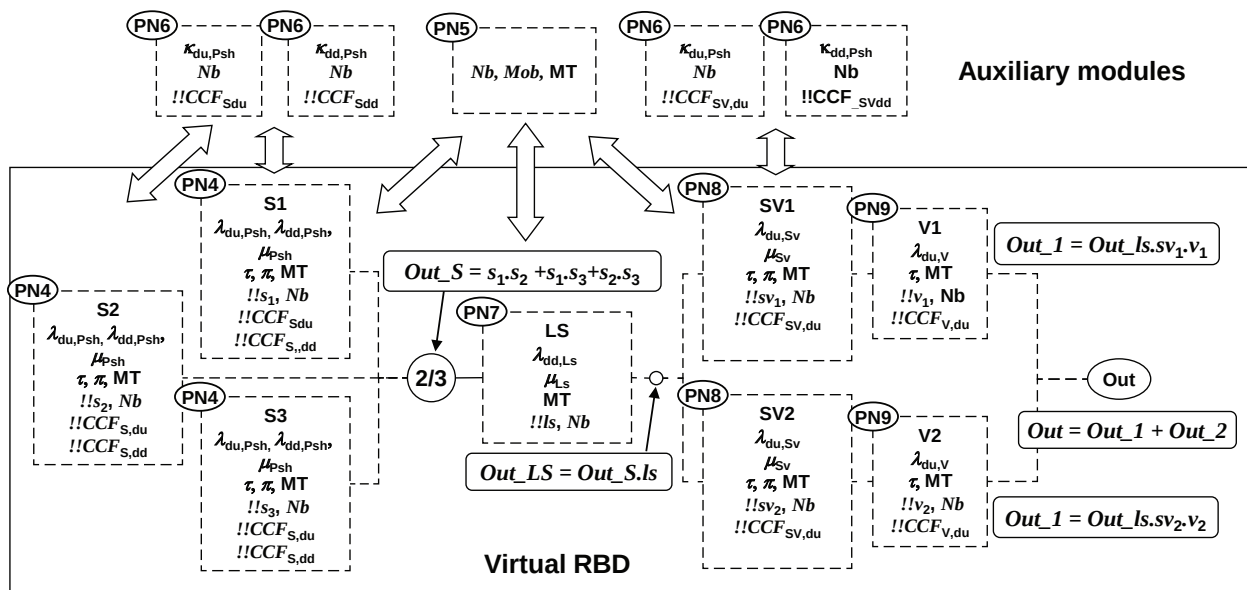


Figure 88 — TA3-5 PN modelling

The above sub-PN modules in [Figure 87](#) can be used in order to model the TA3-5 typical application which implies the modelling of:

- the three sensors with dangerous detected, dangerous undetected failures, common cause failures and non instantaneous periodical tests (PN4);
- logic solver with dangerous detected failure and common cause failure (PN7);
- the solenoid valves with dangerous undetected failures, common cause failure and non instantaneous periodical tests (PN8);
- the isolation valves with dangerous undetected failures, common cause failure and instantaneous periodical tests and instantaneous repairs (PN9).
- This gives the model presented in [Figure 88](#) which is split into two parts:
- the virtual reliability block diagram which models the overall behaviour of TA3-5 (which implements modules of types PN4, PN8 and PN9);
- the auxiliary modules which model the maintenance team mobilization (PN5 module) and the common cause failures impacting the modules of the virtual RBD (PN6 modules).

The links between the virtual RBD and the auxiliary modules is done through variables (e.g. CCF_{Sdd} , CCF_{Sdu} , etc.) and through the repeated place MT.

This Petri net can be used to perform several calculations in order to see the impact of the various parameters. They are summarized in [Table 14](#).

Table 14 — TA3-5 probabilistic results with regard to 4 sets of parameters

N°	Parameters	Average unavailability $\bar{U}(0,T)$ & 90 % confidence interval	Average unavailability u_{psh} of the 2oo3 of sensors and 90 % confidence interval
1	Mob = 0 Periodical tests duration = 0	10^4 [9,37E-04 – 1,02E-03]	1,12E-04 [1,02E-04 – 1,22E-04]
2	Mob = 0 Periodical tests duration = 2	1,14E-03 [1,10E-03 – 1,17E-03]	3,27E-04 [3,16E-04 – 3,39E-04]
3	Mob = 720 h Periodical tests durations = 0	3,95E-03 [3,90E-03 – 4,00E-03]	2,55E-03 [2,53E-03 – 2,58E-03]
4	Mob = 720 h Periodical tests durations = 2	4,09E-03 [4,04E-03 – 4,14E-03]	2,74E-03 [2,71E-03 – 2,76E-03]

The case number one models a single maintenance team without mobilization delay and instantaneous periodical tests. Compared to TA3-1, the dangerous detected failures are also considered. Those new assumptions lead to an average unavailability of $u_{psh} = 1,12 \cdot 10^{-4}$ for the sensor part. This is noticeably higher (56 %) compared to the corresponding result found for TA3-1 ($u_{psh} = 7,19 \cdot 10^{-5}$). Nevertheless the average unavailability of TA3-5 is equal to $\bar{U}(0,T) = 9,76 \cdot 10^{-4}$ which is only slightly higher (8 %) compared to $\bar{U}_{du}(\tau) = 9,0 \cdot 10^{-4}$ obtained for TA3-1.

For the case 2, the periodical test durations of the sensors and of the solenoid valves are equal to 2 h during which they are not available to perform their safety functions. The average unavailability of the sensors increases to $u_{psh} = 3,27 \cdot 10^{-4}$, i.e. an increase of 193 % compared to the previous case, and the average unavailability of TA3-5 is equal to $\bar{U}(0,T) = 1,14 \cdot 10^{-3}$, i.e. an increase of 16 % compared to the previous case.

For the case 3 the periodical test durations are equal to 0 but the mobilization of the maintenance team (I.E. the maintenance rig) is equal to one month. The average unavailability of the sensors increases to $u_{psh} = 2,55 \cdot 10^{-3}$, i.e. an increase of more than 2 000 % compared to the previous case, and the average unavailability of TA3-5 is equal to $\bar{U}(0,T) = 3,95 \cdot 10^{-3}$, i.e. an increase of about 300 % compared to the previous case.

For the case 4 the periodical test durations are equal to 2 h and the mobilization of the maintenance team (I.E. the maintenance rig) is equal to one month. The average unavailability of the sensors increases to $u_{psh} = 2,74 \cdot 10^{-3}$, i.e. an increase of 7,5 % compared to the case number 3, and the average unavailability of TA3-5 is equal to $\bar{U}(0,T) = 4,09 \cdot 10^{-3}$, i.e. an increase of about 3,7 % compared to the previous case number 3.

14.4.4 TA3-6: subsea HIPPS and switching from 2oo3 to 1oo2 on sensor fault detection

The three sensors are organized in 2oo3 logic in order to achieve a good compromise between the dangerous failures and the spurious failures. From the safety point of view a 1oo2 is more reliable than a 2oo3. Therefore when a dangerous failure of one sensor is detected one way to “repair” it temporary is to reorganize the remaining sensors in 1oo2 logic when waiting for the maintenance rig to come on location to perform the repairs.

The PN4-bis PN module is similar to PN4 except that the variable Nbs has been added to count the number of dangerous failures which has been detected.

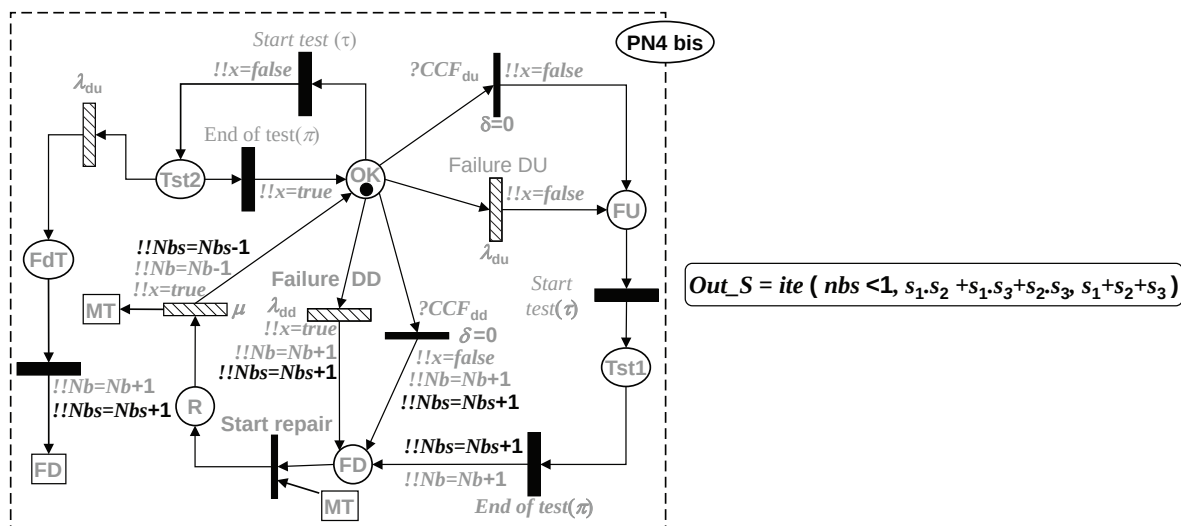


Figure 89 — Module PN4-bis and assertion to switch from 2oo3 to 1oo2

As long as $Nbs = 0$, the sensors works in 2oo3. When one dangerous failure is detected the logic is switched to 1oo2 thanks to the assertion presented on the right hand side of Figure 89. This assertion means:

If $Nb < 1$ (i.e. no failure has occurred)

Then $Out_S = s_1.s_2 + s_1.s_3 + s_2.s_3$ (2oo3 logic)

Else $Out_S = s_1 + s_2 + s_3$ (1oo2 logic)

where s_i is the Boolean variable indicating the state of component S_i

Except the modifications explained above, the PN model presented in Figure 88 can be used as it is.

This leads to the results presented in Table 15.

The average unavailability u_{psh} of the sensors decrease from $2,74 \cdot 10^{-3}$ to $1,35 \cdot 10^{-3}$, i.e. about 50 % compared to the set of parameter n°4 in Table 14 and the average unavailability of TA3-6 decreases from $\bar{U}(0,T) = 4.09 \cdot 10^{-3}$ to $\bar{U}(0,T) = 2.7 \cdot 10^{-3}$, i.e. about 34 %.

Table 15 — TA3-6 probabilistic results (parameter set n°5)

N°	Parameters	Average unavailability $\bar{U}(0,T)$ & 90 % confidence interval	Average unavailability u_{psh} of the 2oo3 of sensors and 90 % confidence interval
5	Mob = 720 h Periodical tests duration = 2 h Switch 2oo3 to 1oo2	2,70E-03 [2,65E-03 - 2.75E-03]	1,35E-03 [1,33E-03 - 1,37E-03]

In order to keep the model simple, the time needed to switch the logic has not been considered but, even if the actual benefit would be reduced in a small proportion, the beneficial impact of this policy is far from negligible. This is typically a behaviour which is not possible to handle properly by using fault trees or formulae.

14.5 Typical application TA4: multiple safety system

14.5.1 TA4: description

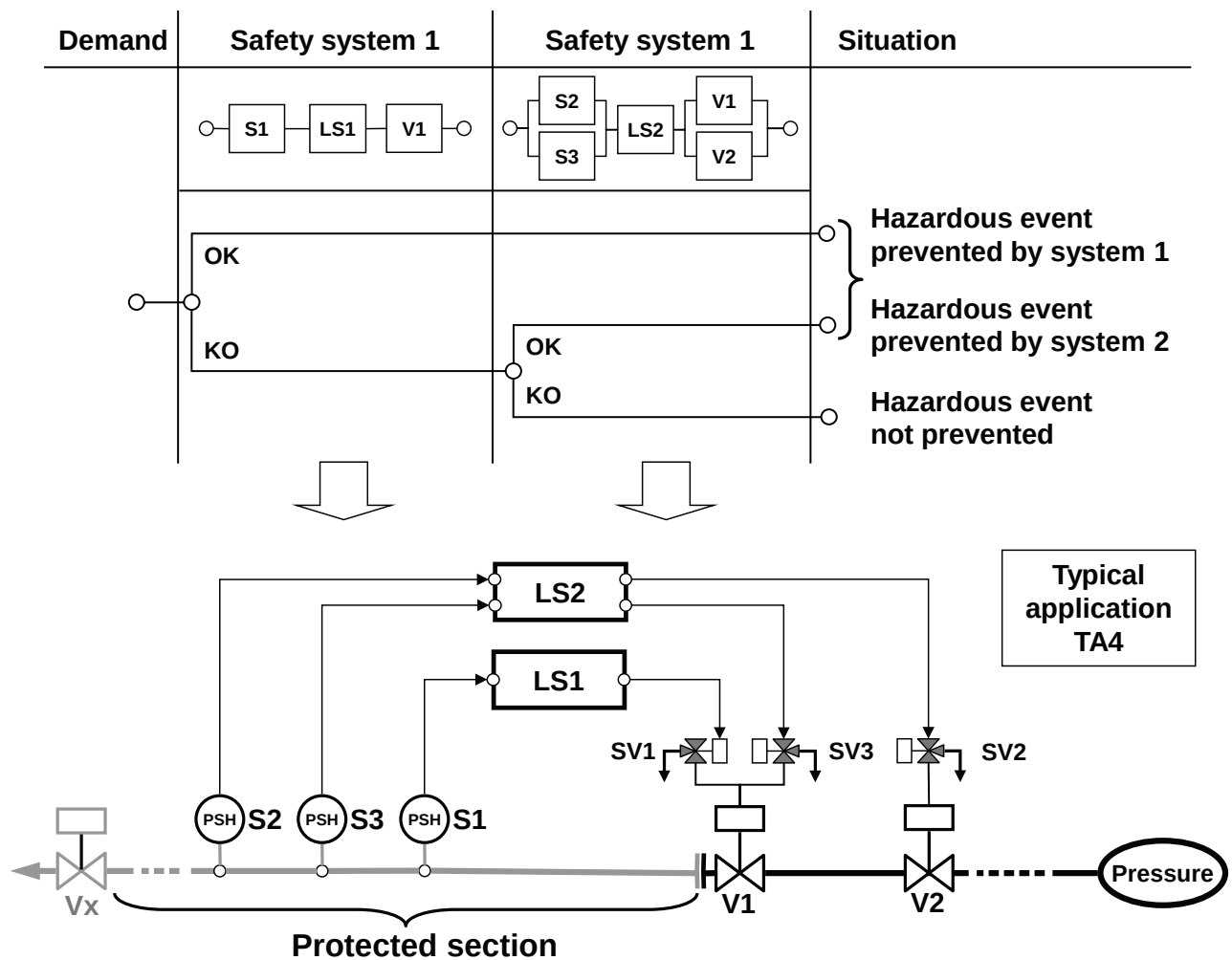


Figure 90 — Simple multiple safety system architecture - Typical application TA4

A simple architecture of a multiple safety instrumented system is illustrated in [Figure 90](#). It comprises two safety instrumented systems working one after the other. The first one may be a safety loop of the BPCS (Basic process control system) and the second one a safety loop of the ESD system (Emergency shut down system) or a HIPPS.

Type of system: pressure protection system (see [Annex A](#)).

Mode of operation: low demand mode.

Layout: multiple safety system comprising two safety loops:

Safety system 1: one pressure sensor (S1), one logic solver (LS1) and one isolation valve (V1) organized in series (similar to the example TA1 with the solenoid valve in addition);

Safety system 2: two redundant sensors (S2 and S3), one single logic solver (LS2) and two redundant valves (V1 and V2) organized in series.

Design philosophy: de-energize to trip.

Hazardous event: over pressurization and burst of the low rating section of the installation.

Functioning description: when the pressure exceeds given thresholds, the sensors (S1, S2 and S3) send signal to the logic solvers (LS1 and LS2) which in turn command the valves (V1 and V2) through their respective solenoid valve (SV1, SV2 and SV3) to close.

14.5.2 TA4: assumptions

Assumptions:

- Periodical tests perfects and performed at the same time.
- Installation stopped during repair of the valves.
- Periodical test durations are negligible
- Dangerous detected and undetected failures of a given component are independent
- Constant dangerous failure rates.
- Components as good as new after repairs.

Reliability parameters:

Same as TA2 (see [14.3.1.1](#)) except for the MRT of the valves which is replaced by a mean time to intervention of 8 h ($\mu = 0,125$).

14.5.3 TA4: analysis

As the installation is stopped during the repair of the valves the risk due to valve faults exists only during the fault detection time (8 h; i.e. $\mu = 0,125$).

When an overpressure occurs, then the multiple safety system works in several steps:

- 1) the sensor (S1) sends a signal to the logic solver (LS1) which in turn commands the valve (V1) through the solenoid valve (SV1) to close.
- 2) if the valve V1 does not close, then the pressure increases until the threshold of the sensors S2 and S3 is trespassed.
- 3) the sensors S2 and S3 send signal to the logic solver (LS2) which in turn commands the valve (V1) through the solenoid valve (SV3) and the valve (V2) through the solenoid valve (SV2) to close.
- 4) if the valves do not close the hazardous event occurs.

Therefore, in case of a demand, the failure of the system 1 provokes a demand on the system 2.

As these safety systems share the valve V1, they are not independent. The solenoid valves commanding V1 have been duplicated in order to decrease the impact of this dependency.

14.5.4 TA4: Probabilistic calculations

14.5.4.1 Analytical formulae

This multiple safety system implies the analysis of 17 minimal cut sets: $LS_2 \bullet S_1$, $LS_1 \bullet LS_2$, $LS_2 \bullet V_1$, $LS_2 \bullet SV_1$, $V_1 \bullet V_2$, $SV_2 \bullet V_1$, $S_1 \bullet S_2 \bullet S_3$, $LS_1 \bullet S_2 \bullet S_3$, $S_2 \bullet S_3 \bullet V_1$, $S_2 \bullet S_3 \bullet SV_1$, $S_1 \bullet SV_3 \bullet V_2$, $LS_1 \bullet SV_3 \bullet V_2$, $SV_1 \bullet SV_3 \bullet V_2$, $S_1 \bullet SV_2 \bullet SV_3$, $LS_1 \bullet SV_2 \bullet SV_3$ and $SV_1 \bullet SV_2 \bullet SV_3$. This would be possible to implement the analytical formulae approach on these minimal cut sets but this would be very tedious. Therefore this is not done in this Technical Report.

14.5.4.2 Fault tree approach

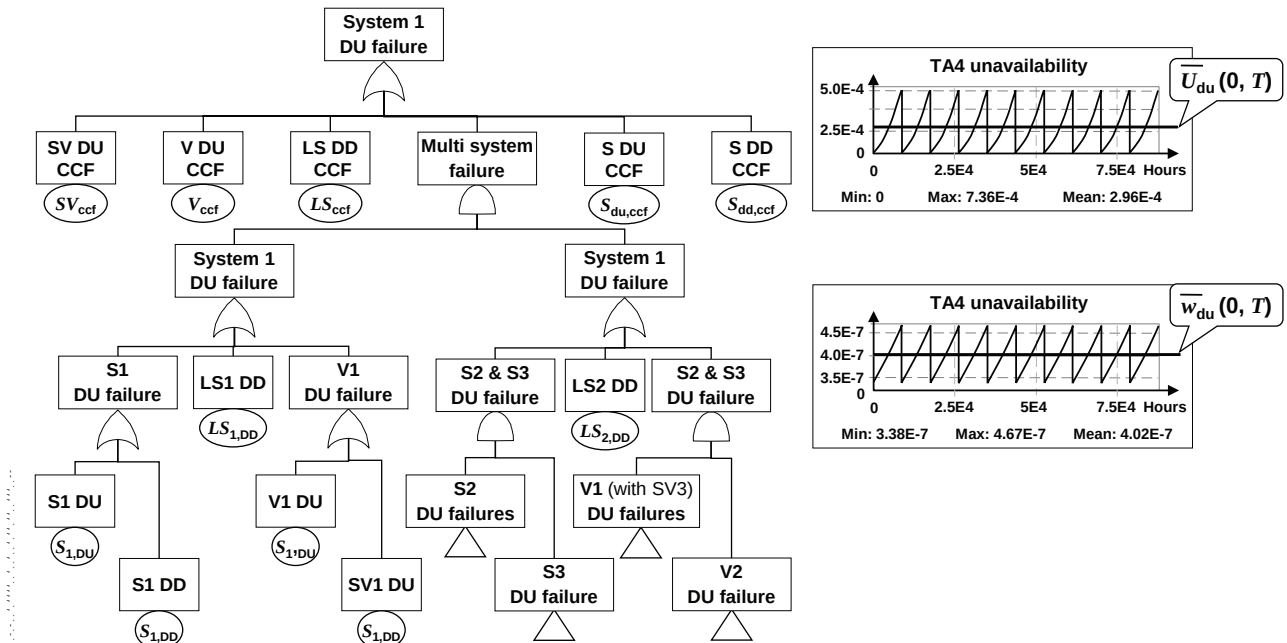


Figure 91 — Fault tree modelling of the typical example TA4

The structure of the fault tree presented in Figure 91 is the following:

- the sub-tree below the logical gate Or1 models the safety system 1,
- the sub-tree below the logical gate Or3 models the safety system 2,
- the redundancy between the safety system with regard to the independent failures is modelled by the logical gate And5,
- the logical gate Or13 at the top of the fault tree gathers the common cause and the independent failures.
- the sub-trees for “S2 failure” and “S3 failure” are similar to the sub-tree for “S1 failure” (gate Or9)
- the sub-tree for “V1 fail to close (with SV3)” and “V2 fail to close” are similar to the sub-tree developed for “V1 fail to close (with SV1)”, (gate Or6).

Using the reliability parameters described above, the average unavailability of TA4 calculated over 10 years is equal to $\bar{U}(0, T) = 2.96 \cdot 10^{-4}$ and the average dangerous failure frequency is equal to $\bar{w}(0, T) = 4.02 \cdot 10^{-7}$

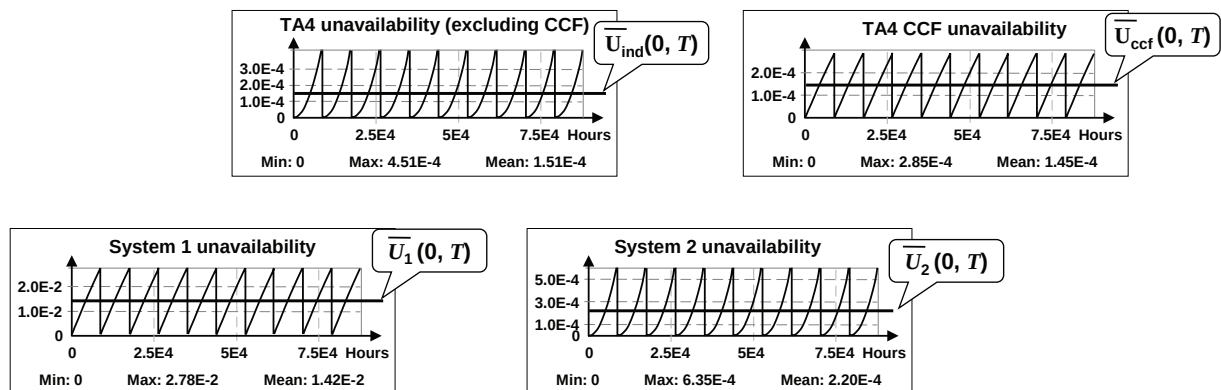


Figure 92 — TA4 unavailability excluding Common Cause Failures and CCF contributions

This fault tree provides also:

- the average unavailability of the safety system 1, $\bar{U}_1(0,T)=1,42 \cdot 10^{-2}$,
- the average unavailability of the safety system 2, $\bar{U}_2(0,T)=2,2 \cdot 10^{-4}$,
- the average unavailability of the multiple safety system without the common cause failures, $\bar{U}_{ind}(0,T)=1,51 \cdot 10^{-4}$,
- the average unavailability due to the common cause failures, $\bar{U}_{ccf}(0,T)=1,45 \cdot 10^{-4}$.

A simplistic calculation which is often performed is to make the product of the average unavailabilities of the different safety systems. In the TA4 case it is obvious that $\bar{U}_{ind}(0,T)$ is very different from the product $\bar{U}_1(0,T) \cdot \bar{U}_2(0,T)=3,12 \cdot 10^{-6}$. In this case, this simplistic calculation leads, for the independent failures, to a result which is 48 times lower than the actual value!! Therefore the contribution of the independent failure appears to be negligible compared to the common cause failure contribution. In fact they actually participate for about half of the result. Finally, and even if the common cause failure limit to some extent the impact of such approximation, the error on the overall unavailability of the multiple system would be of 104 %.

This example shows clearly that the analyst should be very cautious with some popular ideas like that the average unavailability of a multiple system may be calculated by the simple product of the average unavailabilities of its individual safety systems or that the independent failures are negligible and that only the common cause failures should be considered.

14.5.4.3 Petri net and Monte Carlo simulation approach

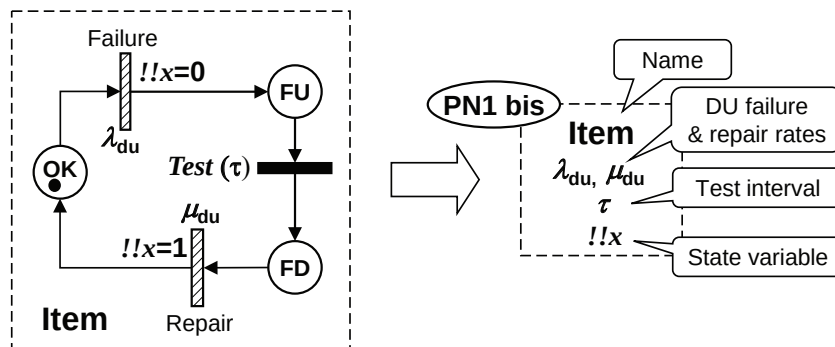


Figure 93 — Sub-PN module PN1 (second version)

The PN presented in Figure 93 is very similar to the module PN1 already described in Figure 49. The difference is that the repair is not instantaneous but modelled by a constant repair rate. Such module describes the behaviour of an item with undetected failures which are periodically tested and repaired with a constant repair rate. According to the simplified assumptions of the typical example PN4, the dangerous undetected failures of the sensors, the solenoid valves and of the valves can be modelled by this type of PN module.

The dangerous detected failures of the sensors and logic solvers can be modelled by the modules of the type PN3 (see Figure 79).

The Figure 94 shows the RBD driven Petri net related to TA4. The following logical equations are needed to complete the model:

- System 1 available: $Out1 = s_{1,du} \cdot s_{1,dd} \cdot ls_{1,dd} \cdot sv_{1,du} \cdot v_{1,du}$
- System 2 available: $Out2 = (s_{2,du} \cdot s_{2,dd} + s_{3,du} \cdot s_{3,dd}) \cdot ls_{2,dd} \cdot (sv_{3,du} \cdot v_{1,du} + sv_{2,du} \cdot v_{1,du})$
- Multiple system available (independent failures): $Out3 = Out1 \cdot Out2$

- Multiple system available (with CCF): $Out4 = Out3 \cdot s_{du_ccf} \cdot s_{dd_ccf} \cdot l_{sdd_ccf} \cdot sv_{du_ccf} \cdot v_{du_ccf}$

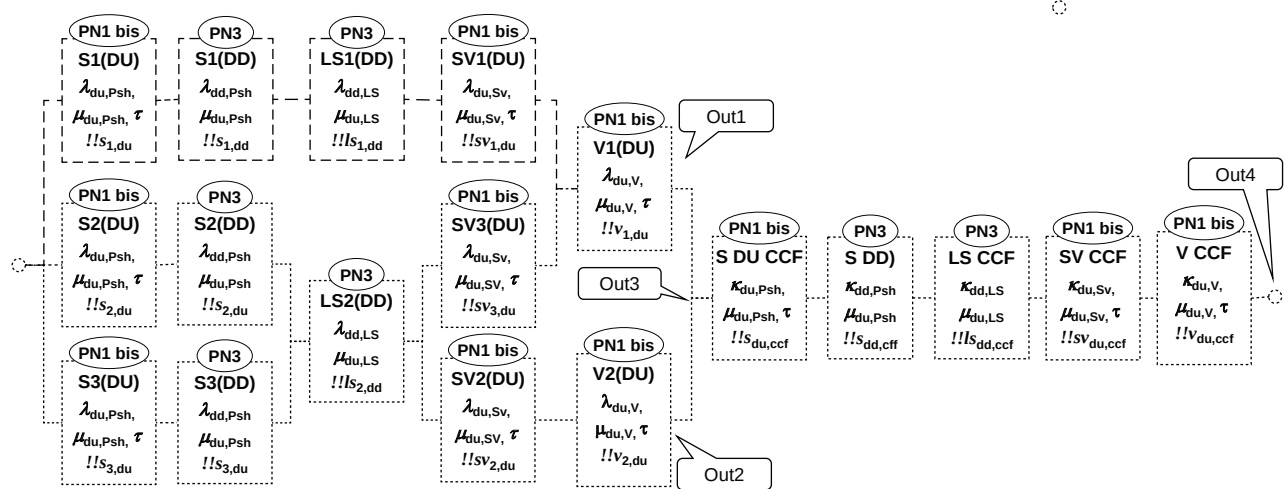


Figure 94 — TA4 PN modelling

The Monte Carlo simulation of this system gives an average unavailability of $\bar{U}_{du}(0,T) = 2.95 \cdot 10^{-4}$ with a 90 % confidence interval equal to $[2,88 \cdot 10^{-4} - 3,02 \cdot 10^{-4}]$. As usual this is very similar to the previous results obtained by fault tree.

The Monte Carlo simulation provides also an average number of failure of 0,035 over 10 years with a 90 % confidence interval equal to $[0,034 \cdot 9 - 0,035 \cdot 5]$. This leads to an average failure frequency of $\bar{w}_{du}(0,T) = 4.0 \cdot 10^{-7}$ with a 90 % confidence interval equal to $[3,98 \cdot 10^{-7} - 4,05 \cdot 10^{-7}]$. This is also very close with the average dangerous failure frequency obtained by fault tree.

Of course with the PN presented in [Figure 94](#), it would be very easy to implement the staggering of the periodical tests of the sensors, solenoid valves and valves as well as partial and full stroking. This does not add more value compared to what has been already described in the previous subclause and is not presented here.

14.6 Typical application TA5: emergency depressurization system (EDP)

14.6.1 TA5: description and basic assumptions

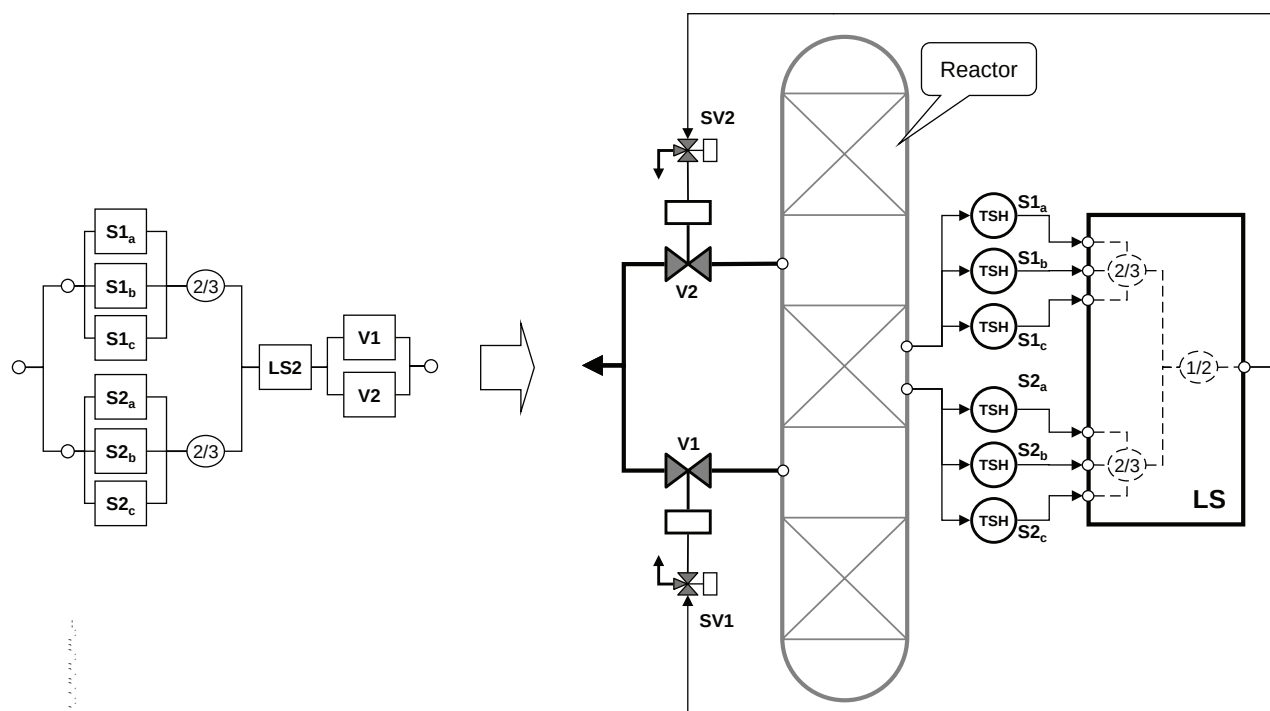


Figure 95 — Emergency depressurization system

Figure 95 shows a simplified emergency depressurization system implemented on a hydrocracking unit which is a classical installation found in the downstream branch of oil and gas industries. The aim of this safety system is to quickly depressurize the reactor when the temperature increases and crosses a given thresholds in order to avoid a runaway of the exothermic chemical reaction. The system itself is presented on the right hand side of the figure and the equivalent reliability block diagram on the left hand side.

Type of system: temperature protection system (see Annex A).

Mode of operation: low demand mode.

Layout: two groups of three temperature sensors organized in 2oo3 (S1a, S1b, S1c) and (S2a, S2b, S2c), one logic solver (LS) and two isolation valve (V1, V2) organized in parallel and piloted by two solenoid valves (SV1, SV2).

Design philosophy: de-energize to trip.

Hazardous event: over pressurization and runaway reaction.

Functioning description: when the temperature exceeds a given threshold, the sensors send signals to the logic solver. When the logic solver receives at least two signals in at least one of the two groups, then it commands the solenoid valves to open. This releases the hydraulic pressure until a pre-established level where the operator can close the valves.

The sensors are located in different parts of the reactor to protect against high temperature.

Assumptions:

- Periodic tests are performed when the reactor is stopped.
- Installation stopped during repair of undetected dangerous failures.

- the 2oo3 logic of a group of sensors is switched to 1oo2 in case of one dangerous detected failure in the group.
- Installation stopped during periodical tests and repair for Logic Solver.
- Dangerous detected and undetected failures of a given component are independent
- Constant dangerous failure rates.
- From the dangerous failure point of view, components are as good as new after repair.
- Failures not covered by periodical tests are never detected and then never repaired.

Table 16 — TA5 basic reliability parameters

Parameter	Component					
	Galvanic isolator	Thermo-couple	Temperature transmitter	Logic solver	Solenoid valve	Isolation valve
Periodical test coverage	99 %	60 %	95 %	90 %	99 %	99 %
Beta factor	2 %	2 %	2 %	1 %	2 %	2 %
DU failure rate h ⁻¹	1,79 10 ⁻⁷	1,28 10 ⁻⁶	2,50 10 ⁻⁷	4,2 10 ⁻⁹	1,95 10 ⁻⁷	2,45 10 ⁻⁷
DD failure rate	—	2,42 10 ⁻⁵	2,34 10 ⁻⁶	1,75 10 ⁻⁶	—	—
Safe failure rate	5,50 10 ⁻⁸		2,30 10 ⁻⁸	6,92 10 ⁻⁶	2,70 10 ⁻⁸	5,00 10 ⁻⁷
Periodical test interval	17 520 h	17 520 h	17 520 h	87 600 h	17 520 h	17 520 h
Maintenance time	48 h	48 h	48 h	48 h	48 h	48 h

It is necessary to split each of the failure rates of the above tables in 4 parts to take into account the periodical test coverage and the beta factor:

Table 17 — TA5 elaborated reliability parameters

Component	Parameter					
	DU failure rate				DD failure rate	
	Covered		Not covered			
	Indep.	CCF	Indep.	CCF	Indep.	CCF
Galvanic isolator	1,73 10 ⁻⁷	3,54 10 ⁻⁹	1,75 10 ⁻⁹	3,57 10 ⁻¹¹	—	—
Thermo-couple	7,5 10 ⁻⁷	1,53 10 ⁻⁸	5,0 10 ⁻⁷	1,02 10 ⁻⁸	2,38 10 ⁻⁵	4,85 10 ⁻⁷
Temperature Transmitter	2,33 10 ⁻⁷	4,75 10 ⁻⁹	1,23 10 ⁻⁸	2,5 10 ⁻¹⁰	2,30 10 ⁻⁶	4,68 10 ⁻⁸
Total sensor	1,16 10⁻⁶	2,36 10⁻⁸	5,14 10⁻⁷	1,05 10⁻⁸	2,60 10⁻⁵	5,32 10⁻⁷
Logic solver	3,77 10⁻⁹		4,19 10⁻¹⁰		1,75 10⁻⁶	
Solenoid valve	1,89 10 ⁻⁷	3,86 10 ⁻⁹	1,91 10 ⁻⁹	3,90 10 ⁻¹¹	—	—
Isolation valve	2,38 10 ⁻⁷	4,85 10 ⁻⁹	2,4 10 ⁻⁹	4,9 10 ⁻¹¹	—	—
Total valve	4,27 10⁻⁷	8,71 10⁻⁹	4,31 10⁻⁹	8,80 10⁻¹¹	—	—

When the galvanic isolator, thermocouple and temperature transmitters are periodically tested at the same time, they can be gathered in a macro-component (sensor) as this has been done in the line “total sensor” of the [Table 17](#).

The same has been done for the solenoid valves and the isolation valves which have been gathered into a macro-component “valve”.

The logic solver being alone, its dangerous failure rates have not been split between independent and common cause failures.

14.6.2 TA5: analysis

The reactor is stopped during the periodical tests, and then the risk disappears during these operations. When a dangerous failure is detected by a periodical test, then it is repaired before restarting the reactor. From modelling point of view this is equivalent to have instantaneous periodical tests and repairs (i.e. MTTS = 0, see [3.1.36](#))

According to the assumptions, the various states of the 2oo3 groups are presented in [Table 18](#).

Table 18 — Various states of the 2oo3 subsystems

N°	DU	DD	n	Voting logic	System state
1	0	0	1	2oo3	OK
2	0	1	3	1oo2	OK
3	0	1 + 1	3	1oo2	<i>safety action</i>
4	1	0	3	2oo3	OK
5	1	1	9	1oo2	OK
6	1	1 + 1	9	1oo2	<i>safety action</i>
7	2	0	3	2oo3	Inhibited
8	2	1	9	1oo2	Inhibited
9	2	1 + 1	9	1oo2	<i>safety action</i>
10	3	0	1	2oo3	inhibited
11	3	1	3	1oo2	inhibited
12	3	1 + 1	3	1oo2	<i>safety action</i>
13	CCF	0	1	2oo3	inhibited
14	CCF	1	3	1oo2	inhibited
15	CCF	1 + 1	3	1oo2	<i>safety action</i>
16		CCF	1	1oo2	<i>safety action</i>

This is a quite complex behaviour: 4 states are OK (1, 2, 4 and 5), 6 states imply a spurious trip of the reactor (3, 6, 9, 12, 15 and 16) and 6 states are dangerous because the safety action is inhibited (7, 8, 10, 11, 13 and 14).

It should be noted that in states 9, 12 and 15 the safety action is restored by dangerous detected failures. This is a typical case of logical incoherence (i.e. the underlying logical equation is not monotonous). Generally, when such behaviour occurs, it is not considered as it is not really usual to count on new failures to increase the safety.

When a dangerous detected failure of a sensor occurs, then it may be repaired without shut down the reactor. This occurs in the states 2, 5, 8, 11, and 14.

When a dangerous detected common cause failure occurs on a 2oo3 group of sensor, then the safety action occurs and the reactor is shut down. Therefore, during the repair of a dangerous detected common cause of sensors and the risk has disappeared.

When a dangerous detected failure of the logic solver occurs then the safety action occurs and the reactor is shut down. Therefore, during the repair of the logic solver, the reactor is shut down, either by periodical test or safety action and the risk has disappeared.

The reliability data kept for the probabilistic calculations are summarized in [Table 19](#).

Table 19 — TA5 reliability parameters used for probabilistic calculations

Component	Parameter					
	DU failure rate				DD failure rate	
	Covered		Not covered			
	Indep.	CCF	Indep.	CCF	Indep.	CCF
Sensor	1,16 10 ⁻⁶	2,36 10 ⁻⁸	5,14 10 ⁻⁷	1,05 10 ⁻⁸ (group of 3 sensors)	2,6 10 ⁻⁵	-
Periodical test interval	17 520 h	17520 h	17 520 h	17 520 h	-	-
MRT	-	-	-	-	48 h	-
Logic solver	3,77 10 ⁻⁹		4,19 10 ⁻¹⁰		-	
Periodical test interval	87 600 h		87 600 h		-	
MRT	-		-		-	
Valve	4,27 10 ⁻⁷	8,71 10 ⁻⁹	4,31 10 ⁻⁹	8,80 10 ⁻¹¹ (two valves)	-	-
Periodical test interval	17 520 h	17 520 h	17 520 h	17 520 h	-	-
MRT	-	-	-	-		

From spurious safety action point of view the reactor is shut down:

- In case of a state of the type 3, 6, 9, 12, 15 or 16 occurs on one of the two 2oo3 groups.
- and, as this is a “de-energize to trip” safety system, each time the energy is lost somewhere:
 - a) loss of power supply, e.g. loss of auxiliary supply (UPS, compressed air, etc.), electric supply;
 - b) rupture of the links between the sensor and the logic solver;
 - c) rupture of the links between the logic solver and the depressurization valve.

14.6.3 TA5: probabilistic calculations

As this safety system operates on demand, its average unavailability (also called PFD_{avg} for a safety instrumented system) is the relevant parameter to calculate.

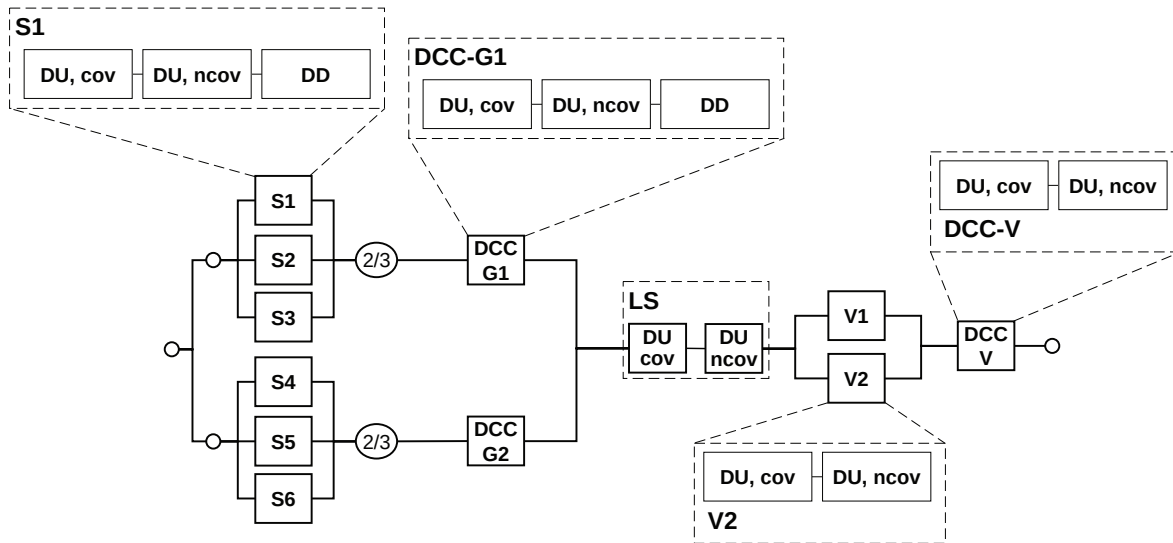


Figure 96 — RBD model for the EDP safety system

Figure 96 presents the reliability block diagram modelling the EDP safety system according to the assumptions described above. Each block is divided in several sub-blocks according to the reliability data to be taken into consideration.

As the probability of common cause failure to more than 2 or 3 similar components is very low (except if explicitly identified), it has been considered, for the sensors, that they are existing only within a given 2oo3 group.

14.6.3.1 Analytical formulae

The reliability block diagram presented in Figure 96 produces 849 minimal cut sets (i.e. 849 combinations of failures leading to the overall failure of the EDP safety system):

a) 4 minimal cut sets of order 1

- 1) $LS_{DU,Cov}$
- 2) $LS_{DU,nCov}$
- 3) $V_{DU,CCF,Cov}$
- 4) $V_{DU,CCF,nCov}$

b) 8 minimal cut sets of order 2

- 1) $G_{1,DU,CCF,Cov}, G_{2,DU,CCF,Cov}$
- 2) $G_{1,DU,CCF,nCov}, G_{2,DU,CCF,Cov}$
- 3) $G_{1,DU,CCF,Cov}, G_{2,DU,CCF,nCov}$
- 4) $G_{1,DU,CCF,nCov}, G_{2,DU,CCF,nCov}$
- 5) $V_{1,DU,Ind,Cov}, V_{2,DU,Ind,Cov}$
- 6) $V_{1,DU,Ind,nCov}, V_{2,DU,Ind,Cov}$
- 7) $V_{1,DU,Ind,Cov}, V_{2,DU,Ind,nCov}$

- 8) $V_{1,DU,Ind,nCov}$ $V_{2,DU,Ind,nCov}$
c) 108 minimal cut sets of order 3
d) 720 minimal cut sets of order 4

It is obvious that without a tool identifying the minimal cut sets and able to process them from a probabilistic point of view, it would be difficult to undertake the calculations by hand of 849 minimal cut sets with analytical formulae.

Among those 849 minimal cut sets, the single and double failures with self-explaining notations have been explicitly identified. The contributions to the average unavailability of some of those minimal cut sets are easy to evaluate (single failures and double failures 1, 4, 5 and 8) because only one periodic test duration should be considered. This is more difficult for some others (double failures 2, 3, 6 and 7) because two periodical test interval durations should be considered.

Therefore only the 4 single failures are considered:

- 1) $LS_{DU,Cov} \Rightarrow \bar{u}_1(0, \tau_{LS}) = 3.77 \cdot 10^{-9} \cdot 87600 / 2 = 1.65 \cdot 10^{-4}$
- 2) $V_{DU,CCF,Cov} \Rightarrow \bar{u}_2(0, \tau_V) = 8.71 \cdot 10^{-9} \cdot 17520 / 2 = 7.63 \cdot 10^{-5}$
- 3) $LS_{DU,nCov} \Rightarrow \bar{u}_3(0, T) = 4.19 \cdot 10^{-10} \cdot 131400 / 2 = 2.75 \cdot 10^{-5}$
- 4) $V_{DU,CCF,nCov} \Rightarrow \bar{u}_4(0, T) = 8.80 \cdot 10^{-11} \cdot 131400 / 2 = 5.78 \cdot 10^{-5}$

The calculations of the average unavailabilities due to the covered dangerous failures are based on the periodical test intervals and the calculations of the average unavailabilities of the not covered dangerous failures are based on 15 years of operation.

One approximation commonly done is to consider that only the single failures contribute to the average unavailability of the safety system and that the other minimal cut sets are negligible. That leads to $\bar{U}(0, T) \approx \bar{u}_1 + \bar{u}_2 + \bar{u}_3 + \bar{u}_4 = 2.75 \cdot 10^{-4}$. This estimation which forgets 99,5 % of the combinations of failures leading to the dangerous failure of the safety system is not conservative. Therefore this raise the question: which is the level of non-conservativeness?

It should also be noted that the sensors have no impact on the above estimation and this raise a second question: which is the usefulness of the redundancy and the sophisticated sensor design?

The answers are given in [14.6.3.2](#) where calculations are performed by using the fault tree approach and provide, with the same reliability data and assumptions, a more accurate result of $\bar{U}(0, 15y) = 3.5 \cdot 10^{-4}$: then the error of the drastically simplified calculation performed just above is greater than 21 %.

Therefore, more minimal cut sets should be considered and analysed in order to estimate the order of magnitude of what is neglected. This can be done by using the formulae developed in the Clause 7.

14.6.3.2 Fault tree approach

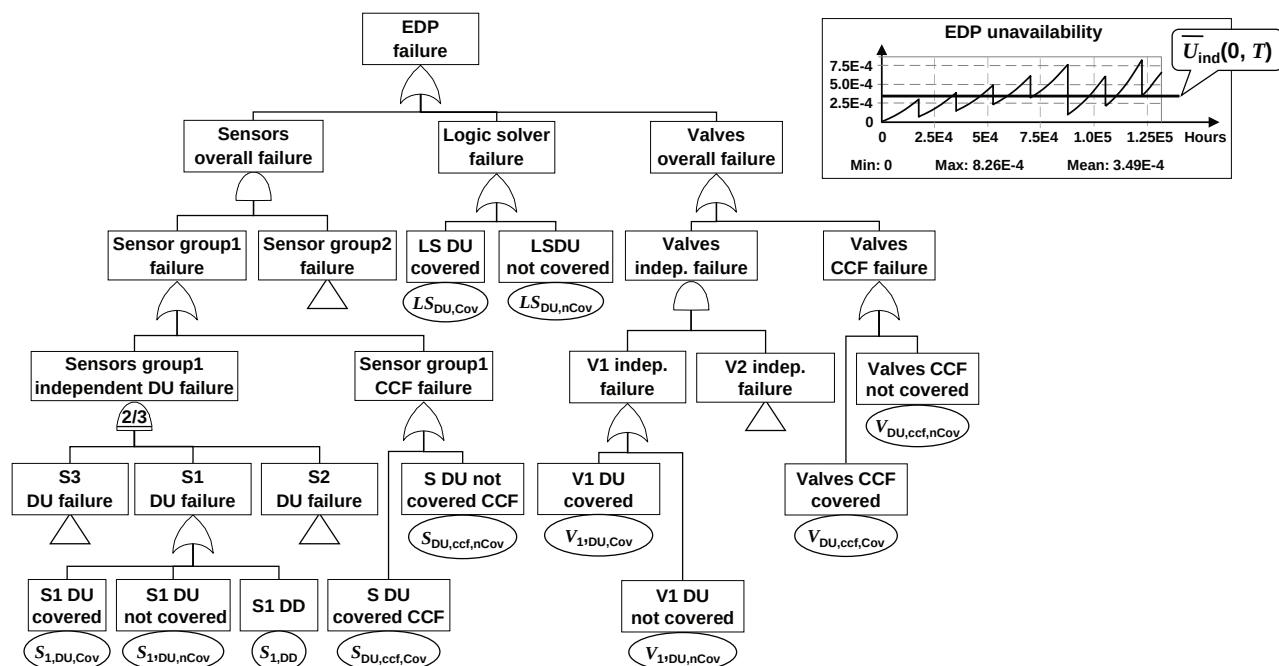


Figure 97 — Fault tree modelling of the EDP system

The fault tree modelling the EDP system failure is represented in Figure 97. The sub-trees represented by the transfer gates (triangle) “S2 failure” and “S3 failure” are similar to the sub-tree “S1 failure (left hand side of the figure). The sub-tree represented by the transfer gate “Sensor group 2 failure” is similar to the sub-tree “sensor group 1 failure” (left hand side of the figure). The sub tree represented by the transfer gate “V2 independent failure” is similar to the sub-tree “V1 independent failure” (right hand side of the figure).

This fault tree model the behaviour of the EDP safety system with the exception of the switch 2oo3 to 1oo2 when one dangerous failure occurs. This is a systemic property which cannot be modelled just by describing individual failures. Nevertheless, this is a conservative approach and the real average unavailability value is lower than the result produced by the above fault tree.

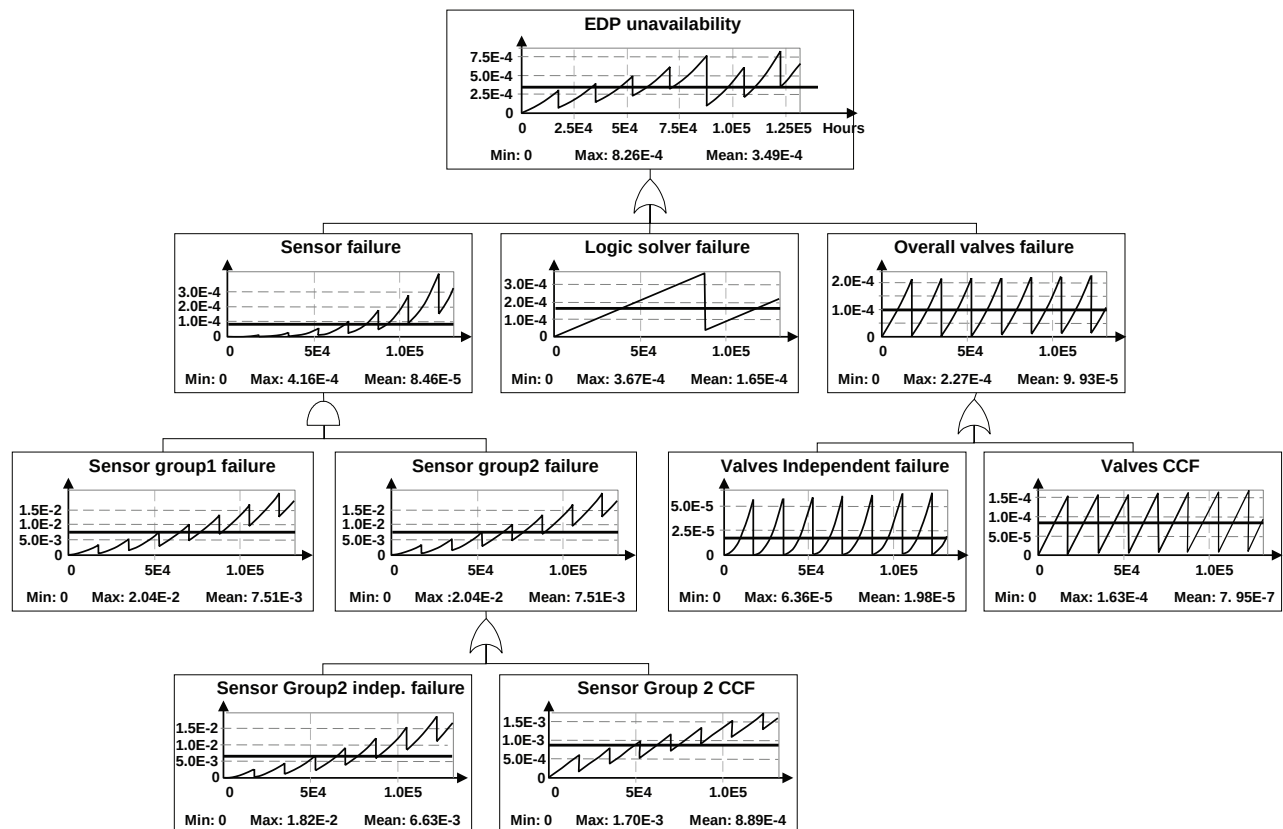


Figure 98 — Illustration of the various contributors to the overall unavailability

The average unavailability of the EDP safety system given by this fault tree and calculated over a period of 15 years is $\bar{U}(0,15y) = 3,5 \cdot 10^{-4}$. Therefore this is in the range of SIL3.

As said above, this result should be conservative. The degree of conservativeness can be evaluated by considering the [Figure 98](#) where the various contributors to the EDP safety system unavailability have been presented. Looking at this figure shows that the main contribution comes from the logic solver and that the valves and the sensors have similar contributions.

[Figure 98](#) shows also that the average unavailability increases when the time increases. This is due to the larger periodical test interval for the logic solver and the dangerous undetected failures not covered by the periodical tests. Nevertheless, the whole unavailability saw-tooth curve remains in the range of the SIL3 requirements all over the 15 years. This is a permanent SIL3 system over this period.

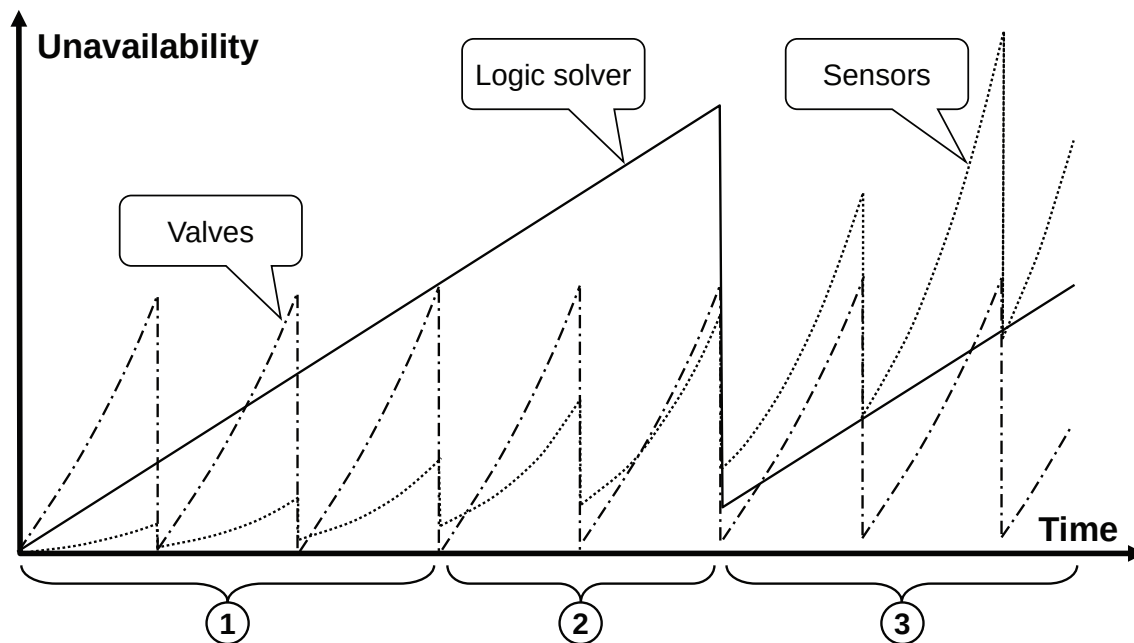


Figure 99 — Contribution to EDP unavailability

The various contributions to the unavailability of the EDP safety system are summarized in [Figure 99](#). At least three zones can be identified:

- 1) the valves are the main contributors;
- 2) the logic solver becomes the main contributor;
- 3) the sensors become the main contributors.

These very different behaviours cannot be identified, just by considering the PFD_{avg} over the 15 years of operation. With this figure it becomes obvious that, even if 20 % accuracy is considered to be sufficient enough, the result obtained just by considering the single failures in [14.6.3.1](#) is no longer relevant in zone 3.

14.6.3.3 Multi-phase Markovian approach

The EDP safety system has too many states for the Markov graph to be built by hand.

14.6.3.4 Petri net and Monte Carlo simulation approach

According to the analysis of the states of the 2003 groups of sensors, the reactor is stopped when two dangerous detected failures occurs within the same 2003 group. This feature has not been modelled in the fault tree developed above because this is a systemic behaviour which cannot be modelled only by considerations at component level. The Petri net models have not such limitation and, therefore a relevant PN model can be built to evaluate the impact of such a feature on the average unavailability of the EDP safety system.

The PN in [Figure 100](#) models the three kinds of sensor failures: independent covered dangerous failures on the left hand side, independent covered dangerous failures in the middle and dangerous detected failures on the right hand side.

The main novelty of this PN is the mechanism implemented to ensure that the repaired sensors wait until all the repairs have been done before to be restarted again:

- When one sensor fails, the variable Nb_{du} is updated to $Nb_{du}+1$ (then, Nb_{du} counts the number of dangerous undetected failures at each instant).

- When the periodical test is performed, then the reactor is stopped and the repair can start.
- When the repair is finished then Nb_{du} is updated to $Nb_{du}-1$.

Then the repaired sensor waits for the reactor restarts until it is put back on line again. This occurs when all repairs have been done (i.e. $Nb_{du} = 0$).

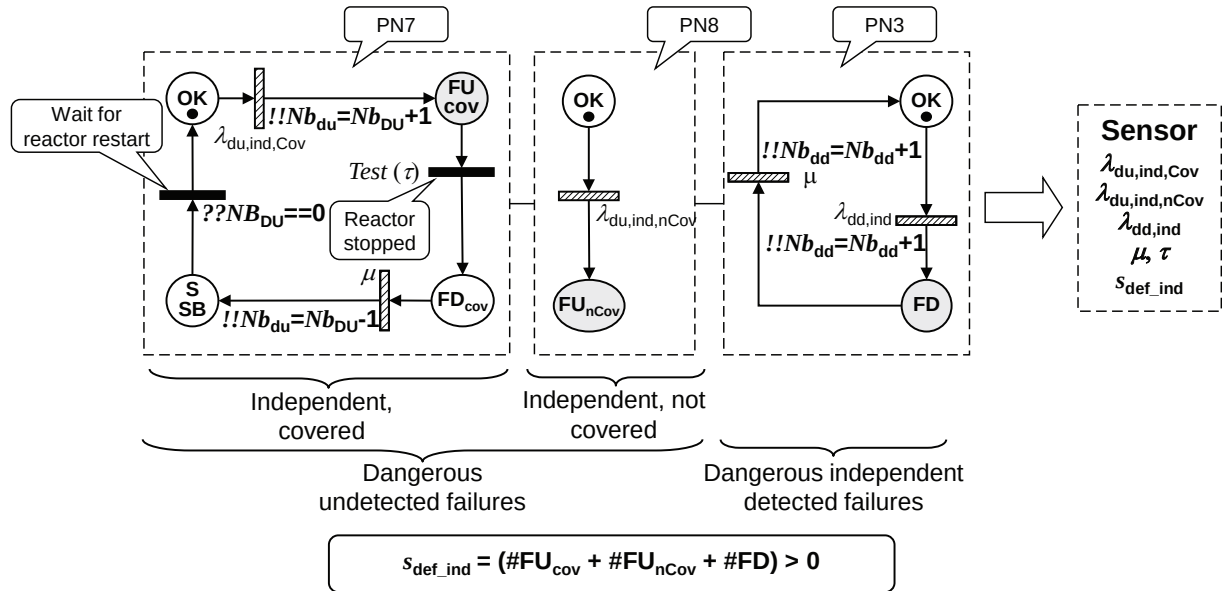


Figure 100 — PN Modelling of the independent failure of a single sensor

From sensor state point of view the three following classes are considered:

- 1) the sensor is in perfect state when there is one token in each of the OK places;
- 2) the sensor is failed and the reactor is running when there is one token in each of the hatched places;
- 3) the reactor is stopped when there is one token in all the other places.

Then only the states of class 2) are considered when establishing the PFD_{avg} of the safety system. This is done through the Boolean variable $S_{ccf} = [(\#FU_{cov} + \#FU_{nCov}) > 0]$ where $\#X$ indicates the marking of the place X.

The CCF modelling is presented in [Figure 101](#). This is similar to the PN presented in [Figure 100](#) except that the part devoted to dangerous detected failures has been removed (because the reactor is stopped when such a failure occur).

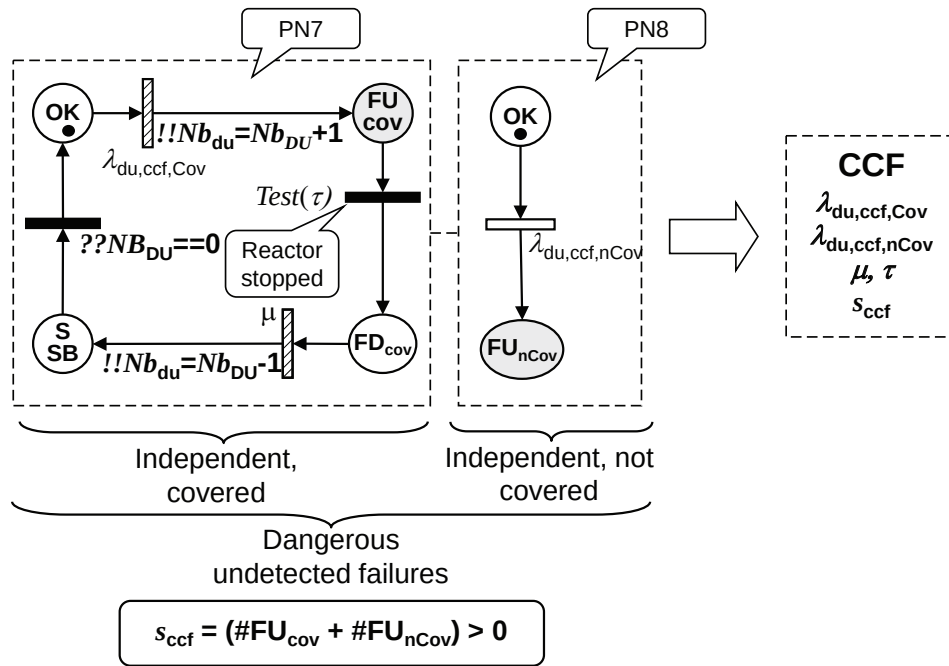


Figure 101 — PN modelling of sensors common cause failures

Figure 102 shows the reliability block diagram used in the background of the PN model. The boxes are populated with sub-PN similar to those presented in Figure 100 and Figure 101. Some formulae have been introduced in order to model the 2oo3 independent failures ($G1_{def_ind}$, $G2_{def_ind}$), the risk related to each sensor groups ($G1_R$, $G2_R$) and the overall risk due to sensors ($Risk$).

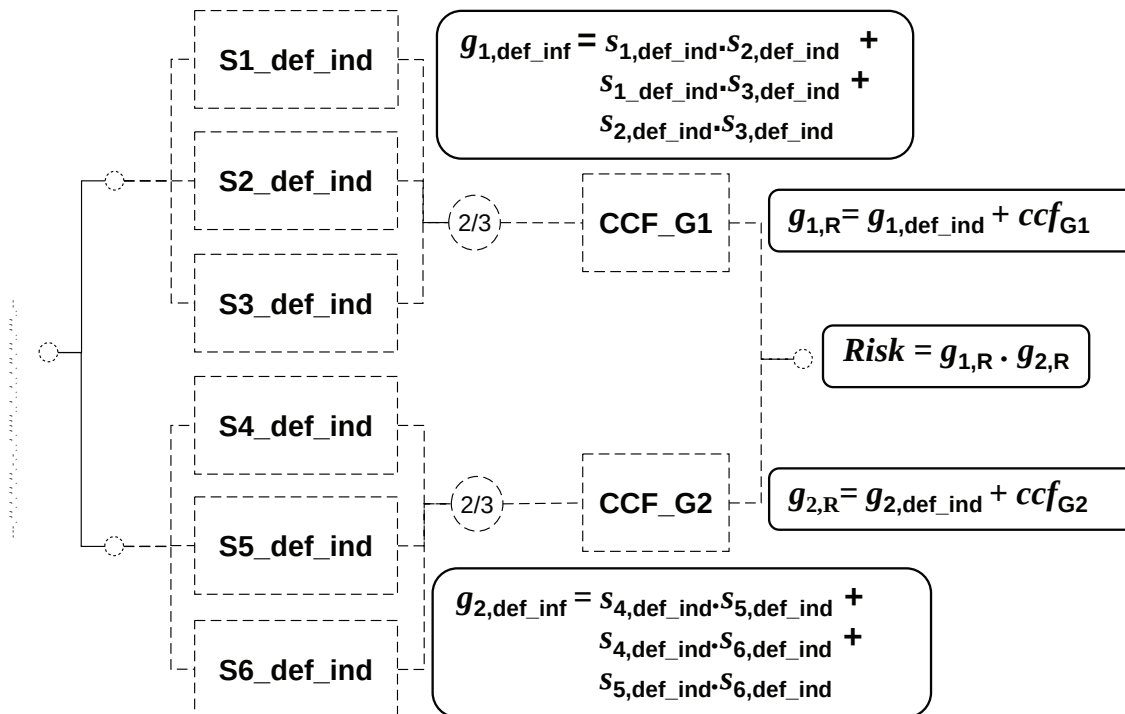


Figure 102 — RBD driven PN modelling the sensors failure

For the moment, this model is equivalent to the fault tree model. The Monte Carlo simulation provides the following results:

- $G1_def_ind = 6,63 \cdot 10^{-3}$ [$6,55 \cdot 10^{-3} - 6,71 \cdot 10^{-3}$] (Fault tree: $6,63 \cdot 10^{-3}$)
- $G2_def_ind = 6,63 \cdot 10^{-3}$ [$6,55 \cdot 10^{-3} - 6,71 \cdot 10^{-3}$] (Fault tree: $6,63 \cdot 10^{-3}$)
- $G1_R = 7,52 \cdot 10^{-3}$ [$7,44 \cdot 10^{-3} - 7,61 \cdot 10^{-3}$] (Fault tree: $7,51 \cdot 10^{-3}$)
- $G2_R = 7,49 \cdot 10^{-3}$ [$7,41 \cdot 10^{-3} - 7,58 \cdot 10^{-3}$] (Fault tree: $7,51 \cdot 10^{-3}$)
- $Risk = 8,47 \cdot 10^{-5}$ [$7,75 \cdot 10^{-5} - 9,18 \cdot 10^{-5}$] (Fault tree: $8,46 \cdot 10^{-5}$)

The results obtained by fault tree and PN are very close.

Now the fact that a double dangerous detected failure within a group of sensor led to the reactor shut down and, then, to the disappearance of the risk can be taken into account. This can be done just by slightly modifying the formula of *Risk*:

$$Risk = G1_R \cdot G2_R \cdot (Nb_{dd_G1} < 2) \cdot (Nb_{dd_G2} < 2)$$

Running the Monte Carlo simulation leads to the following results:

$$Risk = 8,46 \cdot 10^{-5} [7,74 \cdot 10^{-5} - 9,18 \cdot 10^{-5}]$$

As expected, the result is lower than in the previous case but the difference is almost imperceptible. Therefore, in average, this assumption has no impact on the average unavailability of the EDP safety system and the calculations made by fault tree have a very good accuracy.

The models developed just above deal only with the sensor part of the EDP safety system. The whole model could be modelled in the same way but this does not add value to this Technical Report with regard to what is already developed above in this subclause.

14.7 Conclusion about typical applications

The various approaches presented in this Technical Report have been illustrated in [Clause 14](#) and their modelling abilities and limitations have been highlighted. Five typical applications have been analysed in detail with various assumptions in order to encompass a wide variety of typical problems in association with their solutions. This covers most the needs of reliability analysts modelling and calculating the probability of failure of safety systems.

Annex A (informative)

Systems with safety functions

This annex, which presents systems with safety functions, should in principle reflect all such applications for the installations (Level 3) and plants (Level 4) (see A.1.4 in ISO 14224^[15]) and their underlying equipment classes (Level 6) and components (Level 8). Depending on the scope of work for the reliability calculation, building blocks and result parameters may vary between levels 3–8. The table also reflects the equipment classes related to safety systems/components given in Table F.2 of ISO 14224.^[15]

Table A.1 — Safety systems and safety functions for which reliability analysis can be applicable

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components a)	Comments regarding safety system/function and/or useful advice for reliability analysis
A1 – Emergency/ process shutdown			“Loss of containment” would cover the “Safety functions” mentioned.
1. Emergency shutdown (ESD)	Isolation/sectioning (area) Ignition source control Real trip Automatic initiation of ESD according to logic Manual initiation of ESD (incl. push button in field and CAP)	Input devices, Control logic units, Valves, Switchgear (i.e. components like circuit breaker and relays) Equipment classes within rotating (e.g. combustion engines) and mechanical equipment category (e.g. heater and boilers). Power supply: a) Electric: Electric generator, UPS b) Hydraulic power unit c) Air supply equipment Loading arms (LNG), Storage tanks (e.g. high level monitoring), Vessel (e.g. reactors)	Cause and Effect assessment can address dependencies between different systems. Other systems may be an initiator for the ESD system and the ESD system may be a part of other systems. Note: will normally only apply to parts of functions (or elements), ref., OLF GL070 Minimum SIL Table 7 . CAP = Critical Alarm Panel Loading arms may have Emergency Quick Disconnect including hydraulically actuated valves, and such safety functions may be subject to reliability analysis. High temperature protection in reactors beds (refining or petrochemical plant) with potential for runaway may be an issue of reliability concern.

^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in [Annex A](#) of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^{a)})	Comments regarding safety system/function and/or useful advice for reliability analysis
2. Process shutdown (PSD)	Isolation/ sectioning (equipment, process units) Real trip Automatic initiation of PSD according to logic Manual initiation of PSD (incl. push button in field and operator stations)	Input devices, Control logic units, Valves, Switchgear (i.e. components like circuit breaker and relays) Equipment classes within rotating (e.g. combustion engines) and mechanical equipment category (e.g. heater and boilers). Power supply: a) Electric: Electric generator, UPS b) Hydraulic power unit c) Air supply equipment Loading arms (LNG), Storage tanks (e.g. high level monitoring), Vessel (e.g. reactors)	See comment ESD regarding LNG loading arms and reactors.
3. Emergency depressurization (EDP) (Blowdown)	Emergency depressurization Flare ignition	Valves	EDP may be a function within the ESD system or a separate system. Other safety functions can depend on the EDP system.
A2 – Over-pressure protection systems			Safety instrumented system replacing pressure relief functions Vacuum relief system is also a type of pressure protection system.
4. HIPPS (High Integrity Pressure Protection System)	Pressure protection	Input devices, Control logic units, Valves, Switchgear (i.e. components like circuit breaker and relays) Power supply: a) Electric: Electric generator, UPS b) Hydraulic power unit c) Air supply equipment	HIPPS covers also Pipeline Protection System (PPS). See definitions of HIPPS in 3.6 .
^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in Annex A of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.			

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224[15]) (*new equipment class or components ^a)	Comments regarding safety system/function and/or useful advice for reliability analysis
5. Pressure relief system	Pressure protection	Valves (e.g. PSV) *Disc	Non-instrumented (mechanical) secondary pressure protection system (see ISO 10418[33]) are normally not subject to the type of reliability analysis presented in this Technical Report Mechanical system not included in IEC 61508.[2] Reliability calculations may be done on components within the system (e.g. PSV). Table A.68 in ISO 14224[15] describes various types of valves. PVRV (Pressure vacuum relief valve) is a term sometimes used for PSV-vacuum relief.
B – Fire and Gas detection C – Fire water D – Fire fighting			
6. Fire and gas detection	Detection of fire hazards and loss of containment Activation of protection measures	Fire and gas detectors, Control logic units (assumed to cover fire central). Activation of protection functions (equipment), e.g. fire fighting systems. HVAC dampers, ESD functions.	Note: will normally only apply to parts of functions (or elements), ref., OLF GL070 Minimum SIL Table 7
7. Fire water system	Active fire protection	Pumps, valves (incl. deluge valves), Piping, Combustion engines (incl. diesel motors), Electric generators	
8. Fire fighting system	Active fire protection	Valves, Nozzles, Fire-fighting equipment Fire and gas detectors Fire-fighting equipment Foam (AFFF). Equipment class not defined in Table A4 in ISO 14224[15], but in Table F.2. Sprinkler (wet and dry), Deluge. Fire monitors, Water mist, Gaseous agents	

^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in [Annex A](#) of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^{a)})	Comments regarding safety system/function and/or useful advice for reliability analysis
E – Process control			Covered also by A 1 above. Normally not defined as a safety system. A risk reduction factor (RRF) may be accounted for, but then according deterministic criteria given by IEC 61511 ^[3]
9. Process control	Process control, and monitoring	Input devices, Control logic units, Valves, etc.	
10. Chemical injection	Hydrate function Corrosion control		
F – Public alarm			
11. Emergency communication	PA and alarm for emergency announcement	*Public Announcement and Communication System (PACOS) and associated field equipment Input devices, Control logic units Other equipment: Loudspeaker, Radio and Warning lights	Emergency communication is referring to PA and alarm for emergency announcement on an installation upon signal from F&G/ESD. This system will normally be part of the ESD and F&G system. During an emergency situation the PA and alarm system performs an essential safety function by transmitting automatic alarms upon signal from the ESD and/or the F&G system. The main purpose is to alert personnel by generation and broadcasting of emergency alarms (incl. flashing light/ beacon in noisy areas). If the PA and alarm system should fail on demand, the consequences will depend on the scenario. The required reliability/ integrity level for this Safety Instrumented Function (SIF) should be evaluated and assessed as part of IEC 61508 ^[2] / IEC 61511 ^[3] analysis performed for the specific installation.
G – Emergency preparedness			
12. Evacuation system	Evacuate personnel	Evacuation equipment	Lifeboats
^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in Annex A of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.			

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^a)	Comments regarding safety system/function and/or useful advice for reliability analysis
H – Marine equipment			
13. Disconnection system	Move off location in hazardous situation	*Miscellaneous marine disconnection system. In addition process systems and riser equipment are key equipment.	Disconnection system applies for floating production installations such as FPSO, Semi, etc. Hazardous situations could be weather (e.g. typhoon, ice conditions)
14. Station keeping	Collision avoidance	Thrusters Dynamic positioning equipment	Covers Dynamic Positioning (DP) and mooring equipment. DP class requirements (see ISO 19901-7) are based on qualitative requirements (HWFT = Hardware fault tolerance)/Redundancy and not system quantitative requirement (PFD). The international maritime industry does not normally calculate reliability of such systems. ISO 19901 – Part 7 “Station keeping system for floating offshore structure and mobile units”.
15. Loading system	Loss of containment Collision avoidance	*Loading system equipment Dynamic positioning equipment	Different loading systems exist and would imply different analysis approach aligned with such. Different geographical area would affect the need for undertaking detailed analysis.

^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in [Annex A](#) of ISO 14224; b) The indicated equipment class does not exist in ISO 14224, Table A.4.

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^{a)})	Comments regarding safety system/function and/or useful advice for reliability analysis
16. Ballast water	Trim, heeling and emergency pumping (leaks) Typical Safety Instrumented Function (SIF): Start of ballast syst. (de-ballasting and ballasting by pumping + operation of valves). Ballasting by gravity filling (valve operation) Emergency stop of ballasting (pumps/valves). Close watertight doors (to ensure Hull integrity/ buoyancy) Prevent door closing in the event of object/person in door opening (i.e. personnel safety)	Valves, Input devices, Control logic units (incl. relays), UPS (incl. Circuit breaker/ Contactor), Power transformer, Switchgear/ switchboards, Pumps, Hydraulic power units	Safety strip for watertight doors is a SIF that has been introduced due to a fatal accident related to operation of the self-closing watertight doors on a Norwegian floating production platform.
I – Electrical and Telecommunications			
17. Uninterruptable Power Supply (UPS)	Power to PA, F&G, EDP and ESD	UPS	
18. Emergency power and lighting	Charging UPS Ensure safe escape and evacuation Emergency power supply to well barriers during drilling/ well intervention operations (mud pumps, etc.)	Combustion engines, Electric generator (emergency generator driven by combustion engine for emergency power)	
19. Telecommunications	Transmission of safety function signals (PPS, HIPPS, ESD, etc.).	*Telecommunications system includes transmission systems (Fibre optical cable, Radio link or Satellite) and network equipment.	This system will interface to the PSD, ESD, HIPPS and the fire and gas systems. Telecommunication is a vital part of safety systems where initiator and final elements are on different installations (e.g. pipeline protection systems). It is important to evaluate telecommunication system to obtain a robust and safe design.
J – Other utilities			
^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in Annex A of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.			

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^{a)})	Comments regarding safety system/function and/or useful advice for reliability analysis
20. Flare system	Safe disposal of gas Flare igniting system Automatic ignition of flare upon EDP	*Flare ignition system	HIPPS systems may impact the design of flare system. The ISO 23251 ^[31] allows that the design load from disposal system (flare or vent stack system) takes credit from the favourable response or HIPPS, operator intervention or basic process control. The disposal system performance as a whole may have to be assessed taking into account the likelihood of the overpressure contingencies and the reliability of the multiple layers of protection that reduce or eliminate individual relief loads.
21. Natural ventilation and HVAC	Ensure minimum ventilation rate to avoid gas build up. Overpressure against classified areas Shutdown of dampers upon gas detection Prevent ingress of gas through main air intake (upon gas detection) Close fire damper to critical room Cooling of critical rooms/ equipment (by DX units, fin-fan coolers, fan coils, etc.) Smoke control	Input devices, Control logic units,, Valves, Switchgear (i.e. components like circuit breaker and relays) Power supply: a) Electric: Electric generator, UPS b) Hydraulic power unit c) Air supply equipment	
22. Materials handling	Lift MOB boat Safety functions related to minimizing risk while use of crane	Crane	Materials handling (crane, drawworks) may represent risk requiring reliability analysis. Man-machine interfaces
K – Drilling and Wells			Including subsea drilling and subsea well completion Note that “Drilling”, “Well completion (downhole)”, “Mechanical” and “Well intervention” are ISO 14224 ^[15] Equipment categories (see Table A.4).
^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in Annex A of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.			

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^{a)})	Comments regarding safety system/function and/or useful advice for reliability analysis
23. Drilling	Well integrity Dropped object protection Workplace safety	Examples of drilling equipment classes are: — Blowout preventer (BOP) — Mud pumps — Diverter	The importance of the Blowout Preventer (BOP) as well barrier function implies that the BOP is subject to specific reliability analysis, and related to the overall drilling/intervention and/or facility risk analysis. Likewise for the diverter system. Reliability analysis of operations representing potential dropped object risk for personnel at rig floor or subsea equipment may require the need to do reliability analysis of such rig/drilling equipment. Undertaking such analyses should benefit from the guidance given in this Technical Report and applied as appropriate with its analysis context and decision-support.
24. Well completion (down-hole)	Well integrity	DHSV (SCSSV) Electrical submersible pumps (ESP) Packers	Subsea, onshore or offshore completed wells ESP is not a safety system, but could be part of a reliability analysis addressing safety. DHSV/SCSSV: This valve is normally part of a safety system/safety functions for which reliability analysis can be applicable. E.g. it could be part of the ESD safety loop (isolation of well).
25. Well completion (surface)	Well integrity	Xmas tree (topsides/onshore)	Note that dry Xmas tree belongs to equipment category “Mechanical” in ISO 14224 ^[15] (Table A.4). Reliability analysis may be done to assess test interval for Xmas tree valves, and also for DHSV
^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in Annex A of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.			

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^{a)})	Comments regarding safety system/function and/or useful advice for reliability analysis
26. Well intervention	ESD, PSD functions Well integrity and isolation	Various BOPs and control systems Various other pressure-control equipment and systems	Workover control system (WOCS) is part of overall well intervention systems Note that this covers well interventions using a top-side BOP only, otherwise see item 29 below for subsea well intervention which covers WOCS for such.
27. Rig disconnection	Move rig off hazardous location	Drilling and completion riser	The rig disconnection system may be related to the ESD system on the facility. See also Subsea well intervention.
L – Subsea			Note that “Subsea production” is an ISO 14224 ^[15] Equipment category (see Table A.4). “Subsea pipeline”, “Subsea well intervention” and “Saturation diving” reflects in this Table A.1 subsea related systems
28. Subsea production	Subsea isolation: — Topside initiated subsea isolation (ESD functions), — Topside initiated PSD functions — Subsea initiated PSD functions — Subsea HIPPS Subsea isolation of well, manifold, flowline and pipeline, subsea processing, etc.	Subsea Xmas tree Subsea isolation equipment, e.g. by use of SSIV (Subsea Isolation Valve) or subsea HIPPS valves. Subsea production control Other subsea equipment classes	See also ESD function above. OLF GL 070 presents calculation example for ESD, and PSD for Subsea production, but these calculations are subject to somewhat uncertainty Relations to API RP 170, Subsea HIPPS and ISO 13628–6 ^[56] , production control should also be noted. ISO 13628–1 ^[55] and other ISO 13628 standards may be relevant as well.
^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in Annex A of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.			

Table A.1 (continued)

System	Safety function (examples)	Equipment class (Level 6/ ISO 14224 ^[15]) (*new equipment class or components ^{a)})	Comments regarding safety system/function and/or useful advice for reliability analysis
29. Subsea pipeline	Subsea isolation: — Topside initiated sub-sea isolation (ESD functions), — Topside initiated PSD functions — Subsea initiated PSD functions — Subsea HIPPS Subsea isolation of flowline and pipeline.	Flowlines Pipelines Subsea isolation equipment, e.g. by use of SSIV (Subsea Isolation Valve) or subsea HIPPS valves.	Subsea infield flowlines, intrafield pipelines and export pipeline are here covered. Note that onshore pipelines may have similar issues. ISO 16708 ^[52] should also be used as appropriate, and ISO 13628-1 ^[55] and other ISO 13628 standards may be relevant as well. Note that ISO 13623 ^[53] (e.g. 12.4) addresses functional testing of pipeline equipment.
30. Subsea Well intervention	ESD, PSD and EQD functions Well integrity and isolation Workover Riser monitoring systems and operation on dynamic positioning Workover Riser weak-link in combination with drill compensator function	Various subsea well intervention equipment and well control systems; *a) Riserless well intervention *b) Open water intervention *c) Thru-BOP/Drilling Riser intervention	Workover control system (WOCS) is part of subsea well intervention Subsea production control system is similar to Workover Control system, however used differently and consequence of failure and failure modes are very different. ISO 13628-7 ^[57] covers Safety and Reliability on Workover Control Systems. See also API RP 17G. ^[60]
31. Saturation diving	Divers environmental control systems	*a) Life support systems *b) Pressure control systems	Diving systems are always subject to international class societies (DNV, Lloyds etc.). This covers all safety systems in traditional diving equipment. The latest generation of equipment using PLC technology has been subject to IEC 61508 ^[2] and SIL calculation accordingly. The life support system equipment includes hyperbaric chamber and oxygen bottle.

^a Please note this difference which is useful for input to next revision and relations to ISO 14224 (at equipment class level or below): a) The equipment class does exist in Table A.4 in ISO 14224, but no example is further given in Annex A of ISO 14224;b) The indicated equipment class does not exist in ISO 14224, Table A.4.

Annex B (informative)

State analysis and failure classification

B.1 On-demand mode of operation

Figure B.1 gives a synthesis of the different classes of states that may be considered when analysing a safety system operating on demand mode:

- nominal state;
- class of states less safe than the nominal state (“less safe” and “critical with regards to safety action inhibition” states);
- class of states safer than the nominal state (“safer” and “critical with regards to spurious actions” states);
- class of state where the safety action is completely inhibited (“KO”);
- the state where the hazardous event has occurred;
- the state where the spurious action has occurred.

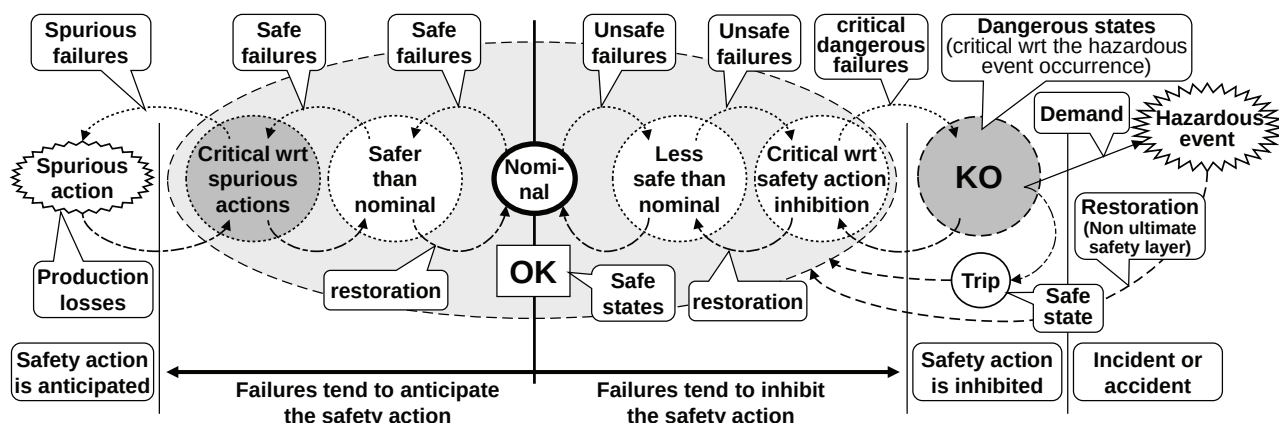


Figure B.1 — Synthesis of states and failures classification for on demand mode safety system

The OK class gathers, in the big ellipse in the middle of Figure B.1, all the classes of states where the safety system is able to operate properly on demand. All these states are “safe” because the safety action is going to be successful on demand.

The KO class gathers, in the grey circle at the right hand side of Figure B.1, all the states where the safety system is completely inhibited. All the states within this class are “dangerous” (or “unsafe”). The taxonomy of the failures shown on this figure is described in B.3.

The “Trip” class at the right hand side of Figure B.1 represents a safe state where the safety action is no longer needed because, e.g. the process has been stopped.

NOTE 1 All the classes of states presented on Figure B.1 do not necessarily exist. For the simplest safety systems, the nominal state may be also critical with regards to safety or spurious actions.

NOTE 2 An unsafe failure does not necessarily lead to a dangerous state. It may only lead to a state where the probability of success, in case of demand, is lower than in the nominal state.

NOTE 3 Industrial safety systems are generally organized in several safety layers acting in sequence (multiple safety system). In this case, if the safety system in [Figure B.1](#) does not constitute the ultimate safety layer, a demand occurring in the state KO will not lead to a “hazardous event” but to a demand on the succeeding safety layer which would, normally trigger the safety action. This will reveal the states KO as well as the failure which has lead to the demand: therefore they can be restored before a “hazardous event” actually occurs. If the safety system is the ultimate safety layer, a demand occurring in the state KO immediately leads to the “hazardous event”: therefore, no restorations can be undertaken.

NOTE 4 The “KO” and “hazardous event” classes are closer and closer in time when the demand rate increases. The limit is when the safety system operates in purely continuous mode where they are merged in a single class (see [Figure B.3](#)).

NOTE 5 The “safer than nominal” class is safer from the considered safety function point of view but degraded from spurious action point of view (which may have in turn a side effect on safety).

NOTE 6 Transitions are potentially possible from each class of states to each other. For the sake of simplicity only the most illustrative have been kept on [Figure B.1](#).

NOTE 7 In de-energize to trip design, all the failures able to lead to a spurious action should be considered (see [5.3](#)). This may imply other components than those considered within the dangerous failure analysis. The safe failure rates should be evaluated accordingly.

[Figure B.1](#) brings to light three critical state classes:

- the states of the class KO (on the right hand side) which are critical with regard to the occurrence of the hazardous event. They are distant from the hazardous event by only one event (i.e. the demand of the safety action).
- the states of the class “critical with regard to the safety action inhibition” (on the right hand side). They are distant from the KO state by only one event (i.e. the dangerous critical failures).
- the states of the class “critical with regard to the spurious actions” (on the left hand side). They are distant from the spurious action by only one transition (i.e. the spurious failures).

[Figure B.2](#) illustrates on a typical simple 2oo3 system what is presented in a general way on [Figure B.1](#). Each of the three components may have 3 states (W: working, D: dangerous undetected fault, S: safe fault). Then, for example, at the system level, the state 2WD means: 2 component working and 1 dangerously faulty. The safe failures are immediately revealed and repairs start immediately; the dangerous failures are periodically tested and are not repairable between tests.

This figure shows that all the classes presented in [Figure B.1](#) do not necessarily exist, that they are interlinked and that some states can belong to several classes (e.g. the state “WDS” is critical both from dangerous and spurious failures point of views. The state “2WS” illustrate also a state which is safer than the nominal state because, due to the safe fault, the 2oo3 has been switched to 1oo2.

Due to the common cause failures the system may jump directly to the class KO or generate a spurious action. Therefore all the states in the class OK are also critical states both from the inhibition of the safety action and of spurious failures point of views. CCFs have been represented only for the nominal state 3W to simplify [Figure B.2](#).

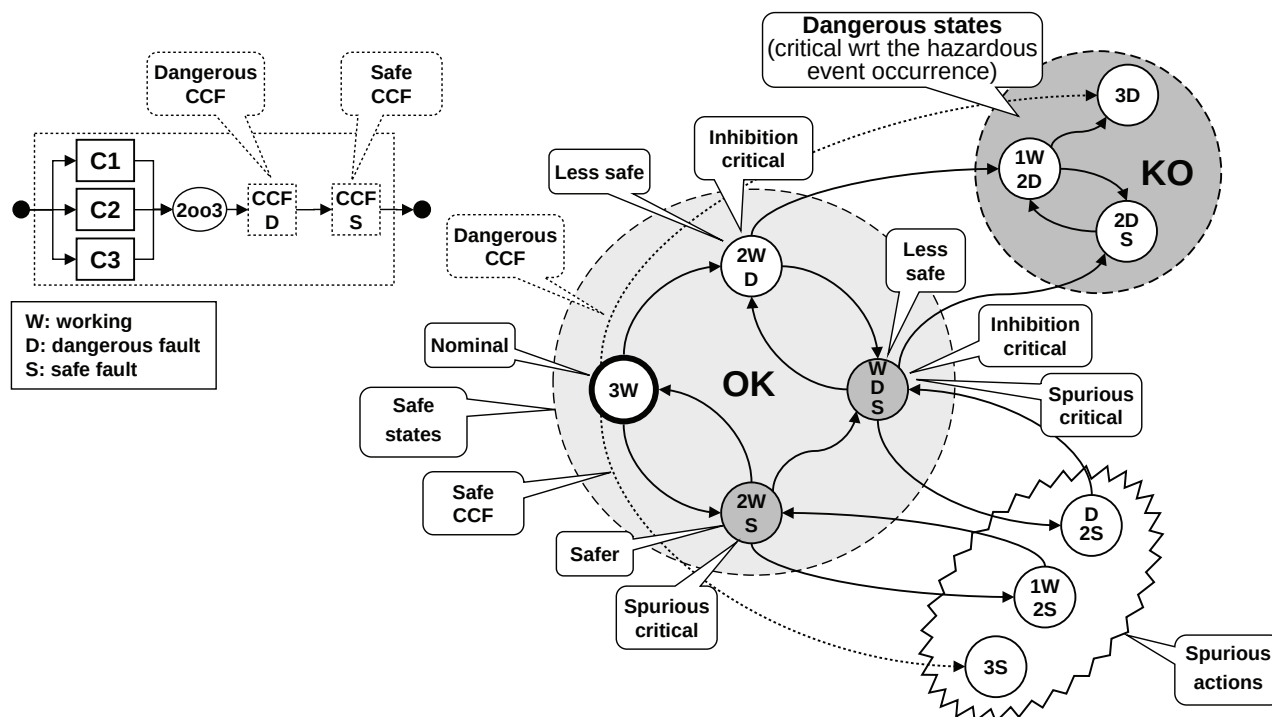


Figure B.2 — Example of the classes of states for a 2oo3 system

B.2 Continuous mode of operation

Figure B.3 gives a synthesis of the different classes of states that may be considered when analysing a safety system operating on continuous mode.

These classes are similar with the classes identified for safety systems working on demand mode. The main difference is on the right hand side of the Figure B.3 when a critical dangerous failure occurs. Depending on the kinetics of the process more or less time may be available before the hazardous event occurs. This may be used to detect the critical failure and the small circle represents the state when this detection is possible. If the critical dangerous failure is detected quickly, the protected system can be tripped before the occurrence of the hazardous event. If it is not detected then the hazardous event actually occurs and the “KO” and “hazardous event” classes are merged in a single class. The taxonomy of the failures shown on this figure is described in B.3.

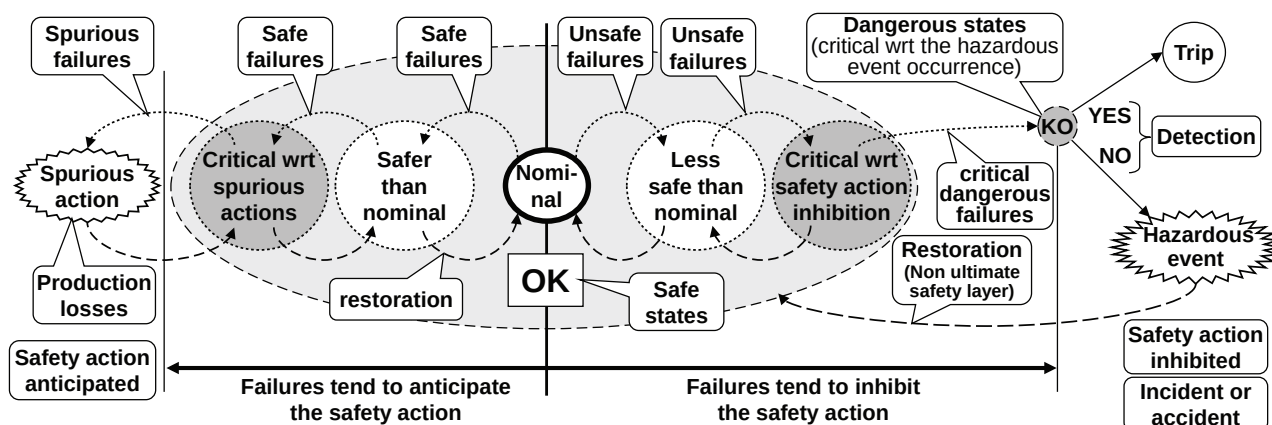


Figure B.3 — Synthesis of states and failures taxonomy for one continuous mode safety systems

NOTE 1 Industrial safety systems are generally organized in several safety layers acting in sequence (multiple safety system). Therefore, if the safety system in [Figure B.3](#) does not constitute the ultimate safety layer, the critical failure leading to the “hazardous event” generates a demand toward a succeeding safety layer which can trip the installation to a safe state. In this case the critical dangerous failure can be restored before an accident actually occurs. This restoration is not possible if it is the ultimate safety layer because the accident occurs at once.

NOTE 2 The “safer than nominal” class is safer from safety function point of view but degraded from spurious action point of view (which may have in turn a side effect on safety).

NOTE 3 Transitions are potentially possible from each class of states to each other. For the sake of simplicity only the most illustrative have been kept in [Figure B.3](#).

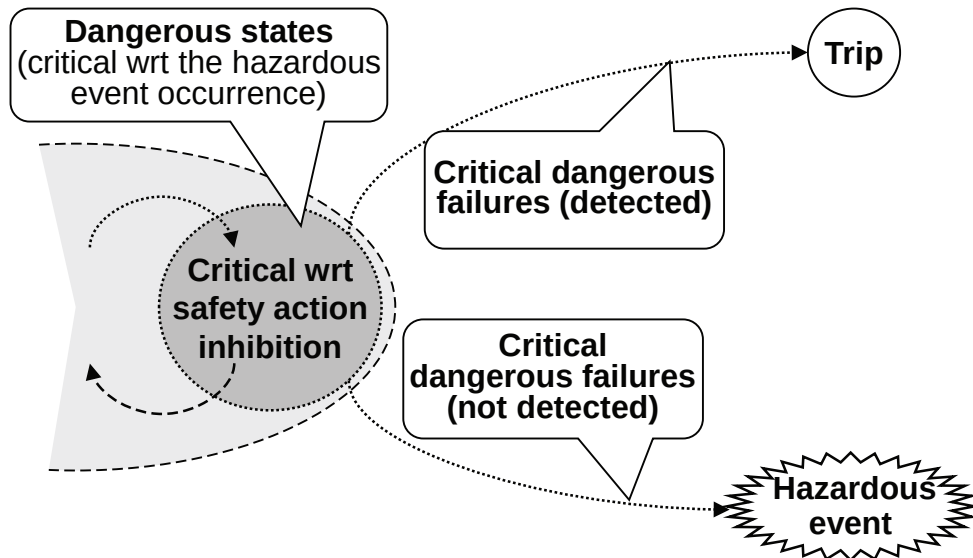


Figure B.4 — Equivalent representation of the states in left hand side of [Figure B.3](#)

The [Figure B.3](#) exhibits the same critical states classes as in [Figure B.1](#) but the class of the states KO (on the right hand side in dotted line) now encompasses states with no real existence as they immediately give either:

- a real trip of the protected system (if they are detected); or
- an hazardous event (if they are not detected).

Therefore the state KO can be removed of [Figure B.3](#) to obtain [Figure B.4](#). Then the states which are critical with regards to the inhibition of the safety action are also critical with regards to the hazardous event.

B.3 Failure classification

The failures may be classified in several ways as illustrated in [Figure B.5](#), [Figure B.6](#) and [Figure B.7](#).

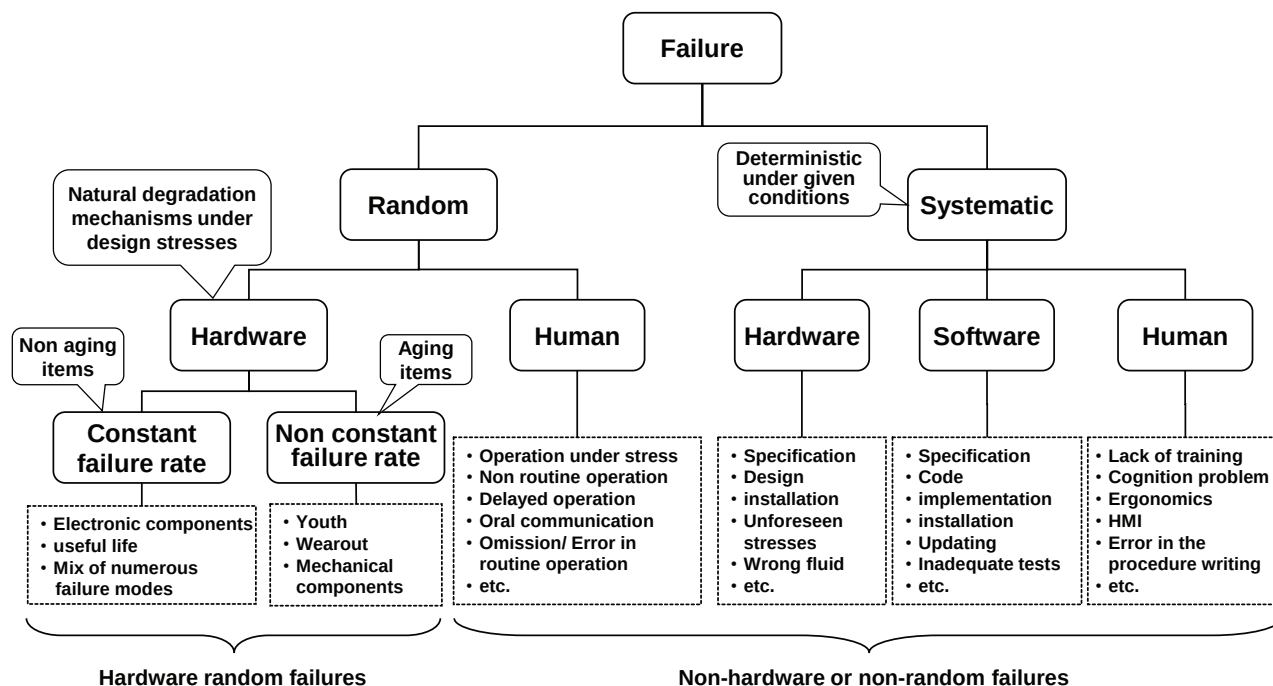


Figure B.5 — Random versus systematic failures

The classification illustrated in Figure B.5 split the failures between those who occur automatically (i.e. in a deterministic way) when some conditions are met (systematic failures, see 3.2.17) and those which occur only with a given probability (random failures, see 3.2.16). This Technical Report mainly deals with random hardware failures and random human failures which may be, to some extent, handled in the same way. There is no specific development in this Technical Report for systematic failures which are covered to the extent that they are part of the failure rates elaborated from the field feedback. See ISO 14224[15], Annex B for some advices on this topic.

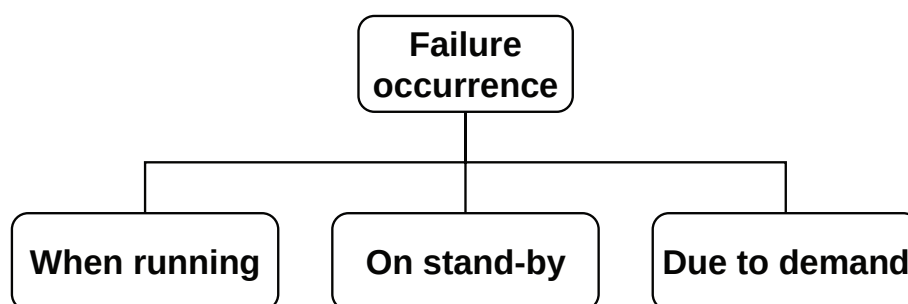


Figure B.6 — Classification according to the state of the failing item

Figure B.6 classifies the failures according to the state of the system when the failure occurs. Failures can occur when the system is running (e.g. continuous mode safety system), on stand by (e.g. on demand mode of operation safety system) or when changing of state due to a demand.

Figure B.7 classify the failures according to the way that they are detected. They may reveal themselves or not and when they do not reveal themselves, diagnostic tests or periodic tests can be performed.

The diagnostic test frequency is generally high enough to consider that the detected failures are discovered almost immediately and can be merged with the ordinary immediately revealed failures as defined in Note 1 to entry of 3.2.10.

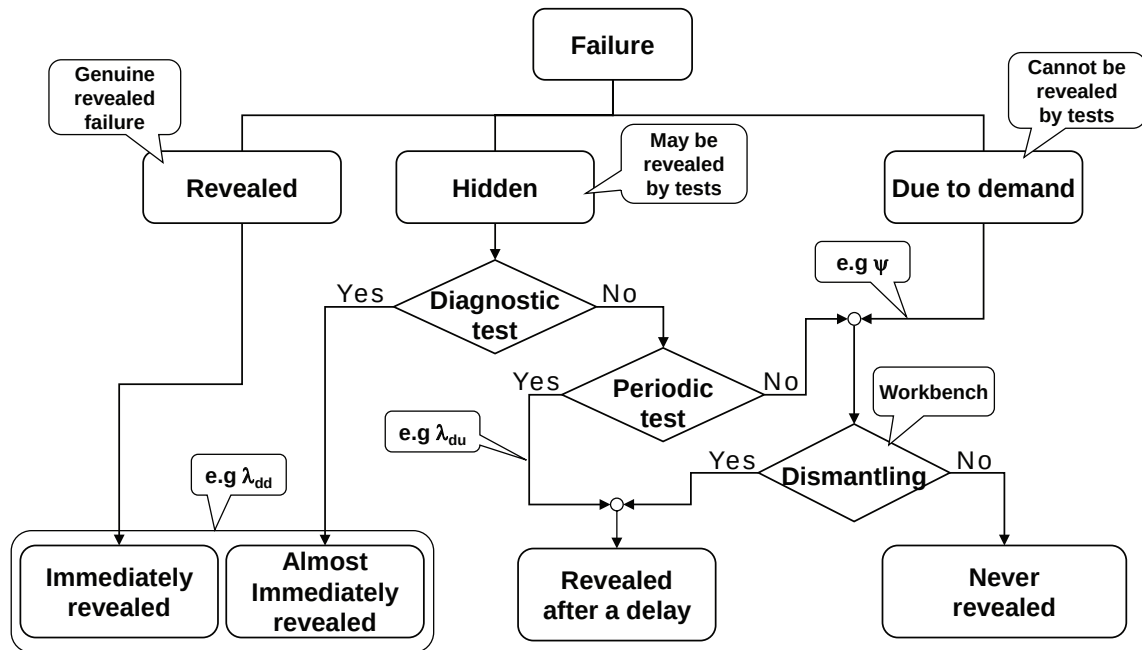


Figure B.7 — Classification according to the detection means

Sometimes there is a residual part of the failures which cannot be revealed by these means (e.g. the failures occurring due to the change of state of a safety system following a demand of the safety action). In this case the item may be dismantled on a workbench to bring the incipient failures to light in order to prevent them to propagating to a complete failure. If this is not sufficient or possible, the component should be periodically replaced by a new one.

Annex C (informative)

Relationship between failure rate, conditional and unconditional failure intensities and failure frequency

In this annex the overall system parameters are noted by $\Lambda_{eq}(t)$ (system failure rate), $\Lambda_v(t)$ (conditional failure intensity), $w(t)$ (unconditional failure intensity) and $f(t)$ (failure density). These parameters are often mixed up because their definitions seems close to each others and because they have close numerical values in simple cases. Understanding the differences between these parameters is very useful to perform relevant reliability, availability and frequency calculations. It should be noted that the same relationships exist between $\Lambda_{eq}(t)$, $f(t)$, $R(t)$ on one side and $\Lambda_v(t)$, $w(t)$, $A(t)$ on the other side.

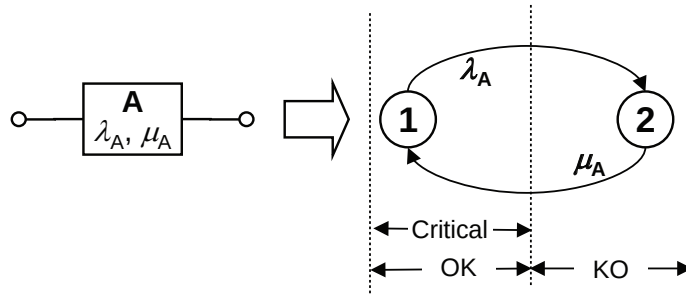


Figure C.1 — Single repairable component

Figure C.1 represents the simplest case made of a single repairable component (or a series of single repairable components) with a failure rate λ_A and a restoration rate μ_A . The Markov graph (see Clause 9) on the right hand side represents the behaviour of this simple system. The two states have been split into two classes:

- OK: the system is working properly (state 1). This is also a critical state as only one failure transition is needed to reach the failed state.
- KO: the system is failed (state 2).

With such a simple case, the analytical formulae are well known:

$$\begin{aligned} R(t) &= P_1(t) = \exp(-\lambda_A t) & A(t) &= \{\lambda_A + \mu_A \cdot \exp[-(\lambda_A + \mu_A)t]\} / (\lambda_A + \mu_A) \\ f(t) &= P_1(t) \cdot \lambda_A = \lambda_A \cdot \exp(-\lambda_A t) & w(t) &= A(t) \cdot \lambda_A \\ \Lambda_{eq}(t) &= f(t) / R(t) = \lambda_A & \Lambda_v(t) &= w(t) / A(t) = \lambda_A \end{aligned}$$

Typical results are presented in Figure C.2:

- The conditional failure intensity $\Lambda_v(t)$ and the failure rate $\Lambda_{eq}(t)$ which are constant and equal to the component failure rate λ_A .
- The unconditional failure intensity $w(t)$ is lower than $\Lambda_v(t)$ because $A(t)$ is lower than 1.
- The failure density $f(t)$ is decreasing and goes to 0 when t goes to infinity.

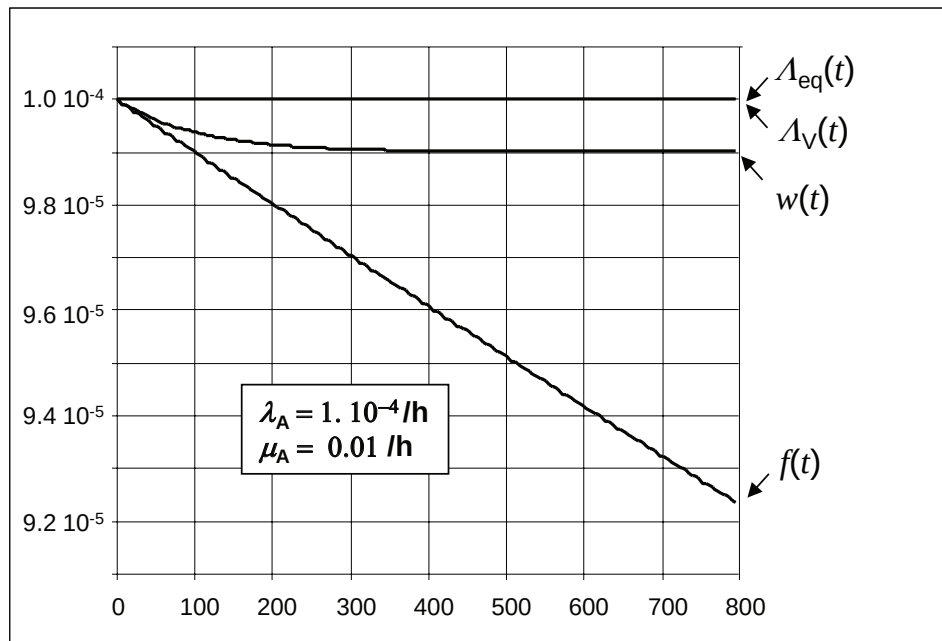


Figure C.2 — Comparison of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_v(t)$ related to Figure C.1

Figure C.3 represents the simplest case made of a single periodically tested repairable component (or a series of single periodically tested repairable components) with a repair rate μ_A . The Markov graph (see Clause 9) on the right hand side represents the behaviour of this simple system.

The three states have been split into two classes:

- OK: the system is working properly (state 1). This is also a critical state as only one failure transition is needed to reach the failed state.
- KO: this class comprise two state:
 - 2: the system is failed undetected
 - 3: the system is under repair

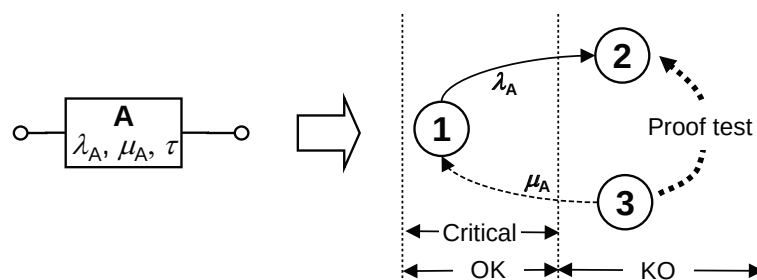


Figure C.3 — Single periodically tested repairable component

Figure C.4 illustrates typical results. Like in the non-periodically tested case, $\Lambda_v(t)$ and $\Lambda_{eq}(t)$ are equal to λ_A and $f(t)$ goes to 0 when t goes to infinity. The unconditional failure intensity $w(t)$ is a saw-teeth curve with an average $\bar{w}(t)$ over $[0, t]$ reaching an asymptotic value.

Those periodically tested components are often approximated by using the model presented in Figure C.1 with $\mu_A = 2/\tau$ where τ is the test interval. This lead to the same $\Lambda_{eq}(t)$ and $\Lambda_v(t)$ and to $w_{approx}(t)$ which is an approximation of $w(t)$. Figure C.4 shows that, with our example, this approximation is conservative.

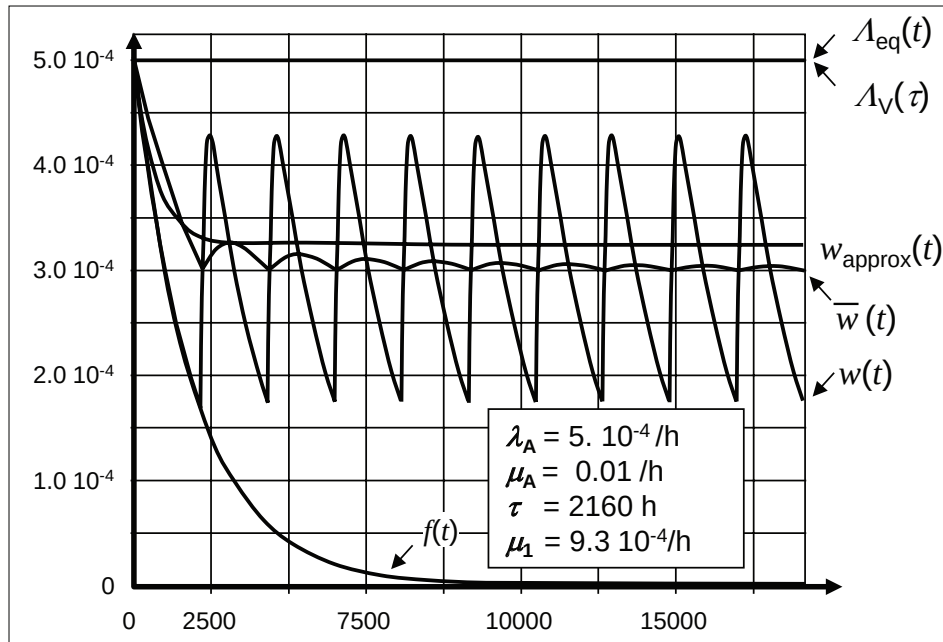


Figure C.4 — Comparison of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_v(t)$ and $w_{approx}(t)$ related to [Figure C.3](#)

[Figure C.5](#) represents a simple system made of 2 identical repairable components with the same failure and repair rates (λ , μ). The Markov graph (see [Clause 9](#)) on the right hand side represents the behaviour of this simple system where the similar states have been aggregated in order to simplify the model.

The states have been split between two classes:

- OK: it comprises two states where the system is working properly
 - 1: not critical (more than one failure to get the KO state)
 - 2: critical (only one transition to get the KO state)
- KO: the system is failed

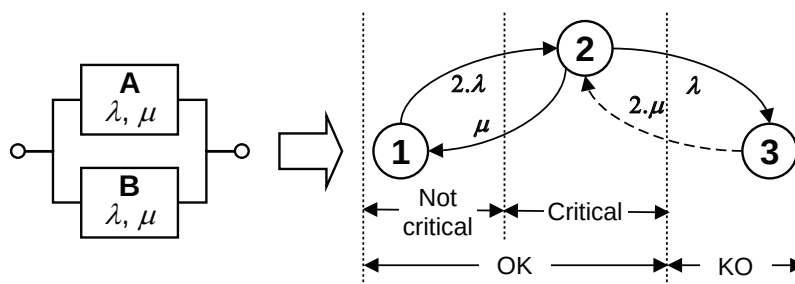


Figure C.5 — Simple system with two identical repairable components

When the KO state is not repaired (i.e. when the transition in dotted line is removed and the state 3 becomes an absorbing state) this graph allows calculating the following parameters:

- The reliability $R(t) = P_1(t) + P_2(t)$ where $P_1(t)$, respectively $P_2(t)$, is the probability to be in state 1, respectively in state 2.
- The failure density $f(t) = \lambda P_2(t)$.

The equivalent failure rate $\Lambda_{eq}(t) = f(t) / R(t)$.

When the KO state is repaired (i.e. the transition in dotted line is applied and the state 3 is a repairable state) this graph allows calculating the following parameters:

- The availability $A(t) = P_1(t) + P_2(t)$ where $P_1(t)$, respectively $P_2(t)$, is the probability to be in state 1, respectively in state 2.
- The conditional failure intensity $w(t) = \lambda P_2(t)$.
- The equivalent failure rate $\Lambda_V(t) = w(t) / A(t)$.

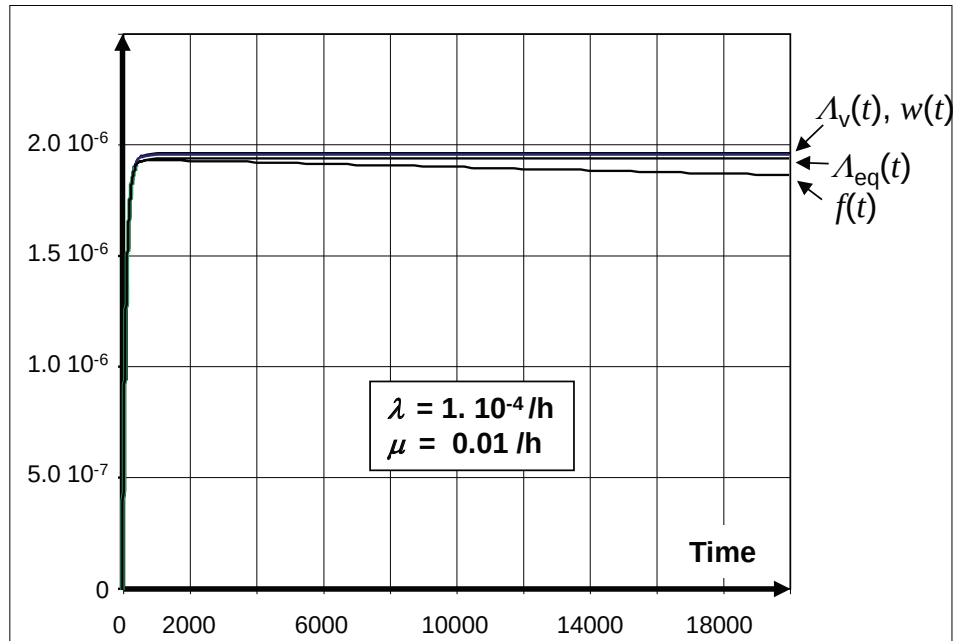


Figure C.6 — Comparison of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_V(t)$ related to Figure C.5

Therefore this is the structure of the Markov graph which decide if either $R(t)$, $f(t)$ or $\Lambda_{eq}(t)$ are calculated or $A(t)$, $w(t)$ and $\Lambda_V(t)$.

This example is typical of systems where the component failures are quickly detected and repaired. In this case $\Lambda_V(t)$, $\Lambda_{eq}(t)$ and $w(t)$ have very close numerical values even on the long term for which they have almost the same asymptotic values. Note that these approximations are robust event if the MTTRes is not so short (here 50 h are needed to repair the KO state).

On the contrary, $f(t)$ has the same numerical value only on the short-term because it goes to 0 when t goes to infinity. Therefore it can be used as an approximation of $\Lambda_{eq}(t)$ or $w(t)$ only on the short-term when it is $\ll 1$.

Figure C.7 shows another example: the individual components are not repaired alone and the repair occurs only when both A and B have failed (repair rate μ_S).

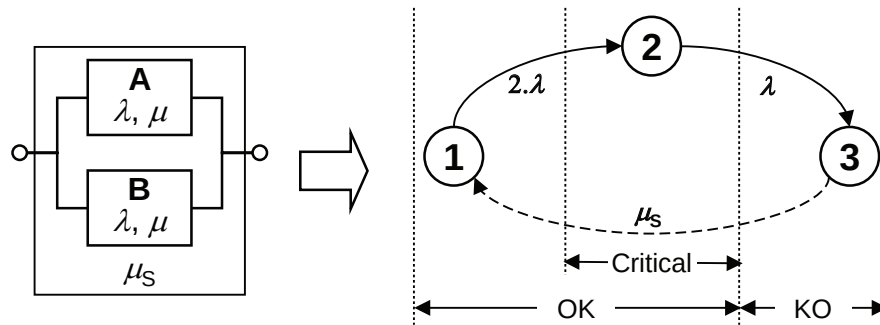


Figure C.7 — Simple system repaired after whole failure

Figure C.8 shows the same parameters as in Figure C.6 but the behaviour is not the same and the parameters have similar numerical values only on the short-term.

If $\Lambda_V(t)$ and $w(t)$ remains close each other, the values of $f(t)$, $\Lambda_{eq}(t)$ and $w(t)$ diverge from each others when the time is increasing:

- $f(t)$ increases first and then decreases to its asymptotic value, 0, after having reached a maximum value.
- $\Lambda_{eq}(t)$ reaches an asymptotic value but it is clearly greater than these of $\Lambda_V(t)$ and $w(t)$.
- $\Lambda_V(t)$ and $w(t)$ have very close numerical value even on the long term because the asymptotic value of $A(t)$ remains close to 1.

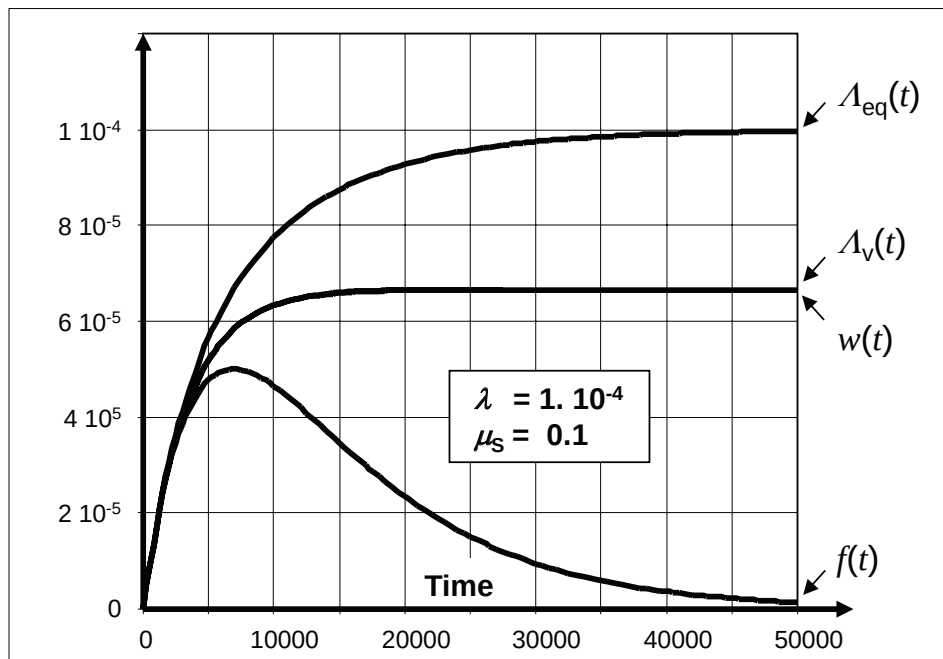


Figure C.8 — Comparison of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_V(t)$ related to Figure C.7 (quick repair)

Therefore in this case $\Lambda_V(t)$ and $w(t)$ does not constitute conservative estimates of the equivalent failure rate $\Lambda_{eq}(t)$ of the system. This is typical of systems quickly repaired as a whole but with no individual repairs of the components failures.

When the MTTRes of the KO state increases, the asymptotic value of $A(t)$ decreases and cannot be assimilated to 1. Then $\Lambda_V(t)$ and $w(t)$, beyond the short-term, have different numerical values and reach different asymptotic values. This general case is illustrated by Figure C.9 where $f(t)$, $\Lambda_{eq}(t)$, $w(t)$ and $\Lambda_V(t)$ are really different and should not be mixed up.

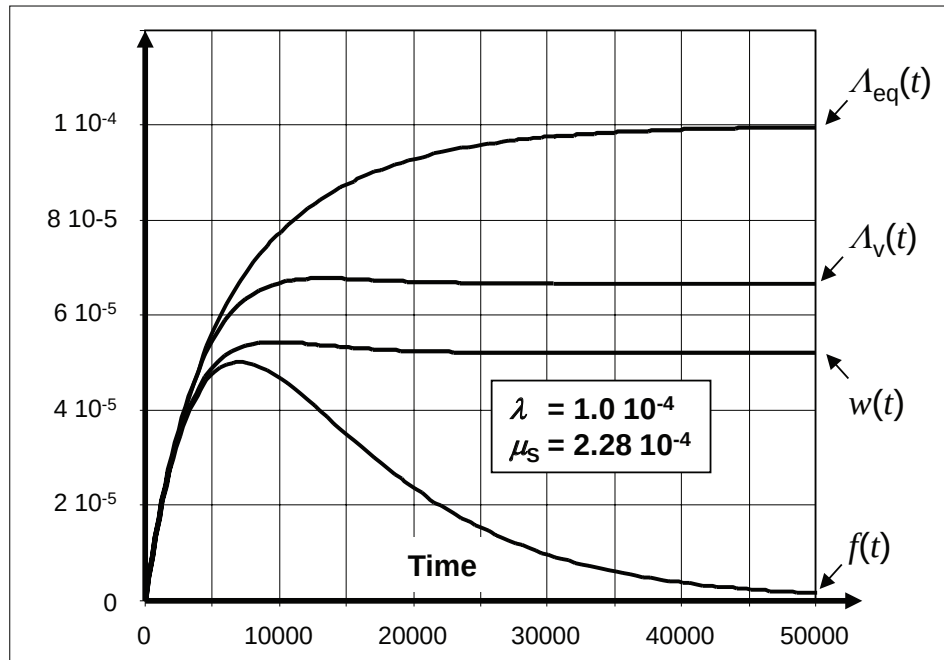


Figure C.9 — Comparison of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_v(t)$ related to Figure C.7 (long repair)

The above examples were dealing with immediately revealed failure (even if the last one can be considered as an approximation of a periodically tested system with a periodic test interval of $\tau = 2/\mu_S = 8\,760$ h).

Figure C.10 illustrates a simple periodically tested system where the redundant components are tested at the same time. The model on the right hand side is a multiphase Markov model which models $A(t)$ when the repair transitions in dotted lines are considered and $R(t)$ otherwise. The impact of periodic tests is represented by bold dotted arrows: the state 2 at the end of a periodic test interval gives the state 3 at the beginning of the following test interval, the state 4 gives state 5 and state 6 gives also state 5.

Due to the periodic tests, the behaviours of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_v(t)$ are of course rather different of the three cases depicted above. This is presented in Figure C.11.

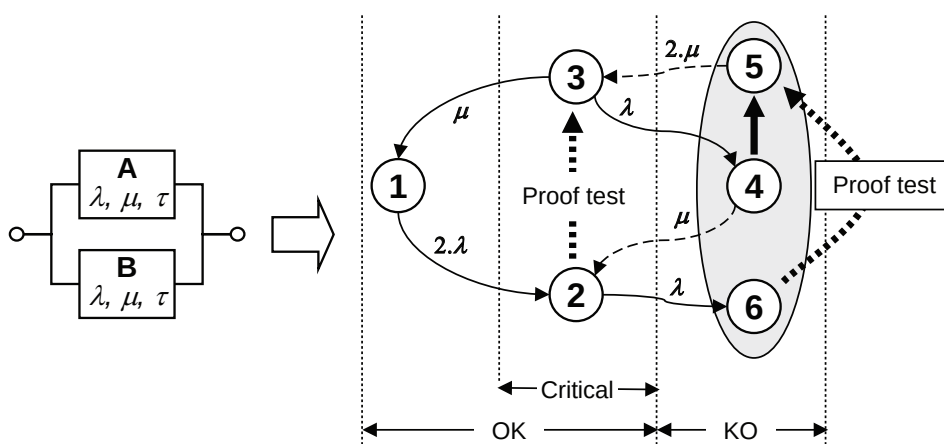


Figure C.10 — Simple system with periodically tested components

Those curves are far from simple and, contrarily as above with immediately revealed failures, have no asymptotic values. It is necessary to consider their averages to make some asymptotic trends happening.

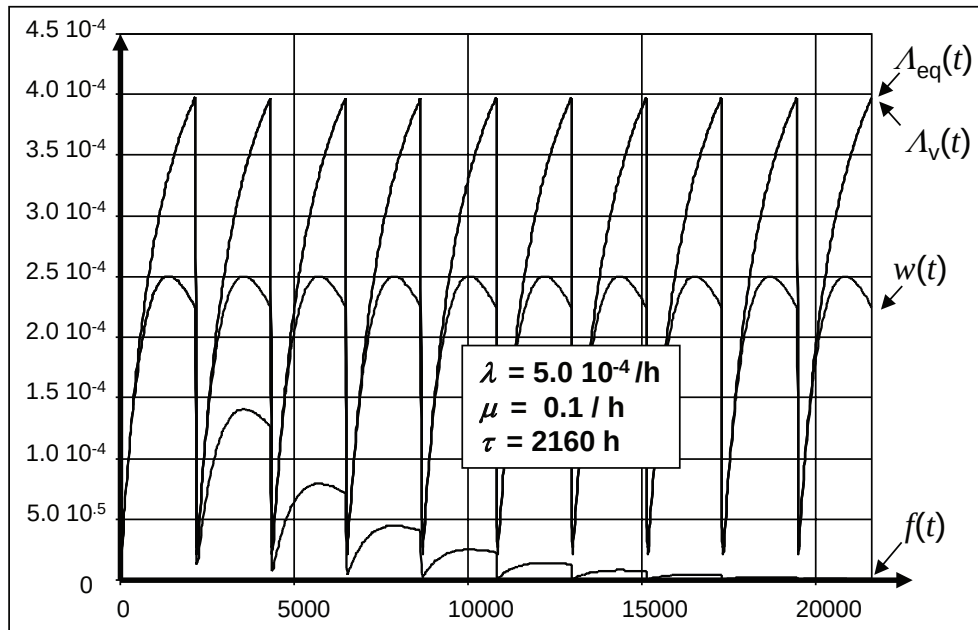


Figure C.11 — Comparison of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_v(t)$ related to Figure C.10

In Figure C.12 the various parameters have been averaged on $[0, t]$. The shapes of the curves are more similar to those established in Figure C.8 or Figure C.9. Again asymptotic values are reached when t goes to infinity. With the numerical values which have been used, $\bar{\Lambda}_{eq}(t)$ and $\bar{\Lambda}_v(t)$ have close numerical values. The asymptotic value $\bar{w}(t)$ is lower than the previous ones and again $\bar{f}(t)$ goes to 0 when t goes to infinity.

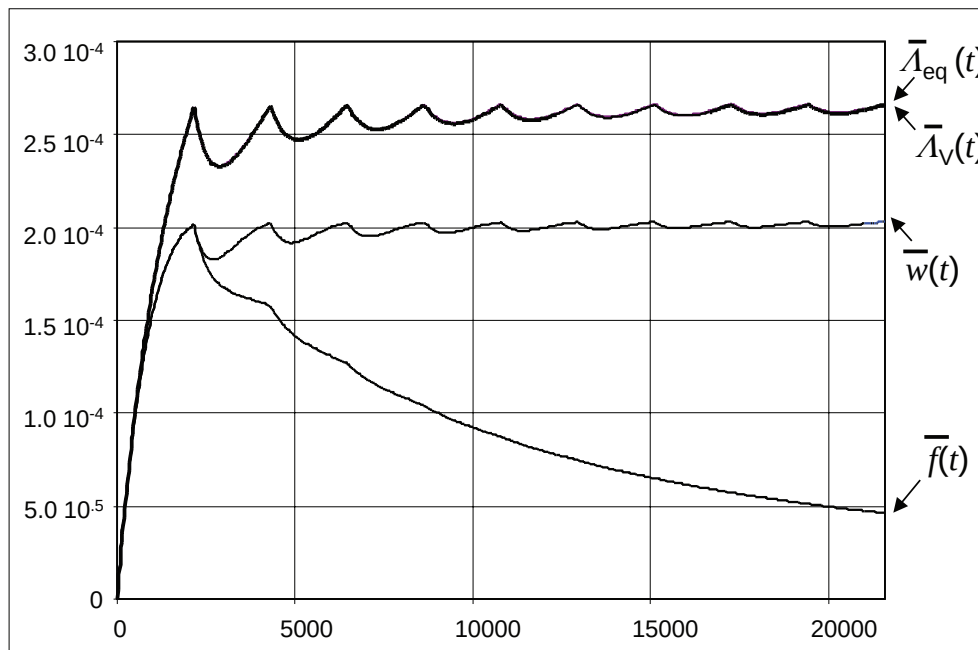


Figure C.12 — Comparison of the averages of $f(t)$, $\Lambda_{eq}(t)$, $w(t)$, $\Lambda_v(t)$ related to Figure C.10

After analysing these six different cases it is possible to draw general conclusions:

- In the general case $\Lambda_v(t)$, $\Lambda_{eq}(t)$ and $w(t)$ have similar values only on the very short term.
- $\Lambda_v(t)$, $\Lambda_{eq}(t)$ and $w(t)$ or their averages reaches asymptotic values which, in some cases, can be used as estimations of each others.

- The asymptotic value of $f(t)$ is 0 and useless for any approximation.
- $\Lambda_V(t)$ is always a conservative estimation of $w(t)$. When dealing with safety, $A(t)$ is close to 1 and therefore the numerical values of $\Lambda_V(t)$ and $w(t)$ are very close.
- $\Lambda_V(t)$ is a good approximation of $\Lambda_{eq}(t)$ when all the failures are quickly detected and repaired. In other cases $\Lambda_V(t)$ is not always a conservative estimation of $\Lambda_{eq}(t)$.

INFO BOX: The very conclusion of this annex is that it is necessary to be very cautious when performing approximate calculations in order to avoid mixed-up between $f(t)$, $\Lambda_{eq}(t)$, $w(t)$ and $\Lambda_V(t)$ and of their averages over a given period.

Annex D (informative)

Broad models for demand mode (reactive) safety systems

D.1 Dangerous failures analysis

[Figure D.1](#) broadly illustrates the intrinsic behaviour of an “on demand” mode safety system considered as an individual item.

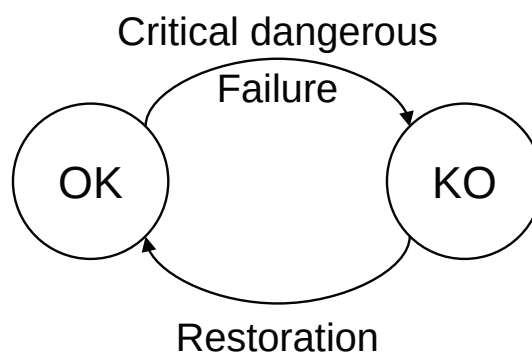


Figure D.1 — Broad available model of an “on demand” safety system

Such a safety system evolves between one class of states (OK) where it is able to perform its safety function and a class of states (KO) where it is no longer able to perform its safety function. Both OK and KO classes may comprise several states.

From a safety point of view the system is “available” when it is in class OK and “unavailable” when it is in class KO. The transition between OK and KO is critical (cf. [Figure B.1](#)) and characterizes the “critical” dangerous failure. This may be a dangerous detected failure as well as a dangerous undetected failure.

When a failure occurs, it is repaired only after having been detected. This may be done immediately or after a periodic test has been performed. This is summarized on [Figure D.1](#) by the term “restoration” (see definition number [3.1.32](#) and [Figure 5](#)): this includes the detection time, the time before that the repair starts, the repair time itself and the time before the system is put on line again. If the protected installation is stopped as soon as the failure has been detected, then the risk of having a hazardous event disappears and it is only necessary to take into account the detection time and the time to intervention (i.e. the time needed to stop the installation after the failure detection) in the calculations. Due to the potential impact on the failure probability, this should be stated in the operating procedures and taken into account in the calculations.

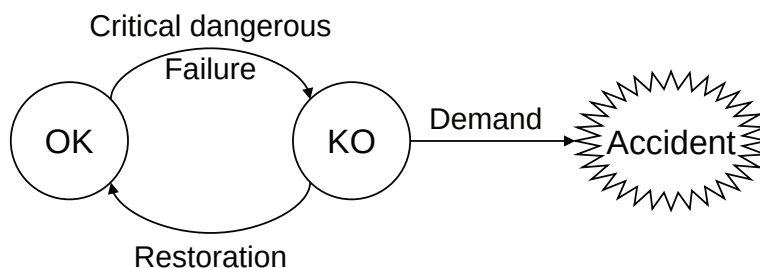


Figure D.2 — Broad model of an “on demand” safety system (ultimate safety system)

A “curative” (or “reactive”) safety system achieves its safety action when it is demanded (e.g. when a given physical parameter exceeds a given threshold) and this yield to two different cases to be considered when a demand occurs during its restoration time (in particular before the failure has been detected):

- The safety system is the ultimate safety system and an accident occurs. This is modelled in [Figure D.2](#) by the transition from “KO” to “accident”.
- The safety system is not the ultimate safety system and a demand is generated toward the succeeding safety system. In this case the safety system restoration can start or continue. This is modelled in [Figure D.3](#).

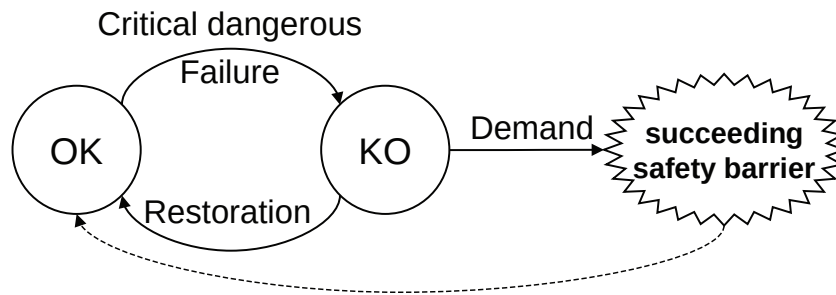


Figure D.3 — Broad model of an “on demand” safety system (non ultimate safety system)

[Figure D.2](#) and [Figure D.3](#) show that the probability of accident (or a demand on the following safety layer):

- increases when:
 - the probability of the critical dangerous failure increases;
 - the probability of the demand increases;
- decreases when the probability of restoration increases.

A rough reasoning allows evaluating the number of accident (or demands on the succeeding safety layer) due to dangerous failures from [Figure D.2](#) and [Figure D.3](#).

If MTBF is the mean time between critical dangerous failures of the safety system then the number of critical dangerous failure of the safety system over the period $[0, T]$ is: $N_f(T) \oplus \frac{T}{MTBF}$.

If MTTRes is the mean time to restore (i.e. the mean time to jump from KO to OK) then, on the long term, the average unavailability of the safety system is given by: $\bar{U} = \frac{MTTRes}{MTBF}$.

Then the mean number of critical failures of the on demand safety system is: $N_f(T) = \bar{U} \cdot \frac{T}{MTTRes}$.

Each time the safety system enters in the KO states there is a probability ζ that a demand occurs when the restoration is not finished yet. Then the number of hazardous event (or accident) will be: $N_a(t) = \zeta \cdot N_f(t)$.

The repair time distribution is not known but two opposite hypotheses to evaluate ζ can be considered:

- 1) The restoration rate μ is constant ($MTTRes = 1/\mu$) The probability to have a demand during the restoration time is equal to: $\zeta^{(1)} = \frac{\lambda_d}{\lambda_d + \mu} = \frac{\lambda_d}{\lambda_d + \frac{1}{MTTRes}}$.
- 2) The restoration time is constant. The probability to have a demand during the restoration time is equal to: $\zeta^{(2)} = 1 - e^{-\lambda_d \cdot MTTR}$.

In the above formulae, λ_d is the demand rate of the safety system. As $\zeta^{(2)}$ is higher than $\zeta^{(1)}$ the hypothesis number 2 is conservative but, when $\lambda_d \cdot MTTRes \ll 1$, $\zeta^{(1)}$ and $\zeta^{(2)}$ have close numerical values.

Therefore the expected number of accidents (or hazardous events, or demands on the succeeding safety layer), $N_a(T)$ is:

$$N_a^{(1)}(T) \approx \zeta^{(1)} N_f(T) = \bar{U} \cdot \frac{\lambda_d \cdot \text{MTTRes}}{\lambda_d \cdot \text{MTTRes} + 1} \frac{T}{\text{MTTRes}} = \bar{U} \cdot \frac{\lambda_d \cdot \text{MTTRes}}{\lambda_d \cdot \text{MTTRes} + 1} \quad \text{or} \quad N_a^{(2)}(T) \approx \zeta^{(2)} N_f(T) = \bar{U} \cdot (1 - e^{-\lambda_d \cdot \text{MTTRes}}) \frac{T}{\text{MTTRes}}$$

When the mean time to demand ($\text{MTTD} = 1/\lambda_d$) is large compared to the MTTRes then $\lambda_d \cdot \text{MTTRes} \ll 1$. In this case $\lambda_d \cdot \text{MTTRes} + 1$ approximately 1 and $(1 - e^{-\lambda_d \cdot \text{MTTRes}}) \approx \lambda_d \cdot \text{MTTRes}$. Therefore when the probability to have a demand during the restoration time is low, both formulae give the same result.

$$N_a(T) \approx \bar{U} \cdot \lambda_d T \quad (\text{D.1})$$

INFO BOX: Provided that the probability of a demand during the restoration of the safety system is sufficiently low (i.e. $\text{MTTD} \gg \text{MTTRes}$), the expected number of accidents due to undetected dangerous failures can be evaluated simply by the average unavailability (i.e. PFD_{avg}) multiplied by the demand rate λ_d and the length T of the period under interest.

$N_a(T)$ is basically an expected number of accidents but when $N_a(T) \ll 1$ it can be interpreted as the probability of accident over $[0, T]$: $N_a(T) \ll 1 \Rightarrow P_a(T) \approx N_a(T)$.

$N_a(T)$ leads directly to the average accident frequency: $\bar{\Phi}_a = \bar{\Phi}_a(T) = \frac{N_a(T)}{T} \approx \bar{U} \cdot \lambda_d = \text{PFD}_{\text{avg}} \cdot \lambda_d$

$\bar{\Phi}_a(T)$ is an average failure frequency, i.e. it is of the same nature as a PFH. This allows making the link between the PFD_{avg} and PFH parameters for systems working in low demand mode.

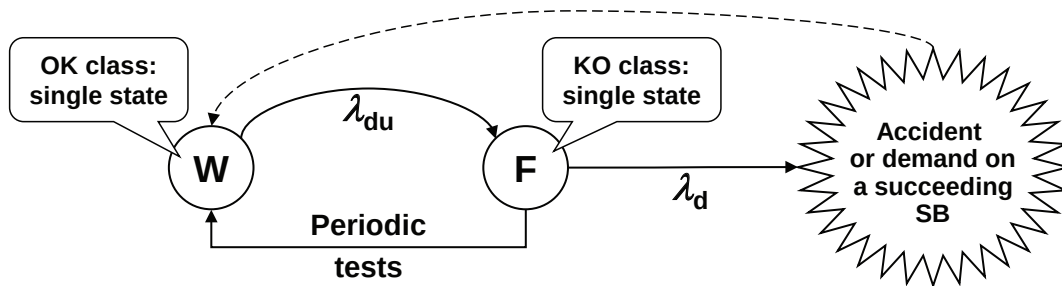


Figure D.4 — Simple periodically tested “on demand” safety system

The approximation provided by Formula D.1 can be analysed more in depth for the dangerous undetected failures through the study of the very simple periodically tested safety system presented in [Figure D.4](#). This is a safety system made of components in series and tested at the same time. It is very simple with only two states (working -W- and failed -F-) but sufficient to illustrate the probabilistic behaviour. It has been treated by using a multiphase Markov model (see [9.2](#) and [Annex L](#)).

The parameters are the following:

- λ_{du} is the “critical undetected dangerous failure rate” and in this case $\text{MTTF} = 1/\lambda_{du}$.
- τ is the test interval and μ is the “repair” rate.
- λ_d is the “demand rate” and in this case $\text{MTTD} = 1/\lambda_d$.

When an undetected failure occurs is discovered by a periodic test then, in average, this failure has occurred at half the test interval (then $\text{MFDT} = \tau / 2$). Then the overall restoration time (MTTRes) is equal to $\tau / 2 + \text{MRT}$ where MRT is the mean repair time after the failure has been detected. Finally $\text{MTTRes} = \tau / 2 + 1/\mu$.

The plot in [Figure D.5](#) shows the comparison between the evaluation of the accident probability by the analytical formula $P_a(t) \equiv P_{a,du}(T) \approx \bar{U}_{du}(T) \cdot \lambda_d T$ and the exact result from Markovian software.

Calculations have been performed with the following hypothesis:

- Duration $T = 50\,000$ h.
- Test interval $\tau = 10\,000$ h.
- Repair rate $\mu = 1/\text{MRT} = 0,1 \text{ h}^{-1}$; therefore $\text{MRT} \ll \tau/2$ and the MTTRes is very close to $5\,000$ h.
- λ_{du} is equal to 10^{-5} h^{-1} and 10^{-6} h^{-1} .

In [Figure D.5](#) the ratio $\text{MTTD} / \text{MTTRes}$ has been used on abscissa axis and the ratio between the approximation and the exact result on the ordinate axis.

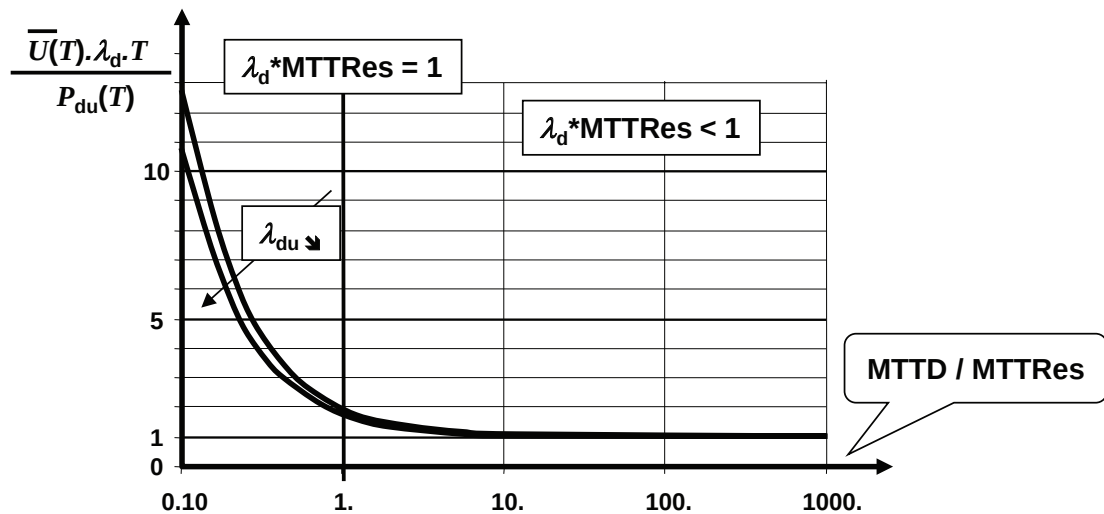


Figure D.5 — Ratio “approximation / exact result” as function of $\text{MTTD} / \text{MTTRes}$

This figure shows that the approximation is very good as long as the MTTD is larger than 10 times the MTTRes , i.e. Five times the test interval. When the MTTD decreases the approximation becomes more conservative and less realistic.

INFO BOX: This is not a sound demonstration but it shows that when the demand and the test frequencies become of the same order then calculating the probability of accident by simplified calculations becomes questionable.

When the frequency of the demand becomes very high (i.e. when the high demand or continuous mode are reached) both $\zeta^{(1)}$ and $\zeta^{(2)}$ tends toward 1 and then the number of accident tends toward:

$$N_{a,du}(T) \approx \bar{U}_{du} \frac{T}{MTTRes} = \frac{T}{MTBF}$$

When $MTTRes \ll MTTF$, which is normally the case for a safety system:

$$N_{a,du}(T) \approx \frac{T}{MTBF} = \frac{T}{MTTF + MTTRes} \approx \frac{T}{MTTF} = \lambda_{du} T \approx F(T)$$

where $F(T)$ is the unreliability of the safety system over $[0, T]$.

INFO BOX: The limit of the demand frequency is the continuous mode of operation. In this case, the accident probability is given by the safety system “unreliability”. Therefore, as shown by the formula just above, the unreliability $F(T)$ may be used to assess the accident probability when the demand frequency increases.

In [Figure D.6](#) the ratio $MTTD / MTTRes$ has been used on abscissa axis and the ratio between the approximation and the exact result on the ordinate axis. As expected, the approximation is very good when $MTTD$ is smaller than 10 % of the $MTTRes$, i.e. is smaller than 5 % of the test interval. When the $MTTD$ increases the approximation becomes more and more conservative and less and less realistic.

This is not a sound demonstration but this shows again that when the demand frequency and the test frequency become of the same order then the results provided by simplified calculations becomes questionable.

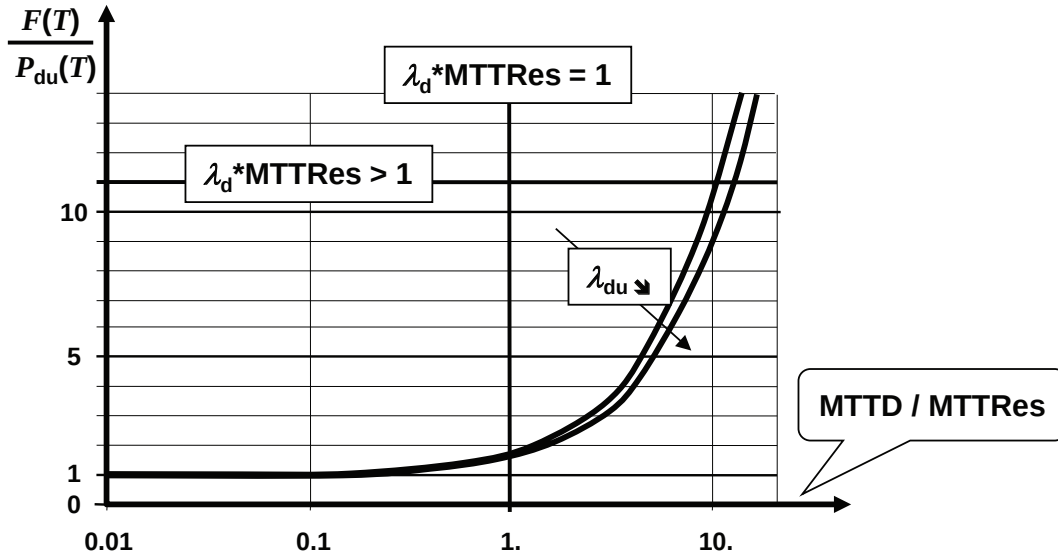


Figure D.6 — Ratio “approximation / exact result” as function of $MTTD / MTTRes$

Conclusion: when the ratio $MTTD / MTTRes$ is large (>10 in our example) $\bar{U}_{du}(T) \cdot \lambda_d T$ is a pretty good approximation of the probability of hazardous event due to dangerous undetected failures $P_{a,du}(T)$.

When it is small ($<0,1$ in our example) $N_{a,du}(T) \oplus \frac{T}{MTBF}$ or $F(T)$ are pretty good approximations of $P_{a,du}(T)$.

In between (0,1 to 10 in our example) both approximations are overpredicting the hazardous event frequency. Therefore the safety systems operating in the “on demand” mode operation should be split into three categories:

- 1) “low demand mode” safety systems;
- 2) “medium demand mode” safety systems;

3) “high demand mode” safety systems.

Usual approximations works for the types 1 and 3 but should be considered cautiously for the type 2.

D.2 Spurious failure analysis

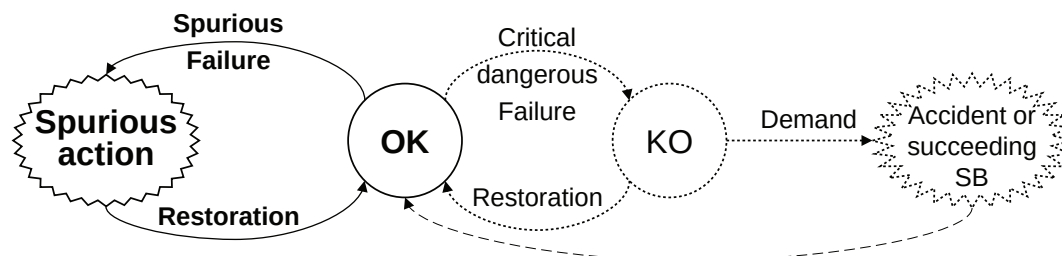


Figure D.7 — Spurious failure modelling

In the engineer collective unconscious it is a good practice to maximize the safe failures to ensure the safety. This is even required by some international standards (e.g. IEC 61508[2]). Unfortunately this is useful only if the safe failures do not provoke too much undesirable spurious failures. The number of accident occurring when restarting the installations would certainly not incite to consider the spurious failure as welcome failures. Of course a spurious action is less dangerous than a critical dangerous failure but too much spurious trips may also lead to dangerous situations. Therefore the approach based on maximizing the safe failures should be considered only when a safe state really exists and that spurious actions are not harmful.

A safe failure is really safe only when it occurs between states belonging to the class OK but may become dangerous if it induces a spurious safety action. Note that an individual safe failure may be safe alone and spurious in conjunction with another failure previously occurred. For example, in the case of a “2 out of 3” voting system the safe failure of a component is only safe when it fails first. If it fails when another component has already had a safe failure, then this generates a spurious action and the safe failure becomes a spurious failure. This swap between safe and spurious for a given failure mode of a given component is a “systemic effect which cannot be handle properly at the component level.

Figure D.7 shows that when the spurious failure probability increases the time spent in the state “spurious action” increases and the time spent in OK and KO decreases. If the time spent in KO decreases, the probability of accident also decreases. Reciprocally, when the critical dangerous failure probability increases the time spent in the state KO increases and the time spent in “spurious action” and OK decreases. If the time spent in OK decreases, the probability of spurious actions also decreases. This highlights an important property of safety systems: increasing safety tends to increase actions trip frequency and, reciprocally, decreasing spurious action frequency tends to decrease safety.

Reducing the probability of accident just by increasing the number of spurious failures (i.e. of safe failures) is not a sound approach because the perfect safety of the installation would be reached when it is stopped all the time due to spurious failures. This can be used only to some extent and a safety system should be designed in order to reach a good balance between safety and spurious actions. Even when spurious actions have no safety issues, the production of the installation is likely to be severely affected if spurious trips are not properly considered.

The “de-energize to trip” design is likely to increase the probability of spurious failure (5.3). A lot of components belonging to the safety system but without dangerous failures may be able to participate in spurious actions. They should be considered (see 5.3) in the spurious failure analysis.

Therefore both accidents and spurious actions or spurious trips should be analysed together in order to avoid that the safety team increases the spurious failure probability while the production team tries to reduce it. The safe failures of components should not be increased blindly without measuring the impact from spurious failures on the overall safety system.

D.3 Average unavailability as a safety indicator

The average unavailability (so-called PFD_{avg}) of a safety system is accepted as a key safety indicator worldwide. It is obviously necessary but is it sufficient?

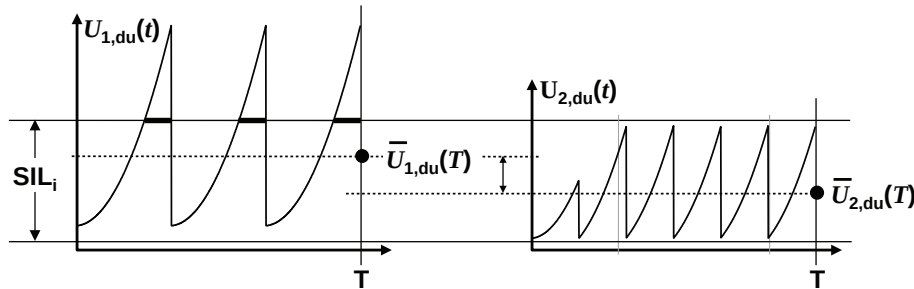


Figure D.8 — Typical saw tooth curves related to low demand mode safety systems

Figure D.8 shows two typical curves related to low demand mode safety systems involving periodically tested components: the unavailability increases between the tests with maximums just prior to the tests; it decreases almost to 0 just after tests. This gives a typical saw tooth shape. The two saw tooth curves in Figure D.8 are related to the same safety system: all the tests are performed at the same times on the left hand sides while they have been staggered on the right hand side. According to the average unavailabilities, these two cases are considered to be equivalent (i.e. they have the same SIL). Obviously they are not identical from a risk point of view: the system on the left hand side makes excursions above the upper SIL boundary about 1/3 of the time while the system on the right hand side remains all the time below this upper SIL boundary. These two cases are perhaps similar for the designer seated in his office but they are certainly different for the operator living on the installation: there are many more risky periods with the system on the left than with the system on the right.

Figure D.9 shows the typical unavailability $U_{du}(t)$ for a safety system involving residual dangerous failures which cannot be detected by periodic tests. In this case, the average unavailability $\bar{U}_{du}(T)$ increases all the time when T increases.

In the example, $U_{du}(t)$ is permanently below the SIL_i boundary over $[0, t_1]$, swinging around the SIL_i boundary over $[t_1, t_2]$ and $[t_2, t_3]$, and above the upper SIL_{i-1} boundary (i.e. in the SIL_{i-1} zone) over $[t_2, t_3, T]$. Beyond t_2 and t_3 , $U_{du}(t)$ is even permanently higher than $\bar{U}_{du}(T)$. From the operator point of view $\bar{U}_{du}(T)$ obviously underestimates the risk for the period $[t_2, t_3, T]$ and this is even worse for $[t_2, t_3, T]$.

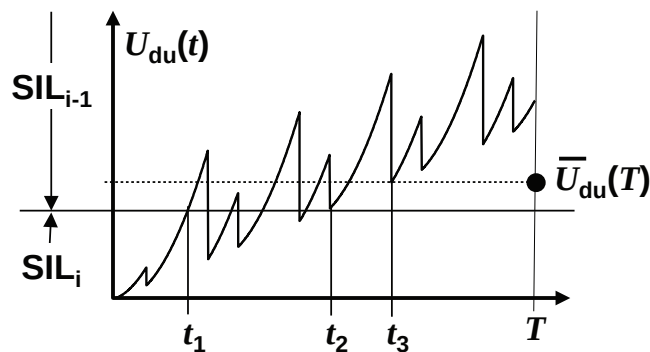


Figure D.9 — Safety system with dangerous undetected failures not detected by test

Therefore if the average unavailability (i.e. PFD_{avg}) is necessary to qualify globally the safety level provided by a safety system, it is not sufficient to represent the risks actually borne by the operators: even if the requirements are satisfied from average point of view, the instantaneous unavailability should also be considered and too large excursions avoided or limited in time (e.g. less than 10 % of the time).

Annex E (informative)

Continuous mode (preventive) safety systems

E.1 Dangerous failures analysis

A preventive safety system achieves its safety action continuously. This is illustrated in [Figure E.1](#) where the model is simpler than [Figure D.2](#). This models a continuous mode safety system which is the ultimate protection before an accident occurs. Note that the detection of the critical dangerous failure before that the accident occurs (see [Figure B.3](#) and [B.4](#)) is not considered to be possible in this simple model.

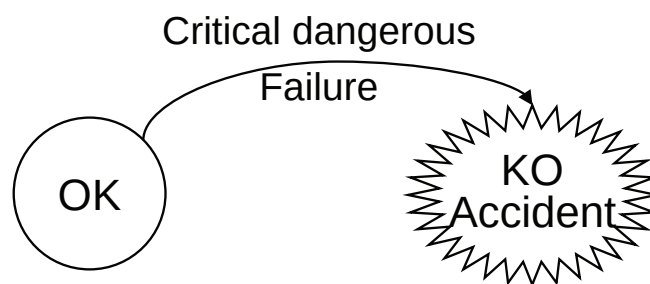


Figure E.1 — Broad model of an ultimate safety system working in continuous mode

Again two classes of state OK and KO are found but the accident occurs as soon as the safety system enters in KO and the hazardous event frequency is equal to the failure frequency of the safety system: $\bar{\Phi}_a(T) = \bar{w}(T)$

The probability of accident over a given interval $[0, t]$ is simply the unreliability of the safety system and [Figure E.1](#) is, in fact, a “reliability” model (cf. [Clause 9](#)) where the state KO is an absorbing state because there is no restoration after an accident. This model is valid if the critical dangerous failure is detected or not detected by periodic tests. If Λ_{df} is the equivalent dangerous failure rate of this system then $P_a(T) = F(T) = 1 - e^{-\Lambda_{df}T} \approx \Lambda_{df}T$ if $\Lambda_{df}T \ll 1$. Therefore Λ_{df} also becomes the accident rate. The equivalent dangerous failure rate it is not necessarily constant for actual safety systems and in general cases time-dependent $\Lambda_{df}(t)$ have to be considered.

INFO BOX: As $F(t)$ is the approximation found for “on demand” mode safety systems when the MTTD becomes short compared to the MTTRes, the “continuous” mode appears to be the limit of the “on demand mode” when MTTD tends to zero.

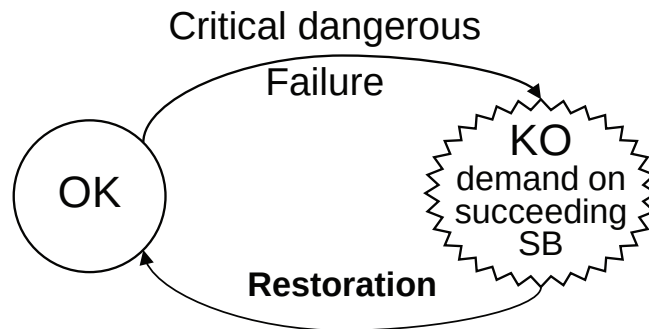


Figure E.2 — Broad model of a safety system working in continuous mode (non ultimate safety system)

When the continuous mode safety system is not the ultimate safety system, it only generates a demand on another safety system (operating “on demand”) when it fails. Therefore, if the succeeding safety system reacts properly the accident is avoided. Consequently, as shown in [Figure E.2](#), it is essential to consider the restoration of the dangerous critical failures to calculate the probability of a demand on the following safety system. Most of the control systems are of this type. They operate in continuous mode and when they fail, they generate a demand on the emergency shut down system or the process shut down systems.

The probability to have a demand over $[0, T]$ on the succeeding safety system is just the same as above, i.e. the safety system unreliability $F(T)$. It should be calculated by using the model in [Figure E.1](#) (i.e. the possible restoration is not taken into account).

The model in [Figure E.2](#) is an availability model (the KO state can be restored). The notions of MTTF, MTTRes and MTBF of the safety system introduced in the previous chapter can also be used here.

Over a given period $[0, T]$ the number of arrivals in KO (i.e. the number of demands N_d on the succeeding safety system) is simply: $N_d(T) \oplus \frac{T}{\text{MTBF}}$.

If the $\bar{w}_{df}(T)$ is the average unconditional failure frequency of the safety system: $N_d(T) = \bar{w}_{df}(T) \cdot T$. Finally, by comparison with the previous formula $\bar{w}_{df}(t) \oplus \frac{1}{\text{MTBF}}$ is obtained.

INFO BOX: The average frequency of dangerous failure $\bar{w}_{df}(T)$, also named PFH, is the key parameter to characterize the functional safety of systems working in continuous mode of operation (see IEC 61508[2], IEC 61511[3], etc.).

The above formula cannot be applied to the model in [Figure E.1](#) because there is no restoration arc and therefore the MTTRes is infinite as well as the MTBF. Nevertheless, this system will be restored in the reality. As this restoration time is not known it is possible to use $\bar{w}_{df}(t) \oplus \frac{1}{\text{MTTF}}$ which is a conservative estimation.

Finally: $N_d(T) = \bar{w}_{df}(T) \cdot T = \frac{T}{\text{MTTF}}$.

In the very simple case of [Figure E.1](#), $\bar{w}_{df} \approx \frac{1}{\text{MTTF}} = \lambda_{df} = \frac{\lambda_{df} T}{T} \approx \frac{F(T)}{T}$ is obtained. This formula remains valid when the critical failure rate is not constant but the conditions of validations, $F(T) \ll 1$ should not be forgotten as this formula goes to 0 when T becomes large.

E.2 Spurious failure analysis

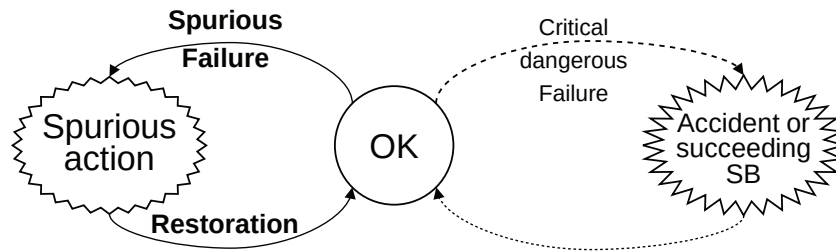


Figure E.3 — Spurious failure modelling

The spurious failures can be analysed exactly in the same way as for “continuous mode” safety systems. Here also, increasing frequency of safe failures (e.g. at component level) increases the spurious failures frequency (at safety system level). In turn this decreases the probability of accident. At the end, the probability of accident vanishes because the process is stopped all the time: this should be thoroughly considered.

Annex F (informative)

Multi-layers safety systems/multiple safety systems

F.1 Dangerous failure analysis

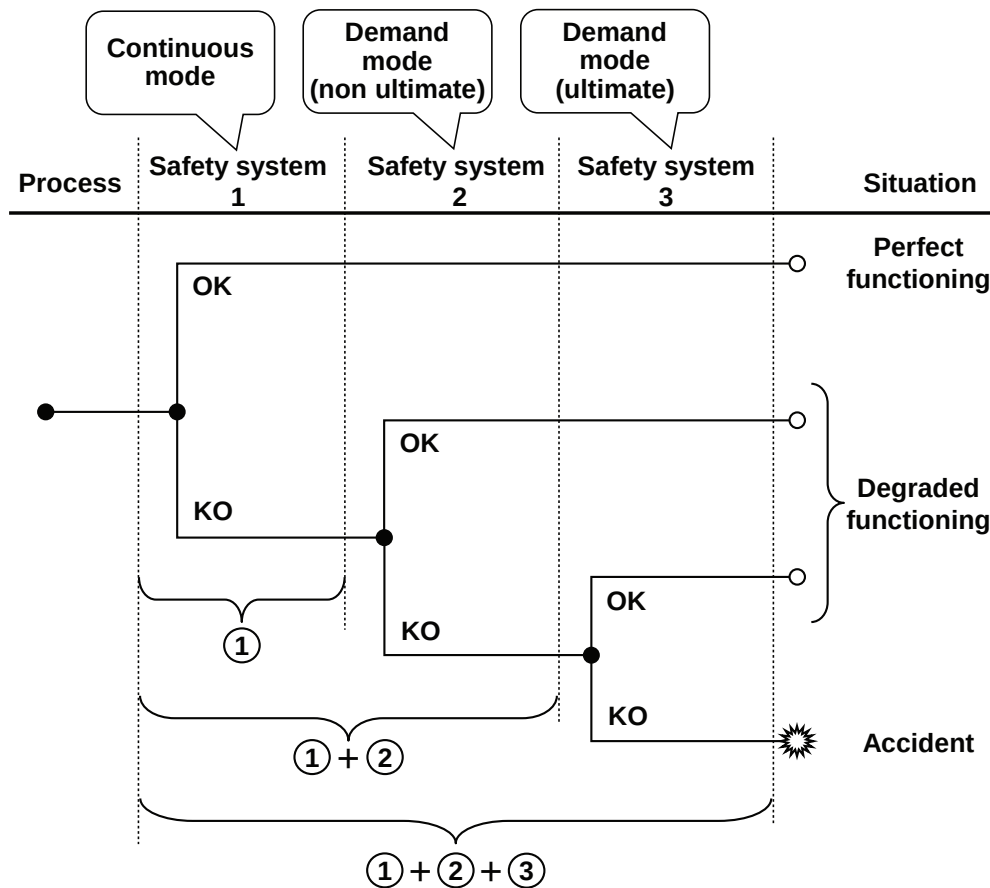


Figure F.1 — Dangerous failures of three safety systems running in sequence

Figure F.1 illustrates a multiple safety system made of 3 individual safety systems working in sequence. An event tree^[8] is used to model the functioning and dysfunctioning of these individual safety systems. Each of them may be also considered as being a safety layer (as per the LOPA^[9] terminology). This is a very general configuration encountered in oil and gas industries where two safety layers are commonly implemented (e.g. process control + emergency shut down) and sometime three (e.g. process control + emergency shut down + HIPS):

- The first safety layer is a “non ultimate continuous mode” safety system.
- The 2nd safety layer operates on demand when the first layer fails. It is a “non ultimate on demand mode” safety system.
- The 3rd safety layer operates on demand as ultimate safety barrier when both the first and second layers have failed.

In addition, [Figure E.1](#) shows that:

- the overall safety system gathering layers 1 and 2 operates in continuous mode;
- the overall safety system gathering layers 1, 2 and 3 operate in continuous mode.

Then the “on demand mode” is only relevant to analyse individual intermediate safety layers. From an accident probability point of view the continuous mode of operation should always be considered.

A very commonly used approach for calculating the probability P_a of the sequence leading to the accident is to estimate a probability of failure for each safety system and to multiply these probabilities together. For example $P_a = p_1.p_2.p_3$, where p_1 , p_2 and p_3 represent the respective probabilities of failures of the safety systems 1, 2 and 3. The promoters of this kind of calculation claim that this is valid when the safety layers are independent. Unfortunately this clause of independence is hardly ever fulfilled in the oil and gas industries due to the following reasons:

- 1) The common cause failures existing between the safety layers.
- 2) The final elements (e.g. valves) which are often shared by several layers (e.g. common process shut down and emergency shut down valves).
- 3) The impact of the systemic dependencies introduced by the periodic tests.

None of the above difficulties can be handled properly just by considering the safety layers individually and systemic approach considering all the safety systems as a whole need to be implemented. Whereas problems 1 and 2 are well known, the third one is generally ignored because it is outside the common knowledge of engineers. In fact the formula $P_a = p_1.p_2.p_3$ is only valid to combine constant probabilities, p , or probabilities at time t , $p(t)$, provided that they are independent. Average probabilities - which from mathematical point of view are “mathematical expectations” and not “probabilities” - can definitively not be handled in this way and integral calculations are needed. When the safety system comprises periodically tested components, it is not possible to split this integral into parts related to the individual systems and it should be calculated as a whole. The result may be significantly higher than what is obtained with the simple multiplication which is not conservative. This may be observed when LOPA is implemented.

Moreover, the non-conservativeness increases when the redundancy increases and, therefore, the approximation become worse when the probability decreases. This is misleading because it is exactly the opposite behaviour of the usual approximations used in reliability engineering where the approximations are conservative and become better when the probability decreases. An example may be found in IEC 61511-3[3] Ed.2.

INFO BOX: Hence, when dealing with periodic tests and when simply multiplying the branch probabilities together, the results may become non-conservative.

F.2 Spurious failure analysis

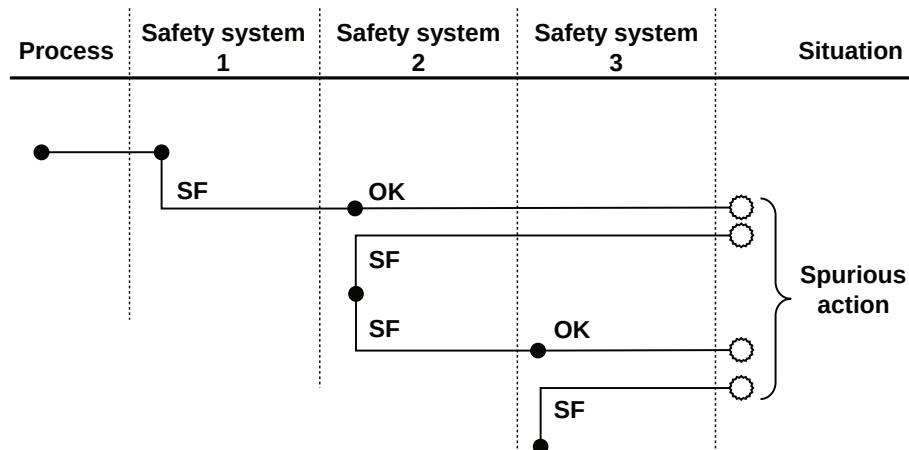


Figure F.2 — Spurious failures of three safety system running in sequence

Each safety system may produce two kinds of spurious failures:

- 1) The action corresponding to the OK branch is spuriously triggered.
- 2) The demand on the following safety layer is spuriously triggered.

In the first case the spurious action is straightforward whereas in the second case the spurious action is triggered through the scheduled reaction of the following safety systems (which normally react properly to a demand).

In our example, the first safety system can only generate a spurious demand on the second safety system. The second safety system may generate directly a spurious action as well as a demand on the third safety system. The third safety system may only produce a spurious action. This is presented in [Figure F.2](#) where SF stands for “spurious failure”.

The propagation of a spurious failure from a given safety system to the protected installation depends on the design of the succeeding safety systems. Then, similarly to the critical dangerous failures the spurious failures should be thoroughly analysed and a systemic analysis is needed.

NOTE In de-energize to trip design, all the failures able to lead to a spurious action should be considered (see [5.3](#)). This may imply other components than those considered within the dangerous failure analysis.

Annex G (informative)

Common cause failures

NOTE This annex develops [5.4.2](#) which is to be read first.

G.1 Small similar components number

G.1.1 Beta (β) factor approach

As said in [5.4.2](#), it was believed at the beginning that increasing redundancy allows reducing the probability of failure as low as it was wanted. At the beginning of the 1980s and after long discussions, the shortcomings of such an approach have finally been recognized in the nuclear field and the so-called β factor[20] approach has been introduced as a safeguard.

The principle is very simple. It consists in splitting the failure rate of the individual component in two parts:

- independent failure: $(1-\beta).\lambda$
- common cause failure: $\beta.\lambda$

Therefore the β factor is only the percentage of the failure rate which is reputed to be related to CCF. It is generally considered that the multiple failures due to the β factor occur simultaneously or in a short period of time. This is a conservative hypothesis which is implemented, for example, in IEC 61508[2].

Of course the main difficulty is to evaluate the value of this β factor because it is a safeguard (i.e. an artefact) rather than a well identified failure mode, that CCF are not very frequent and that they are not necessarily identified as such when they occur. This is certainly a domain needing improvement in reliability data collection and also to be better in capturing failure causes (cf. ISO 14224[15], Table B.3).

Nevertheless β factor ranging from 1 % to 10 % are often used in probabilistic calculations. The literature about CCF is rather sparse but some indications are provided for safety instrumented systems in standards dealing with functional safety (IEC 61508[2]). Ten % is certainly conservative in most of the cases, 1 % is perhaps more realistic, 5 % seems to be a good safeguard. Anyway, a thorough analysis of the potential CCF should be performed and a report documented to explain the choice of the value of β factor.

The β factor approach is interesting because it is easy to implement in the calculation models.

G.1.2 The PDS extension of the Beta (β) factor model

G.1.2.1 Basic model

The traditional way of accounting for common cause failures (CCF) in the process industry has been the β -factor model. In this model it is assumed that a certain fraction of the failures (equal to β) are common cause, i.e. failures that will cause all the redundant components to fail simultaneously or within a short time interval. One problem with this approach is that for *any* M-out-of-N (MooN) voting the rate of common cause failures is the same, regardless of the configuration. If λ_{DU} is the component failure rate, the MooN voted system has a common cause failure contribution equal to $\beta.\lambda_{DU}$. Hence, this approach does not explicitly distinguish between different voting logics, and the same result is obtained e.g. for 1oo2, 1oo3 and 2oo3 voted systems.

In the PDS method^[13], an extended version of the β -factor model that distinguishes between different types of voting is proposed. Here, the rate of common cause failures explicitly depends on the configuration, and the beta-factor of a MooN voting logic may be expressed as:

$$\beta(\text{MooN}) = \beta \cdot C_{\text{MooN}}, (M < N),$$

C_{MooN} is then a modification factor for various voting configurations, and β is the factor which applies for a 1oo2 voting. This means that if each of the N redundant modules has a failure rate λ_{DU} , then the MooN configuration will have a system failure rate due to CCF that equals:

$$C_{\text{MooN}} \cdot \beta \cdot \lambda_{\text{DU}}$$

NOTE A MooN voting ($M < N$) means that at least M of the N redundant modules give a shutdown signal for a shutdown to be activated.

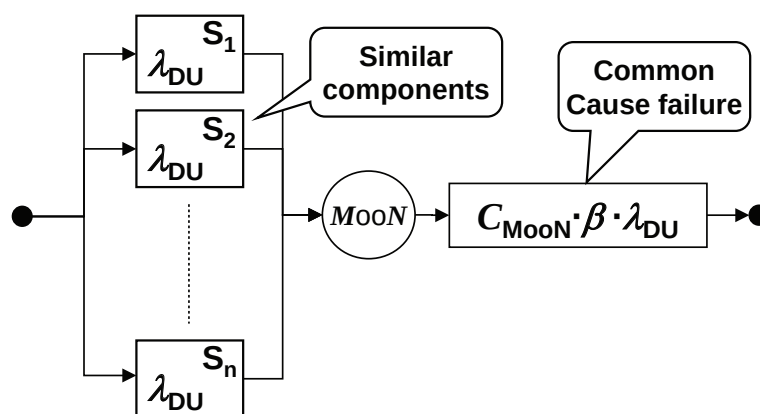


Figure G.1 — PDS CCF model

By using this model, the parameter β is maintained as an essential parameter whose interpretation is now entirely related to a duplicated system. Further, note that the effect of voting is introduced as a *separate* factor, C_{MooN} , independent of β . This makes the model easy to use in practice.

G.1.2.2 Determining values for the C_{MooN} factors

There is of course no definite choice of the values for the C_{MooN} values. It may also be argued that the C_{MooN} factors should differ for different types of equipment. Experience data (reliability data) on common cause failures are in general scarce and data which explicitly differentiate between different voting logics even more scarce. The factors therefore should be chosen based on the limited data available combined with expert judgements. When determining the C_{MooN} values, it is also important to ensure that the effect of added redundancy is not overestimated. Empiric results from a study made in Swedish Power plants, [SKI Technical report NR 91:6. CCF analysis of high redundancy systems, safety/relief valve data analysis and reference BWR application. Stockholm 1992] and subsequent analysis of data from this study indicates the following:

- Given a failure of two redundant components, the likelihood of having a simultaneous failure of a third added component may be higher than 0,3, sometimes as high as 0,5.
- When introducing more and more components it appears that the effect of added redundancy decreases as the number of components increases.
- For systems where the number N of parallel components are high (say more than 7-8 components) the likelihood of having N simultaneous failures seem to be higher than having $N-1$ (and sometimes also $N-2$) components failing.

Assuming that the above observations can be generalized, the following assumptions are made when establishing values for the C_{MooN} factors:

- Given a common cause failure of two redundant components, then the probability of a third similar component also to fail due to the same cause will be 50 %.
- To cater for the (observed) reduced effect of added redundancy, the following has been assumed:
 - when going from 3 to 4 components then the probability of the fourth component also failing will be 60 %
 - When going from 4 to 5 components then the probability of the fifth component also failing will be 70 %, etc.
 - Finally when having 7 or more components the assumed effect of adding one more component is negligible, i.e. it is said that if 7 components have already failed simultaneously, then the likelihood of the 8th component also failing is 100 %.

Then [Table G.1](#) of C_{MooN} for some typical voting configurations is obtained.

Table G.1 — Suggested C_{MooN} factors for different voting logics

$M \setminus N$	$N = 2$	$N = 3$	$N = 4$	$N = 5$	$N = 6$
$M = 1$	$C_{1oo2} = 1,0$	$C_{1oo3} = 0,5$	$C_{1oo4} = 0,3$	$C_{1oo5} = 0,21$	$C_{1oo6} = 0,17$
$M = 2$	-	$C_{2oo3} = 2,0$	$C_{2oo4} = 1,1$	$C_{2oo5} = 0,7$	$C_{2oo6} = 0,4$
$M = 3$	-	-	$C_{3oo4} = 2,9$	$C_{3oo5} = 1,8$	$C_{3oo6} = 1,1$
$M = 4$	-	-	-	$C_{4oo5} = 3,7$	$C_{4oo6} = 2,4$
$M = 5$	-	-	-	-	$C_{5oo6} = 4,3$

It should be pointed out that the above figures are suggested values based on the listed assumptions. C_{1oo5} and C_{1oo6} have been given with two decimals in order to be able to distinguish the two configurations. Based on the general formulae for calculating C_{MooN} -factors as described in the next section, the user can modify these factors based on experience and knowledge. For example, there might be arguments for using more or less conservative values for special equipment types.

G.2 Large similar components number (shock model)

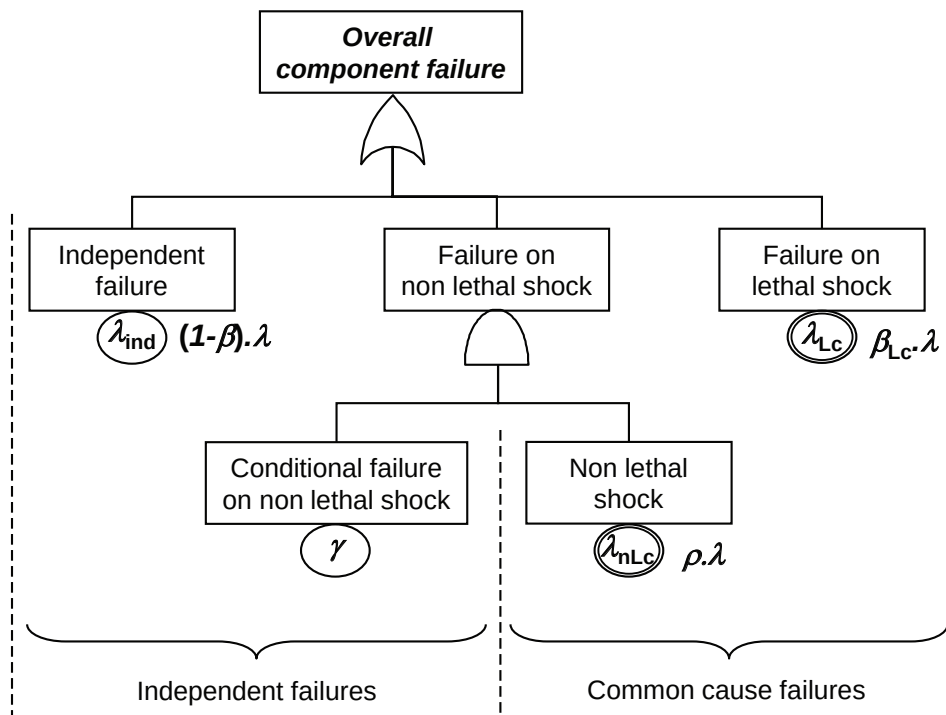


Figure G.2 — Fault tree illustrating the shock model

If numerous double failures and a few triple failures have been observed from field feedback, it is not confirmed that quadruple failure have been observed. No multiple failures beyond four have ever been reported. The probability of multiple dependent failures due to implicit CCF decreases when the number of affected components increases and if the β factor model is realistic for double failures it is slightly pessimistic for triple failures and it becomes too pessimistic for quadruple failures and beyond.

For example it is necessary that all the production wells be closed in case of a “blocked outlet” on a production platform. An implicit CCF may prevent the closure of 2, 3 or even perhaps 4 wells, but surely not the closure of 200 wells (otherwise the CCF would be explicit and should be analysed as such). Therefore the β factor model is ineffective for a big number of affected components.

Several models have been developed to deal with this difficulty but most of them require too many parameters (e.g. multiple Greek letters or α -models) to be practically used. The exception is the so-called “binomial failure rate” (“shock” model) approach which has been introduced by Vesely in 1977 and improved by Atwood in 1986^[21], (see Figure G.2). The principle is that when a CCF occurs, it provokes on the affected components a *shock* which may be lethal or not. The lethal shock is similar to the β factor model when the non-lethal shock has only a given probability to fail the affected components. The probability of having k failures due to the non-lethal shock is then binomially distributed.

As shown in Figure G.2, like the β factor model, the shock model is very easy to implement in probability calculation models. It is a good compromise as only 3 basic parameters are needed.

- λ_{Lc} lethal shock failure rate;
- λ_{nLc} non-lethal shock failure rate;
- γ conditional probability of failure of a component given a non-lethal shock.

Figure G.2 shows how the shock model is implemented for a single component:

- independent failure (λ_{ind});

- failure on non lethal shock (λ_{nLc}, γ);
- failures on lethal shock (λ_{Lc}).

If a collection of N similar component is considered and $\lambda_{ccf}t \ll 1$ is obtained, the failure rate due common cause failure can be written as: $\lambda_{ccf} \approx \lambda_{Lc} + C_N^2 \gamma^2 \lambda_{nLc} + C_N^3 \gamma^3 \lambda_{nLc} + C_N^4 \gamma^4 \lambda_{nLc} + C_N^5 \gamma^5 \lambda_{nLc} + \dots$

In the above formula, $C_N^k \gamma^k \lambda_{nLc}$ is the failure rate of k component failures due to the occurrence of a non lethal shock.

Similarly with the β factor model, it can be written:

- $\lambda_{Lc} = \beta_{Lc} \lambda$ and
- $(C_N^2 \gamma^2 + C_N^3 \gamma^3 + C_N^4 \gamma^4 + C_N^5 \gamma^5 + \dots) \lambda_{nLc} = \beta_{nLc} \lambda$

Finally the following results are obtained:

- overall failure rate: $\lambda = \lambda_{ind} + \lambda_{Lc} + \gamma \lambda_{nLc}$
- link with conventional β factor: $\beta = \beta_{Lc} + \beta_{nLc}$
- failure rate on lethal shock: $\lambda_{Lc} = \beta_{Lc} \lambda$
- failure rate on non lethal shock: $\lambda_{nLc} = \frac{\beta_{nLc}}{(C_N^2 \gamma^2 + C_N^3 \gamma^3 + C_N^4 \gamma^4 + C_N^5 \gamma^5 + \dots)} \lambda$ which can be written $\lambda_{nLc} = \rho \lambda$
- independent failure rate: $\lambda_{ind} = \lambda[1 - (\beta_{Lc} + \gamma \cdot \rho)]$

The problem is now to evaluate the values λ_{ind} , λ_{Lc} , λ_{nLc} and γ . Reference[21] provides some information about the evaluation of the shock model parameters but the engineering judgement can be used when no data are available.

A pragmatic way to do that is to align the shock model to produce the same result as the conventional β factor model by considering that common cause cannot produce multiple failures beyond triple failures.

This leads to: $\rho = \frac{\beta_{nLc}}{(C_N^2 \gamma^2 + C_N^3 \gamma^3)}$

Multiple failures beyond triple failures being negligible, then the quadruple failures due to a non lethal shock are certainly lower more than 10 times of the double failures due to the same non lethal shock.

This implies $C_N^2 \gamma^2 > 10 \cdot C_N^4 \gamma^4$ which leads to $\gamma < \sqrt{\frac{C_N^2}{10 \cdot C_N^4}}$ and therefore taking $\gamma \approx \sqrt{\frac{C_N^2}{10 \cdot C_N^4}}$ is certainly a conservative hypothesis.

Finally, for a set of N similar components, the model can be implemented in the following way:

- estimation of the lethal part of the beta factor: β_{Lc}
- estimation of the non-lethal part of the beta factor: $\beta_{nLc} = \beta - \beta_{Lc}$
- estimation of the lethal failure rate $\lambda_{Lc} = \beta_{Lc} \lambda$
- estimation of $\gamma \approx \sqrt{\frac{C_N^2}{10 \cdot C_N^4}}$
- estimation of $\rho = \frac{\beta_{nLc}}{(C_N^2 \gamma^2 + C_N^3 \gamma^3)}$
- estimation of $\lambda_{nLc} = \rho \lambda$
- estimation of $\lambda_{ind} = \lambda[1 - (\beta_{Lc} + \gamma \cdot \rho)]$

The above approach estimates the relevant parameters from the number N of components and the classical β factor and its repartition between lethal and non lethal parts. When dealing with a big number

of components (e.g. a great number of wells to close in case of blocked outlet) this provides conservative results where double and triple failures are the main contributors and where the unrealistic multiple failures are not completely neglected.

Further detail about shock model can be found in references [20] and [21]

G.3 Systematic failures considerations

Systematic failures are another kind of dependent failures. As defined in 0, this is the combination between a pre-existing unknown fault and a given external condition which lead to a systematic failure. This arises in a deterministic way as soon as the given condition occurs as illustrated in Figure G.3.

Depending of the safety system under study the related cause may never, always or randomly happens (e.g. probability p in Figure G.3):

- 1) When the cause is outside the operating range of the safety system, it never happens and, therefore, the related systematic failure never happens.
- 2) When the cause belongs to the routine operating range of the safety system, it occurs very often and the related systematic failure can be quickly detected and, usually, removed by a relevant design modification.
- 3) When the cause belongs to the non-routine operating range, it occurs randomly according to the probability to run the safety system in this non-routine range. In this case, even that it is deterministic in nature, the related systematic failure occurs randomly with the same probabilistic distribution of its related cause which, itself, follows the probabilistic distribution of the non-routine operations initiating this cause.

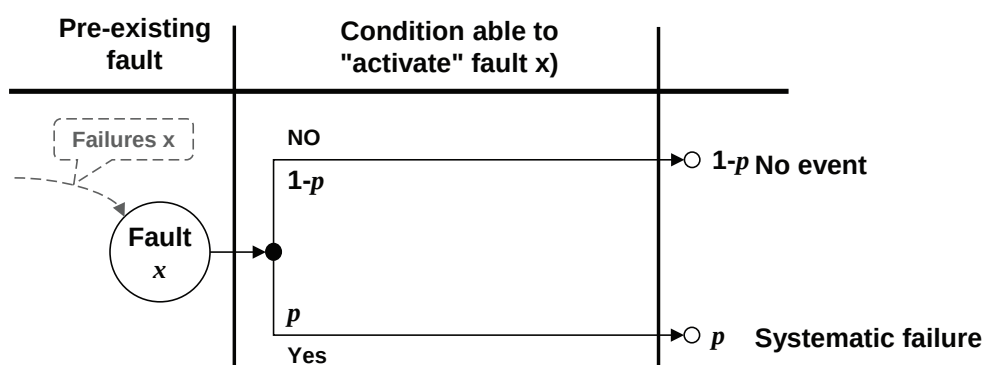


Figure G.3 — Illustration of a systematic failure

The case 1 does not matter as long as the operating range is not modified. The analyst should have in mind that a safety system which is free of systematic failures when used under given condition may experience such systematic failure when used in different conditions.

The case 2 is the easiest to handle because it can be detected during the commissioning stage and corrected (e.g. modification of the design of hardware items or debugging of software) before the safety system is really used.

The case number 3 is the most difficult to handle because it may occurs in situations with very low probabilities (i.e. not detectable during the commissioning phase). If it is rather unusual for hardware items this becomes an increasing problem with the software where remaining bugs are likely to lead to such situations. It is possible to consider the compound event "random condition + pre-existing systematic fault" as a single random event called "systematic failure" (see Figure G.3). This is used, for example, for some software reliability approaches. Provided that probabilistic data are available, such compound random events may be processed as other random event with the approaches described in this Technical Report.

It should be noted that systematic failures are not necessarily identified as such when a reliability data collection is performed. Therefore the failure rates elaborated from the field feedback contain a part of systematic failures. It should also be noted that the systematic failures are a kind of common cause failures (in fact the most perfect ones!). Therefore they can be covered by the common cause failure approaches.

www.TeraStandard.com

Annex H (informative)

The human factor

H.1 Introduction

It is outside the scope of this Technical Report to deal in detail with the human factor. An International Standard like IEC 62508 provides valuable information on this subject. Therefore only important topics are pointed out. ISO 20815:2008^[16], I.10 refers to additional guidance for reducing human error. References^{[36],[37],[38],[39],[40],[41],[42],[43]} and^[44] provide also further guidance.

Human error modelling depends on how the error is classified. Error is generally considered to be of three types:

- Memory and attention lapses and action slips – here the plan of action is satisfactory but action deviates from intention in some unintentional way. Slips and lapses generally occur in routine tasks with operators who know the process well and are experienced in their work. For example highly automated action routines can be applied in the wrong situation due to attention lapses. Inattention and overattention are potential problems leading to omissions, repetition and routines applied in the wrong situation or at the wrong time.
- Mistakes in planning actions – inadvertent human errors that occur when the elements of a task are being considered by the operator. They are decisions to carry out (or omit) actions that are subsequently found to be wrong, although at the time the operator believed them to be correct.
- Rule violations – Deliberate deviations from the rules, procedures, instructions and regulations, such as omitting a step in a procedure in order to get the job done more quickly

In addition, three levels of human performance can be considered:

- Skill based – predetermined actions carried out in a certain order where performance is controlled by stored patterns of behaviour triggered by stimuli in the environment. The operator is essentially running on automatic and carrying out routine actions. Routine monitoring and detection tasks can be included here. Such a level of operating is primarily concerned with slips and lapses.
- Rule based – actions (or inaction) following well defined mentally stored or available rules (if X then do Y) and involving conscious effort in applying them. This is planned-for-problem solving where solutions are governed by stored rules. Such a level of operating is primarily concerned with mistakes, for example misclassification of a situation leading to the application of the wrong rule.
- Knowledge based – here action follows problem solving in novel situations where no predetermined rules exist and depends instead upon stored knowledge and analytical processes to determine what rules apply. Again, such a level of operating is concerned with mistakes. Limited mental resources and lack of knowledge can cause errors in determining the appropriate rule.

These levels are hierarchical. Performance at the knowledge-based level will call upon rule-based procedures which will call upon skill-based operations to execute the procedure. Human performance is most reliable at the skill-based level and least at the knowledge-based level.

H.2 Quantification

In human reliability analysis the ultimate purpose is to reduce the presence of factors which increase the chance of human error at these different levels of performance and optimize the factors which increase the chance of success. These performance shaping factors exist in equipment, man-machine

interface and task design, training and experience, procedures, information, communication, stress and fatigue, workload, and underlying organizational and management factors. Other factors to be taken into account are dependencies, for example when the same operator carries out the same task on otherwise independent systems.

The following rules of thumb are here suggested with respect to the need for probabilistic quantification of PFD_{avg} (or PFH):

- The contribution from human errors is not included in the quantification if the safety instrumented function (SIF) is activated automatically without any required operator intervention (e.g. a process shutdown function (PSD) activated automatically upon high temperature).
- The contribution from human errors should be considered included in the quantification also for a fully automatic SIF, if the function is frequently inhibited or bypassed during special operational modes. An example can be a low pressure alarm (PALL) function on the inlet piping that frequently needs to be inhibited during start-up.
- The contribution from human errors should be included in the quantification if a person/operator is an active element in the execution of the SIF. For example, an operator may be expected to initiate a valve closure (shutdown) or valve opening (blow down) upon an alarm from the safety instrumented system (SIS).

Hence, human errors should be considered in the reliability calculations if human tasks are likely to have a direct influence on the performance of the SIF. In addition the probability of human error depends very much on the situation: low for routine operations, high for operations under stress and very high for emergency situations.

At present there is a lack of measured human error data. For this reason there have been a wealth of methods derived which compensate for this. In many cases the analyst will look for actual data specific to the task or similar task, but mostly this has to be supplemented with a structured judgement method. Generic data, distilled from actual data, plays an important role in probability estimates. There are several methods to execute a Human Reliability Assessment (HRA). These include THERP (Technique for Human Error Rate Prediction), SLIM (Success likelihood Index Method), and HEART (Human error assessment reduction technique) which are probably the best known methods. Most of the methods have components which are based on data obtained in real-life situations. So, when HEPs from different methods are compared they should show some agreement for specific types of task. Important parameters are:

- Whether the task is skill-based (operator is essentially running on automatic and carrying out routine actions), rule-based (actions or inaction following well defined mentally stored or available rules and involving conscious effort in applying them) or knowledge based (action follows problem solving in novel situations where no predetermined rules exist and depends instead upon stored knowledge and analytical processes). Quantification is generally in the order of 10^{-4} , 10^{-2} and 10^{-1} respectively for the operator action for these three kinds of task.
- Whether the task is regularly performed and commonplace, more complex with less time available where the operator may make a mistake, or where the operator may omit to carry out actions for tasks with dependence on situational cues and/or dependence upon memory the quantification is generally in the order of 10^{-4} , 10^{-3} and 10^{-2} respectively for the operator action. Complex unfamiliar tasks under (time) stress typically acquire a value of between 10^{-1} to 1 per demand on the operator.

Absolute Probability Judgement (APJ), where the expert or group of experts estimate an HEP, frequently uses nominal probability values in look-up tables and is probably the simplest of all approaches. Another method, that of Paired Comparisons (PC), makes use of the fact that experts find it easier to make comparative than absolute judgements.

For example, a look-up table may give generic tasks and associated human error probabilities. An example for 2 very different task conditions is given in [Table H.1](#) (HEART method):

Table H.1 — Human error probability

Generic task	Nominal human reliability (5th –95th percentile bounds)
Restore or shift a system to original or new state following procedures, with some checking.	0,003 (0,000 8 – 0,007)
Totally unfamiliar task, performed at speed with no real idea of likely consequences.	0,55 (0,35 – 0,97)

The following two sections give examples of approaches to human failure.

H.3 Signal detection theory

Signal detection theory (SDT) is a way to model human performance in uncertain or ambiguous situations and is a good example of quantification of human performance. The model can be used for tasks requiring decision making (yes/no decisions) which may occur in:

- Emergency situations where e.g. the operator has to diagnose the causes of alarms.
- Inspection and checking where it needs to be decided whether a component of the system is in an acceptable state or not.
- Continuous monitoring of a system state where the operator is watching out for signals of deviations e.g. process control room tasks, drilling.

The model applies where signals may not conclusively indicate a given condition. In many cases evidence (signals) of a particular condition may occur in the presence of “noise”. It may not be easy to determine the difference between signal and noise if the distribution of “evidence” of what constitutes a signal overlaps with the distribution of “evidence” of what constitutes noise. E.g. there may be conflicting evidence or weak evidence as to whether a component is mounted the correct way up. Here there are four classes of response to the evidence:

- Hit – the operator correctly detects the signal (e.g. on checking the equipment the operator identifies that the component was incorrectly mounted when it was incorrectly mounted).

Hit probability = Number of hits/number of signal events.

- Miss – the operator missed the signal (e.g. on checking the equipment the operator did not identify that the component was incorrectly mounted).

Miss probability = Number of missed signals/number of signal events.

- False alarm – the operator identifies a noise event as a signal (e.g. on checking the equipment although the component is correctly mounted the operator reports that it is incorrect).

False alarm probability = Number of reported signals / Number of noise events.

- Correct reject – the operator identifies a noise event as noise (e.g. on checking the equipment the operator correctly finds no evidence to suggest the component is incorrectly mounted).

Correct accept probability = Number of events reported as noise / number of noise events.

Hits and correct rejects are good. A graph of the probability of hits against the probability of false alarms indicates the operator’s performance where optimum performance is $p_{Hit} = 1$, $p_{False Alarm} = 0$.

Because of signal-noise overlap not all responses can be good. The quality of the signal and the experience of the operator will affect the ability of the operator to make hits and correct rejects. However, in addition, a criterion will be selected by the operator to separate the distributions which, depending on its position along the line of evidence, will affect the frequency of false alarms and hits which will either increase or

decrease together. For example an operator may be biased towards false alarms if not missing signals is critical or towards misses if false alarms are considered unacceptable. Over time an operator may adjust their response frequency to match encountered conditions such that in high reliability systems where signals are unexpected there can be a bias towards rejecting the signal when it does occur.

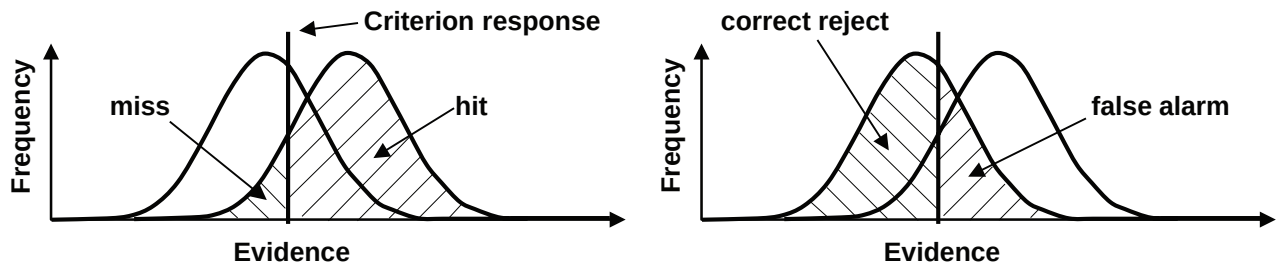


Figure H.1 — Overlapping noise (left) and signal (right) distributions of probability of evidence when there is a signal event (top) or noise event (below)

These concepts of human performance which split human actions up into different types are important for linking the human performance into the system. Ultimately the human being should do something involving detection, diagnosis and response. The indicators (displays, signals) are in the technical system. The perception of the displays and signals is in the human component. The internal response of the human is diagnosis and selection of the procedural rule to apply. The subsequent action responses of the person are through physical interaction with the technical system.

NOTE Even though an operator makes a correct diagnosis and selects the right rule to apply, the action can still be incorrectly carried out.

H.4 A Paired Comparisons example

In the operation of an emergency system in response to alarms, the uncertainty in the signal can affect the probability of the activation of the safety system. Indicators of shallow gas blow out are a good example. Not only are signals unreliable but drilling personnel can have very little time to make a decision.

The probabilities of different signals triggering a response can be calculated in a number of different ways, e.g.:

- Absolute probability judgement (APJ) which might consider different nominal values depending upon the nature of the task and perhaps come up with a nominal value of 0,5 for failing to respond to a shallow gas indicator (whether or not it is a real signal).
- The method of Paired Comparisons, which compares the different signals with one another to determine which is considered more likely to be acted upon or ignored.

The last method is developed as a worked example below. It provides a scale of values which can be calibrated if two known data points are available, preferably around the upper and lower bounds of the scale. As well as the calibration points, it is essential that the assessor gather data on the relevant signals and ask experts to compare one with another in pairs.

Considering, for example, signals which may indicate shallow gas (drilling parameters, mud parameters, measurement while drilling parameters, electric logging etc.) the experts (drill crew) should be able to answer to questions like: “What indicator of the two is the least likely to generate a control response: Increasing mud return flow or Increase in resistivity”

Every item (e.g. shallow gas indicator) is paired with every other one and assigned a row and column in a table. For each pair the preferred item is specified and then the number of times as item is preferred is totalled for each column (see [Table H.2](#)).

Table H.2 — Record form for a single evaluator

	A	B	C	D	E	F	G	H	I	J
A	.	A	A	A	A	A	A	A	A	A
B	A	.	B	B	E	B	G	B	B	J
C	A	B	.	D	C	C	C	C	C	C
D	A	B	D	.	E	D	D	D	D	J
E	A	E	C	E	.	E	E	E	E	E
F	A	B	C	D	E	.	F	F	I	F
G	A	G	C	D	E	F	.	G	G	G
H	A	B	C	D	E	F	G	.	I	J
I	A	B	C	D	E	I	G	I	.	J
J	A	J	C	J	E	F	G	J	J	.
Choice score (S)	9	5	6	5	7	3	4	0	2	4

This is repeated for the sample of experts and the choice scores totalled and averaged and then divided by $n-1$ (where n is the number of items) to obtain the probability that the item was preferred (p). To avoid zeros 0.5 is added to each S value. The p is then converted to a z score, where the z value for a score is the number of standard deviations from the mean in the normal distribution: $z = \frac{S - m}{\sigma}$ where S is a score, m is the mean and σ is the standard deviation.

Table H.3 — Standard conversion table

z	0	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0	0.5000	0.5040	0.5080	0.5120	0.5160	0.5199	0.5239	0.5279	0.5319	0.5359
0.1	0.5398	0.5438	0.5478	0.5517	0.5557	0.5596	0.5636	0.5675	0.5714	0.5753
0.2	0.5793	0.5832	0.5871	0.5910	0.5948	0.5987	0.6026	0.6064	0.6103	0.6141
0.3	0.6179	0.6217	0.6255	0.6293	0.6331	0.6368	0.6406	0.6443	0.6480	0.6517
0.4	0.6554	0.6591	0.6628	0.6664	0.6700	0.6736	0.6772	0.6808	0.6844	0.6879
0.5	0.6915	0.6950	0.6985	0.7019	0.7054	0.7088	0.7123	0.7157	0.7190	0.7224
0.6	0.7257	0.7291	0.7324	0.7357	0.7389	0.7422	0.7454	0.7486	0.7517	0.7549
0.7	0.7580	0.7611	0.7642	0.7673	0.7704	0.7734	0.7764	0.7794	0.7823	0.7852
0.8	0.7881	0.7910	0.7939	0.7967	0.7995	0.8023	0.8051	0.8078	0.8106	0.8133
0.9	0.8159	0.8186	0.8212	0.8238	0.8364	0.8289	0.8315	0.8340	0.8365	0.8389
1	0.8413	0.8438	0.8461	0.8485	0.8508	0.8531	0.8554	0.8577	0.8599	0.8621
1.1	0.8643	0.8665	0.8686	0.8708	0.8729	0.8749	0.8770	0.8790	0.8810	0.8830
1.2	0.8849	0.8869	0.8888	0.8907	0.8925	0.8944	0.8962	0.8980	0.8997	0.9015
1.3	0.9032	0.9049	0.9066	0.9082	0.9099	0.9115	0.9131	0.9147	0.9162	0.9177
1.4	0.9192	0.9207	0.9222	0.9236	0.9251	0.9265	0.9279	0.9292	0.9306	0.9319
1.5	0.9332	0.9345	0.9357	0.9370	0.9382	0.9394	0.9406	0.9418	0.9429	0.9441

z	0	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
1.6	0.9452	0.9463	0.9474	0.9484	0.9495	0.9505	0.9515	0.9525	0.9535	0.9545
1.7	0.9554	0.9564	0.9573	0.9582	0.9591	0.9599	0.9608	0.9616	0.9625	0.9633
1.8	0.9641	0.9649	0.9656	0.9664	0.9671	0.9678	0.9686	0.9693	0.9699	0.9706
1.9	0.9713	0.9719	0.9726	0.9732	0.9738	0.9744	0.9750	0.9756	0.9761	0.9767
2.0	0.9772	0.9772	0.9783	0.9798	0.9793	0.9898	0.9803	0.9808	0.9812	0.9817
2.1	0.9821	0.9826	0.9830	0.9834	0.9838	0.9842	0.9846	0.9850	0.9854	0.9857
2.2	0.9861	0.9864	0.9868	0.9871	0.9875	0.9878	0.9881	0.9884	0.9887	0.9890
2.3	0.9893	0.9896	0.9898	0.9901	0.9904	0.9906	0.9909	0.9911	0.9913	0.9916
2.4	0.9918	0.9920	0.9922	0.9925	0.9927	0.9929	0.9931	0.9932	0.9934	0.9936
2.5	0.9938	0.9940	0.9941	0.9943	0.9945	0.9946	0.9948	0.9949	0.9951	0.9952
2.6	0.9953	0.9955	0.9956	0.9957	0.9959	0.9960	0.9961	0.9962	0.9963	0.9964
2.7	0.9965	0.0066	0.9967	0.9968	0.9969	0.9970	0.9972	0.9972	0.9973	0.9974
2.8	0.9974	0.9975	0.9976	0.9977	0.9977	0.9978	0.9979	0.9979	0.9980	0.9981
2.9	0.9981	0.9982	0.9982	0.9983	0.9984	0.9984	0.9985	0.9985	0.9986	0.9986
3.0	0.9987	0.9987	0.9987	0.9988	0.9988	0.9989	0.9989	0.9989	0.9990	0.9990

By using the areas underneath normal distribution curve, the probabilities of different outcomes are used to determine the number of standard deviations the score is from the mean. The p is converted to a z score using a standard conversion Table H.3. When p is less than 0.5, use $1-p$ to find negative z values.

The z scores are then recalibrated as a 0-100 scale where 0 and 100 are used for the extreme values.

Table H.4 — Paired comparisons arranged in ascending order of S scores

	H	I	F	J	G	D	B	C	E	A
Choice (S) score	0	2	3	4	4	5	5	6	7	9
Adjusted score S' (add 0.5)	0,5	2,5	3,5	4,5	4,5	5,5	5,5	6,5	7,5	9,5
Probability $P' = (S'/n)$	0,05	0,25	0,35	0,45	0,45	0,55	0,55	0,65	0,75	0,95
z'	-1,64	-0,67	-0,39	-0,13	-0,13	+0,13	+0,13	+0,39	+0,67	+1,64
Scaling	0	30	38	46	46	54	54	62	70	100

In the example the HEP would be the probability that the indicator is not responded to. Following the scale conversion, two seed HEPs are used to calculate the other values. Ideally these are high and low values (e.g. A and H for our example). What this means is that two of the item values have known HEPs and these can be used to calibrate the scale where:

$\text{Log}_{10} \text{HEP} = a.S + b$ where S is the scale value and a and b are constants. This rests on the premise that there is an exponential relationship between the scale values and the absolute probability of failure.

Suppose indicator A has a known HEP of 0,5 per demand on operator and H a known HEP of 10^{-3} per demand on operator as the chance of failing to respond to the indicator.

Then $\text{Log}_{10} 0,5 = (a \times 100) + b$ and $\text{Log}_{10} 10^{-3} = (a \times 0) + b$ which is solved as: $b = -3$, $a = 0,033\ 01$

So the value of D , for example, would be calculated as $\text{Log}_{10} \text{HEP}(D) = 0,033\ 01 \times 54 - 3 = -1,217\ 46$

So $\text{HEP}(D) = 0,0606$

In such a way a scale of probabilities for a set of items (e.g. A to H in our example) can be derived.

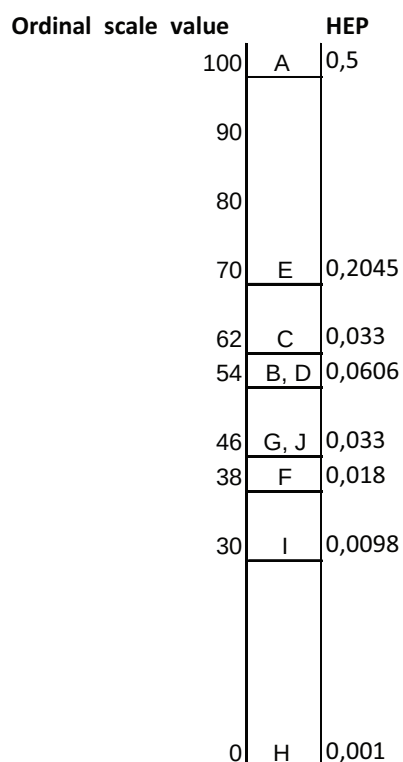


Figure H.2 — Scale of probability from the score of the pairing method

Annex I (informative)

Analytical formulae

I.1 Analytical formulae development (low demand mode)

This annex is based on the approach developed in 6.2.3. It aims to detail the establishment of the formulae related to single, double or triple failures of safety system operating in low demand mode. Those formulae are valid only if the assumptions described in Clause 7 are met: low probability of failure, quick repairs, constant failure and repair rates, test duration small compared to the test interval, etc.

NOTE If the protected installation is stopped during tests and repairs then MRT should be replaced by the time elapsing between the detection of the fault and the stopping of the installation, and the test duration should be taken equal to 0.

I.2 Single failure analysis

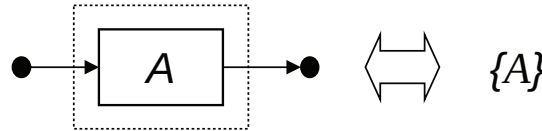


Figure I.1 — Single component and single failure

The reliability block diagram in Figure I.1 represents a single element and the corresponding single failure (A = “component A failed”) which may be a part and a minimal cut set of a complete safety system.

I.2.1 Dangerous undetected failures

If the protected system is stopped during testing and repair of the protection systems, the risk of accident is absent during the corresponding downtime and very simple formulae can be established for the dangerous undetected failures:

- Unreliability due to the single failure: $F_{du}(t) \approx \lambda_{du}t$
- Unavailability due to the single failure: $U_{du}(t) \approx \lambda_{du}(t \text{ modulo } \tau)$

NOTE ($t \text{ modulo } \tau$) represents the remaining part of the division of t by τ (i.e. the time remaining until the next test).

In order to simplify the notations, “unreliability” means “unreliability due to single failure” and “unavailability” means “unavailability due to single failure” in the remaining part of this annex.

Thanks to the hypothesis “as good as new after a test”, the unavailability is similar from a test interval to another and it is only necessary to develop the calculations over a test interval of length τ .

This leads to:

- Average unavailability over a periodic test interval: $\bar{U}_{du} \approx \frac{1}{\tau} \int_0^\tau \lambda_{du} \delta \cdot d\delta = \frac{1}{\tau} \frac{\lambda_{du} \tau^2}{2} = \frac{\lambda_{du} \tau}{2}$
- Number of accident over the time interval $[0, \tau]$ is equal to $N_{a,du}(\tau) \approx \bar{U}_{du} \cdot \lambda_d \tau = \frac{\lambda_{du} \lambda_d \tau^2}{2}$ (according to Equation D.1)

- Hazardous event probability over $[0, \tau]$: $P_{a,du}(\tau) \approx N_{a,du}(\tau) = \frac{\lambda_{du}\lambda_d\tau^2}{2}$ (according the underlying hypothesis $\lambda_d\tau \ll 1$).
- Hazardous event frequency: $\bar{\Phi}_{a,du}(\tau) = \frac{N_{a,du}(\tau)}{\tau} \approx \bar{U}_{du} \cdot \lambda_d = \frac{\lambda_{du}\lambda_d\tau}{2}$.

If the protected system is not stopped during the repair of the safety system then a hazardous event may arise if a demand occurs during the repair. The probability of a repair is equal to $P_r(\tau) \approx \lambda_{du} \cdot \tau$ and the duration of a repair lasts $MRT_{du} = 1/\mu_{du}$. Then:

- Expected time of unavailability due to repair: $\Theta_r(\tau) \approx \lambda_{du} \cdot \tau / \mu_{du}$
- Unavailability due to repair: $\bar{U}_r \approx \frac{\Theta_r(\tau)}{\tau} = \frac{\lambda_{du}}{\mu_{du}}$
- Number of accident over the time interval $[0, \tau]$: $N_{a,r}(\tau) \approx \bar{U}_r \cdot \lambda_d \tau \approx \frac{\lambda_{du}\lambda_d\tau}{\mu_{du}}$ (according to Equation D.1)
- Hazardous event probability over $[0, \tau]$: $P_{a,r}(\tau) \approx N_{a,r}(\tau) \approx \frac{\lambda_{du}\lambda_d\tau}{\mu_{du}}$ (according the underlying hypothesis $\lambda_d\tau \ll 1$)
- Hazardous event frequency: $\bar{\Phi}_{a,r} \equiv \bar{\Phi}_{a,r}(\tau) = \frac{N_{a,r}(\tau)}{\tau} \approx \bar{U}_r \cdot \lambda_d = \frac{\lambda_{du}\lambda_d}{\mu_{du}}$

NOTE The formula of the instantaneous unavailability $U_{du}(t)$ has not been established because it is rather complicated when repair is considered.

When a test is performed, there is a probability to fail the safety system due to the test itself just because of the sudden change of the state of some component. This is a failure occurring on demand as per 3.2.13. Such a failure has to be repaired and the same repair duration MRT_{du} as for failures occurred within the test interval is considered. Therefore, if γ is the probability of failure due to the test:

- Unavailability due to the tests: $\bar{U}_\gamma \approx \gamma \frac{\Theta_r(\tau)}{\tau} = \gamma \frac{1}{\mu_{du}\tau}$
- Number of accident over $[0, \tau]$: $N_{a,\gamma} \equiv N_{a,\gamma}(\tau) \approx \bar{U}_\gamma \cdot \lambda_d \tau \approx \frac{\gamma \lambda_d}{\mu_{du}}$ (according to Equation D.1)
- Hazardous event probability over $[0, \tau]$: $P_{a,\gamma} \equiv P_{a,\gamma}(\tau) \approx N_{a,\gamma}(\tau) \approx \frac{\gamma \lambda_d}{\mu_{du}}$ (according the underlying hypothesis $\lambda_d\tau \ll 1$)
- Hazardous event frequency: $\bar{\Phi}_{a,\gamma}(\tau) = \frac{N_{a,\gamma}(\tau)}{\tau} \approx \bar{U}_\gamma \cdot \lambda_d = \frac{\gamma \lambda_d}{\mu_{du}\tau}$

Another kind of failure occurring on demand can occur when a demand of the safety action occurs. The difference with γ is that it cannot be repaired before that the demand occurs. Some of them may be detected by, e.g. dismantling the device on a workbench to inspect its internals. Others cannot be detected, e.g. mechanical blockage due to a sudden change of the state. As above such failure may be modelled by a simple not time-dependent probability ψ .

- Unavailability due the demands themselves: $\bar{U}_\psi = \psi$
- Number of accident over $[0, \tau]$: $N_{a,\psi}(\tau) = \bar{U}_\psi \lambda_d \tau \approx \psi \lambda_d \tau$
- Hazardous event probability over $[0, \tau]$: $P_\psi(\tau) \approx N_{a,\psi}(\tau) \approx \psi \lambda_d \tau$
- Hazardous event frequency: $\bar{\Phi}_{a,\psi}(\tau) = \frac{N_{a,\psi}(\tau)}{\tau} \approx \bar{U}_\psi \cdot \lambda_d = \psi \lambda_d$

The protected system is not necessarily stopped during the tests. Then, if the safety system is unavailable for its safety function (e.g. if it is disconnected) during the test, this is necessary to properly consider that in the calculations. If π is test duration, then:

- Unavailability due to test duration: $\bar{U}_{\pi} = \frac{\pi}{\tau}$
- Number of accident over $[0, \tau]$: $N_{a,\pi} \equiv N_{a,\pi}(\tau) \approx \bar{U}_{\pi} \cdot \lambda_d \tau \approx \lambda_d \pi$ (according to Equation D.1)
- Hazardous event probability over $[0, \tau]$: $P_{\pi} \approx N_{a,\pi} \approx \lambda_d \cdot \pi$
- Hazardous event frequency: $\bar{\Phi}_{a,\pi}(\tau) = \frac{N_{a,\pi}(\tau)}{\tau} \approx \bar{U}_{\pi} \cdot \lambda_d = \lambda_d \frac{\pi}{\tau}$

If the safety system is disconnected to perform the test its reconnection can be forgotten when the test is finished (human error). If nothing is scheduled to detect this kind of error, then it will remain unavailable until the next test (i.e. during τ). If ω is the probability of this human error then:

- Unavailability due to human error: $\bar{U}_{\omega} = \omega \frac{\tau}{\tau} = \omega$
- Number of accident over $[0, \tau]$: $N_{a,\omega}(\tau) = \bar{U}_{\omega} \cdot \lambda_d \tau \approx \omega \lambda_d \tau$
- Hazardous event probability over $[0, \tau]$: $P_{\omega}(\tau) \approx N_{a,\omega}(\tau) \approx \omega \lambda_d \tau$
- Hazardous event frequency: $\bar{\Phi}_{a,\omega}(\tau) = \frac{N_{a,\omega}(\tau)}{\tau} \approx \bar{U}_{\omega} \cdot \lambda_d = \omega \lambda_d$

Other parameters can be introduced accordingly when it is needed (e.g. no detection of the failure, repair when no failure is present, etc.).

Gathering all the results above gives:

- Average unavailability: $\bar{U}_{du}(\tau) \approx \frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi$
- Number of accidents over $[0, \tau]$: $N_{a,du}(\tau) \approx (\frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi) \lambda_d \cdot \tau$
- Hazardous event probability over $[0, \tau]$:

$$P_{a,du}(\tau) \approx N_{a,du}(\tau) \approx (\frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi) \lambda_d \cdot \tau$$

- Average accident frequency over $[0, \tau]$:

$$\bar{\Phi}_{a,du}(\tau) = \frac{N_{a,du}(\tau)}{\tau} \approx (\frac{\lambda_{du} \cdot \tau}{2} + \frac{\lambda_{du}}{\mu_{du}} + \frac{\gamma}{\mu_{du} \cdot \tau} + \frac{\pi}{\tau} + \omega + \psi) \lambda_d$$

I.2.2 Dangerous immediately revealed/detected failures

Revealed failures, as well as failures detected by some diagnostic tests, are known immediately or almost immediately after they occur. Then, restoration normally begins quickly and, then, $1/\mu_{dd} = \text{MTTRes}_{dd} \approx \text{MRT}_{dd}$.

This case is very similar to [Figure D.4](#) which has been adapted in [Figure I.2](#):

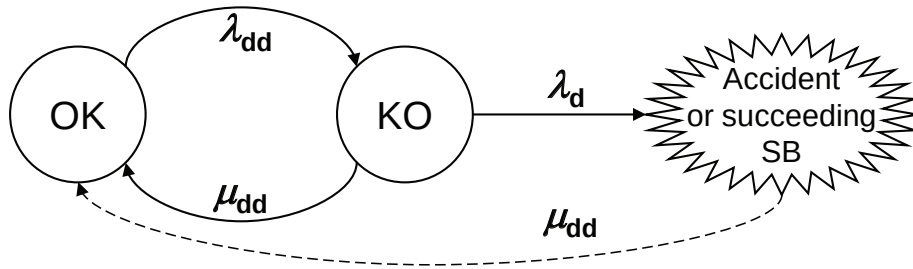


Figure I.2 — Dangerous detected failures modelling

Here the restoration time is short, i.e. the ratio $\rho = \text{MTTD} / \text{MTTRes}_{dd}$ is high and the approximations work well:

- Safety system unreliability: $F_{dd}(t) \approx \lambda_{dd}t$
- Safety system unavailability: $U_{dd}(t) = \frac{\lambda_{dd}}{\lambda_{dd} + \mu_{dd}} [1 - e^{-(\lambda_{dd} + \mu_{dd})t}]$
- Safety system average unavailability: $\bar{U}_{dd} \approx \frac{\lambda_{dd}}{\lambda_{dd} + \mu_{dd}}$ (asymptotic value which is reached after 2 or 3 MTTRes_{dd})
- Number of accidents over $[0, \tau]$: $N_{a,dd}(T) \approx \bar{U}_{dd} \cdot \lambda_d \cdot T = \frac{\lambda_{dd}}{\lambda_{dd} + \mu_{dd}} \lambda_d \cdot T$
- Hazardous event probability: $P_{a,dd}(T) \approx N_{a,dd}(T) \approx \bar{U}_{dd} \cdot \lambda_d \cdot T = \frac{\lambda_{dd}}{\lambda_{dd} + \mu_{dd}} \lambda_d \cdot T$
- Hazardous event frequency: $\Phi_{a,dd}(T) = \frac{N_{a,dd}(T)}{T} \approx \frac{\lambda_{dd} \cdot \lambda_d}{\lambda_{dd} + \mu_{dd}}$

From a theoretical point of view, this model belongs to the “completely and quickly repairable” systems, meaning that it spends most of its time in the state OK. When it jumps to KO, it stays there for a very short time ($\approx 1/\mu_{dd}$) and KO is almost an “instantaneous” state. When it jumps out of KO, it comes back to OK with the probability $\frac{\mu_{dd}}{\lambda_d + \mu_{dd}}$ and to the accident (or demand) state with the probability $\frac{\lambda_d}{\lambda_d + \mu_{dd}}$.

Therefore the instantaneous state KO can be removed to obtain the equivalent graph in [Figure I.3](#) where $\Lambda_d \approx \lambda_{dd} \frac{\lambda_d}{\lambda_d + \mu_{dd}}$ is the equivalent accident (or demand) rate of the safety system.

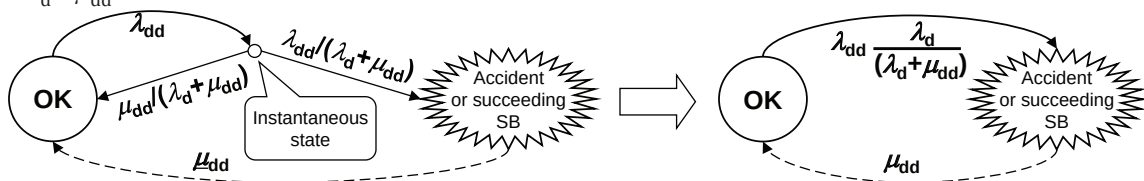


Figure I.3 — Equivalent model for dangerous detected failures

From the equivalent failure rate, the probability of accident over $[0, T]$ is simply the unreliability provided by this model: $P_{a,dd}(T) = F(T) \approx \Lambda_d T \approx \lambda_{dd} \frac{\lambda_d}{\lambda_d + \mu_{dd}} T$.

This is the same result as above found in another way.

I.2.3 Spurious failures

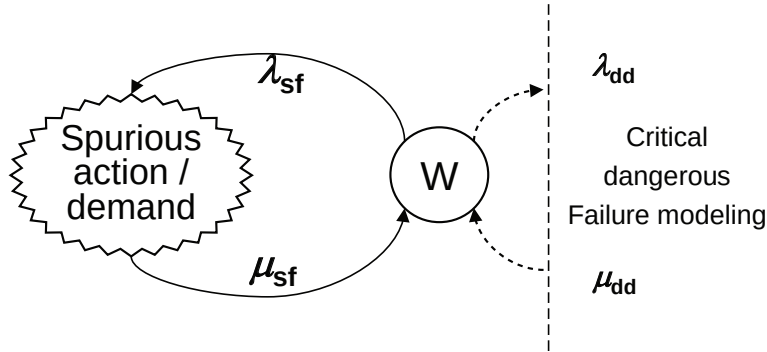


Figure I.4 — Model for spurious failures

Figure I.4 models the behaviour of a single component from spurious failure point of view. The spurious failure rate is equal to the safe failure rate λ_{sf} of the component and the repair rate is equal to the repair rate μ_{sf} of the component. Such failure is normally immediately revealed: $MFDt = 0$ and $MTTRes_{sf} = MRT_{sf}$.

NOTE In “de-energize to trip” design, all the failures able to lead to a spurious action should be considered (see 5.3). This may imply other failures than those considered within the dangerous failure analysis (e.g. the links between the components). The safe failure rate λ_{sf} should be evaluated accordingly

The part related to the dangerous failure is not considered in the calculations because this simplifies the calculations and provide conservative results.

From Figure I.4 the following is obtained:

- Probability of spurious failure (i.e. “unreliability” with regard to spurious failures):

$$F_{st}(t) = F_{sf}(t) \approx \lambda_{sf} \cdot t$$

- Mean time between spurious failure: $MTBF_{st} = MTBF_{sf} = MTTF_{sf} + MTTRes_{sf} = \frac{1}{\lambda_{sf}} + \frac{1}{\mu_{sf}}$
- Number of spurious failure over $[0, T]$: $N_{st}(T) = N_{sf}(T) = \frac{T}{MTBF_{sf}}$
- Average spurious failure frequency is given by:

$$\bar{\Phi}_{st}(T) = \bar{\Phi}_{sf}(T) = \frac{N_{sf}(T)}{T} \approx \frac{1}{MTBF_{sf}} = \frac{1}{MTTF_{sf} + MTTRes_{sf}} = \frac{\lambda_{sf} \mu_{sf}}{\lambda_{sf} + \mu_{sf}}$$

When $F_{sf}(T) \ll 1$, it can be used to provide a conservative value of the average spurious failure frequency over $[0, T]$: $\bar{\Phi}_{st}(T) = \bar{\Phi}_{sf}(T) \geq \frac{1}{MTTF_{sf}} = \frac{\lambda_{sf} T}{T} \approx \frac{F_{sf}(T)}{T}$. The approximation $F_{sf}(T)/T$ cannot be used when T becomes large at it goes to 0.

I.3 Double failure analysis

Figure I.5 presents RBD comprising two redundant blocks and the corresponding minimal cut set which is a double failure.

Even for a rather simple system like this one, the number of combinations drastically increases (especially if the extra parameters like π , ω ... introduced in 1.2.1, are considered). In addition, the tests intervals are not necessarily identical and even if they are identical they can be performed not at the same time. Therefore it is almost impossible to treat all the possible cases. The aim of this part of the annex is to give some advice and to encourage the analysts to understand and develop their own formulae rather than to provide a catalogue of ready-to-use formulae.

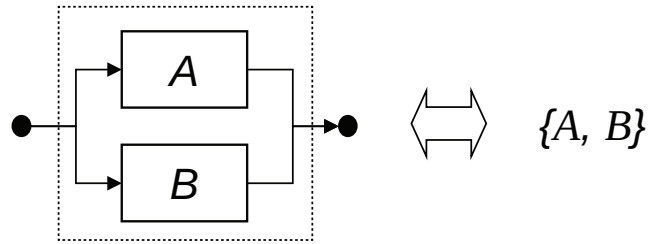


Figure I.5 — Redundant components and double failure

I.3.1 System unavailability

The system presented on [Figure I.5](#) is failed when both A and B are failed, i.e. when they are unavailable at the same time: Then its unavailability is given by:

$$U_{AB}(t) = U_A(t) \cdot U_B(t) = [U_{A,du}(t) + U_{A,dd}(t)] [U_{B,du}(t) + U_{B,dd}(t)]$$

This formula comprises three types of terms:

- 1) one term related to dangerous detected failures: $U_{A,dd}(t) \cdot U_{B,dd}(t)$
- 2) one term related to dangerous undetected failures: $U_{A,du}(t) \cdot U_{B,du}(t)$
- 3) two terms related to both type of failures: $U_{A,dd}(t) \cdot U_{B,du}(t), U_{A,du}(t) \cdot U_{B,dd}(t)$

As seen in I.2.2 $U_{A,dd}(t)$ and $U_{B,dd}(t)$ converge quickly toward asymptotic values which are also averages values: $\bar{U}_{A,dd} \approx \frac{\lambda_{A,dd}}{\lambda_{A,dd} + \mu_{A,dd}} = u_A$, $\bar{U}_{B,dd} \approx \frac{\lambda_{B,dd}}{\lambda_{B,dd} + \mu_{B,dd}} = u_B$. Therefore the above terms can be expressed as:

- 1) term related to dangerous detected failures: $U_1(t) = u_A \cdot u_B$
- 2) term related to dangerous undetected failures: $U_2(t) = U_{A,du}(t) \cdot U_{B,du}(t)$
- 3) terms related to both type of failures: $U_{3,A}(t) = u_A \cdot U_{B,du}(t), U_{3,B}(t) = U_{A,du}(t) \cdot u_B$

Gathering all the results obtained over a given duration T lead to:

- Hazardous event probability: $P_a(\tau) \approx N_a(T) \approx [\bar{U}_1(T) + \bar{U}_2(T) + \bar{U}_{3,A}(T) + \bar{U}_{3,B}(T)] \lambda_d \cdot \tau = \bar{U}_{AB} \cdot \lambda_d \cdot \tau$
- Hazardous event frequency: $\bar{\Phi}_a(T) \approx \bar{U}_{AB} \cdot \lambda_d$

The term $U_1(t)$ and the terms $U_{3,A}(t)$ and $U_{3,B}(t)$ can be easily evaluated from the results obtained in the single failure analysis (cf. [I.2](#)).

Then only the term $U_2(t) = U_{A,du}(t) \cdot U_{B,du}(t)$ is new and is analysed below.

I.3.2 Dangerous undetected failures

I.3.2.1 Simultaneous tests

The simplest case is presented on [Figure I.6](#) where the two components are tested at the same time with the same test interval.

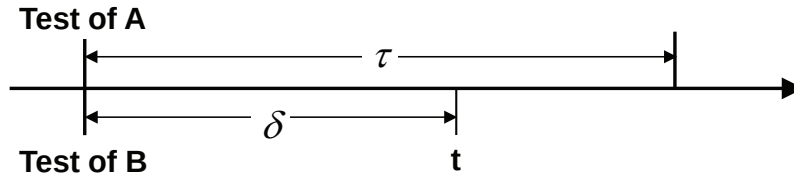


Figure I.6 — Two components tested at the same time

Within a test interval, the failure may occur in four different ways. Therefore the unavailability due to dangerous undetected failures has 4 different terms:

- 1) A and B are failed undetected: $\bar{U}_2^1(\tau)$
- 2) A and B are under repair after a test: $\bar{U}_2^2(\tau)$
- 3) B fails during the repair of A: $\bar{U}_2^3(\tau)$
- 4) A fails during the repair of B: $\bar{U}_2^4(\tau)$

At the end the average unavailability of the system is obtained as:

$$\bar{U}_2(\tau) = \bar{U}_2^1(\tau) + \bar{U}_2^2(\tau) + \bar{U}_2^3(\tau) + \bar{U}_2^4(\tau)$$

These four terms are evaluated hereafter.

I.3.2.2 A and B are failed undetected

If $\delta = t \bmod \tau$ (i.e. δ equal t modulo τ):

- unavailability: $U_2^1(t) = U_{A,du}(\delta) \cdot U_{B,du}(\delta) \approx \lambda_{A,du} \delta \cdot \lambda_{B,du} \delta = \lambda_{A,du} \cdot \lambda_{B,du} \delta^2$
- average unavailability: $\bar{U}_2^1(\tau) \approx \frac{\int_0^\tau \lambda_{A,du} \cdot \lambda_{B,du} \delta^2 d\delta}{\tau} = \frac{\lambda_{A,du} \cdot \lambda_{B,du} \tau^2}{3}$

It is observed that $\bar{U}_2^1 \oplus \frac{4}{3} \bar{U}_{A,du} \cdot \bar{U}_{B,du}$

INFO BOX: This means that the multiplication of the individual average unavailabilities of A and B does not provide a conservative estimation of the average unavailability of the system AB. Nevertheless these kinds of calculations are often observed.

When the components are tested at the same time their availabilities are good (just after a test) and bad (just before test) at the same time. This introduces a correlation between the component availabilities which is a systemic dependency. The coefficient 4/3 which is greater than 1, only appears when establishing the formula for the whole system. It is the manifestation of this “systemic dependency” existing between the components A and B because of the periodic tests.

I.3.2.2.1 A and B under repair after a test

In this case:

- unavailability: $U_2^2(\delta) \approx (\lambda_{A,du} \cdot \tau) \cdot e^{-\mu_{A,du} \cdot \delta} (\lambda_{B,du} \cdot \tau) \cdot e^{-\mu_{B,du} \cdot \delta} = \lambda_{A,du} \cdot \lambda_{B,du} \cdot \tau^2 \cdot e^{-(\mu_{A,du} + \mu_{B,du}) \cdot \delta}$
- average unavailability over $[0, \tau]$: $\bar{U}_2^2(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\mu_{A,du} + \mu_{B,du}} \tau$

NOTE The implicit hypothesis is that A and B has their own repair team. This may be an optimistic hypothesis. If there is a single repair team, the repair policy should be modelled (e.g. component A repaired first) or an equivalent repair rate ($\mu_{AB,du}$, μ_A , μ_B) should be evaluated.

I.3.2.2.2 One component fails during the repair of the other

B fails during the repair of A gives:

- $f_2^3(\delta) \approx (\lambda_{A,du} \cdot \tau) \cdot e^{-\mu_{A,du} \delta} \cdot \lambda_{B,du} e^{-\lambda_{B,du} \delta} = \lambda_{A,du} \cdot \tau \cdot \lambda_{B,du} e^{-(\lambda_{B,du} + \mu_{A,du}) \delta}$ is the probability of failure of B between δ and $\delta + d\delta$ when A is under repair
- Unavailability: $U_2^3(t) = \int_0^t f_2^3(\delta) d\delta \approx (\lambda_{A,du} \cdot \tau) \cdot \frac{\lambda_{B,du}}{\lambda_{B,du} + \mu_{A,du}} (1 - e^{-(\lambda_{B,du} + \mu_{A,du}) t})$
- Average unavailability: $\bar{U}_2^3(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\lambda_{B,du} + \mu_{A,du}} \tau$ when considering $e^{-\mu_{A,du} \tau} \approx 1$ and $1/\mu_{A,du} \ll \tau$.
- A fails during the repair of B gives:
- Unavailability: $U_2^4(t) \approx (\lambda_{B,du} \cdot \tau) \cdot \frac{\lambda_{A,du}}{\lambda_{A,du} + \mu_{B,du}} (1 - e^{-(\lambda_{A,du} + \mu_{B,du}) t})$
- Average unavailability: $\bar{U}_2^4(\tau) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\lambda_{A,du} + \mu_{B,du}} \tau$ when considering $e^{-\mu_{B,du} \tau} \approx 1$ and $1/\mu_{B,du} \ll \tau$.

I.3.2.3 Non-simultaneous tests

Figure I.7 shows an example where the components have different tests intervals. Therefore they are not tested at the same time. With such a test pattern, simultaneous tests occurs only when $\theta_A + n \cdot \tau_A = \theta_B + m \cdot \tau_B$ where m and n are integers and θ_A and θ_B the length of the first test intervals (if they are different from the others). This can occur often, rarely or never (every 4 tests of A and 7 test of B in Figure I.7). It is the work of the analyst to perform the analysis of this pattern: on Figure I.7 the time intervals 2 to 10 start with the test of one component and end with the test of the other; intervals 1 and 4 start after the tests of the two components. The interval 11 ends with the tests of the two components.

For the example presented in Figure I.7, the renewal period (i.e. the duration after which the same pattern of intervals comes back) is $T_R = 4 \cdot \tau_A = 7 \cdot \tau_B$. This implies that it is necessary to establish the unavailability formulae for the 11 different intervals comprised within T_R . This can be done just by using formulae similar to those established above. Finally the average unavailability can be calculated

over the renewal period by: $\bar{U}_{du}(T_R) = \frac{\sum_i T_i \cdot \bar{U}_{du,i}(T_i)}{T_R}$ where $T_R = \sum_i T_i$ is the renewal duration and T_i is the length of the interval i .

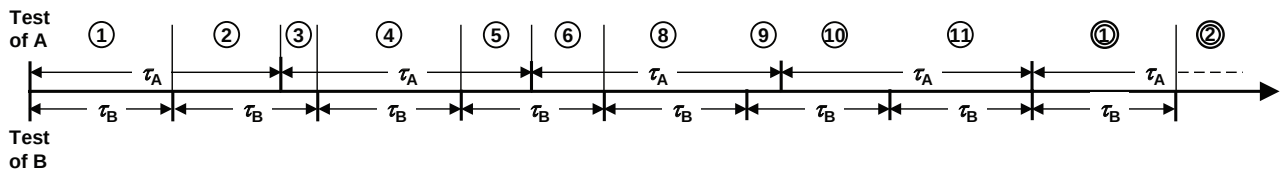


Figure I.7 — Example of components not simultaneously tested

A difficulty arises when there is no renewal period or when it is larger than the period of interest. In these cases, all the intervals over the period of interest have to be identified and considered for the calculations.

Figure I.8 provides the example of the special case where the test intervals have the same duration (τ).

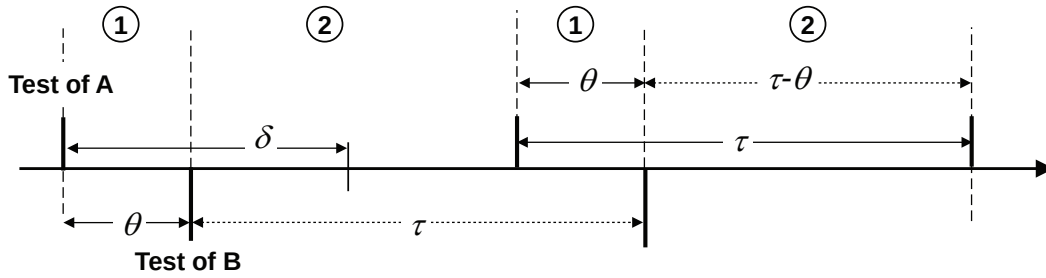


Figure I.8 — Example of similar test intervals

In this case, all the time intervals begin after the test of one component and end when the other is tested. There is no renewal period but the test pattern comprises only the two types of intervals identified as 1 and 2 on [Figure I.8](#). In addition, the test interval τ is the sum of one interval 1 and one interval 2. On the long range (i.e. after 2 or 3 test interval) the average unavailabilities of such a system converges toward asymptotic values in each of the intervals (see the NOTE hereafter).

A test duration beginning by a test of A comprises one interval of type 1 (duration θ) and one interval of type 2 (duration $\tau - \theta$). A test duration τ beginning by a test of B comprises one interval of type 2 (duration $\tau - \theta$) and one interval of type 1 (duration θ). Therefore the asymptotic average unavailability over the test duration is the same either it begins by a test of A or a test of B. It can be established by establishing the asymptotic average unavailabilities of intervals 1 and 2 and: $\bar{U}_{AB,du}(\tau) = \bar{U}_{AB,du}^{(1)}(\theta) + \bar{U}_{AB,du}^{(2)}(\tau - \theta)$

NOTE The asymptotic value $\bar{U}^{as}(T) = \lim_{t \rightarrow \infty} \bar{U}(t, t+T)$ of the average unavailability cannot be mixed up with the asymptotic value $U^{as} = \lim_{t \rightarrow \infty} U(t)$ of the instantaneous unavailabilities. U^{as} exists only where a steady-state is reached and there is no steady-state in this example.

Interval of type 1:

Within an interval of type 1, the instant δ , $\delta \in [0, \theta]$, measured from the last test of A is also the instant $\tau - \theta + \delta$ measured from the last test of B. Then for the test interval of type 1:

- Unavailability: $U_{AB,du}^{(1)}(t) \approx \lambda_{A,du} \delta \cdot \lambda_{B,du} (\delta + \tau - \theta) = \lambda_{A,du} \cdot \lambda_{B,du} \delta (\delta + \tau - \theta)$
- Average unavailability: $\bar{U}_{AB,du}^{(1)}(\theta) \approx \frac{\int_0^\theta \lambda_{A,du} \cdot \lambda_{B,du} (\delta^2 + \delta(\tau - \theta)) d\delta}{\tau} = \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\tau} \left(\frac{\theta^3}{3} + \frac{\theta^2(\tau - \theta)}{2} \right)$

Interval of type 2:

Within an interval of type 2, the instant δ , $\delta \in [\theta, \tau]$, measured from the last test of A is also the instant $\delta - \theta$ measured from the last test of B. Then for the test interval of type 2:

- Unavailability: $U_{AB,du}^{(2)}(t) \approx \lambda_{A,du} \cdot \lambda_{B,du} \delta (\delta - \theta)$
- Average unavailability: $\bar{U}_{AB,du}^{(2)}(\tau - \theta) \approx \frac{\int_\theta^\tau \lambda_{A,du} \cdot \lambda_{B,du} \delta (\delta - \theta) d\delta}{\tau} = \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{\tau} \left(\frac{\tau^3}{3} - \frac{\tau^2\theta}{2} - \frac{\theta^3}{3} + \frac{\theta^3}{2} \right)$

Interval of type 1 + interval of type 2:

Gathering the results for type 1 and type 2 gives the average unavailability for a whole test interval of duration equal to τ : $\bar{U}_{AB,du}(\tau) = \bar{U}_{AB,du}^{(1)}(\theta) + \bar{U}_{AB,du}^{(2)}(\tau - \theta) \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du}}{6} (2\tau^2 - 3\tau\theta + 3\theta^2)$

When τ has a fixed value, the average unavailability is a function of θ . The analysis of the derivative of $\frac{d}{d\theta} (2\tau^2 - 3\tau\theta + 3\theta^2) = 6\theta - 3\tau$ shows that:

- $\bar{U}_{AB,du}(\tau)$ is minimum when $\theta = \tau/2$: $\bar{U}_{AB,du}(\tau) = \lambda_{A,du} \cdot \lambda_{B,du} \frac{5}{24} \tau^2$

- $\bar{U}_{AB,du}(\tau)$ is maximum when $\theta = 0$: $\bar{U}_{AB,du}(\tau) = \lambda_{A,du} \cdot \lambda_{B,du} \frac{1}{3} \tau^2$.

INFO BOX: Therefore the average unavailability decreases with 37,5 % when going from simultaneous tests to perfectly staggered tests. Whether this “systemic” improvement is significant or not will depend on the application.

This analysis should be completed with the unavailabilities due to repairs but when a failure is discovered by a test several operational procedures can be considered:

- nothing is done on the other component which will be tested at its next planned test;
- the other component is immediately tested and the protected equipment is stopped if both components are failed (in this case the risk of accident during repair is removed);
- the other component is immediately tested in order to verify that a common cause failure has not occurred.

The formulae corresponding to all these cases are not established but this can be done in a similar way to what has been described above.

I.3.3 Double dangerous detected / immediately revealed failures

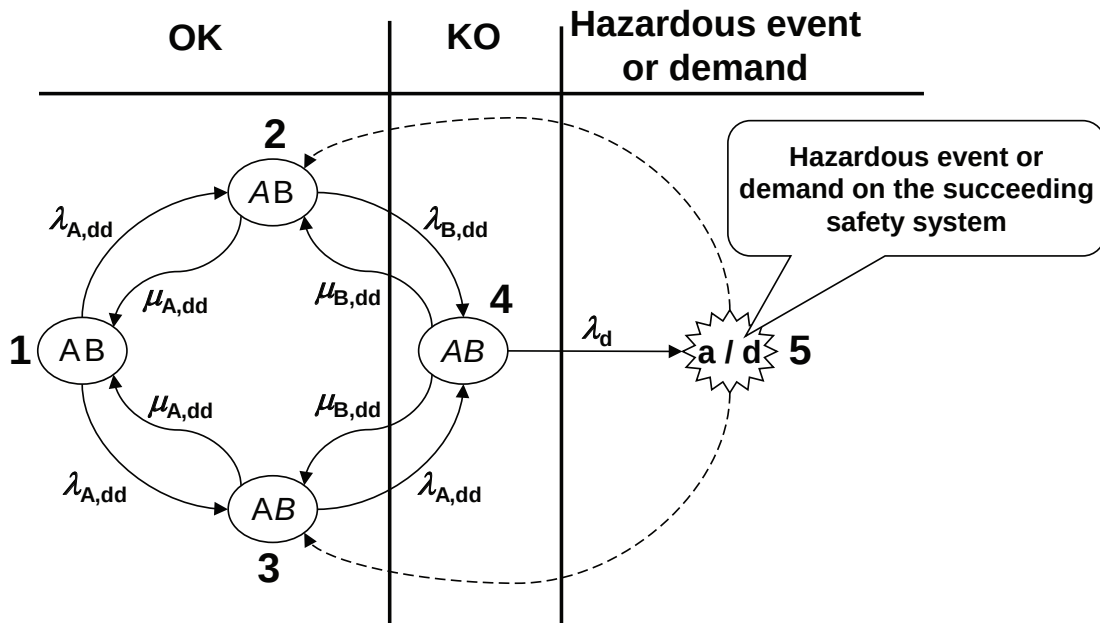


Figure I.9 — Dangerous detected failures model for a two-component minimal cut set

Figure I.9 illustrates the case when both A and B have only dangerous detected failures (see the Markovian modelling in Clause 9). The system failure occurs when both A and B are failed (state n° 4) and if a demand occurs in this state then a hazardous event (ultimate safety system) or a demand on the following safety layer (non-ultimate safety system) occur.

In case of dangerous detected failures, the detection is normally quasi instantaneous and repair can start quickly. Therefore the times to restore are small compared to the times to fail and, compared to the state 1, the states 2, 3 and 4 are quasi instantaneous. Then for the overall safety system:

- Equivalent failure rate: $\Lambda_{AB,dd} \approx \lambda_{A,dd} \frac{\lambda_{B,dd}}{\lambda_{B,dd} + \mu_{A,dd}} + \lambda_{B,dd} \frac{\lambda_{A,dd}}{\lambda_{A,dd} + \mu_{B,dd}}$

- Critical failure rate (see [Annex B](#)): $\Lambda_{AB,dd}$
- Mean time do fail (MTTF): $1 / \Lambda_{AB,dd}$
- Mean time to restore (MTTRes): $\frac{1}{\mu_{A,dd} + \mu_{B,dd}}$
- Average unavailability: $\bar{U}_{AB,dd} = \frac{MTTRes}{MTTRes + MTTF} \approx \frac{\Lambda_{AB,dd}}{\Lambda_{AB,dd} + \mu_{A,dd} + \mu_{B,dd}} \approx \frac{\Lambda_{AB,dd}}{\mu_{A,dd} + \mu_{B,dd}}$
- Probability of hazardous event (accident) over $[0, T]$: $P_{a,dd}(T) \approx N_{A,dd}(T) \approx \bar{U}_{AB}(T) \cdot \lambda_d \cdot T$ (or probability of a demand on the succeeding safety layer):
- Hazardous event (or accident) frequency: $\bar{\Phi}_{a,dd}(T) \approx \bar{U}_{AB}(T) \cdot \lambda_d \approx \frac{\Lambda_{AB,dd} \lambda_d}{\mu_{A,dd} + \mu_{B,dd}}$

[Figure I.10](#) illustrates a model equivalent to the model of [Figure I.9](#) where the equivalent failure rate $\Lambda_{AB,dd}$ models the jump from the OK class to the state AB. The jump from AB state to the OK class is given by $\mu_{A,dd} + \mu_{B,dd}$ as there is two possibilities to repair the state AB.

Compared to the time spent in the OK class - $1/\Lambda_{AB,dd}$ -, the time spent in AB is very short - $1/(\mu_{A,dd} + \mu_{B,dd})$ - Then when the system jumps to AB its goes out this state very quickly with a probability $(\mu_{A,dd} + \mu_{B,dd})/(\mu_{A,dd} + \mu_{B,dd} + \lambda_d)$ to come back to OK and a probability $\lambda_d/(\mu_{A,dd} + \mu_{B,dd} + \lambda_d)$ to reach the hazardous event state.

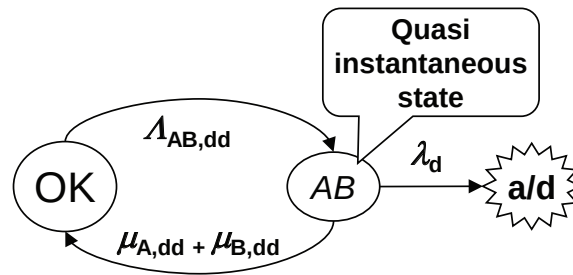


Figure I.10 — Equivalent model for immediately revealed failures

Then, the accident rate can be found from the Markov model in [Figure I.10](#) and:

- Hazardous event rate: $\Lambda_{a,dd} \approx \Lambda_{AB,dd} \frac{\lambda_d}{\mu_{A,dd} + \mu_{B,dd} + \lambda_d} \approx \Lambda_{AB,dd} \frac{\lambda_d}{\mu_{A,dd} + \mu_{B,dd}}$
- Probability of accident: $P_{a,dd}(T) \approx \Lambda_{a,dd} \cdot t = \frac{\Lambda_{AB,dd}}{\mu_{A,dd} + \mu_{B,dd}} \lambda_d \cdot T$

Therefore $P_{a,dd}(T) \approx \bar{U}_{AB,dd} \lambda_d \cdot T$ which is the same formula as this established just above by a different approach. All these approximations are valid when the failure rates and the demand rate are negligible compared to the restoration rates: $\Lambda_{AB,dd} + \mu_{A,dd} + \mu_{B,dd} \approx \lambda_d + \mu_{A,dd} + \mu_{B,dd} = \mu_{A,dd} + \mu_{B,dd}$.

I.3.4 Spurious failures

The number of safe failures needed to lead to a spurious action depends on the design of the safety system from which the minimal cut set $\{A, B\}$ is coming. If $\{A, B\}$ is coming from a 1oo2 architectures, then each safe failure of A or B is leading to a spurious action. But if $\{A, B\}$ is coming from 2oo3, 3oo4, etc. architecture then a spurious action occurs when both A and B have a safe failure.

[Figure I.11](#) describes the case where a safe failure from any of the components A or B is able to produce a spurious action (i.e. $\{A, B\}$ is coming from a 1oo2).

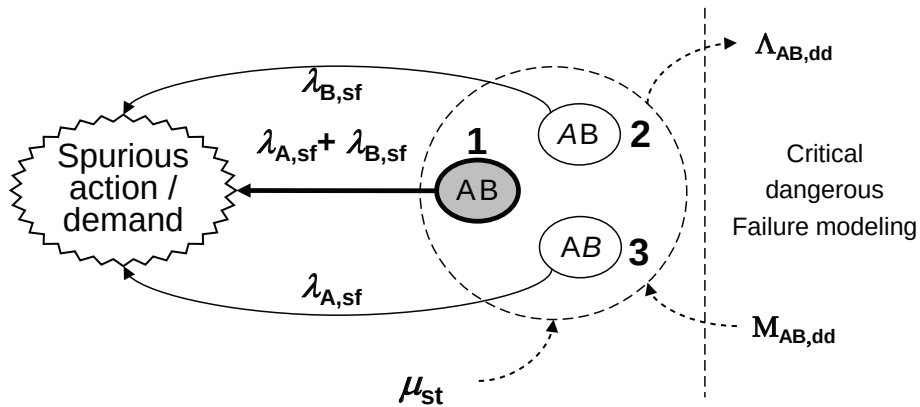


Figure I.11 — Spurious failure model - case 1002 -

Under the same hypothesis as in the previous subclause (i.e. the restoration times are short), state 2 and 3 are quasi instantaneous and it is necessary to consider only state 1:

- Spurious failure rate: $\Lambda_{AB,st}^{1002} \approx \lambda_{A,sf} + \lambda_{B,sf}$

This is similar to the same case already analysed in [I.2](#).

[Figure I.12](#) illustrates the case where $\{A, B\}$ is coming from a 2 out of 3 voting architecture. Then two safe failures (e.g. both A and B) are needed to trigger a spurious safety action. The establishment of the spurious failure rate is similar to what has been described above for the equivalent dangerous failure rate $\Lambda_{AB,dd}$ from [Figure I.9](#) or [Figure I.10](#). Then, directly from [Figure I.12](#):

- Spurious failure rate: $\Lambda_{AB,st}^{2003} \approx \lambda_{A,sf} \frac{\lambda_{B,sf}}{\lambda_{B,sf} + \mu_{A,sf}} + \lambda_{B,sf} \frac{\lambda_{A,sf}}{\lambda_{A,sf} + \mu_{B,sf}}$

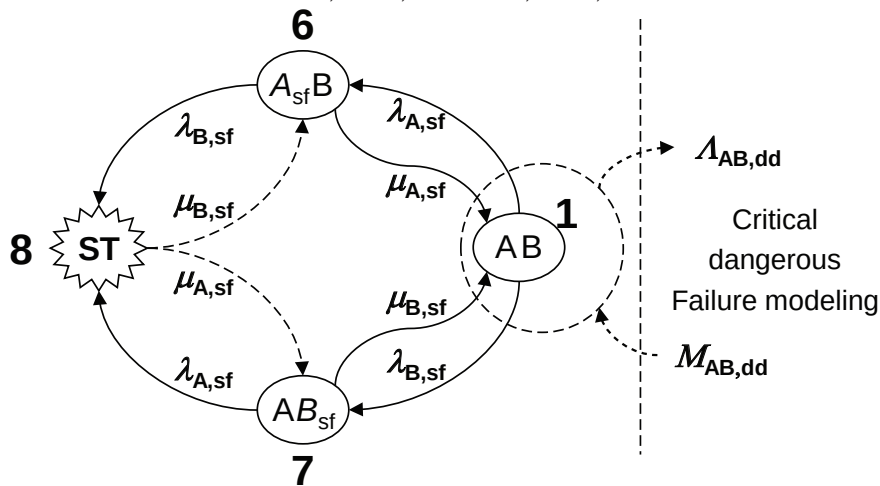


Figure I.12 — Spurious failure model - case 2003 -

In both cases (1002 and 2003):

- Probability of spurious failure over $[0, t]$ is: $P_{AB,st}(t) \approx 1 - e^{-\Lambda_{AB,st} \cdot t} \approx \Lambda_{AB,st} t$ when $\Lambda_{AB,st} t \ll 1$.

NOTE When the λ_{sf} are high then, the probability of spurious action may be very high and even equal to 1 if several spurious actions are likely to occur in a given period. In this case, the above probability is useless but can be replaced by the frequency of spurious failure.

From [Figure I.12](#):

- Mean time to spurious failures: $MTTF_{st}^{(2)} = \frac{1}{\lambda_{AB,st}^{(2)}}$
- Mean time to repair the spurious failures: $MTTRes_{st}^{(2)} \approx \frac{1}{\mu_{A,sf} + \mu_{B,sf}}$
- Number of spurious failures: $N_{st}^{(2)}(T) \approx \frac{T}{MTBF_{st}^{(2)}} = \frac{T}{MTTF_{st}^{(2)} + MTTRes_{st}^{(2)}}$
- Average spurious failure frequency: $\bar{\phi}_{st}^{(2)} = \frac{1}{MTBF_{st}^{(2)}}$

I.3.5 Combination of dangerous detected and undetected failures

As shown in [I.3.1](#), the unavailability due to the combination of a dangerous undetected failure of A with a dangerous detected failure of B is given by $U_{A,du}(t) \cdot u_B$ and the unavailability due to a dangerous detected failure of A with a dangerous undetected failure of B is given by $u_A \cdot U_{B,du}(t)$.

u_B , u_A are constant values and $U_{A,du}(t)$ and $U_{B,du}(t)$ are related to the instantaneous unavailabilities of single component. Therefore the results established in [I.2](#) for single failures are available to perform these unavailability calculations and no more development is done in this Technical Report.

I.4 Triple failure analysis

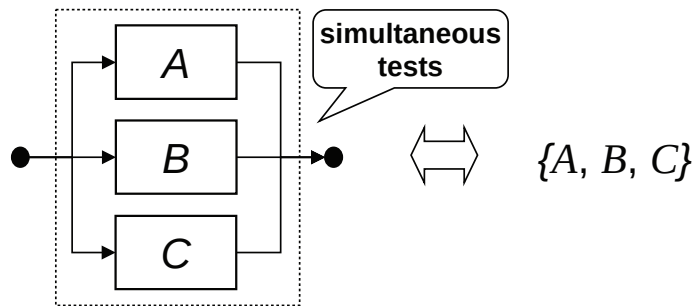


Figure I.13 — Redundant components and triple failures

The analysis which was not simple for double failures becomes very complicated for triple failures because the number of situations increases more or less exponentially with the order of the minimal cut sets under consideration. Nevertheless this analysis can be performed by skilled reliability engineers applying the principles developed above for single and double failures.

Developing in this Technical Report the complete general analysis of triple failures would lead to a great number of pages covered with complicated formulae. This would not be very effective as their meanings would likely be not really understood by users not skilled in probabilistic calculations.

Therefore only the more simplistic situation is analysed: three components tested at the same time and with only the basic parameters defined in [I.2](#). None of the extra parameters (γ , ω , π , test staggering, etc.) are considered hereafter.

The basic reliability parameters are the following:

- $\lambda_{A,du}$, $\lambda_{B,du}$, $\lambda_{C,du}$: dangerous undetected failure rates of components A, B and C (i.e. hidden failures detected by periodic tests);
- $\lambda_{A,dd}$, $\lambda_{B,dd}$, $\lambda_{C,dd}$: dangerous detected failure rates of components A, B and C (i.e. immediately revealed failures or failures detected by diagnostic tests);
- $\mu_{A,du}$, $\mu_{B,du}$, $\mu_{C,du}$: repair rates of dangerous undetected failure of components A, B and C;
- μ_{dd} : restoration rate of immediately revealed dangerous failure ($1/\mu_{dd} = MTTRes_{dd}$);

- τ : test interval of components A, B and C tested at the same time.

As said above, this part of the annex deals only with the simplest case of minimal cut sets of order three and the reader should be aware that this constitutes a simplification which may be non conservative when the assumptions are not fulfilled. The aim of this part of the annex is only to provide some guidance and to encourage the analysts to understand and develop their own formulae rather than to provide a catalogue of ready-to-use formulae.

I.4.1 Dangerous undetected failures

The simplest case is presented on [Figure I.14](#) where the three components are tested at the same time with the same test interval. With this hypothesis the average unavailability of the safety system is also the average unavailability calculated over a test interval.

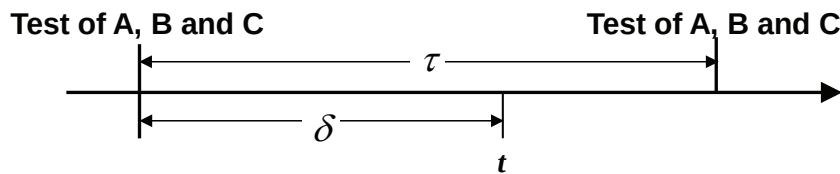


Figure I.14 — Three components tested at the same time

In the following, focus is only on the dangerous undetected failures where 8 cases can be identified:

- 1) A, B and C failed undetected: $\bar{U}_3^1(\tau)$
- 2) A, B and C under repair just after a test: $\bar{U}_3^2(\tau)$
- 3) One component failing during the repair of the two others: $\bar{U}_3^3(\tau), \bar{U}_3^4(\tau), \bar{U}_3^5(\tau)$
- 4) Two component failing during the repair of the third one: $\bar{U}_3^6(\tau), \bar{U}_3^7(\tau), \bar{U}_3^8(\tau)$

The lower index 3 has been kept (e.g. $\bar{U}_3^1(\tau)$) to remain homogeneous with the analysis of double failures. At the end, the total average unavailability of the system is obtained as the sum of the individual unavailabilities of the above seven case: $\bar{U}_3(\tau) = \bar{U}_3^1(\tau) + \bar{U}_3^2(\tau) + \bar{U}_3^3(\tau) + \dots + \bar{U}_3^7(\tau)$

Only the two first cases are analysed in this Technical Report but the principles are the same for the other ones.

I.4.1.1 A, B and C failed undetected

- Unavailability: $U_3^1(t) = U_{A,du}(\delta) \cdot U_{B,du}(\delta) \cdot U_{C,du}(\delta)$ where $\delta = t \bmod \tau$
- Unavailability (approximation): $U_3^1(t) \approx \lambda_{A,du} \delta \cdot \lambda_{B,du} \delta \cdot \lambda_{C,du} \delta = \lambda_{A,du} \cdot \lambda_{B,du} \cdot \lambda_{C,du} \delta^3$
- Average unavailability: $\bar{U}_3^1(\tau) \approx \frac{\int_0^\tau \lambda_{A,du} \cdot \lambda_{B,du} \cdot \lambda_{C,du} \delta^3 d\delta}{\tau} = \frac{\lambda_{A,du} \cdot \lambda_{B,du} \cdot \lambda_{C,du} \tau^3}{4}$

As for the double failures, it is observed that:

$$\bar{U}_3^1(\tau) \approx \frac{8}{4} \bar{U}_{A,du}(\tau) \bar{U}_{B,du}(\tau) \bar{U}_{C,du}(\tau) = 2 \bar{U}_{A,du}(\tau) \bar{U}_{B,du}(\tau) \bar{U}_{C,du}(\tau)$$

INFO BOX: Again this means that the multiplication of the individual average unavailabilities of A, B and C does not provide a conservative estimation of the average unavailability of the system ABC. Unfortunately these kinds of calculations are often observed.

The coefficient $8/4 = 2$ which is greater than 1, only appears when establishing the formula for the whole system. It is the manifestation of the “systemic dependency” introduced between the components A and B because of the periodic tests. This is worse than for double failures and the coefficient increases when the order of the minimal cut sets increases.

I.4.1.2 A, B and C under repair just after the test

- Unavailability: $U_3^2(\delta) \approx (\lambda_{A,du} \cdot \tau) \cdot e^{-\mu_{A,du} \cdot \delta} (\lambda_{B,du} \cdot \tau) e^{-\mu_{B,du} \cdot \delta} (\lambda_{C,du} \cdot \tau) e^{-\mu_{C,du} \cdot \delta} = \lambda_{A,du} \cdot \lambda_{B,du} \cdot \lambda_{C,du} \cdot \tau^3 e^{-(\mu_{A,du} + \mu_{B,du} + \mu_{C,du}) \cdot \delta}$
- Average unavailability: $\bar{U}_3^2 \approx \frac{\lambda_{A,du} \cdot \lambda_{B,du} \cdot \lambda_{C,du}}{\mu_{A,du} + \mu_{B,du} + \mu_{C,du}} \tau^2$

NOTE The implicit hypothesis is that A, B and C have their own repair teams. This may be an optimistic hypothesis. If there is a single repair team for the three components then the repair policy should be modelled (e.g. component A repaired first or component with the higher repair rate repaired first) or an equivalent repair rate ($\mu_{ABC,du} < \mu_{A,du} + \mu_{B,du} + \mu_{C,du}$) should be evaluated.

As said above the other 6 cases may be established by using the same principles.

I.4.2 Triple dangerous detected/immediately revealed failures

The principle is similar as for double failures in [I.3.3](#) but there are 7 OK states (ABC, ABC, ABC, ABC, ABC, ABC, ABC) instead of 3.

The formulae can be established by using the same principles but they are more complicated. No more developments are done in this Technical Report.

I.4.3 Spurious failures

Similarly to double failures, the analysis of the spurious failure should know from which architecture the minimal cut set come from. According to the logic of the system under study 1, 2 or 3 individual safe failures may be needed to trigger a spurious failure.

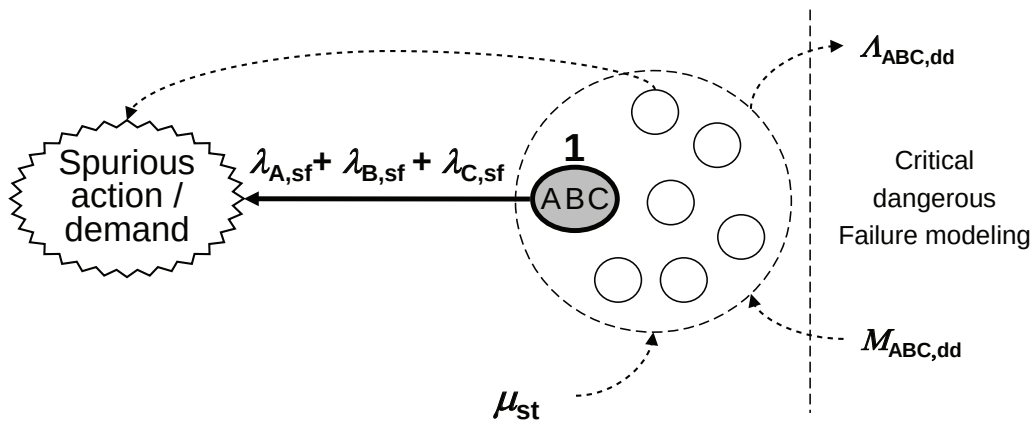


Figure I.15 — Spurious failure model for 3 component minimal cut set (1003 logic)

Figure I.15 is related to a safety system implementing a 1003 voting logic where all single safe failures imply a spurious action of the safety system. The corresponding spurious failure rate is then

$$\Lambda_{ABC,st} \approx \lambda_{A,s} + \lambda_{B,s} + \lambda_{C,s}$$

Other logics (e.g. 2004, 3005, etc.) should be analysed on a case by case basis as this has been done for the 2003 configuration in [I.3](#).

I.5 Example of analytical formulae implementation: the PDS method

I.5.1 Introduction

As shown in [Clause 5](#), predicting the future reliability performance of SIS is a challenging subject full of pitfalls and conflicting interests. On the one hand, the manufacturer of SIS equipment obviously wants to demonstrate reliability estimates that are as good as possible. On the other hand, plant operators would like to have reliability estimates that corresponds to the actual performance in field. Questions are therefore frequently asked on how the reliability predictions may be performed and what a reliability analysis should include.

The PDS²⁾ method has been developed to answer these questions in relation to the oil and gas industries. It is usable in the case where the restrictive assumptions described in [Clause 7](#) are fulfilled and that analytical formulae can be implemented. The method comprises two main parts:

- A *method handbook*[\[13\]](#) primarily used to quantify the safety unavailability of SIS and the loss of production due to spurious actions.
- A *data handbook*[\[18\]](#) that supplements the method with relevant input data.

The main objectives of the PDS method are consistent with the scope of this Technical Report and principally with the framework and assumptions presented in [Clause 7](#):

- present simple calculation formulae together with reliability data relevant for the oil and gas industries;
- provide a new common cause failures (CCF) model that accounts for different types of voting configurations;
- describe the treatment of some non-hardware or non-random failures (see the note hereafter) in reliability analyses;
- describe the treatment of downtime contribution from repair and regular periodic tests, in light of strategies for degraded operation.

NOTE The failures other than “*random hardware failures*” are named “*systematic*” failure in the PDS method. They are different from how the systematic failures are defined in this Technical Report (see definition number [3.2.17](#)). The purpose of this part of the annex is to give a brief summary of some of the mentioned features. The treatment of CCFs in the PDS method is presented in [G.1.2](#). The failure classification in the PDS method is different from the failure classification given in [B.3](#).

I.5.2 Main features of the PDS method

I.5.2.1 PDS performance measures for loss of safety

In PDS three main contributions to the SIS dangerous failures unavailability identified:

- 1) *Unavailability due to dangerous undetected (DU) failures* during the fault detection time (noted PFD in the PDS method), i.e. unavailability caused by dangerous failures not revealed by automatic self-test (i.e. diagnostic test) and that are *detectable* only during periodic testing *or* upon a demand.

NOTE 1 PFD_{avg} is used instead of PFD in IEC 61508,[\[2\]](#) IEC 61511[\[3\]](#) and this Technical Report (see definition [3.1.16](#)).

- 2) *Unavailability due to test-independent failures* (noted P_{TIF} by PDS). This does not include the time-dependent failures and is equivalent to the factor ψ (probability of failure due to demand) introduced in [Clause 7](#).

2) PDS is the Norwegian acronym for “reliability of computer-based safety systems”.

3) *Unavailability due to known or planned downtime* (noted DTU by PDS). This unavailability is caused by components either taken out for repair or for testing/maintenance. This downtime unavailability can be split in two main contributors:

- i) The *known* unavailability due to dangerous (D) failures where the failed component should be repaired (cf. 7.3 and I.2.1). The average period of unavailability due to these events equals the mean repairing time (MRT) (see note), i.e. the time elapsing from the failure is detected until the situation is restored.

NOTE 2 The parameter MRT is not differentiated with the MTTR in the PDS approach.

NOTE 3 ISO 14224[15], C.2.3.2 describes the 'technical unavailability' and how its corrective maintenance causes unavailability (downtime).

- ii) The *planned* (and known) unavailability due to the downtime/inhibition time during periodic testing (see 7.3 and I.2.1) and/or preventive maintenance.

NOTE 4 A component may be available for its safety function during the periodic tests (e.g. a valve is available during a partial of a full stroking).

NOTE 5 ISO 14224[15], C.2.3.2 describes the 'operational unavailability' and how its preventive maintenance (PM) causes unavailability (downtime). If a failure is detected during PM this will give technical unavailability, unless small repair or depending on the operator CMMIS reporting routines

The sum of these three unavailability measures provides a measure of the *SIS* unavailability (PFD_{avg}) which is called critical safety unavailability (CSU) in the PDS approach. It is also referred to as loss of safety.

The contribution to loss of safety from failures in category 3) (3a and 3b) is influenced by the operating philosophy. Temporary compensating measures may, for example, be introduced while a component is down for maintenance, testing or repair. If the component is considered too critical to continue to operate the protected installation (e.g. a critical shutdown valve in single configuration), the production may simply be, for example, shut down during the restoration and testing period. Hence, *DTU* should be treated separately and not as part of category 1) or 2).

The contributions 3a) and 3b) are usually small compared to the contribution from failures in category 1). That is, usually $MRT \ll \tau$. This is, however, not always the case; e.g. for subsea equipment in offshore production, the MRT could be rather long. Category 3b) can often be considered the least critical, as this represents a truly planned unavailability of the safety system and since testing and maintenance is often performed during planned shutdown periods.

I.5.2.2 Modelling of common cause failures (CCF) in PDS

For modelling of common cause failures in PDS, reference is made to the separate description given in G.1.2.

I.5.3 Summary of PDS calculation formulae

The strong underlying assumptions of the PDS method allow relatively simple calculation formula, both for *SIS* operating in the low demand mode and in the continuous / high demand modes. This section gives a brief summary of the formulae implemented in PDS for low demand systems. Details about developing such formulae are available in Clause 7 and above in this Annex.

I.5.3.1 Formulae for PFD

The analytical formulae for PFD are summarized in Table I.1. As seen in the table, the contribution is split into two parts: the contribution from independent failures and the contribution from common cause failures (CCFs). All these formulae are established for similar components and synchronous tests (i.e. all the components are tested at the same time).

Table I.1 — Summary of analytical formulae for PFD

Voting logic	PFD calculation formulae	
	Common cause contribution	Contribution from independent failures
1oo1	-	$\lambda_{DU} \cdot \tau / 2$ (I.2.1)
1oo2	$\beta \cdot \lambda_{DU} \cdot \tau / 2$ (I.2.1)	+ $[\lambda_{DU} \cdot \tau]^2 / 3$ (I.3.2.1)
2oo2	-	$2 \cdot \lambda_{DU} \cdot \tau / 2$ (I.2.1 with 2 possibilities)
1oo3	$C_{1oo3} \cdot \beta \cdot \lambda_{DU} \cdot \tau / 2$ (I.2.1)	+ $[\lambda_{DU} \cdot \tau]^3 / 4$ (I.4.1)
2oo3	$C_{2oo3} \cdot \beta \cdot \lambda_{DU} \cdot \tau / 2$ (I.2.1)	+ $[\lambda_{DU} \cdot \tau]^2$ (I.3.2.1 applied with 3 possibilities of failures)
3oo3	-	$3 \cdot \lambda_{DU} \cdot \tau / 2$ (I.2.1 with 3 possibilities of failures)
1ooN; $N = 2, 3, \dots$	$C_{1ooN} \cdot \beta \cdot \lambda_{DU} \cdot \tau / 2$ (I.2.1)	+ $\frac{1}{N+1} \cdot (\lambda_{DU} \cdot \tau)^N$ (generalization from I.2.1, I.3.1 and I.4.1)
MooN, $M < N$; $N = 2, 3, \dots$	$C_{MooN} \cdot \beta \cdot \lambda_{DU} \cdot \tau / 2$ (I.2.1)	+ $\frac{N!}{(N-M+2)! \cdot (M-1)!} \cdot (\lambda_{DU} \cdot \tau)^{N-M+1}$ (generalization from I.2.1, I.3.1 and I.4.1 and common Boolean algebra)
NooN; $N = 1, 2, 3, \dots$	-	$N \cdot \lambda_{DU} \cdot \tau / 2$ (generalization of I.2.1)

When using the formulae of Table I.1 or any other formulae from Clause 7 or this Annex, it is important that the analyst is aware of its application area as well as any limitations on the use of the formulae. The main assumptions underlying the analytical formulae are described in 7.2.

I.5.3.2 Formulae for P_{TIF}

Table I.2 — Formulae for P_{TIF} , various voting logics

Voting logic	TIF contribution to CSU for MooN voting
1oo1	P_{TIF}
1oo2	$\beta \cdot P_{TIF}$
MooN, $M < N$	$C_{MooN} \cdot \beta \cdot P_{TIF}$
NooN, ($N = 1, 2, \dots$)	$N \cdot P_{TIF}$

For a single component, P_{TIF} expresses the likelihood of a component, to fail on demand (irrespective of the interval of periodic testing) due to a failure which is caused by the demand itself (e.g. rupture of the shaft of a valve, of the spring of a relay ...). For redundant components, voted MooN ($M < N$), the PDS approach considers that the main contribution is due to the common cause failures. Then, the TIF contribution to loss of safety is given by the general formula: $C_{MooN} \cdot \beta \cdot P_{TIF}$, where C_{MooN} is a modification factor when calculating the common cause contribution from various voting configurations, (see G.1.2). The formulae for calculation of P_{TIF} for different voting logics are summarized in Table I.2.

I.5.3.3 Formulae for DTU

The contribution from DTU (category 3a and 3b) is split into two parts:

- 1) The downtime related to repair of (dangerous) failures. The average duration of this period is the mean repairing time (MRT);
- 2) The planned downtime (or inhibition time) resulting from planned activities such as testing and preventive maintenance.

As discussed in previous sections, the contribution from downtime unavailability will depend on the operational philosophy and the configuration of the process as well as the SIS itself. Further, authority and company requirements saying that compensating measures should be introduced upon degradation of a critical safety function will also affect the operational philosophy. Hence, which formulae to apply, will depend on several factors. Below, some approximate formulae and the corresponding assumptions underlying these formulae are given for downtime unavailability due to repair of dangerous failures (DTU_R). For a discussion of formulae for downtime unavailability due to periodic tests (DTU_T), more information is available in [7.3](#) and [I.2.1](#) or [I.2.2](#).

The approximate formulae for the downtime unavailability due to repair, here referred to as DTU_R , are comparable to the PFD formulae presented above. However, given that a dangerous failure has occurred and has been detected, the average “known” unavailability period is MRT (or $1/\mu$ if μ is the repair rate, see [7.3](#) and [I.2.1](#) or [I.2.2](#)) rather than $\tau/2$.

In the following formulae it is considered that the fault detection time of dangerous detected failures is negligible then the MRT of *all dangerous* failures can be included and for a single component, the related downtime unavailability is then approximately $\lambda_D \cdot MRT$. Whether it is correct to treat dangerous detected (DD) failures detected during normal operation in the same way as DU failures revealed during a periodic test or upon a true demand is a matter of discussion. It should be noted that IEC 61508[2] makes the distinction between the MRT of the DD failures and the MTTRes of the DU failures. With the above hypothesis, $MTTRes \approx MRT$ are obtained for the dangerous detected failures.

When establishing formulae for DTU_R , the operational philosophy should be specified. Here, three possible operational/repair philosophies are considered:

- 1) *Always shut down (while repairing)*. This (extreme) philosophy may apply for the most critical safety functions, and means that production is shut down (even for redundant systems) whenever at least one component of the safety function experiences a dangerous failure. In such case, the risk disappears during repair and there will be no contribution to the DTU_R . Nevertheless there will be a contribution to loss of production.
- 2) *Degraded operation if possible; otherwise shutdown*. This may be the most common philosophy. If all redundant components have a dangerous failure there will be a shutdown, otherwise there will be degraded operation. If there is a single *D* failure in a *2oo3* voting, it should be specified whether the degraded operation is a *1oo2* or a *2oo2* voting logic. Note that if a *2oo3* voting degrades to a *1oo2* configuration, the safety performance actually improves, and no degradation term should be added.
- 3) *Always continue production, even with no protection*. This is another extreme philosophy where all the (redundant) components have experienced a dangerous failure, but production is continued during the repair/restoration period even with no protection available.

Observe that the above list is not complete since alternative operational philosophies can be foreseen (“combinations” of the above). Also note that the possibility of introducing compensating measures has not been included in this discussion. [Table I.3](#) presents DTU_R formulae for three common configurations for the two operational philosophies that may give DTU_R contributions.

When the operational/repair philosophy for safe (detected and undetected) failures are specified, similar DTU_R formulae as shown in [Table I.3](#) can be established also for these failure types. If the same repair philosophy applies for all critical failures, the approximate DTU_R formulae as given in [Table I.3](#) can be applied, simply replacing λ_D with λ_{crit} and by replacing the r out of m logics by $(m-r+1)$ out of m logics.

Similar principles as stipulated above can be applied when establishing formulae for downtime unavailability due to testing (DTU_T). [Clause 7](#) and [Annex I](#) may be applied and details are given in reference.^[1] Example of application of the PDS method can be found in reference.^[13]

Table I.3 — Formulae for DTU_R for some voting logics and operational philosophies

Initial voting logic	Failure type	Contribution to DTU_R for different operational/repair philosophies ^a	
		Degraded operation	Operation with no protection
1001	Single failure	N/A	$\lambda_D \cdot MRT = \lambda_D / \mu_D$ (1.2.1)
1002	Single failure	Degraded operation with 1001: $2 \cdot \lambda_D \cdot MRT \cdot \lambda_{DU} \cdot \tau / 2$ (1.3.2.2.2)	N/A
	Both components fail	N/A	$\beta \cdot \lambda_D \cdot MRT$
2003	Single failure	Degraded operation with 2002: ^b $3 \cdot \lambda_D \cdot MRT \cdot 2 \cdot \lambda_{DU} \cdot \tau / 2$ (1.3.2.2.2 with 3 possibilities of single failures)	N/A
	Two components fail	Degraded operation with 1001: $(C_{2003} - C_{1003}) \cdot \beta \cdot \lambda_D \cdot MRT \cdot \lambda_{DU} \cdot \tau / 2$	N/A
	All three components fail	N/A	$C_{1003} \cdot \beta \cdot \lambda_D \cdot MRT$ (1.2.1)
^a The formulae provided here do not distinguish between the $MRT/MTTRes$ for one or two (three) components.			
^b Degradation to a 1002 voting logic gives no contribution to the DTU, since a 1002 voting actually gives increased safety as compared to a 2003 voting logic.			

Annex J (informative)

Sequential modelling

The reliability block diagrams^[4] or the fault trees^[5] described in [Clause 8](#) and [Annex K](#) are focused on individual safety systems but in reality, single safety systems rarely work alone and, when a demand occurs (the so-called “initiating event”) several safety systems or safety layers run in sequence to prevent the occurrence of the related hazardous events. This has been already discussed in [Annex F](#).

The first attempt to deal with such sequences of operations has been published in 1971 by the Danish Atomic Energy Commission^[10]: this is the so-called “*Cause-Consequence Diagram*” model. Few years after, the well known “*event tree*”^[8] approach appears in the famous “*Rasmussen report*” (Wash 1400) where it is implemented to evaluate the probability of accident of nuclear power plants. At the present time the last version of this kind of approach is the “*Layer of Protection Analysis*” (LOPA)^[9].

All these models may be represented by fault trees and calculated by using fault trees but some precaution may be needed (see here after for non-coherence and sequential aspects).

[Figure J.1](#) gives an example of cause-consequence diagram. From the *initiating* event, at the top of the diagram, the diagram is developed to describe what happens when succeeding operations (e.g. the steps of a procedure or demands on safety systems) are successful or not. This is done by using *Yes/No gates* which are the core graphical element of this approach. As example, [Figure J.1](#) describes what happens during the progress of the procedure performed by the drillers when a *kick* occurs. During the drilling operations this arises when the weight of the mud column is no longer sufficient to counterbalance the pressure at the well bottom. This is the beginning of a well leak which is likely to escalate to a blow out if relevant operations are not started immediately. The cause-consequence diagram in [Figure J.1](#) describes 5 of the operations which have to be performed to come back to a safe situation in the case of subsea well drilled from a dynamic positioning drilling rig.

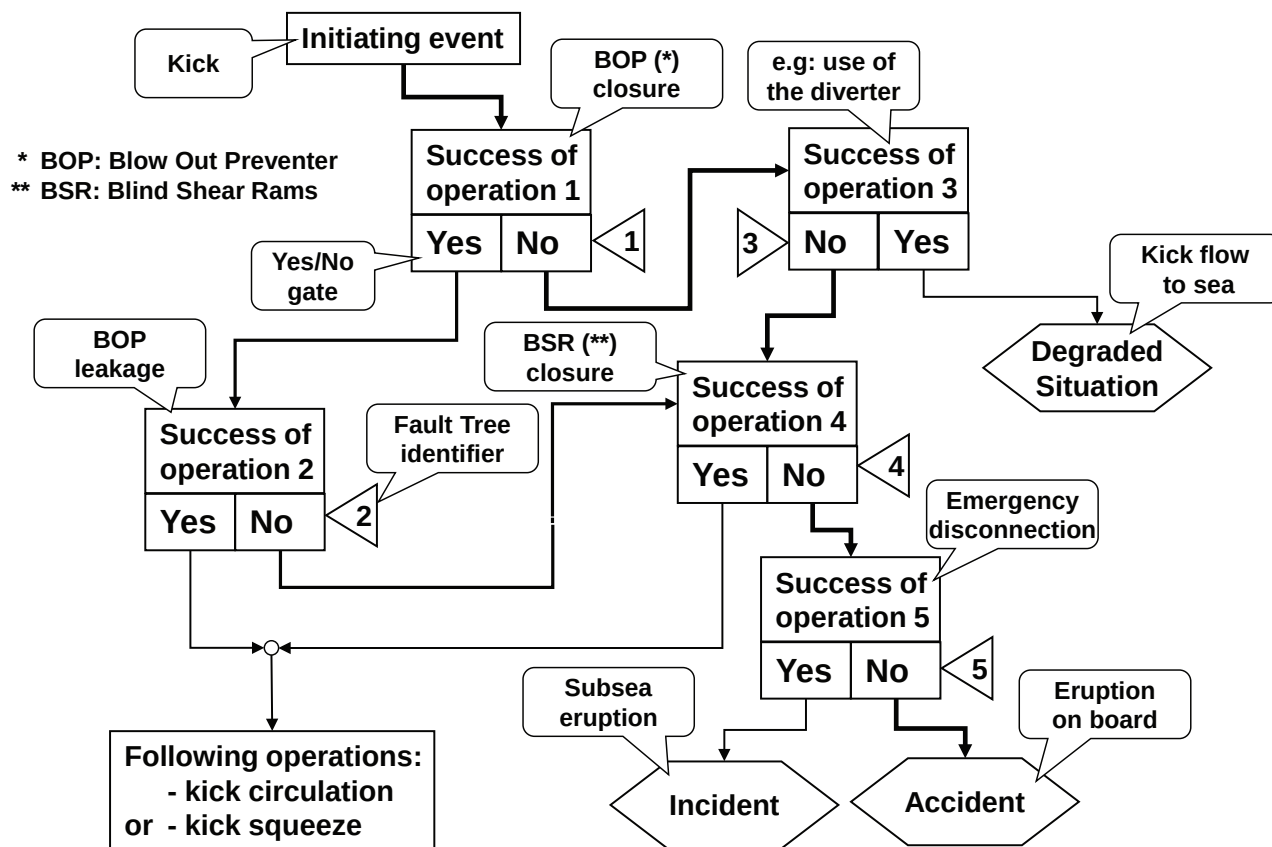
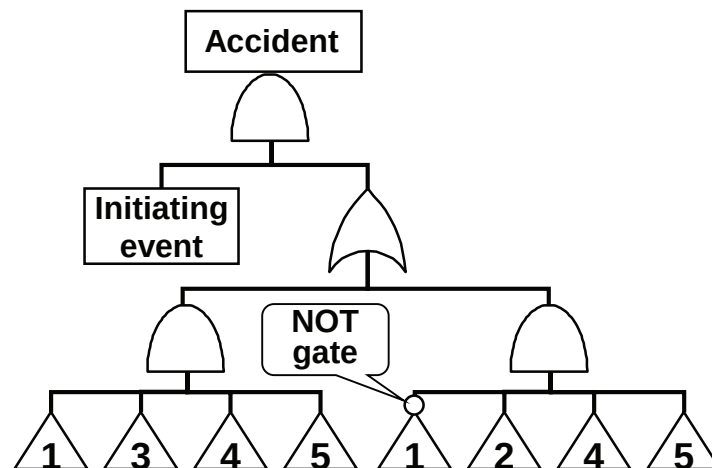


Figure J.1 — Example of cause-consequence diagram

According to the success or the failure of each succeeding operations (involving both operator and safety system actions), several sequences can be identified which lead either toward a normal situation (circulation or squeeze of the erupting flow), a degraded situation (flow sent to the sea through the diverter device), financial losses (riser disconnection) or an accident (eruption on board of the drilling vessel).

The cause-consequence diagram is really an analysis tool which helps the analyst to identify step by step the sequences of events leading to situations ranging from safe, degraded and hazardous situations. The bold line on the figure shows one of the sequences leading to the accident. This constitutes the *consequence* part of the model. The connection with the *cause* part is made by the small numbered triangles which refer to fault trees describing how the operation (or the safety system) related to a given Yes/No gate fails. Therefore, beyond the sequence identification, this diagram is also a method to build big fault trees as this has been done in [Figure J.2](#) which models the two sequences leading to an accident.



NOTE A "NOT" gate has been used in [Figure J.2](#) to represent the success of operation 1.

Figure J.2 — Global fault tree describing the sequences leading to an accident

Another important use of the cause-consequence diagram approach is to help to develop event trees. The event tree analysis is described in detail in IEC 62502[8].

[Figure J.3](#) represents the event tree which can be developed from [Figure J.1](#). In spite it looks very different it is completely equivalent to this cause-consequence diagram. Again the starting point is the initiating event. It is represented on the left hand side of the event tree. Then the various operations (or safety systems) are arranged in columns and two branches are used for each of them to graphically represent their success (Yes) or their failure (No). This allows drawing several event sequences leading to different outcomes represented at the very right hand side of the event tree. They are, of course, the same as those described in [Figure J.1](#).

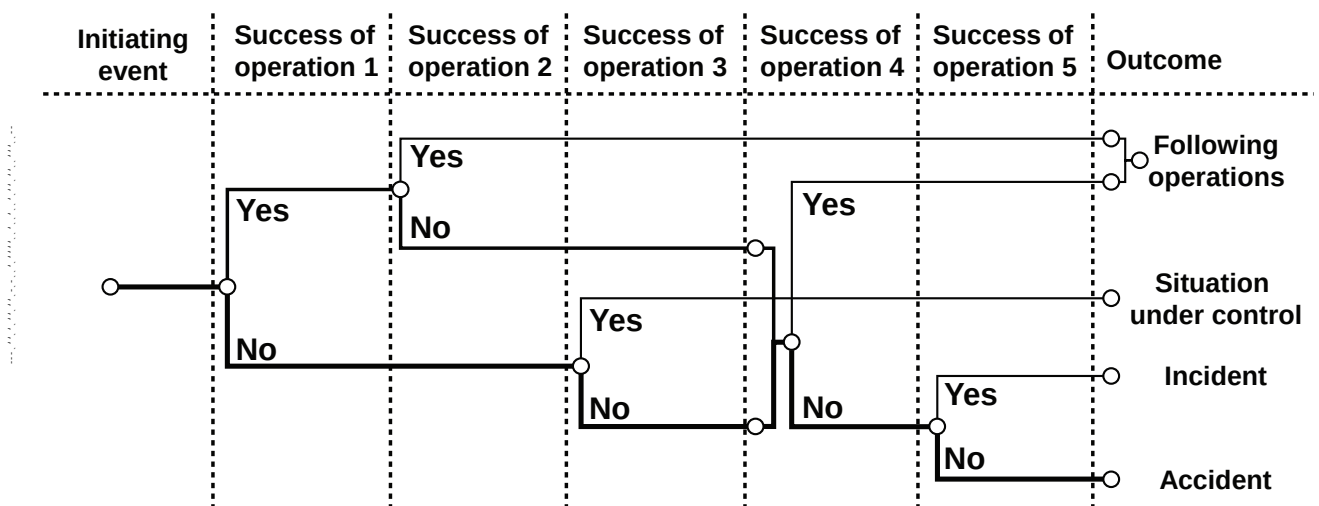


Figure J.3 — Equivalent event tree

Contrarily to the cause-consequence diagrams the event trees does not really help the analyst to build them: several attempts may be needed before finding the good column order resulting in an event tree as simple as possible. Without the preliminary drawing of the equivalent cause-consequence diagram, the example illustrated in [Figure J.3](#) would not have been so easy to build from scratch!

The last *metamorphosis* of the event tree approach is the LOPA[9] which is now widely used to analyse the multiple safety systems. [Figure J.4](#) shows such a LOPA performed for a set of three protection layers.

Each protection layer (i.e. safety system) reduce the probability (or frequency) of occurrence of the hazardous event as this is illustrated by the arrows decreasing in size.

When, like in [Figure J.4](#), each protection layer is able to prevent alone the occurrence of the hazardous event the corresponding event trees is simpler than in the general case. Nevertheless in oil and gas industries the situation is often different as several partial safety loops works in parallel or in sequence to prevent the hazardous event occurrence. This is more complicated than the example in [Figure J.4](#) and closer to what has been described in [Figure J.3](#).

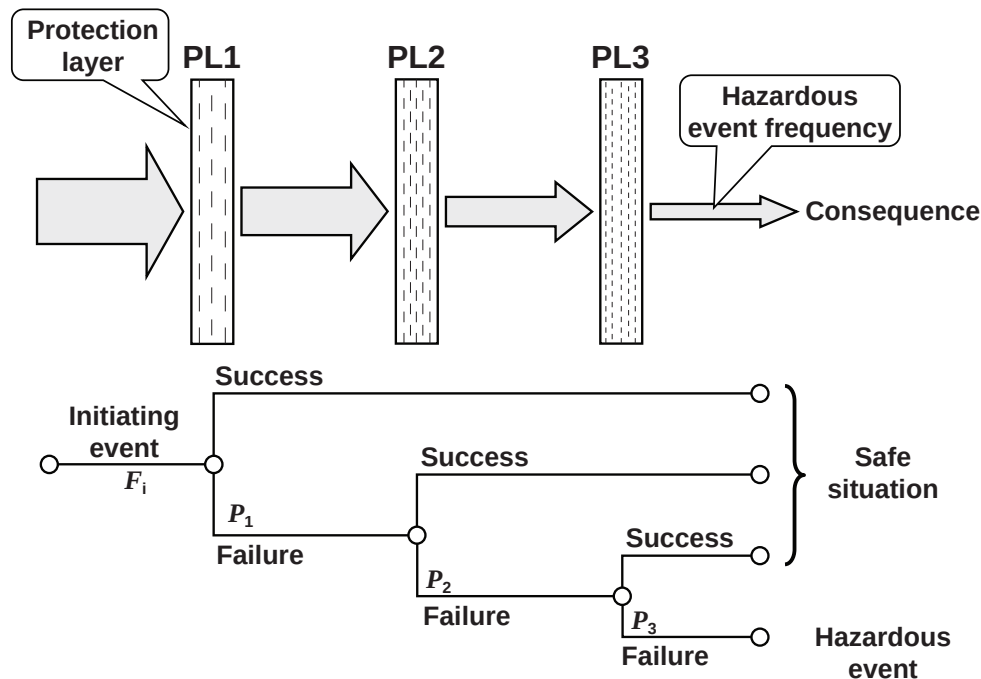


Figure J.4 — Example of LOPA

Simplistic event tree or LOPA approaches suggest calculating the probability of a given sequences just by multiplying the probability of failure of each branch: e.g. $P_1.P_2.P_3$ for the probability of the sequence leading to the hazardous event in [Figure J.4](#) and $F_i.P_1.P_2.P_3$ for the frequency of the hazardous event due to this sequence. This approach assumes that all of the branches are independent of each other therefore when there are heavy dependencies between the multiple systems (e.g. common cause failures, shared components or systemic effects due to periodic tests), the quantification have to be performed using more detailed analysis. For example, in [Figure J.2](#) the two sequences leading to the accident are not independent because they share the sub fault trees 4 and 5 and simple calculations are likely to lead to unexpected results. Besides, IEC 62502 Ed. 1[8] encourages performing the calculations through the equivalent fault tree related to the considered sequence. This approach is fully in line with the cause-consequence approach described above.

The fault tree of [Figure J.2](#) can be improved by isolating the common parts between the sequences as this has been done on the left hand side of [Figure J.5](#) where the common parts appear only once. Nevertheless other dependencies may exist between the sub fault trees (e.g. shared components between the systems used to achieve operations 1 and 4). This approach allows making them appearing just by developing the sub fault trees at a relevant level of detail.

As mentioned above, NOT gates have to be used to model the success of operations (or safety systems). This seems to be a “small” detail but, in fact, this completely changes the nature of the fault tree under study. As a matter of fact this yield to non *monotone* logical functions and so-called *incoherent* fault trees (e.g. the system is failed by a further repair or repaired by a further failure).

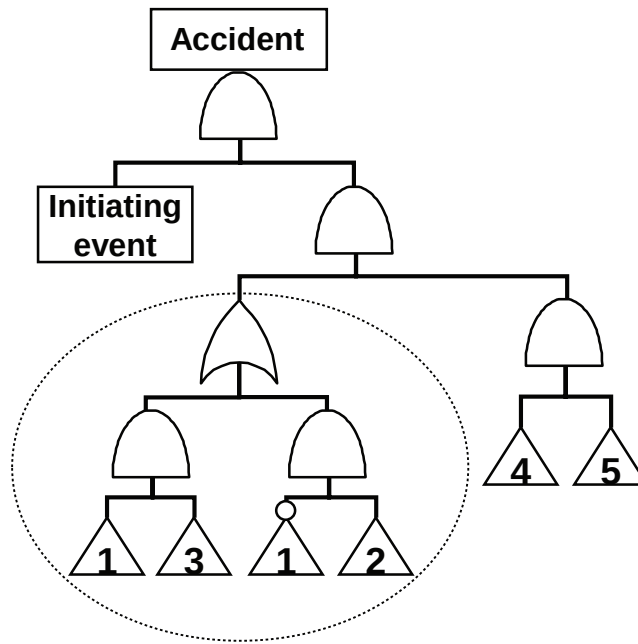


Figure J.5 — Improved fault trees from Figure J.2

The analysis of the small fault tree presented in Figure J.6 is interesting to show the difficulty:

- when A and C are failed and B is functioning, the system S works;
- in this situation, if A is repaired, then the system S fails!

This is an example of non coherent fault tree: the repair of A can leads to the failure of S.

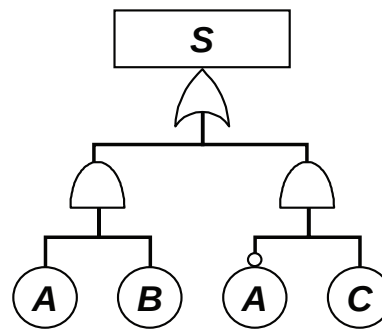


Figure J.6 — Example of non coherent fault tree

The analysis of this fault tree through the ordinary minimal cut set approach would lead to two failure combinations: $A \bullet B$ and $\neg A \bullet C$. This trivial result is not correct as a third combination exists, i.e. $B \bullet C$. This last combination cannot be found by applying the minimal cut sets approach. The *minimal cut set* concept is no longer valid and should be superseded by the *prime implicants* concept.

Nevertheless if $\neg A$ represent the success of a safety item, its probability is very close to 1 and there is $\neg A \approx \Omega$. (Ω being the event which is certain). Then $\neg A \bullet C \approx \Omega \bullet C = C$ and, as $B \bullet C + C = C$, the three combinations above are reduced to only two: $A \bullet B$ and C . Therefore $\neg A$ can be removed from the fault tree which becomes coherent again. This principle can be implemented for the part of the fault tree presented in Figure J.5 surrounded by the dotted ellipse if $P(\neg "1")$ is close to 1.

Nevertheless a problem really arises when A and $\neg A$ in Figure J.6 or the events "1" and $\neg "1"$ in Figure J.7 represent events which can both occurs with a high probability (e.g. day and night, operating phases, etc.). In this case the above approximation cannot be used and the non coherency has to be handled

properly. Fortunately the present time fault tree algorithms (e.g. based on the *binary decision diagrams* approach[51]) are able deal with this situation. It is beyond the scope of this Technical Report to go in more mathematical details but the reliability analysts should be aware of this difficulty when handling sequences involving success branches (i.e. mainly for intermediate sequences between the perfect functioning of the various operations and the occurrence of the hazardous event).

Another problem arises with the sequential aspects. The fault trees presented above (Figure J.7) have an AND gate at the top with the initiating event on the left hand side and a sub-tree modelling the failures of the various safety barriers on the right hand side. The hazardous event will occur only if the various safety barriers have failed before the initiating occurs. Therefore this logic gate should be a sequential AND gate. Handling such sequential AND gate which does not belong to the Boolean algebra field would need developments which are beyond the scope of this Technical Report. Therefore the ordinary fault tree calculation may be used only in the case where the initiating event is modelled by a constant probability of failure or a constant frequency of failure.

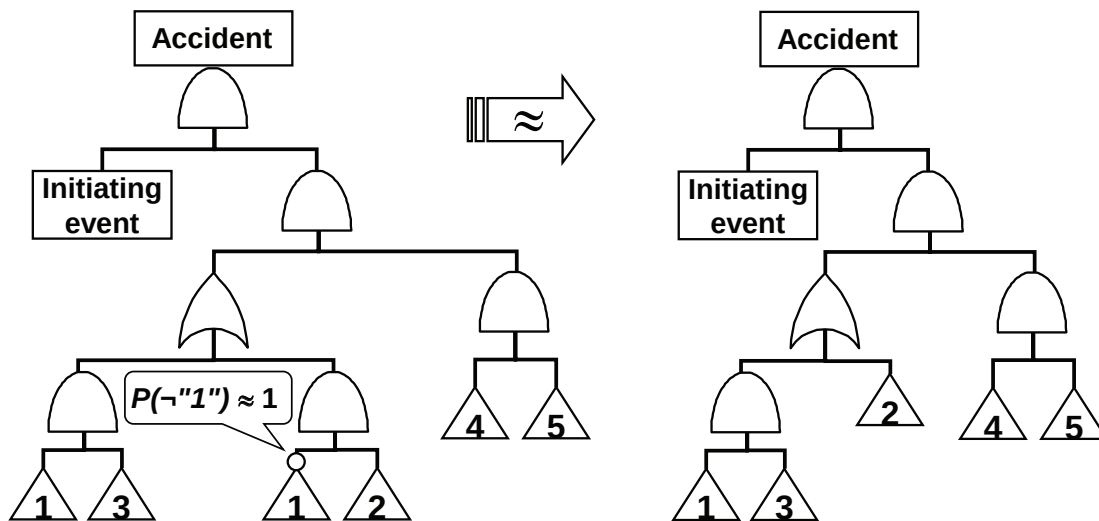


Figure J.7 — Approximation for non coherent fault trees

As conclusion, the cause-consequence diagrams, event trees, LOPA belong to the Boolean models and can be handled through the fault tree approach as described in 8.5 to 8.6. The event trees are more widely used than cause-consequence diagrams. The latter have been somewhat forgotten even though they provide an effective help to build event sequences: the reliability engineers ought to (re)discover them.

Annex K (informative)

Overview of calculations with Boolean models

K.1 Overview of calculations with Boolean models

The simple fault tree (FT) on the right hand side of [Figure K.1](#) is used in order to explain the basic probabilistic calculations which can be performed by using Boolean models. It is equivalent to the RBD represented on the left hand side.

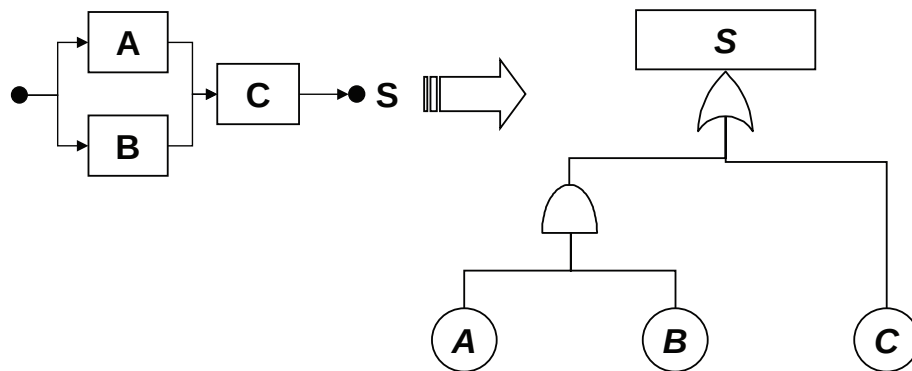


Figure K.1 — Simple fault tree

In the following part the following notations are used:

- X : component X is failed
- $\bar{X}$: component X is not failed

The logical equations represented by these two Boolean models are the following:

- FT : $S = A \bullet B + C$
- RBD : $\bar{S} = (\bar{A} + \bar{B}) \bullet \bar{C}$

The RBD and the FT are equivalent because $\bar{\bar{S}} = S$. The RBD is the “dual” of the FT and reciprocally. This can be demonstrated by applying the Boolean properties:

$$\bar{\bar{S}} = \overline{(\bar{A} + \bar{B}) \bullet \bar{C}} = \overline{(\bar{A} + \bar{B})} + \bar{\bar{C}} = \bar{\bar{A}} \bullet \bar{\bar{B}} + C = A \bullet B + C = S$$

The link between a Boolean equation and probabilistic calculations is made through two simple basic properties:

$$P(X+Y) = P(X)+P(Y)-P(X \bullet Y) \text{ where } P(.) \text{ is the probability function}$$

$$P(X \bullet Y) = P(X).P(Y) \text{ when } X \text{ and } Y \text{ are independent (i.e. behaves independently from each others)}$$

The last hypothesis is very important as the components are never completely independent. It is the work of the analyst to determine if the dependencies are reasonably low to be neglected or not.

Considering that A , B , and C are reasonably independent gives:

$$P(S) = P(A \bullet B + C) = P(A \bullet B) + P(C) - P(A \bullet B \bullet C) = P(C) + P(A) \cdot P(B) - P(A) \cdot P(B) \cdot P(C)$$

For more convenience this can be written $P_S = P_C + P_A \cdot P_B - P_A \cdot P_B \cdot P_C$

Probabilistic calculations based on Boolean equations are basically time independent. They are only related to the failure architecture of the system and not at all to the possible time-dependent behaviours of its components.

Nevertheless if the components behave independently from each others the above formula can still be used: $P_S(t) = P_C(t) + P_A(t) \cdot P_B(t) - P_A(t) \cdot P_B(t) \cdot P_C(t)$.

This formula is valid only if $P_A(t)$ does not depends of events arising on B and C over $[0, t]$, that $P_B(t)$ does not depends of events arising on A and C over $[0, t]$ and that $P_C(t)$ does not depends of events arising on A and B over $[0, t]$.

The above formula is a particular case of the Sylvester - Poincaré formula (also known as the *inclusion-exclusion* principle^[22]) established for two minimal cut sets. When n minimal cut sets have been

identified the logical function can be written as $S = \bigcup_{i=1}^n M_{cs_i}$ and the full Poincaré formula is given by:

$$P_S = P(\bigcup_{i=1}^n M_{cs_i}) = \sum_{i=1}^n P(M_{cs_i}) - \sum_{i < j} P(M_{cs_i} \cdot M_{cs_j}) + \sum_{i < j < k} P(M_{cs_i} \cdot M_{cs_j} \cdot M_{cs_k}) - \dots$$

The number of minimal cut sets increases exponentially with n and there is a combinatory explosion of the number of terms of the Poincaré formula. Therefore, most of the time, only the 1st term is calculated:

$$P_S < \sum_{i=1}^n P(M_{cs_i})$$

It provides a good conservative estimation when the probabilities are low:

$$P_S \approx \sum_{i=1}^n P(M_{cs_i}) \text{ if } P_S < 0.1$$

This is normally the case for safety systems and the approximation may be used by hand when the number of minimal cut sets is not too high. Nevertheless very effective algorithms (e.g. Binary Decision Diagrams^[51]) are now available which allow fast and exact calculations on industrial size systems.

Theoretically, such calculations are valid only for the unavailability $U_S(t)$, provided that there is as many repair teams as components. Unreliability calculations and periodic tests introduce subtle dependencies between the components. Then the formula cannot be used straightforwardly and precautions are needed to avoid obtaining erroneous results. This is explained below.

K.2 Unavailability calculations

The principle of unavailability calculation is shown on [Figure 32](#). A , B and C being independent, the component unavailabilities $U_A(t_1)$, $U_B(t_1)$ and $U_C(t_1)$ can be calculated at time t_1 independently from each others. Then they can be combined through the fault tree according to the rules described above to calculate $U_S(t_1)$. Doing that for a relevant number of time values over a period of interest gives the whole curve $U_S(t)$ over this period of time.

Calculating the average unavailability $\bar{U}_S(T)$ over the period $[0, T]$ of the system modelled in [Figure K.2](#) is not straightforward because it cannot be obtained by combining the average component unavailabilities by the same formula as this used for actual probabilities. This is not valid in the general case: $\bar{U}_S(T) \neq \bar{U}_C(T) + \bar{U}_A(T) \cdot \bar{U}_B(T) - \bar{U}_A(T) \bar{U}_B(T) \bar{U}_C(T)$

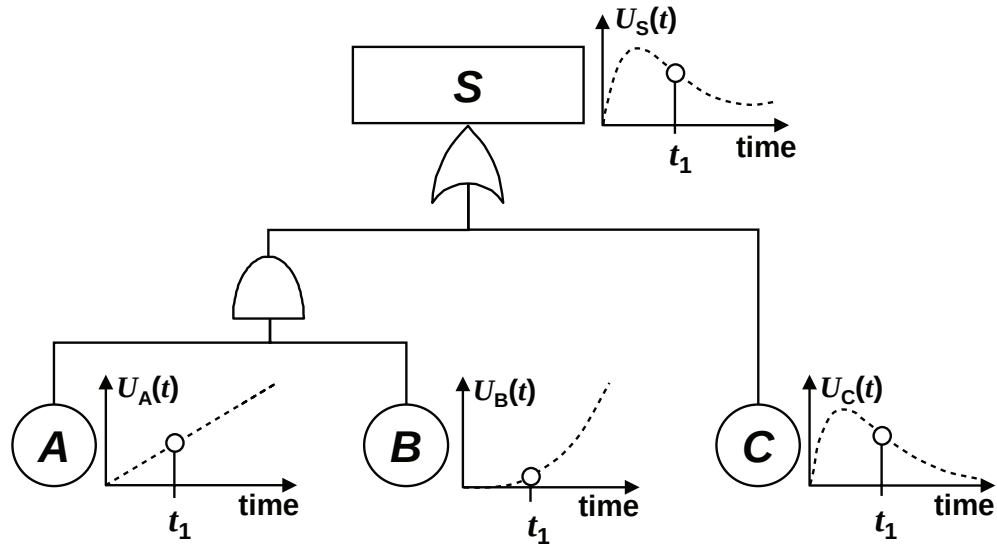


Figure K.2 — Unavailability calculation

Fortunately the calculation of the average unavailability is possible just by calculating the average value of the curve $U_S(t)$:

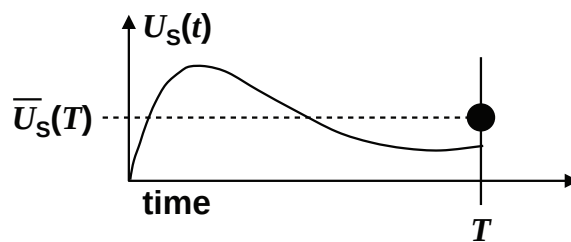


Figure K.3 — Average unavailability as average of the instantaneous unavailability curve

This case is encountered when periodic tests are performed to detect and repair the dangerous undetected failures. An example is shown in [Figure K.4](#).

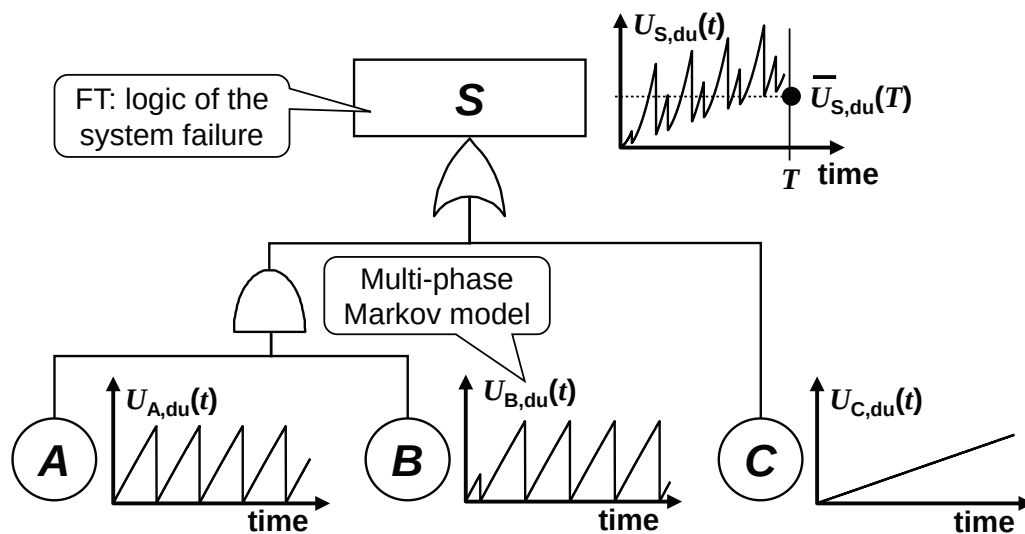


Figure K.4 — FT driven Markov process modelling a safety system with periodically tested components

This is an example of fault tree driven Markov process where a “mixed” approach is undertaken: the logics is described by the fault tree and the component behaviour is modelled either by analytical formula, $U_C(t)$, or by a Markovian multi-phase models, $U_A(t)$ or $U_B(t)$, (see [Clause 9](#)).

The tests of A and B have been staggered and component C is not repairable in order to show the impact of these characteristics on the overall unavailability. Then $U_S(t)$ never go back to zero and $\bar{U}_S(T)$ increases when T increases. Therefore, the average unavailability is not a good indicator of the safety provided by this system. The resulting system unavailability is a characteristic saw tooth curve which shows that the probability of failure:

- is higher before tests than after tests;
- becomes permanently higher than the average after some times (see the right hand side of $U_S(t)$ on [Figure K.4](#)).

The acceptability for the operators of those high risk values is questioning and this should be considered when designing the safety system (see [D.3](#))

Even for a very simple safety system with only three components, the overall behaviour can be rather complicated and the average unavailability difficult to be evaluated just by using analytical formulae. Therefore, the use of systemic models is encouraged even to handle small safety systems.

INFO BOX: FT calculations based on average unavailabilities of periodically tested components (e.g. with PFD_{avg} of the form $\lambda\tau / 2$) leads to non conservative results. The non conservativeness increases when the redundancy increases. That means, for example, that the safety system is likely to be qualified for a higher SIL than it is really. Some available software packages perform these kinds of calculations. This is a part of the duty of the analyst to be cautious and to verify that software he/she uses implement sound algorithms.

The only case where components average unavailabilities can be used to perform relevant calculations is illustrated in [Figure K.5](#): this is when these average unavailabilities are the asymptotic values of the unavailabilities of the components. This is obviously not the case when periodic tests are considered.

In [Figure K.5](#) the unavailabilities $U_A(t)$, $U_B(t)$ and $U_C(t)$ converge toward asymptotic values when t becomes large enough.

Then it is possible to write: $\bar{U}_S(T) \approx \lim_{t \rightarrow \infty} U_S(t) \approx U_S^{as} = U_C^{as} + U_A^{as} U_B^{as} - U_A^{as} U_B^{as} U_C^{as}$

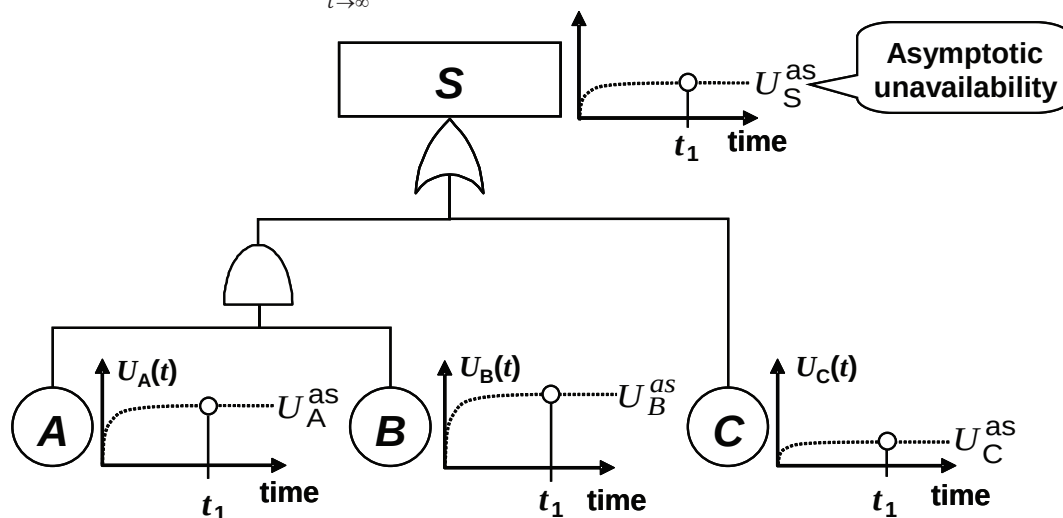


Figure K.5 — Average unavailability as asymptotic value of the instantaneous unavailability

This behaviour is typical of systems with completely and quickly repairable failures, i.e. system of which the failures are quickly detected and repaired. When these hypotheses are met, the asymptotic values are reached rather quickly after a duration equal 2 or 3 times time the larger MTTRes of the components. This approximation is conservative and only one calculation of $U_S(t)$ is needed. It works well provided that t is large enough to ensure that all the asymptotic values have been reached.

INFO BOX: This is, in principle, the case of the dangerous detected failures as well as immediately revealed safe failures. This is of the duty of the analyst to verify that the hypotheses related to this particular case are met and that the calculations can be performed in this way.

This is obviously not the case of dangerous undetected failures where the asymptotic unavailability values do not exist.

K.3 Unreliability and frequency calculations

Evaluating the reliability with Boolean models is not really easy when dealing with repairable systems. That conceals a subtle difficulty that analysts are often not aware about and which is illustrated in [Figure K.6](#). In this Figure, A and B are repairable components.

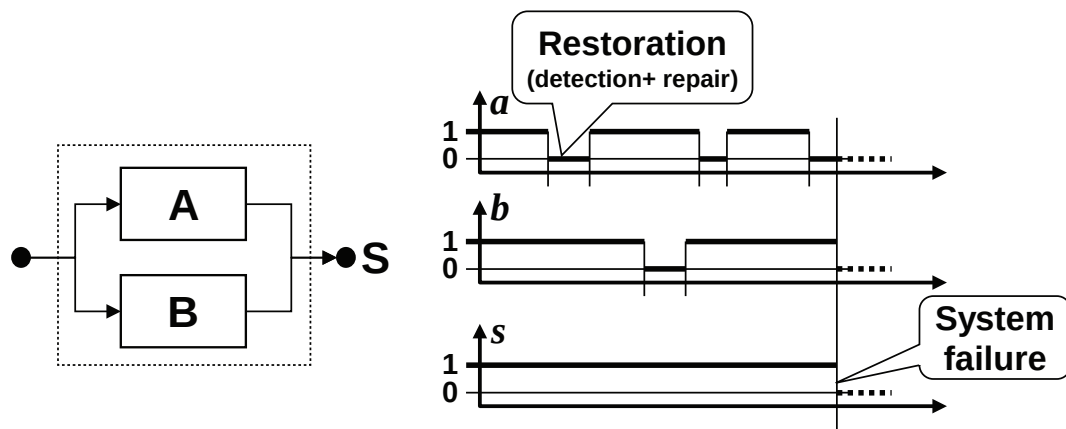


Figure K.6 — Links between component and overall system failures

The whole system fails when both A and B are failed at the same time: B fails when A is under restoration (case illustrated in [Figure K.6](#)) or A fails when B is under restoration. As reliability $R_S(t)$ of the system S is the probability that it does not fail before t , from mathematical point of view its repair is excluded from the calculations. Therefore the repair of A (respectively B) which would also repair S after it has failed is also excluded. Then the failures of A and B can be repaired as long as they don't occur at the same time. When this occurs, the repairs are excluded of the reliability calculation because the repair of one of them would repair the whole system.

To summarize, a failure of A (respectively B) is repairable only if it does not fail the overall system. In [Figure K.6](#), the failure of B fails the system because A was already failed. From this point, A and B are no longer repairable.

Therefore, when the components are repairable, the reliability (as well as the unreliability) calculations induce dependencies between all the components. This is a systemic dependency which prohibits calculating the system unreliability simply by combining the individual unreliabilities of the components.

Several approximations (and the related algorithms) are available for calculating the unreliability from a fault-tree, e.g. Murchland and Vesely asymptotic approximations[23].

The Murchland approximation is based on the fact that the expected number of failure $N_S(t)$ over $[0, t]$ provides a conservative approximation of the probability to have one or more failures i.e. an approximation

of the unreliability $F_S(t)$. This approximation is pretty good when $N_S(t)$ is small compared to 1. This is generally verified when dealing with safety problems. Then $F_S(t) \approx N_S(t) = \bar{w}_S(t) \cdot t$ where $w_S(t)$ is the average unconditional failure intensity which is also the average failure frequency of the system over $[0, t]$.

Coming back to the definition of $\bar{w}_S(t) = \frac{\int_0^t w_S(\delta) d\delta}{t}$ allows to calculate the unreliability from the instantaneous failure frequency $w_S(t)$: $F_S(t) \approx \int_0^t w_S(\delta) d\delta$. Finally the reliability is given by $R_S(t) \approx 1 - \int_0^t w_S(\delta) d\delta$.

The Vesely approximation is based on the idea to replace the true failure rate Λ_S of the system by the conditional failure intensity Λ_V also called “Vesely failure rate”. Then $R_S(t) = e^{-\int_0^t \Lambda_S(\delta) d\delta} \approx e^{-\int_0^t \Lambda_V(\delta) d\delta}$.

This can also be written: $R_S(t) \approx e^{-\int_0^t \Lambda_V(\delta) d\delta} = e^{-\bar{\Lambda}_V(t) \cdot t}$. When $\bar{\Lambda}_V(t) \cdot t \ll 1$, then $R_S(t) \approx \bar{\Lambda}_V(t) \cdot t$ is obtained. In this case the availability $A_S(t)$ is close to 1 and $\Lambda_V(t) = \frac{w_S(t)}{A_S(t)} \approx w_S(t)$. Therefore $\bar{w}_S(t) \approx \bar{\Lambda}_V(t)$ and it is found that $F_S(t) = 1 - R_S(t) \approx 1 - \bar{w}_S(t) \cdot t$ which is again the Murchland approximation.

In the case of failures quickly detected and repaired Λ_V reaches quickly an asymptotic value Λ_V^{as} and this lead to the “Vesely asymptotic” approximation of the reliability: $R(t) \approx e^{-\Lambda_V^{as} \cdot t}$. This approximation is good for large values of t (i.e. when the Vesely failure rate has converged toward its asymptotic value)

Other approximation exists (e.g. Barlow-Proschan). Among all available approximations, the Vesely approximation seems to be the best one.

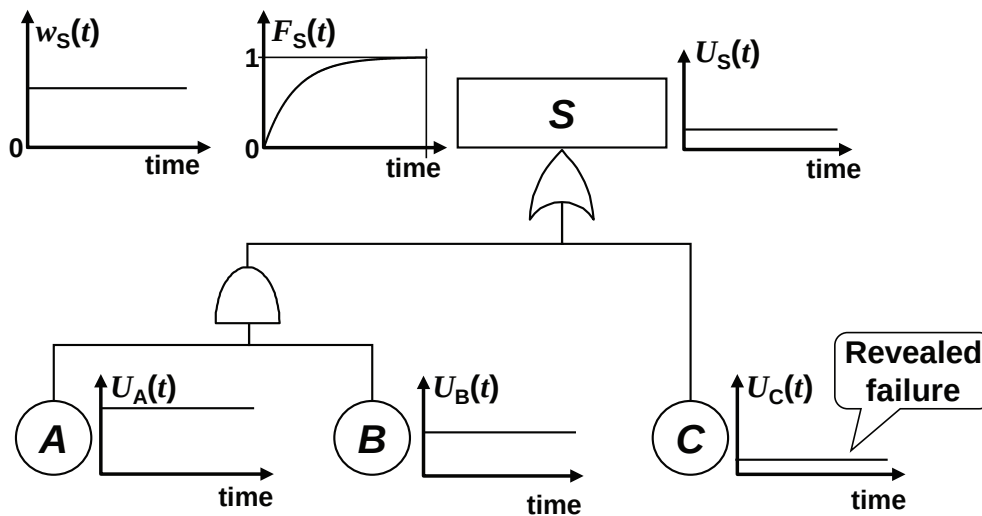


Figure K.7 — Unreliability calculations (revealed and quickly repaired failures)

Beyond its usefulness for the Murchland approximation, the failure frequency $w_S(t)$ is needed to evaluate the PFH (i.e. the average failure frequency as per IEC 61508[2]) of the safety system. The calculation of $w_S(t)$ implies the calculation of the “Birnbaum importance factor” which itself is obtained from the combination of two different conditional probabilities. This is rather difficult by hand but easy by using a fault tree package implementing the relevant algorithms.

As mentioned above (see Figure K.5) the system unavailability converges toward an asymptotic value when the failures are quickly revealed (or detected) and quickly repaired. This is the same for the failure rate $\Lambda_S(t)$, the Vesely failure rate $\Lambda_V(t)$ and the unconditional failure intensity $w_S(t)$ of the safety system converges quickly toward an asymptotic constant value. Figure K.7 shows the evolution of $w_S(t)$, $U_S(t)$ and $F_S(t)$ of such a system.

In this case the availability of the safety system always remains close to 1 and the numerical values of $\Lambda_S, \Lambda_S^{as}, w_S^{as}$ are close to each others: $\Lambda_S \approx \Lambda_S^{as} = \frac{w_S^{as}}{\Lambda_S^{as}} \approx w_S^{as}$. This is why they are commonly mixed-up. 0 provides clarification about the differences between these parameters.

The safety system unreliability $F_S(t)$ presented in [Figure K.7](#) is calculated by using the Vesely approximation: $F_S(t) \approx 1 - e^{-\Lambda_S \cdot t}$

[Figure K.8](#) is similar to [Figure K.7](#) but for periodically tested failures (e.g. dangerous undetected failures). The shapes of the unavailability $U_S(t)$ and of $w_S(t)$ are very different than in the case just above ([Figure K.7](#)). In return, as $F_S(t)$ is a non-decreasing function in any cases its shape is similar in [Figure K.8](#) and [Figure K.7](#).

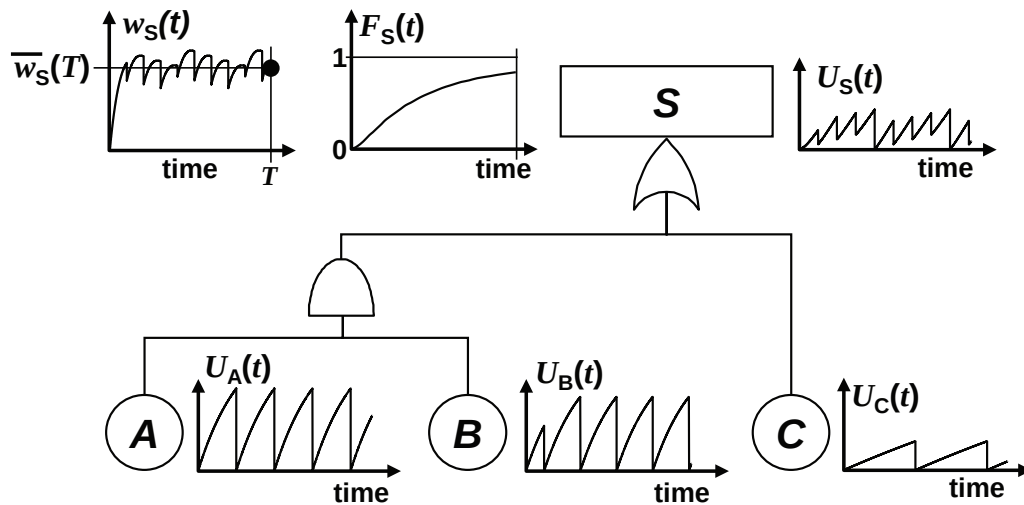


Figure K.8 — Unreliability calculations (periodically tested failures)

In the case of periodically tested components, the equivalent failure rate of the safety system is no longer constant but it can be approximated by the Vesely failure rate $\Lambda_S(t) \approx \Lambda_V(t) = \frac{w_S(t)}{\Lambda_S(t)}$ and $F_S(t)$

calculated by the Vesely approximation $F_S(t) \approx 1 - e^{-\int_0^t \Lambda_V(\delta) d\delta} = 1 - e^{-\bar{\Lambda}_V(t) \cdot t}$

When $\bar{\Lambda}_V(t) \cdot t \ll 1$, the Murchland approximation: $F_S(t) \approx 1 - \bar{w}(t) \cdot t$ can also be used.

The average unconditional failure intensity $\bar{w}_S(t)$ is also the average failure frequency, i.e. the safety system PFH according to international standards dealing with functional safety (e.g. IEC 61508[2] or IEC 61511[3]).

Safe or dangerous immediately revealed failures can be modelled like in [Figure K.7](#) when safe or dangerous undetected failures can be modelled like in [Figure K.8](#). Mixing the two cases is easy when using fault trees. This allows to model complex safety system and provides all probabilistic parameters needed: unavailability, average unavailability, unreliability and failure frequency.

K.4 Hazardous event probability calculation

The probability and the average frequency of hazardous event (e.g. accidents) related to safety systems working in continuous mode are simply the unreliability and the average failure frequency of the related safety system. Therefore it can be calculated just by applying the results presented above in [K.3](#) and illustrated by [Figure K.7](#).

For the safety systems working in demand mode it is more complicated but the approximations presented above can be used provided that MTTD (mean time to demand = $1/\lambda_d$) is large compared to the mean

time to restoration (MTTResS) of the whole safety system. This is the duty of the analyst to verify that this condition applies. In this case the expected number of accident (or of a demand on the succeeding protection layer) can be calculated as $N_a(T) \approx \bar{U}_S \cdot \lambda_d T$ and $\bar{\Phi}_a(T) \approx \bar{U}_S \cdot \lambda_d$.

The probability of accident (or of a demand on the succeeding protection layer) can also be approximated as $P_a(T) \approx N_a(T) \approx \bar{U}_S \cdot \lambda_d T$.

When the mean restoration time is not negligible, it is essential to estimate the MTTResS in order to apply the formulae proposed in [D.1](#) and to consider the coefficient $\zeta^{(2)}$ in the calculation.

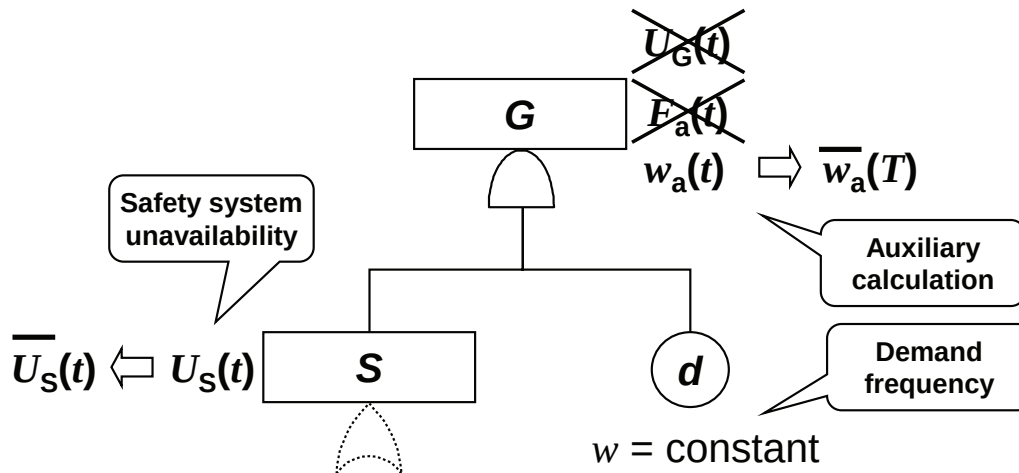


Figure K.9 — Hazardous event probability fault tree modelling

Another possibility is illustrated in [Figure K.9](#) where the safety system failure has been combined with the demand frequency. This approach needs to input an event which is described by its constant frequency rather than its probability of occurrence in the fault tree. Such a calculation provides the average accident frequency $\bar{w}_a(t)$ over $[0, t]$.

The probability of hazardous event $P_a(T)$ over $[0, T]$ could be obtained by integrating $w(t)$ between 0 and T . When $\bar{w}_a(T) \cdot T \ll 1$, then, as usual, $P_a(T) \approx \bar{w}_a(T) \cdot T$.

NOTE A hazardous event occurs only if the safety system is unavailable when the demand occurs. This happens only if the safety system has failed before that the demand occurs. If the safety system fails after that a demand has occurred this does not matter. The model presented in [Figure K.9](#) does not take into account the order of these events. It works only because $w(t) = w = \text{constant}$ and that, in this case, the order does not matter. When $w(t)$ is not constant (e.g. if it is modelled by a sub fault tree) the order of these events should be considered and a sequential AND gate be used. The calculations with such sequential gates are outside of the scope of this Technical Report and the unavailability $U_a(t)$ and the unreliability $F_a(t)$ obtained for the top event of the fault tree in [Figure K.9](#) are meaningless.

Annex L (informative)

Markovian approach

L.1 Introduction and principles

The Markovian approach has been developed at the beginning of the past century and is a popular approach for probabilistic calculations in numerous scientific fields. Its use in reliability engineering is described in IEC 61165[6]. This approach has already been mentioned many times in this Technical Report as illustration in [Clauses 3, 5 and 8](#) as possible fault tree input. It is also the underlying model behind the formulae developed in [Clause 7](#). Further development can be found in reference.[24]

It is a state-transition model and its principle is to identify the various states of the system under interest and to analyse how the system jumps from a state to another state. This allows building graphical models - called Markov graphs - which represent the dynamic behaviour of the system which is modelled.

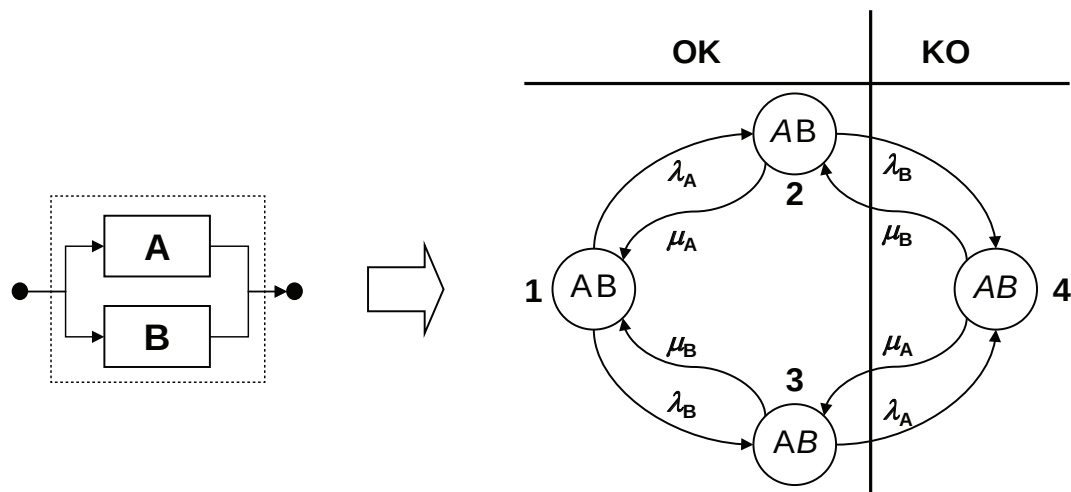


Figure L.1 — Example of simple Markov model

[Figure 36](#) gives an example of simple Markov graph. The states are represented by circles and the jumps (transitions) by arrows.

For example, when the system is in state 1 it may jump to the state 2 with the transition rate λ_A (i.e. the failure rate of component A) or to state 3 with the transition rate λ_B (i.e. the failure rate of component B). From state 2, it may jump to state 4 with the transition rate λ_B or to state 1 with the transition rate μ_A (i.e. the restoration rate of component A). It may jump directly from state 1 to state 4 with the transition rate λ_{CCF} (i.e. a common cause failure), etc.

More generally, the transition between a state i to a state j is characterized by the transition rates $\eta_{ij}(t)$.

$\eta_{ij}(t).dt$ is the conditional probability to jump from i to j within $[t, t+dt]$ provided that the system was in state i at time t . Therefore, failure and restoration or repair rates are only particular cases of transition rates.

The set of the states and transition describe a stochastic process which, when all the $\eta_{ij}(t)$ are constant (i.e. $\eta_{ij}(t) \equiv \eta_{ij}$) is a so-called "homogeneous Markov process". When the transition rates are time dependent, the stochastic process is called "semi-Markov process". Semi-Markov processes are difficult to handle analytically but they can be modelled by using Petri nets (see [Clause 10](#)) and calculated by Monte Carlo simulation (see [Clause 11](#)). Only the homogeneous Markov processes (i.e. the ordinary

Markov processes) which can be analytically handled are considered in the remaining part of this annex. Then only the constant failure rates (i.e. catalectic failures) and constant restoration or repair rates of individual components are modelled in this annex.

[Figure L.1](#) is similar to the simple Markov graph already presented in [Figure I.9](#). It models the small redundant system on the right hand side made of two components A and B which have repairable and immediately revealed failures. Contrary to [Figure I.9](#), only the safety system behaviour has been represented, the demand will be introduced later. The case of periodically tested failure is also described later.

The states have been split into two classes: OK where the system is available to satisfy to its safety function and KO where it is unavailable.

L.2 Probability of the states

The basic idea is to establish the probabilities of the various states at time $t+dt$, from the probabilities of these states at time t .

For example the probability $P_2(t+dt)$ to be in state 2 at time $t+dt$ is equal of the sum of the following probabilities:

- Probability of a jump from 1 to 2: $\lambda_A.P_1(t).dt$
- Probability of a jump from 4 to 2: $\mu_B.P_4(t).dt$
- Probability of no jump outside 2: $[1-(\mu_A+\lambda_B)dt].P_2(t)$

Finally: $P_2(t+dt) = \lambda_A.P_1(t).dt + [1-(\mu_A+\lambda_B)dt].P_2(t) + 0.P_3(t) + \mu_B.P_4(t).dt$

Similarly a set of four equations is obtained:

- $P_1(t+dt) = [1-(\lambda_A+\lambda_B)dt].P_1(t).dt + \mu_A.P_2(t).dt + \mu_B.P_3(t) + 0.P_4(t).dt$
- $P_2(t+dt) = \lambda_A.P_1(t).dt + [1-(\mu_A+\lambda_B)dt].P_2(t) + 0.P_3(t) + \mu_B.P_4(t).dt$
- $P_3(t+dt) = \lambda_B.P_1(t).dt + 0.P_2(t) + [1-(\mu_B+\lambda_A)dt].P_3(t) + \mu_A.P_4(t).dt$
- $P_4(t+dt) = 0.P_1(t).dt + \lambda_B.P_2(t).dt + \lambda_A.P_3(t).dt + [1-(\mu_B+\mu_A)dt].P_4(t).dt$

The above set of equation can be transformed into the following one:

- $[P_1(t+dt)-P_1(t)]/dt = -(\lambda_A+\lambda_B).P_1(t) + \mu_A.P_2(t) + \mu_B.P_3(t) + 0.P_4(t)$
- $[P_2(t+dt)-P_2(t)]/dt = \lambda_A.P_1(t) - (\mu_A+\lambda_B).P_2(t) + 0.P_3(t) + \mu_B.P_4(t)$
- $[P_3(t+dt)-P_3(t)]/dt = \lambda_B.P_1(t) + 0.P_2(t) + (\mu_B+\lambda_A).P_3(t) + \lambda_A.P_4(t)$
- $[P_4(t+dt)-P_4(t)]/dt = 0.P_1(t) + \lambda_B.P_2(t) + \lambda_A.P_3(t) - (\mu_B+\mu_A).P_4(t)$

Finally:

- $dP_1(t)/dt = -(\lambda_A+\lambda_B).P_1(t) + \mu_A.P_2(t) + \mu_B.P_3(t) + 0.P_4(t)$
- $dP_2(t)/dt = \lambda_A.P_1(t) - (\mu_A+\lambda_B).P_2(t) + 0.P_3(t) + \mu_B.P_4(t)$
- $dP_3(t)/dt = \lambda_B.P_1(t) + 0.P_2(t) + (\mu_B+\lambda_A).P_3(t) + \lambda_A.P_4(t)$
- $dP_4(t)/dt = 0.P_1(t) + \lambda_B.P_2(t) + \lambda_A.P_3(t) - (\mu_B+\mu_A).P_4(t)$

This is a set of homogeneous linear differential equations which is very classical in mathematics. It may be solved analytically and powerful algorithms are available to do that. Nevertheless this set of equation is not "free" (i.e. any equation can be established from the others) and an extra equation is needed to solve it. This is simply: $P_1(t) + P_2(t) + P_3(t) + P_4(t) = 1$.

The system of equation can be represented in a conventional matrix form (cf. [Figure L.2](#)) where the vector of the derivatives of the state probabilities appears to be the product of a square matrix multiplied by the vector of the same state probabilities.

The square matrix has as many columns and lines as the number of states of the modelled system and the sum of the coefficients is equal to zero for each lines. This is the characteristic of a so-called Markovian matrix.

$$\begin{bmatrix} \frac{d}{dt} P_1(t) \\ \frac{d}{dt} P_2(t) \\ \frac{d}{dt} P_3(t) \\ \frac{d}{dt} P_4(t) \end{bmatrix} = \begin{bmatrix} (\lambda_A + \lambda_B) & \mu_A & \mu_B & 0 \\ \lambda_A & (\mu_A + \lambda_B) & 0 & \mu_B \\ \lambda_B & 0 & (\mu_B + \lambda_A) & \lambda_A \\ 0 & \lambda_B & \lambda_A & (\mu_B + \mu_A) \end{bmatrix} \times \begin{bmatrix} P_1(t) \\ P_2(t) \\ P_3(t) \\ P_4(t) \end{bmatrix}$$

Figure L.2 — Typical Markovian equation in a matrix form

The particular case described just above can be easily generalized. As previously said the probabilities $P_i(t)$ of state i at time t and from the transition rates $\eta_{i,k}$ can be split into two parts:

- 1) jump toward from another state k to state i : $\sum_{k \neq i} \eta_{k,i} P_k(t) dt$
- 2) stay in i : $[1 - (\sum_{k \neq i} \eta_{i,k}) dt] P_i(t) = [1 - \eta_{i,i} dt] P_i(t)$

And finally: $P_i(t) = [1 - \eta_{i,i} dt] P_i(t) + \sum_{k \neq i} \eta_{k,i} P_k(t) dt$

By simple manipulation the derivative form is obtained: $\frac{dP_i(t)}{dt} = -\eta_{i,i} \frac{dP_i(t)}{dt} + \sum_{k \neq i} \eta_{k,i} \frac{dP_k(t)}{dt}$

As previously shown, the Markov process is mathematically represented by a set of linear differential equations with constant coefficients which may be represented in a matrix form:

$$\frac{d\bar{P}(t)}{dt} = [M] \bar{P}(t)$$

where $[M]$ is the Markovian matrix of the Markov process and $\bar{P}$ the vector of the states probabilities.

This equation is similar to the well known the derivative of a scalar exponential. This can be extended to the matrix form and this leads to the basic Markovian formula:

$$\bar{P}(t) = e^{-t[M]} \bar{P}(0)$$

where $\bar{P}(0)$ is the vector of the initial conditions (e.g. 1 for the probability of the perfect state and 0 for the others).

Nevertheless, the exponential of the matrix has not all the good properties of the ordinary exponential: it is not commutative) but keeps the fundamental property:

$$\bar{P}(t) = e^{-t[M]} \bar{P}(0) = e^{-(t-\theta)[M]} e^{-\theta[M]} = e^{-(t-\theta)[M]} \bar{P}(\theta)$$

Therefore if the states probabilities are calculated for a given time θ , they can be used as new initial conditions for further calculations: the homogeneous Markov process has no memory about how the states have been reached and all the past of the process is summarized by the current state at time t .

This is a very important property which simplifies the calculations and allows handling in a simple way the multi phases Markovian models analysed here after in [L.5](#).

L.3 Sojourn time and probabilistic parameters

When the system enters in the state i , it sojourns there, in average, during $\theta_i = \frac{1}{\eta_{i,i}} = \frac{1}{\sum_{k \neq i} \eta_{i,k}}$. When the

failure or repair rate are constant this lead to the well known equalities: $MTTF = 1/\lambda$ and $MTTRes = 1/\mu$.

If $N_i(T)$ is the number of visits of state i over $[0, T]$, then the cumulated sojourn time $\Theta_i(T)$ in state i over this period is equal to $\theta_i N_i(T)$.

The cumulated sojourn times $\bar{\Theta}(T)$ spent in the various states of the system over a given period $[0, T]$ are easy to calculate by a simple integration of: $\bar{\Theta}(T) = \int_0^T \bar{P}(t) dt$

Therefore the number of visit of a given state i is given by $N_i(T) = \frac{\Theta_i(T)}{\theta_i} = \Theta_i(T) \eta_{i,i}$ and the average frequency of the visits of state i is given by $\bar{\phi}_i(T) = \frac{\Theta_i(T) \eta_{i,i}}{T}$

Then the cumulated sojourn time is useful to compute the frequency of the visits in the various states. It is also useful to calculate the average probabilistic parameters analysed later in this annex.

[Figure L.3](#) illustrates two kinds of Markov graphs of utmost importance for unavailability and reliability calculations:

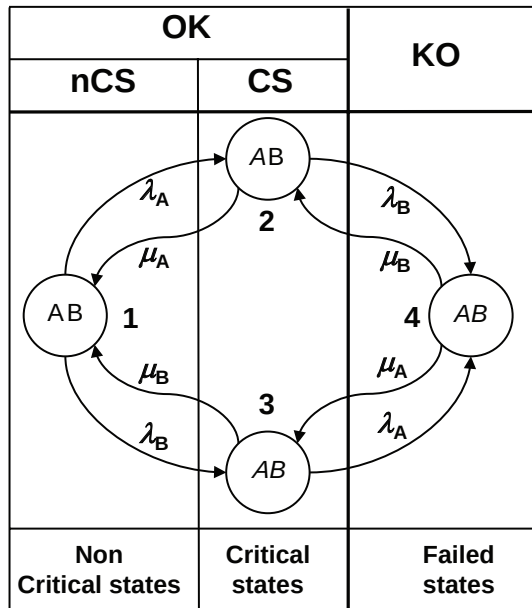
- On the left hand side, it is possible to come back to the class OK after being entered in the class KO. This state can be visited several times: $N_4(t) = \Theta_4(t) \cdot (\mu_A + \mu_B)$. Then $P_4(t) = U_S(t)$ is the system unavailability and $P_1(t) + P_2(t) + P_3(t) = A_S(t)$ is the system availability: this is an availability graph.

More generally with an availability graph $U_S(t) = \sum_{i \in KO} P_i(t)$ and $A_S(t) = \sum_{i \in OK} P_i(t)$ are obtained.

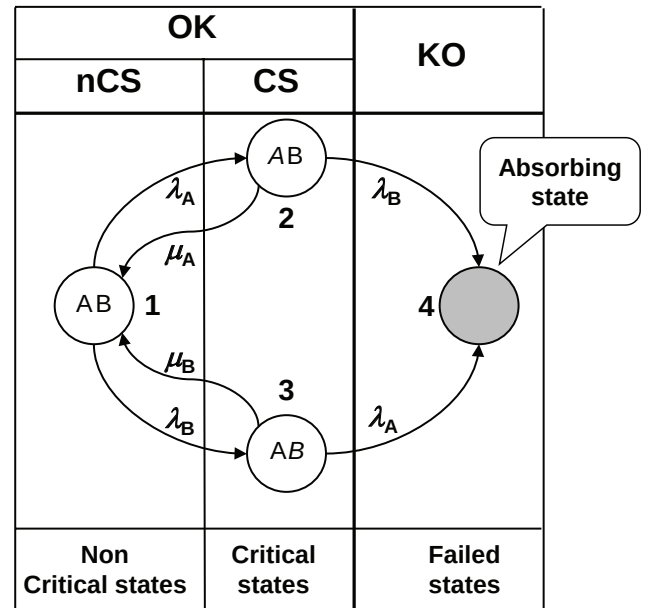
- On the right hand side it is no longer possible to come back to the class OK after being entered in the class KO. $\theta_4 = \infty$ and this state can be visited only once: this is a so-called “absorbing” state. Then $P_4(t) = F_S(t)$ is the system unreliability) and $P_1(t) + P_2(t) + P_3(t) = R_S(t)$ is the system reliability: this is a reliability graph.

More generally with a reliability graph $F_S(t) = \sum_{i \in KO} P_i(t)$ and $R_S(t) = \sum_{i \in OK} P_i(t)$ are obtained.

Transforming an availability graph into a reliability graph is very simple: just remove the transitions from KO to OK and gather the KO states within a single absorbing state.



Markov graph for availability calculation



Markov graph for reliability calculation

Figure L.3 — The three classes of states

Another important topic illustrated in [Figure L.3](#) is that the class OK may be divided into two sub-classes: “non critical” (nCS) and “critical” (CS) sub-classes. For both Markov graphs in [Figure L.3](#) the state 1 is “non critical”, states 2 and 3 are “critical” and state 4 is “failed”.

Any Markov graph comprises a “critical” class but the “non critical class does not necessarily exists. The critical class is very important because it gathers the states which are just at the border between OK and KO and are in direct relationship with the overall failure rate or failure intensities of the system under study.

From the graphs in [Figure L.3](#) it is obvious that the system fails at t if:

- it is in a critical state at t ; and
- has a jump toward KO between t and $t+dt$.

Therefore in both cases:

- Transition rate between OK and KO: $w_S(t) = P_2(t)\lambda_B + P_3(t)\lambda_A$.
- Probability to have the jump from OK and KO during $[t, t+dt]$: $[P_2(t)\lambda_B + P_3(t)\lambda_A]dt = w_S(t)dt$

The parameter $w_S(t)$ is the “unconditional failure intensity” (also called “failure frequency”) of the overall system. This parameter has already been discussed in [Annex K](#). It can be obtained straightforwardly by the Markovian approach: $w_S(t) = \sum_{i \in CS} \eta_{i,KO} P_i(t)$

This is a very interesting parameter which leads directly to the evaluation of the average system failure frequency (i.e. the PFH as per IEC 61508[2]): $\bar{\Phi}_S(T) = \frac{\int_0^T w_S(t)dt}{T} = \bar{w}_S(T)$.

For the availability graphs, $w_S(t)dt$ is linked with the “conditional failure intensity” Λ_V (also called Vesely failure rate) which is the conditional probability of failure per unit of time between t and $t+dt$, given the

system is available at t : $\Lambda_V(t) = w_S(t) / A_S(t)$. Therefore $\Lambda_V(t) = \frac{\sum_{i \in CS} \eta_{i,KO} P_i(t)}{\sum_{j \in OK} \eta_{i,KO} P_j(t)}$

$\Lambda_V(t)$ is the approximation of the equivalent failure rate which has been mentioned in the Boolean approach to perform approximated unreliability calculations.

For the reliability graphs $w_S(t)$ is also the failure density function (pdf) $f_S(t)$ of the system: $f_S(t) = w_S(t)dt$. It is linked with the “equivalent failure rate” of the system $\Lambda_S(t)$ as defined in 3.1.16 by the same formula

$$\text{as above: } \Lambda_S(t) = f_S(t)/R_S(t) \text{ . Therefore: } \Lambda_S(t) = \frac{\sum_{i \in CS} \eta_{i,KO} P_i(t)}{\sum_{j \in OK} \eta_{j,KO} P_j(t)}$$

Then $\Lambda_V(t)$ and $\Lambda_S(t)$ are given by the same formula. This is the nature of the Markov graph which decide which of them is actually calculated. It can also be said that, when a reliability Markov graph is considered, $\Lambda_V(t)$ is also the exact $\Lambda_S(t)$.

When all the failures are completely and quickly repairable the Markov process reaches quickly the equilibrium and all the states probabilities reaches asymptotic values. Then $U_S(t)$, $A_S(t)$, $w_S(t)$ and $\Lambda_V(t)$ reaches also asymptotic values. In this case the Vesely failure rate calculated from an availability graph provides a good approximation for the equivalent failure rate: $\Lambda_V^{as} \approx \Lambda_S^{as}$.

The availability $A_S(t)$ can be calculated directly from the probability to be in class OK and the average availability $\bar{A}_S(t)$ from the cumulated sojourn times spent in OK. In the same way, the unavailability $U_S(t)$ can be calculated directly from the probability to be in class KO and the average unavailability $\bar{U}_S(t)$ (e.g. the PFD_{avg}) from the cumulated sojourn time spent in KO:

- Unavailability: $U_S(t) = \sum_{k \in KO} P_k(t)$,
- Average unavailability: $\bar{U}_S(T) = \frac{\sum_{k \in KO} \Theta_k(T)}{T}$,
- Availability: $A_S(t) = \sum_{k \in OK} P_k(t)$,
- Average availability: $\bar{A}_S(T) = \frac{\sum_{k \in OK} \Theta_k(T)}{T}$.

Then over a given period $[0, T]$ the global MTTF_S(T), the MTTR_S(T) and the MTBF_S(T) of the system can be calculated by the following formulae:

- Mean time to fail over $[0, T]$: $MTTF_S(T) \oplus \frac{\bar{A}_S(T)}{w(T)}$,
- Mean time to restore $[0, T]$: $MTTR_S(T) \oplus \frac{\bar{U}_S(T)}{w(T)}$
- Mean time between failure $[0, T]$: $MTBF_S(T) = MTTF_S(T) + MTTR_S(T)$.

When T is long enough the above values generally converge toward MTTF_S, MTTR_S and MTBF_S.

$\bar{P}(t)$, $\bar{\Theta}(T)$, $\Lambda_S(t)$, $w_S(t)$ and $\Lambda_V(t)$ are the key parameters to perform unavailability, average unavailability, unreliability and frequency calculations. Of course those calculations can be performed by hand but this possible only when the number of states is very small (3 or 4). When the number is greater than that, the use of an adequate software package is useful. Excellent algorithms are now available to do that.

In Figure L.1 the components A and B are restored independently from each other. This is not realistic because there is generally only one repair team. When a component fails it waits until the repair team is available in order to be repaired in turn. In Figure L.4, the first component failed is the first repaired and for modelling that the state 4 has been split in order to differentiate two states: A failed first (state 5) and B failed first (state 4).

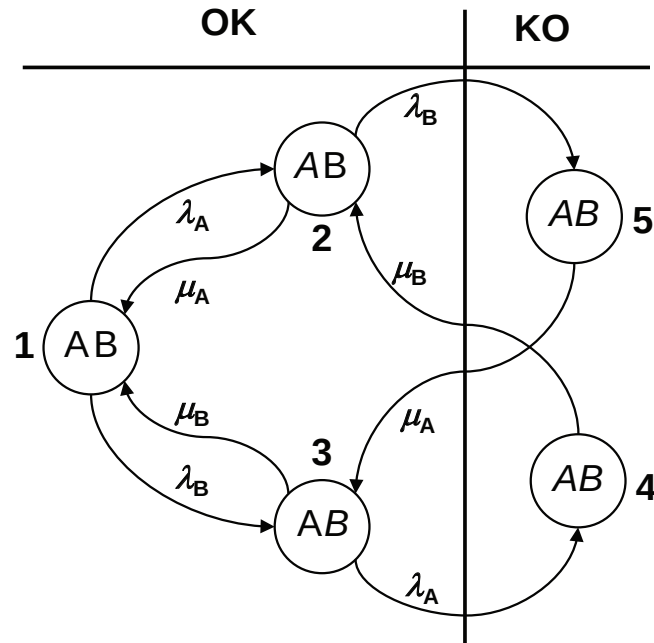


Figure L.4 — Modelling a single repair team (first failed first repaired)

This exhibits the flexibility and the ability of Markov approach to model rather complex system behaviours. Unfortunately the number of states increases exponentially when the number of components of the modelled system increases. Therefore the main limitation of the Markovian approach is due to the resulting combinatory explosion of the states. Using this approach on industrial size systems imply approximations difficult to handle: aggregation of similar states, identification and calculations of the shortest sequences (i.e. the chain of transitions from perfect state to the class KO), etc. In this Technical Report only the use of this approach for components (to be used as fault tree input) or small safety systems is analysed. For larger safety systems the analysts are encouraged to use the mixed fault tree x Markov approach described previously when this is possible or to jump to the Petri nets and the Monte Carlo simulation described below.

L.4 Common cause failure modelling

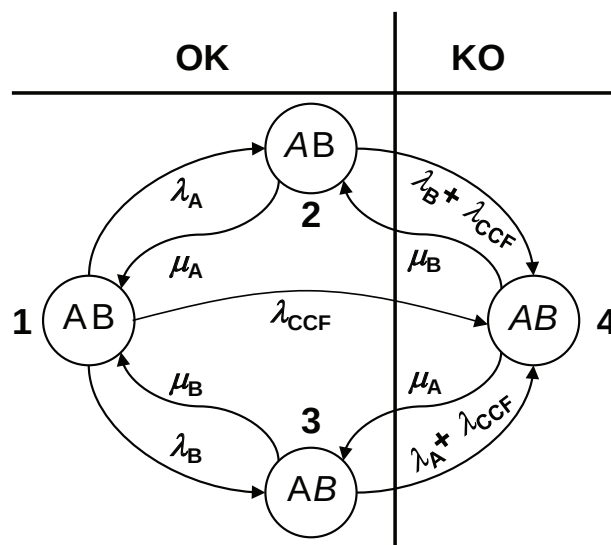


Figure L.5 — Example of common cause failure modelling

The common cause failures can be easily modelled by using the Markovian approach. This is done in [Figure L.5](#):

- The transition from state 1 to state 4 models a CCF which fails both component A and B at the same time. The CCF failure rate is λ_{CCF} .
- λ_A and λ_B are now the independent failure rates of components A and B. Therefore the full failure rates are $\lambda_A + \lambda_{CCF}$ and $\lambda_B + \lambda_{CCF}$. This has been taken into account in the transitions from 2 to 4 and 3 to 4.

The OK state becomes a critical state but this change nothing to the developments described above.

L.5 Multi-phase Markov modelling — Unavailability calculations

If the formulae developed above are relevant for the previous Markov models which comprises only random events (revealed failure and repair) this is definitely not valid when periodically tested failures are going to be modelled because the test interval is deterministic. In this case it is necessary to switch to the “multiphase” Markovian approach to properly deal with the problem.

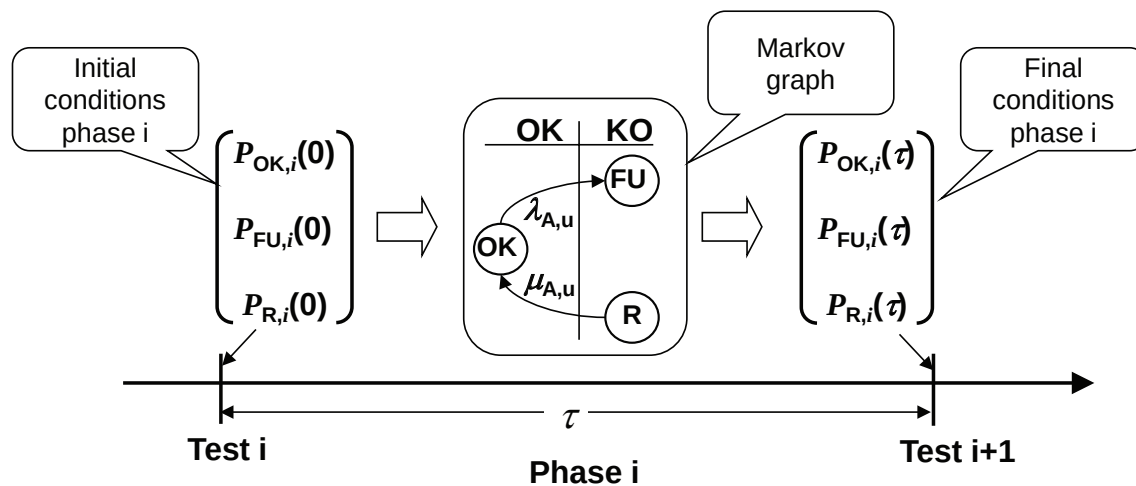


Figure L.6 — Example of simple multiphase Markovian model

The Markov graph presented in [Figure L.6](#) model the behaviour of a periodically tested component between in the phase i which begins just after the test number i and finishes just before the test number $i+1$. From the initial conditions (i.e. the probabilities of the various states at the beginning of phase i) the system behaves as a three states Markovian process during all the duration of phase i and it jumps between the three following states:

- OK: the system is perfect;
- R: the system is under repair;
- FU: the system is failed but the failure is not detect yet.

As shown in [Figure L.6](#), no jump is allowed from FU to R because this failure is not detected yet.

At the end of the phase i , the final conditions are reached (i.e. the probabilities of the various states just before test number $i+1$ is performed). They are calculated from the initial condition as this has been described above for the ordinary Markov processes.

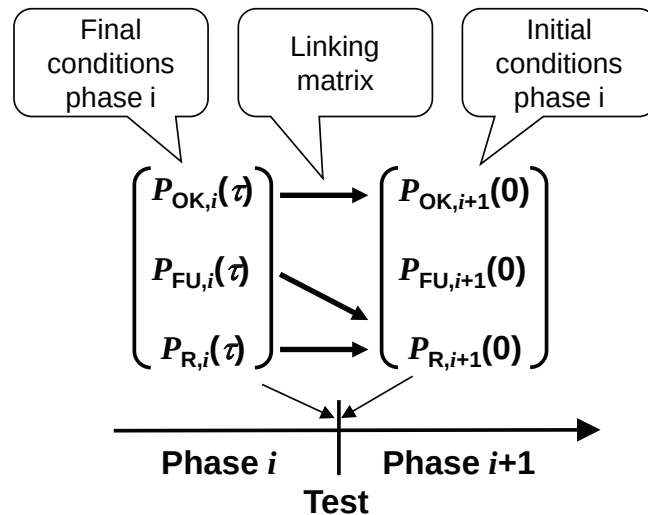


Figure L.7 — Example of linking between phase i and phase i+1

Then the initial conditions of phase $i+1$ can be evaluated. This is done through the linking matrix illustrated graphically in [Figure L.7](#):

- the system remains OK if it is OK;
- repair continue the system is already under repair;
- repair starts if the system is failed.

Therefore the initial conditions at the beginning of phase $i+1$ are calculated by the combination of the final conditions at the end of phase i and the linking matrix. This process can be started at the end of phase 1 and continued along to the relevant number of phase covering the period T under interest.

It should be mentioned that this models works even if the time to repair is not negligible compared to the test interval duration and also if the probability of failure is high during the test interval. It should also be mentioned that this model is valid both for dangerous and safe failures. This is why the notation FU (failed and undetected) is used.

From the graph in [Figure L.6](#) it is obtained:

- Unavailability: $U_S(t) = P_{FU}(t) + P_R(t)$
- Average unavailability over $[0, T]$: $\bar{U}_S(T) = \frac{\Theta_{FU}(T) + \Theta_R(T)}{T}$

The results are illustrated in [Figure L.8](#) where several curves have been plotted. The curve 1 is the classical one (low failure rate, quick restoration). The curves 2, 3 and 4 illustrate cases with a higher failure rate and a restoration rate which decreases. Finally, the curve 5 illustrates the case where the test interval decreases. When the test interval goes to 0, the curve 5 becomes similar to this of an immediately revealed failure.

This model encompasses all the cases and the dangerous detected failures appear only as a particular case of the dangerous undetected failures and it is possible to encompass the low demand mode of operation as well as the high demand mode of operation within the same model.

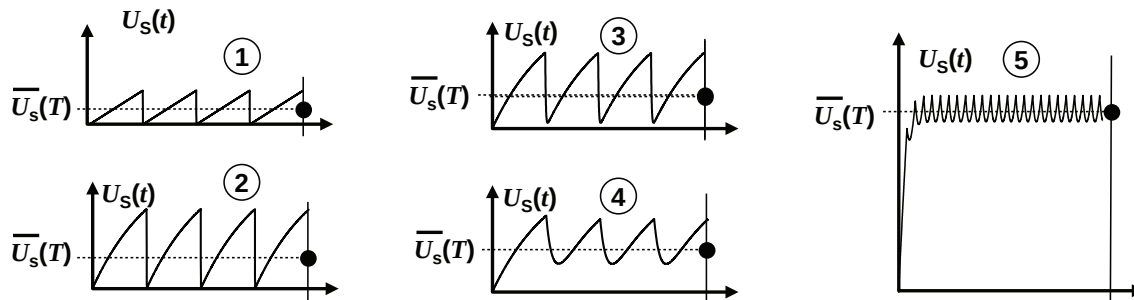


Figure L.8 — Various saw-tooth curves derived from multiphase Markov model

Such models may be used as input in fault tree driven Markov processes as mentioned in Clause 8.

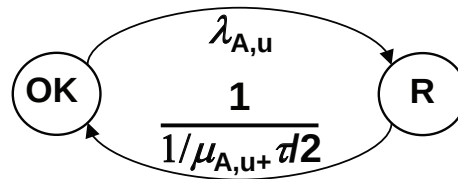


Figure L.9 — Pseudo equivalent Markov graph

The ordinary Markov graph presented on [Figure L.9](#) is often proposed as an equivalent of the multiphase Markovian model of [Figure L.6](#) and [Figure L.7](#). The curves presented in [Figure L.10](#) shows that this simpler model does not behave at all like the actual system but that its asymptotic unavailability is close to the average unavailability of the multiphase model. It is, nevertheless a little bit optimistic.

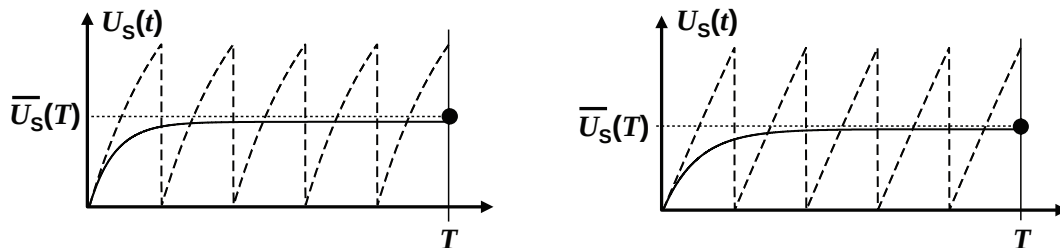


Figure L.10 — Comparison between simple and multiphase Markovian model

It should be noted that to use the correct equivalent restoration rate $(1/\mu_{a,du} + \tau/2)$ it is necessary to evaluate first the mean time to detect the failure by the periodic test. But, the use of some kind of multiphase model is needed to evaluate it. For example, in the case of [Figure L.9](#), the development of the analytical equations related to the multiphase Markov model allows to find that it is close to $\tau/2$. Then there is a kind of vicious circle: the problem should be exactly solved before implementing the approximate model. This may be rather difficult for more complex system and it seems more realistic to use directly a multiphase Markov model when dealing with periodically tested failures.

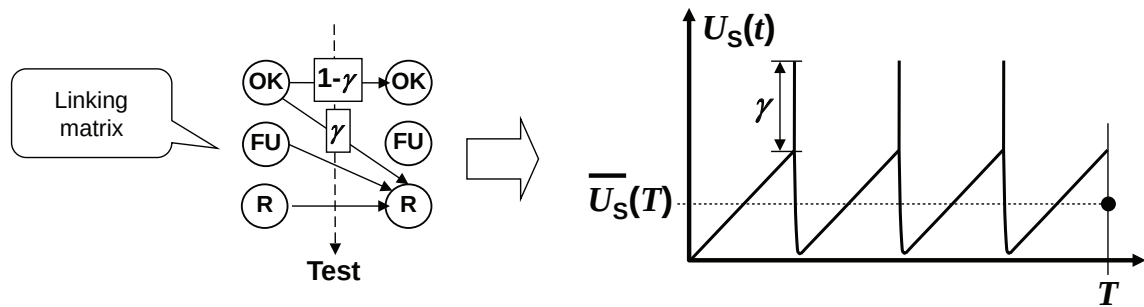


Figure L.11 — Modelling of the failures due to the tests

The model in [Figure L.6](#) and [Figure L.7](#) can be easily adapted to handle other parameters like those mentioned in [7.3.2](#). [Figure L.11](#) illustrates the case when the safety system can be failed due to the test itself (probability γ). For example, the sudden change of state occurring when the test is performed may be harmful (e.g. starting the diesel engine of the emergency power system, moving a mechanical relay, moving a valve, etc.). [Figure L.11](#) shows the jump γ in the system unavailability which occurs when the test is performed. The parameter γ is the probability of failure due to the tests. It was formerly named “probability of failure on demand” but this is used now as a synonym of “average unavailability” (see [3.1.15](#), [3.1.16](#) and [3.2.13](#)). Therefore the analysts should make sure of the real meaning of this term when they meet it.

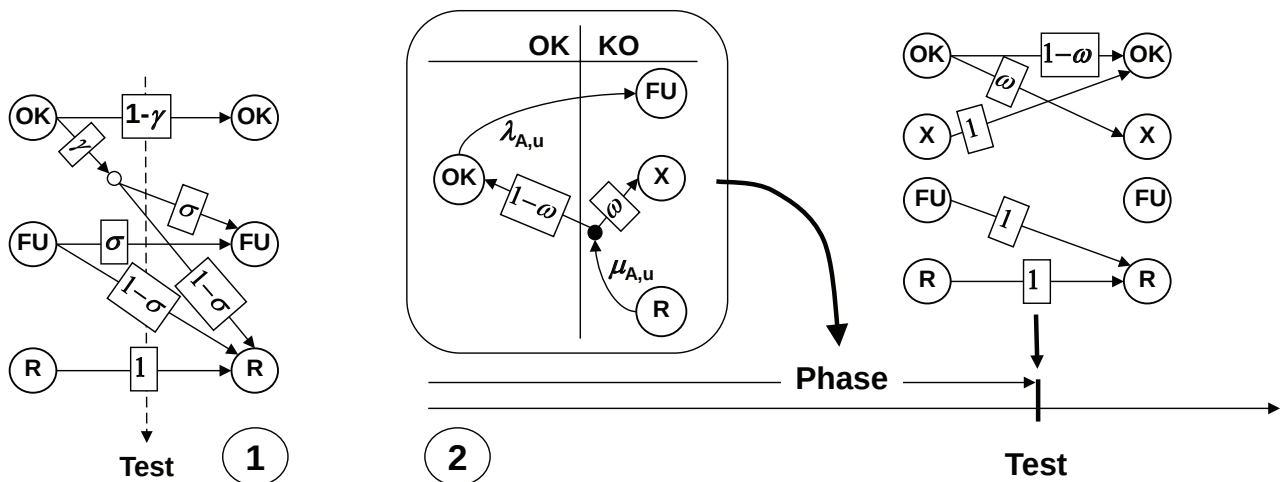


Figure L.12 — Modelling of human error

This kind of modelling is very flexible and can be adapted to various situations. In [Figure L.12](#) several type of human error are modelled. The linking matrix 1 model the probability σ that a failure is not detected by the operators performing the tests. The Markov graph and the linking matrix 2 model the probability ω that the component is not reconnected after a test. Note that in this case, a new state (X) has been introduced to model the reconnection omission. Note also that the same mistake has been modelled at the end of a repair.

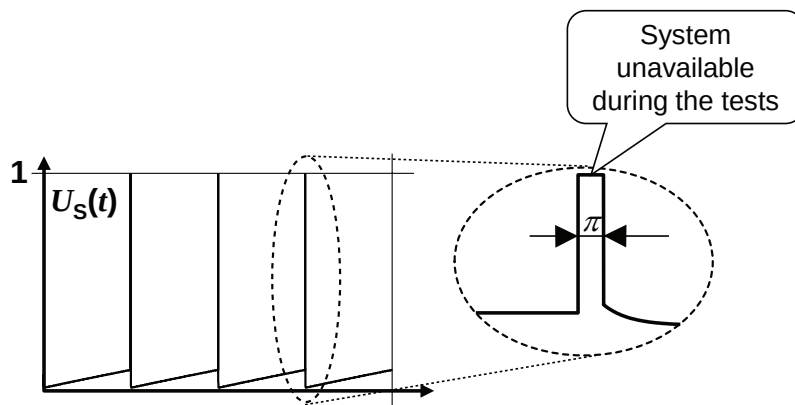


Figure L.13 — Modelling of a periodically tested single component with test duration non equal to zero

When the protected system is not stopped during the tests then it is necessary to take into account the tests durations, especially if the safety system is disconnected to perform the tests. This introduces extra phases (to model the tests themselves) between the test intervals.

Such model leads to curves like in [Figure L.13](#) when the component is unavailable for its function during the tests and this causes a jump up to of 1 of the unavailability. This has been enlarged on the right hand side of the figure. Even if the test duration is small, its contribution to the average unavailability may be not negligible (e.g. $1,4 \cdot 10^{-3}$ for a test lasting 1 h per month).

There is no problem to mix the immediately revealed (detected) and the periodically tested (undetected) failures within the same Markov model. This has been done in [Figure L.14](#). This does not change the principle of the calculations. One state FD has been added to the Markov graph on the left hand side. It belongs to the KO class and it models a failure for which the restoration starts as soon as it is detected. The same state appears in the linking matrix on the right hand side: its probability is just re-transmitted to the next phase without change.

NOTE It has been considered that the system cannot fail of several manners at the same time. If it is possible, extra states and transitions should be introduced

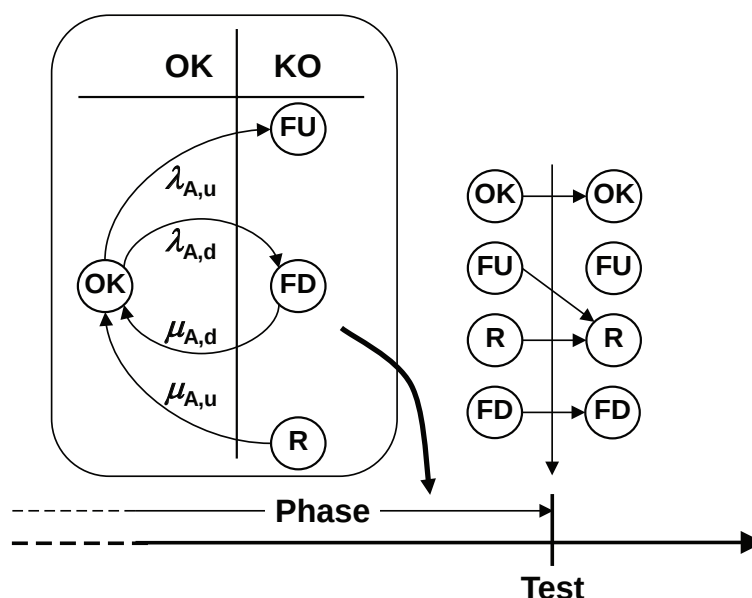


Figure L.14 — Modelling both immediately revealed failures and undetected failures

An example of more complex multiphase model is illustrated in Figure L.15: a system made of two periodically tested redundant components. The components are not tested at the same time and this introduces two types of phases: this beginning with a test of A and this beginning with a test of B. The Markov graph is the same for all the phases but there are two linking matrixes, one for the test of A and another for the test of B.

The length of the phases are not necessarily the same but in Figure L.15 it has been considered that the test are staggered at the optimum (i.e. one component is tested in the middle of the test interval of the other). Parameters like these introduced for the single components (i.e. γ , σ , ω) have not been considered here in order to keep this model as simple as possible.

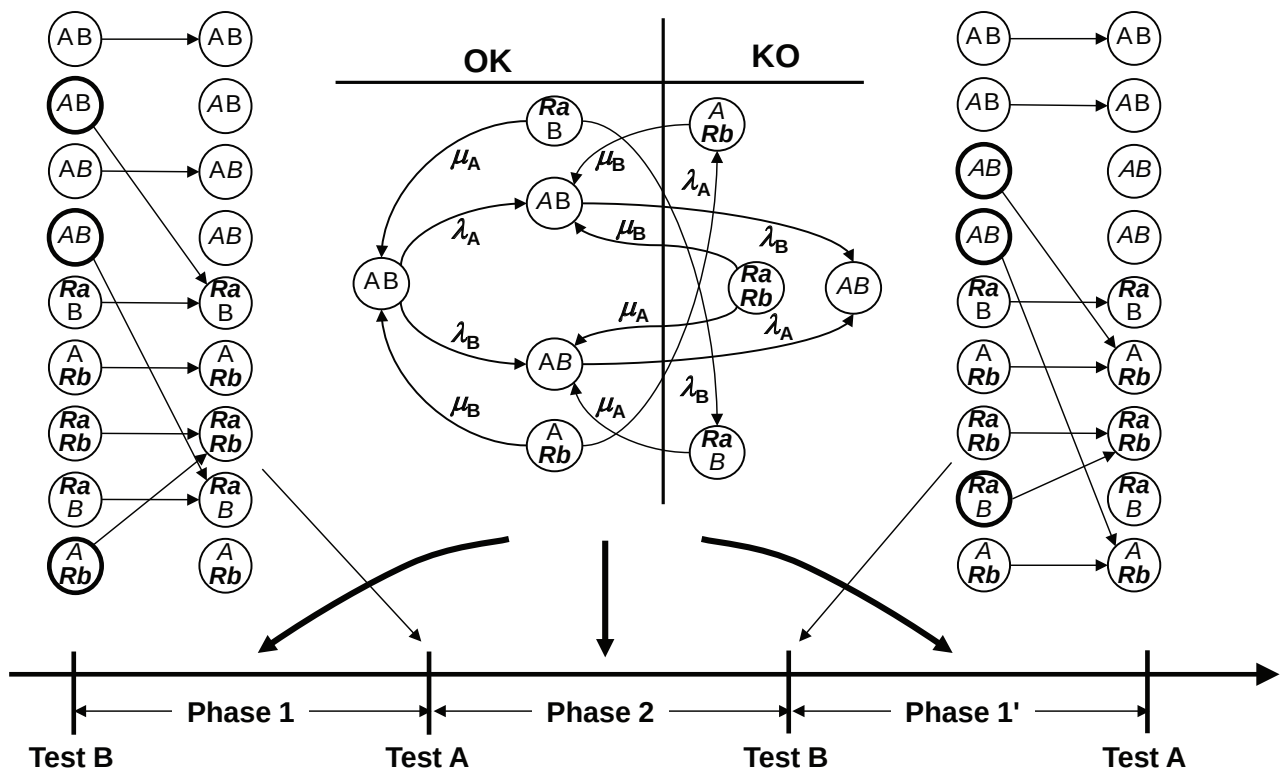


Figure L.15 — Two periodically tested redundant components (with test staggering)

The unavailability of this system is equal to the sum of the probabilities to be in the states of the class KO and the average unavailability the sum of the cumulated sojourn times in the states of the class KO divided by the duration of the period of interest.

L.6 Unreliability and frequency calculations

As explained above in L.3, the Markov models devoted to unreliability calculations are similar to those devoted to unavailability calculation except that jumping out the KO state class is forbidden. Unreliability is focused on the first failure and the calculations discard any restoration of the KO class. Therefore, for reliability calculations it is necessary to identify the states belonging to the KO class and to gather them into a single absorbing state.

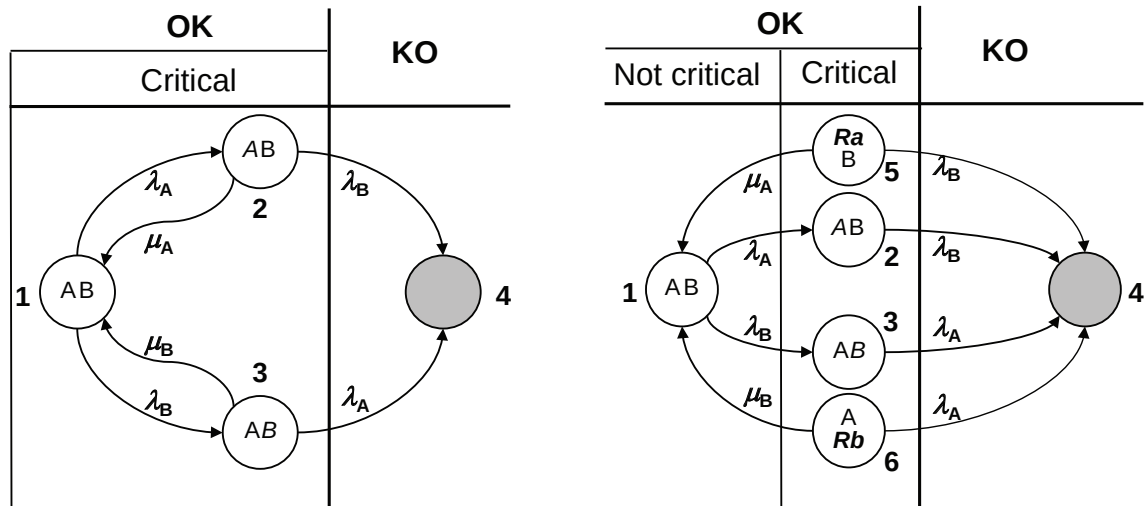


Figure L.16 — Markov graphs for unreliability calculations

Figure L.16 shows (graph on the left) how Figure L.1 and Figure L.4 are modified to perform unreliability calculations. It also shows (graph on the right) how Figure L.15 is be modified. In both cases all states of the KO class have been gathered into a single “absorbing” state (hatched). It is called absorbing because the system stay here after this state has been reached. The system unreliability is given in any cases by: $F_S(t) = P_{KO}(t)$.

Figure L.17 shows the type of unreliability curves which can be obtained from the models presented in Figure L.16.



Figure L.17 — Example of unreliability curves

The frequency of failure can be calculated by the general formula established in 9.1: $w_S(t) = \Lambda_V(t) \cdot A_S(t)$.

In the case of a reliability graph, $w_S(t)$ is also the failure density $f_S(t)$ of the system modelled by the Markov graph. This gives: $f_S(t) = w_S(t) = \Lambda_S(t) \cdot R_S(t)$.

As $R_S(T)$ tends to 0 when t goes to ∞ , $f_S(t)$ continuously decreases when t increases. This is due to the fact that a reliability graphs allows only one system failure. Therefore $f_S(t)$ is meaningful only when $F_S(t) \ll 1$, i.e. $R_S(T) \approx 1$.

The general expression of the unreliability of the system is: $F_S(t) = 1 - e^{-\int_0^t \Lambda_S(\tau) d\tau} = 1 - e^{-\bar{\Lambda}_S t}$

When $F_S(t) \ll 1$:

- $F_S(T) \approx \bar{\Lambda}_S(T) \cdot T$ where $\bar{\Lambda}_S(T)$ is the average of the equivalent failure rate of the system over $[0, T]$.
- $R_S(T) = 1 - F_S(T) \approx 1$.
- $f_S(T) = \Lambda_S(T) R_S(T) \approx \Lambda_S(T)$ and $\bar{f}_S(T) \approx \bar{\Lambda}_S(T)$.

Finally, provided that the system is reliable over $[0, T]$, the frequency of failure can be calculated from the unreliability by the formula: $\bar{w}_S(T) = \bar{f}_S(T) \approx \frac{F_S(T)}{T}$.

This formula is valid in any cases, and in particular for periodically tested systems. It is the duty of the analyst to verify that the hypothesis $F_S(t) \ll 1$ is met because this formula tends to 0 when t increases and may be not conservative.

When all the failures are quickly detected and restored (e.g. safe or dangerous detected failures) the Vesely failure rate $\Lambda_V(t)$ of the availability model (e.g. the graph of [Figure L.1](#)) reaches quickly an asymptotic value Λ_V^{as} which provides a good approximation of the asymptotic value of the equivalent failure rate Λ_S^{as} of the reliability model ([Figure L.16](#) on the left). This allows to calculate $\bar{\Lambda}_S \approx \Lambda_S^{as} \approx \Lambda_V^{as}$ from the availability graph and, therefore, the system unreliability $F_S(t)$. This constant asymptotic failure rate, when it exists, can be calculated directly from the Markov graph.

Let us consider the Markov graph on the left hand side of the [Figure L.16](#): the average sojourn times are $\theta_1 = 1/(\lambda_A + \lambda_B)$ for state 1, $\theta_2 \approx 1/\mu_A$ for state 2 and $\theta_3 \approx 1/\mu_B$ for state 3. When $1/\mu_A$ and $1/\mu_B$ are small compared to $1/(\lambda_A + \lambda_B)$ the states 2 and 3 are quasi instantaneous compared to state 1. Then the equivalent failure rate of the overall system can be directly obtained as shown in [Figure L.18](#).

The system jumps out of state 1 toward state 2 with the transition rate λ_A . Then it immediately jumps to the absorbing state with the probability $\lambda_B/(\lambda_B + \mu_A)$ or comes back to the state 1 with the probability $\mu_A/(\lambda_B + \mu_A)$. Therefore $\lambda_A \lambda_B/(\lambda_B + \mu_A)$ is the transition rate of the sequence 1→2→4. The sum of all the possible sequences from 1 to the absorbing state (i.e. 1→2→4 and 1→3→4) gives an excellent approximation of the equivalent failure rate of the whole system.

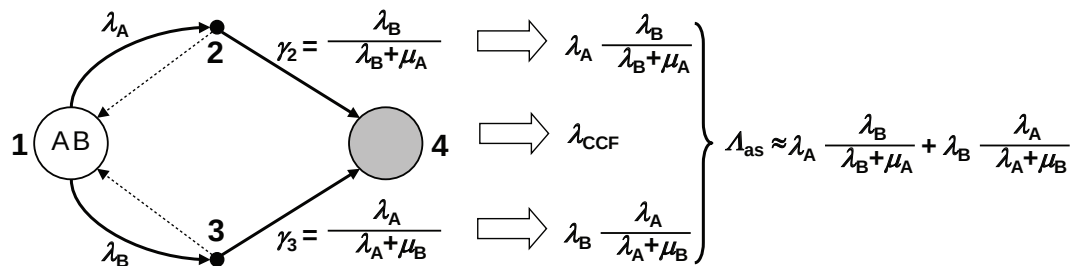


Figure L.18 — Simplified calculation of the equivalent failure rate

The corresponding availability and reliability graphs give the same formula though the reliability graph provides the equivalent failure rate as defined in [3.1.16](#) and the availability graph provides the Vesely failure rate (conditional failure intensity) as defined in [3.1.19](#). The numerical values are close and, as previously said, $\Lambda_S^{as} \approx \Lambda_V^{as}$ when the failures are totally and quickly repairable. This is an important property which has already been used for the approximation of the reliability calculations from Boolean models.

Finally, in this case the $MTTF_S$ can be estimated to $MTTF_S = \frac{1}{\Lambda_S^{as}} \approx \frac{1}{\Lambda_V^{as}}$.

The above formula is valid as soon as the convergence has been obtained, i.e. after 2 of 3 times the higher mean restoration time of the states belonging to the Markov graph.

When the reliability is high over $[0, T]$: $\bar{w}_S(t) \approx \bar{\Lambda}_S \approx \Lambda_S^{as} = \frac{1}{MTTF_S}$.

When there is no asymptotic equivalent failure rate the $MTTF_S$ is a little bit more complicated to calculate and implies the use of the reliability graph. The simplest way is to consider the mean cumulated functioning time $MFT(T)$ over a given period $[0, T]$. This is obviously given by the cumulated time spent in the OK classes:

$$MFT(T) = \sum_{i \in OK} \Theta_i(T)$$

$MFT(T)$ increases when T increases but, as there is an absorbing state, it reaches an asymptotic value when $F_S(T)$ tends toward 1:

$$MTTF = \lim_{T \rightarrow \infty} MFT(T) = \lim_{T \rightarrow \infty} \sum_{i \in OK} \Theta_i(T) \text{ for reliability Markov graphs}$$

L.7 Hazardous event probability calculations

L.7.1 Ultimate safety system

For an ultimate safety system working in continuous mode of operation, a hazardous event occurs as soon as the KO state class is reached. Then KO is an absorbing state and all what has been described in L.6 applies and no more explanations are needed.

The reliability Markov graphs in [Figure L.16](#) discussed above are example of model of safety systems working as ultimate safety barriers in continuous mode of operation.

For an ultimate safety systems working in demand mode, there is a chance to restore the safety system when it is failed before a demand occurs. This demand can be modelled as illustrated by [Figure L.19](#) which represents the [Figure L.1](#) and the [Figure L.14](#) adapted to calculate the probability of a hazardous event for an on demand safety system working as ultimate safety barriers.

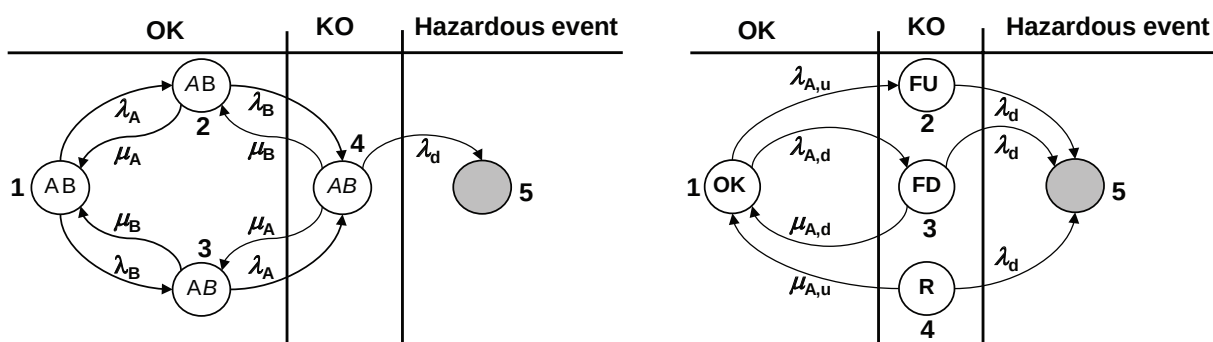


Figure L.19 — Hazardous event modelling in case of ultimate safety system

On these graphs λ_d represent the demand rate.

They are ordinary “reliability” Markov graphs and, again, all what has been described above in [L.6](#) directly apply and no more explanations are needed.

The number of expected hazardous events $N_a(T)$ can be evaluated as explained in [L.7.2](#). If $N_a(T) \ll 1$ then it can be used as an estimation of the probability of accident. This allows using directly the Markov graph established for the safety system itself.

L.7.2 Non-ultimate safety system

When the safety system is not the ultimate safety system a demand is generated toward the following safety system when it fails. This succeeding safety system normally responds properly and the accident is avoided. Therefore, the non-ultimate safety system can be restored and there is no need for an absorbing state as in the previous case. This is true either for safety system working in continuous mode of operation or demand mode of operation.

For safety systems working in continuous mode of operation, a demand on a succeeding safety system occurs as soon as the KO state class is reached. Therefore all what has been described above in [L.1](#) directly apply and no more explanations are needed: the Markov graphs in [L.1](#) can be used to model safety systems working as non ultimate safety barriers in continuous mode of operation.

For safety systems working in demand mode, there is a chance to restore the safety system, when it is failed, before a demand occurs. Two possibilities are available to evaluate the expected number of demand on the succeeding safety barrier:

- 1) Evaluating $MTTRes_s$ of the safety system and using one of the formula established in [D.1](#) for $N_a(T)$;
- 2) Modelling the demand as illustrated in [Figure L.20](#).

Modelling the demand may lead to a complex Markov graph and therefore it is better to consider approach number 1 first. In this case:

- Average unavailability: $\bar{U}_S(T) = \frac{\sum_{k \in KO} \Theta_k(T)}{T}$ (cf. L.3)
- Mean time to restore the system: $MTTR_S(T) = \frac{\bar{U}_S(T)}{w(T)}$ (cf. L.3)

And, the number of demand on the succeeding system is obtained by

- $N_a(T) \approx \bar{U}_S(T) \cdot \lambda_d T$ when $\lambda_d \cdot MTTR_{Res} < 1$
- $N_a(T) \approx \bar{U}_S(T) \cdot \frac{(1 - e^{-\lambda_d \cdot MTTR_S})}{MTTR_S} T$ when $\lambda_d \cdot MTTR_{Res}$ is not small compared to 1.

The demand can be easily modelled. This is done in Figure L.20 which is similar to the Figure L.1 and Figure L.14 adapted to calculate the probability of demand on the succeeding safety system.

Again, λ_d is the demand rate.

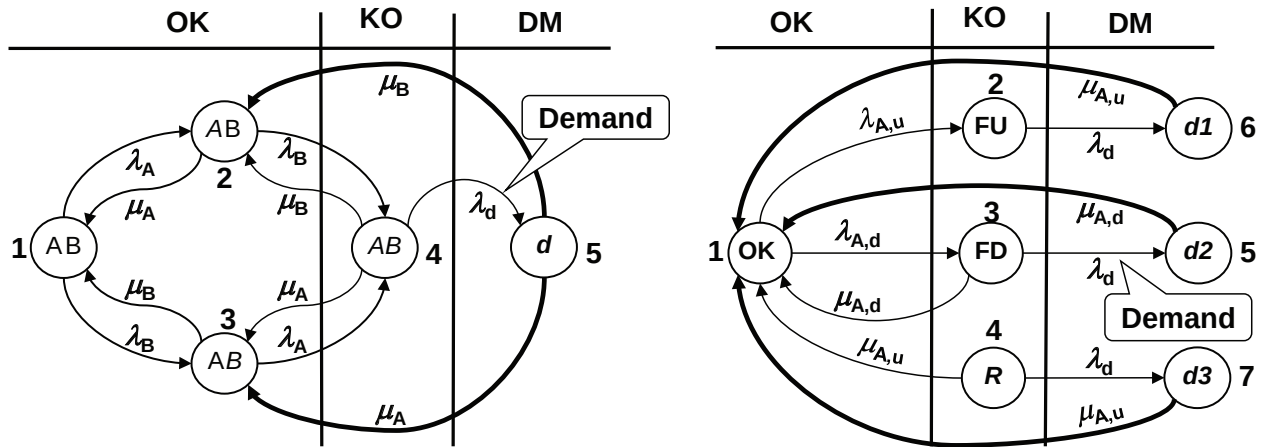


Figure L.20 — Modelling of a demand transferred to a succeeding safety barrier

On the left hand side, when a demand occurs in state 4 this demand is transmitted to the succeeding safety system (i.e. state 5). There is no reason to stop the repairs undertaken in state 4 and they continue in state 5: the transition from 5 to 2 is the same as 4 to 2 and the transition from 5 to 3 is the same as 4 to 3.

On the right hand side this is a little bit more complicated. When a demand occurs in states 2, 3 or 4 this demand is transmitted to the succeeding safety system (i.e. states 5, 6 or 7). For the states 5 and 7 the repairs undertaken in state 3 and 4 go on. Therefore the transition from 5 to 1 have the same transition rates as 3 to 1 and 4 to 1. For the state 6 this demand “reveals” the failure which was hidden in state 2. Then a repair can begin without need to wait for the next test and the transition from 6 to 1 has been introduced with the transition rate $\mu_{A,u}$.

The Markov graphs in Figure L.20 are similar to an “availability” Markov graphs and therefore all what has been described above in this annex applies by switching the role of the class CS to KO and the role of the class KO to DM. The parameters are similar to $\Lambda_S(t), \Lambda_S^{as}, \Lambda_V(t), \Lambda_V^{as}$ and can be manipulated in the same way without additional difficulties. For example:

- Probability to generate a demand toward the succeeding safety system: $P_{dm}(t) = \sum_{k \in DM} P_k(t)$.
- Frequency of demand on the succeeding safety system: $w_{dm}(t) = \sum_{i \in KO, j \in DM} \eta_{i,j} P_i(t)$.
- Average frequency of demand on succeeding safety system: $\bar{\Phi}_{dm}(T) = \bar{w}_{dm}(T) = \frac{\int_0^T w_{dm}(t) dt}{T}$.

L.8 Spurious failure modelling

The spurious failure can be modelled through models of the type “availability” Markov Graph.

Figure L.21 illustrates the example of a system made of two components A and B. On the left hand side, the safe failure of any of them produce a spurious failure (e.g. a spurious trip) of the system when on the right hand side the safe failures of both of them are necessary to produce a spurious failure of the safety system.

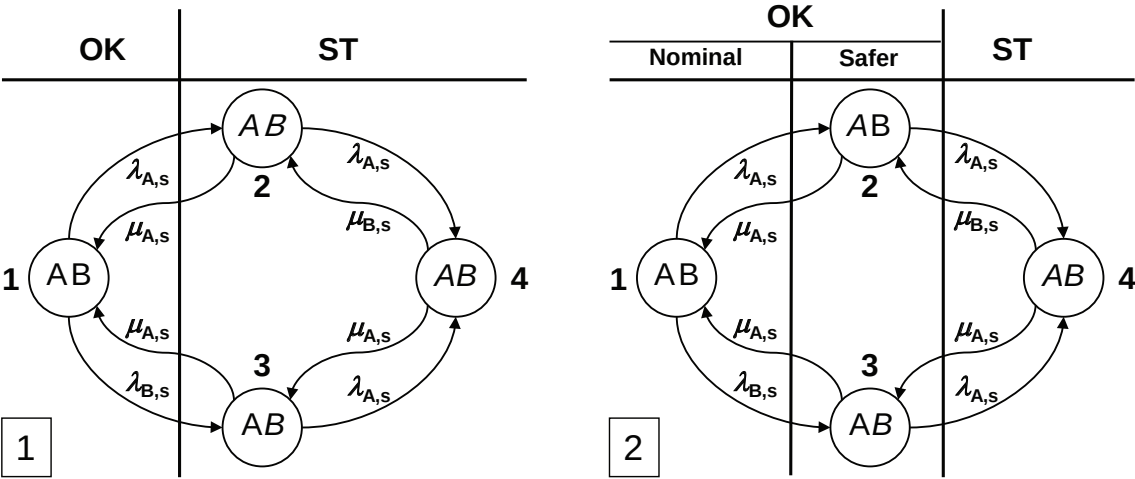


Figure L.21 — Example of spurious failure modelling

Then on the right hand side, a safe failure of A may drive the system into states safer than the nominal state (e.g. from 1 to 2) or to an unsafe state (e.g. from 3 to 4). The effect of a given failure change according to the current state of the system. The nature “safe / unsafe” is not an intrinsic characteristic of the individual components. This is a systemic characteristic and it is the duty of the analyst to identify the spurious failure scenarios.

Annex M (informative)

Petri net modelling

M.1 Basic principle

The basic graphical elements of Petri nets (PN) are represented in [Figure M.1](#) according to IEC 62551[7]. Further development can be found in reference.[25]

A Petri net is split within two different kinds of graphical elements:

- Static elements which are simple drawings.
- Dynamic elements which change according to the state of the modelled system.

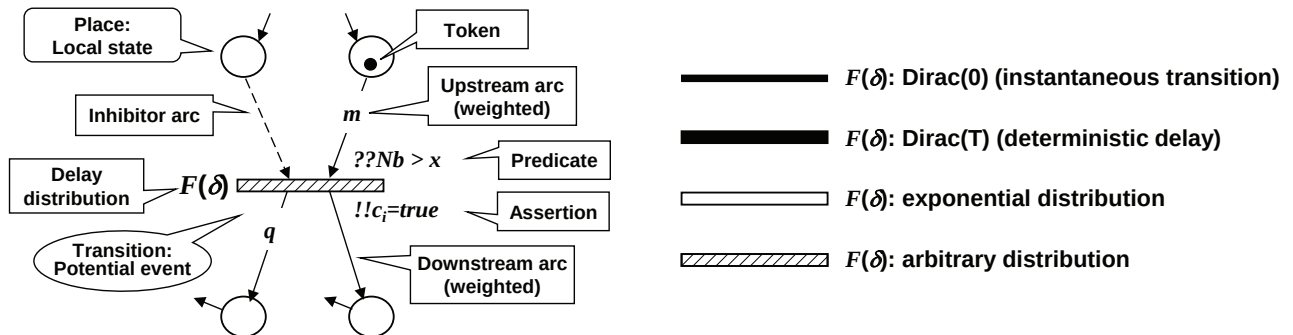


Figure M.1 — Basic graphical elements for Petri nets

The static elements are the following:

- Places (circles) which model potential local states.
- Transitions (rectangles) which models potential events. IEC 62551[7] proposes various drawings in function of the nature of the delay (see on the right hand side of the figure). They update the assertions when they are fired (i.e. when the attached event occurs) and change the marking of the downstream places.
- Upstream arcs linking places to transitions and participating to the validation of transitions. They can be weighted or not. Inhibitor arcs are represented in dotted line.
- Downstream arcs linking transitions to places and indicating the “marking” changes of the places when the transitions are “fired”. They can be weighted or not.

The dynamic elements are the following:

- Tokens (small black circles) which are removed or added in places when transitions are “fired”. The tokens constitute the “marking” of the PN.
- Predicates which participate to the validation of transitions. A predicate is identified by a double question mark $??$. This is any expression which can be “true” or “false” (e.g. $??2 == 5$ is always “false”, $??x > y+z$ is “true” or “false” according to the current values of y and z , etc.).
- Assertions which update variables when transitions are fired. An assertion is identified by a double exclamation mark $!!$. This is a simple equation (e.g. $!!x = y+z$ updates the value of x according to the

current values of y and z). Assertions can be local (i.e. attached to a given transition) or global (i.e. attached to the whole PN).

The “marking” and the values of the variables define the current state of the PN, i.e. the state of the modelled system.

A transition is “valid” (or “enabled”) when all the following conditions are gathered:

- 1) The upstream places have at least a number of token equal to the weight of the corresponding upstream arc.
- 2) The upstream places have a number of token strictly lower than the weight of the corresponding inhibitor arc.
- 3) The predicates are “true”.

When a transition is valid, that means that the corresponding event can occur.

As soon as a transition becomes valid its attached delay distribution is used to evaluate when it is going to be fired. The delay can be deterministic or stochastic. When random delays are considered, the Petri net is named “stochastic Petri net”.

A random delay is evaluated by using random numbers in a so-called Monte Carlo simulation (see [Clause 11](#) and [Annex N](#)).

When the delay has elapsed the transition is “fired” (i.e. the corresponding event occurs) and this results in:

- Removing in upstream places a number of token equal to the weight of the corresponding upstream arcs. (The inhibitor arcs induce no change).
- Adding in downstream places a number of token equal to the weight of the corresponding downstream arcs.
- Updating the variables according to the assertions attached to the transition.
- Updating the variables according to the global assertions.
- Updating all the predicates attached to all the transitions.

Then a new current state is reached: new transitions may become valid, some transitions may be inhibited and some others may don’t change. And so on...

The potential of PN to model dynamic behaviours is virtually endless. This makes them very effective to model the industrial systems and to perform Monte Carlo simulations.

M.2 RBD driven Petri net modelling

As discussed in [6.2.1](#), the safety systems have basically only two states (OK, KO) and therefore using reliability block diagrams (RBD) seems a natural way of modelling. The approach consisting to use an underlying virtual RBD to drive the building of the PN itself is detailed below. It has proven to be very effective because:

- it helps in building the PN;
- it helps to understand the behaviour of the PN;
- it helps to master large models thank to the use of libraries of blocks;
- it allows calculations which are far beyond the capacity of the common RBDs;
- it does not really limit the modelling powerfulness of Petri nets.

The idea is then to use such small Petri nets as a kind of library of blocks within a whole RBD as this has been done in [Figure M.3](#).

The underlying virtual RBD is represented in dotted lines. It is a straightforward re-transcription of the RBD presented in [Figure 18](#):

- The blocks of the RBD have been fulfilled by one of the sub-PN described in [Figure M.2](#):
 - periodically tested component sub-PN₃ for the sensors and the valves;
 - revealed failure sub-PN₁ for the logic solver.
- The logics of the RBD is achieved through four global assertions related to the nodes O₁, O₂, O₃ and “out”.
- These global assertions use the state variables of each component. They are updated each time a local assertions implies a change of their input variables.
- The sensors share the repair team R_s when the valves share the repair team R_v. The repair teams are supposed to be readily available for the logic solvers.

The use of this approach allows developing Petri nets for complex safety systems while keeping the readability of the model and avoiding the mistakes.

The remaining problem is to extract the probabilistic parameters of interest from this model. This work is devoted to the auxiliary Petri net presented in [Figure M.4](#).

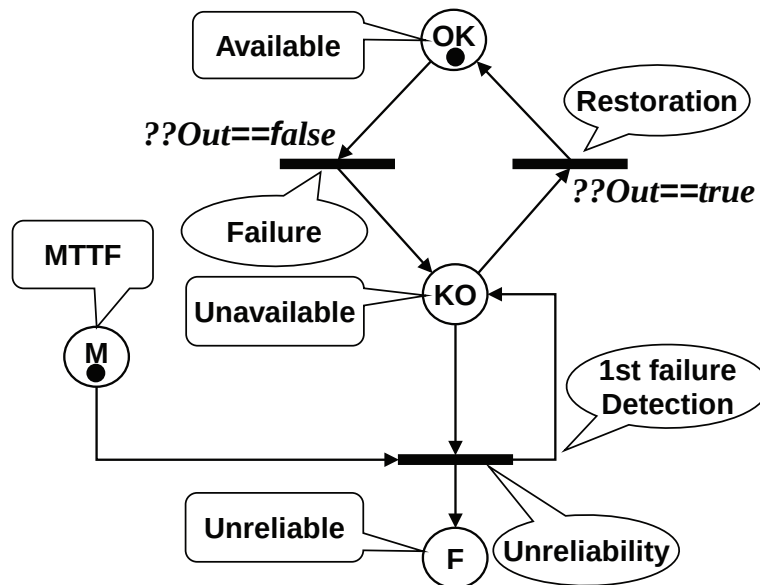


Figure M.4 — Auxiliary PN for Unavailability / Reliability / frequency calculations

The link between this auxiliary sub-PV and the model in [Figure M.3](#) is done through the variable “out” which gives the current state of the overall safety system:

- At the beginning the safety system is OK.
- When the variable “out” becomes “false” then the transition “failure” is fired immediately.
- The token is removed from OK and one token is put in KO indicating that the system is now unavailable.
- The first time that KO is marked, the transition “unreliability” is immediately fired
- This removes the tokens in KO and M and adds one token in KO and F. Then the marking of KO is unchanged.

- When the variable “out” becomes “true” then the transition “restoration” is fired and the token in KO is removed and one token is added in OK. One cycle has been achieved.
- And so on.

It should be noted that the transition “unreliability” is a “one shot” transition: after the first shot, the token in M definitively disappears and the transition will be never fired again.

When animated by a Monte Carlo simulation this model provides all the relevant probabilistic measures which are needed, e.g.:

- Availability $A_S(t)$: place OK marked at t .
- Unavailability $U_S(t)$: place KO marked at t .
- Average availability: mean marking of place OK.
- Average unavailability: mean marking of place KO.
- Unreliability $F_S(T)$: firing frequency of the transition “unreliability”.
- $MTTF_S$: mean marking of place M provided that T is large enough to have at least one failure for each simulated history).
- Failure frequency (not ultimate safety layer): number of firings of the transition “failure” / T .
- Failure frequency (ultimate safety layer): $\approx$ frequency of the transition “unreliability” / $T \approx 1/MTTF_S$.
- MUT_S : mean marking of OK $\times T$ / number of firings of the transition “failure”.
- MDT_S (i.e. $MTTRes_S$): mean marking of KO $\times T$ / number of firings of the transition “restoration”.
- $MTBF_S = MUT_S + MDT_S$.

Therefore, when using Petri net modelling it is useless to distinguish between, low demand, high demand, and continuous mode of operation or spurious failures. The modelling principle is the same in any cases.

M.3 Hazardous event probability calculations

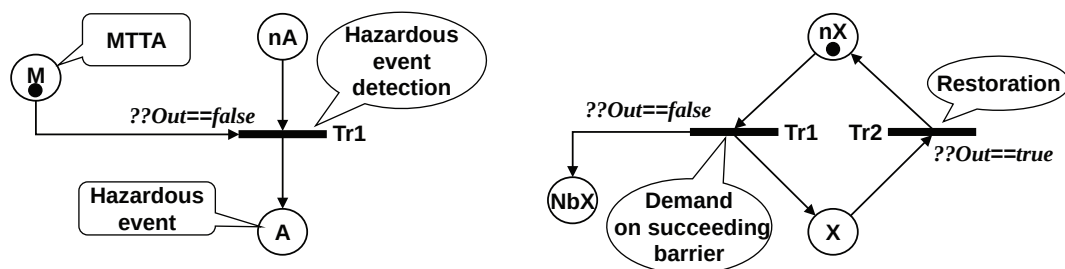


Figure M.5 — Hazardous event modelling in continuous mode of operation

For a safety system running in continuous mode the auxiliary sub-PNs in [Figure M.4](#) can be used as it is but it can be simplified according to [Figure M.5](#).

The sub-PN on the left hand side of [Figure M.5](#) is related to an ultimate safety barrier: the firing frequency of the transition Tr1 provides the statistical estimation of the hazardous event probability. The cumulated time with one token in the place M, gives the mean time to a hazardous event, provided the simulation is performed for a time large enough to have as many hazardous event as simulated histories.

The sub-PN on the right hand side of [Figure M.5](#) is related to a non ultimate safety barrier: the firing frequency of the transition Tr1 provides the statistical estimation of the frequency of the demands

on the succeeding safety barrier. The mean marking of the place NbX gives the mean number of such demands over the time used for the simulation.

For safety systems operating “on demand” mode, three possibilities are available:

- 1) Using the formula $N_a(T) \approx \bar{U}_d \lambda_d T$ established in [D.1](#)
- 2) Evaluating the MTTRes of the safety system and using the complete formula for $N_a(T)$ established in [D.1](#);
- 3) Modelling the demand as illustrated in [Figure M.6](#).

When the probability of hazardous event (or demand on the succeeding safety barrier) is low, the approach number 3 may lead to very long Monte Carlo computations, therefore the approach number 1 or 2 should certainly be considered first.

As said above $\bar{U}_S(T)$ is given by the mean marking of the place KO and MTTRes_S is given by the cumulated time spent in KO divided by the average number of firings of the transition “restoration”:

- Number of hazardous event: $N_{a,S}(T) \approx \bar{U}_S(T) \cdot \lambda_d T$ if $\lambda_d \cdot \text{MTTRes}_S \ll 1$
- Number of hazardous event: $N_{a,S}(T) \approx \bar{U}_S(T) \cdot \frac{(1 - e^{-\lambda_d \cdot \text{MTTRes}_S})}{\text{MTTRes}_S} T$ if $\lambda_d \cdot \text{MTTRes}_S$ is not small compared to 1.

Of course, if it is needed, the demand can be easily modelled as it is illustrated by [Figure M.6](#).

The sub-PN in [Figure M.6](#) models a demand occurring with a demand rate equal to λ_d . When Tr3 is fired then the Boolean variable “Dem” becomes “true”. It remains true during a short while ε (see the explanation below) after which the transition Tr4 is fired. Then “Dem” becomes “false” and the sub-PN is ready to simulate the next demand. Such sub-PN generates demands with a mean time between demands equal to $1/\lambda_d$ as ε is negligible.

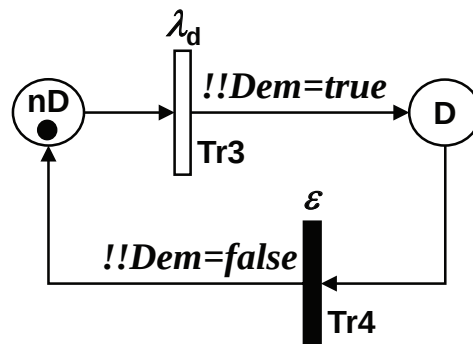


Figure M.6 — Sub-PN for demand modelling

The sub-PNs in [Figure M.7](#) are very similar to those in [Figure M.5](#). The difference is that the transition Tr1 is validated when “Dem” is “true”. The sub-PN in the left hand side is related to the occurrence of a hazardous event when the sub-PN on the right is related to the occurrence of demands on the succeeding safety layer.

Let us consider the transition Tr1 of the sub-PN on the left hand side. When “Out” is “false” and “Dem” becomes “true” this transitions can be fired immediately. But in the sub-PN in [Figure M.6](#), the transition Tr4 can also be fired immediately. There is a conflict between these transitions because if Tr4 is fired first then Tr1 will be inhibited before being fired. Therefore a mechanism is needed to ensure that Tr1 is fired before Tr4 when those transitions are simultaneously valid.

This can be done by giving a higher level of priority to Tr1 than to Tr4, by inhibiting Tr4 as long as Tr1 has not been fired or by choosing the very simple solution implemented here: the delay of Tr4 (i.e. ε) is greater than the delay of Tr1 (i.e. 0). Then Tr1 is always fired before Tr4.

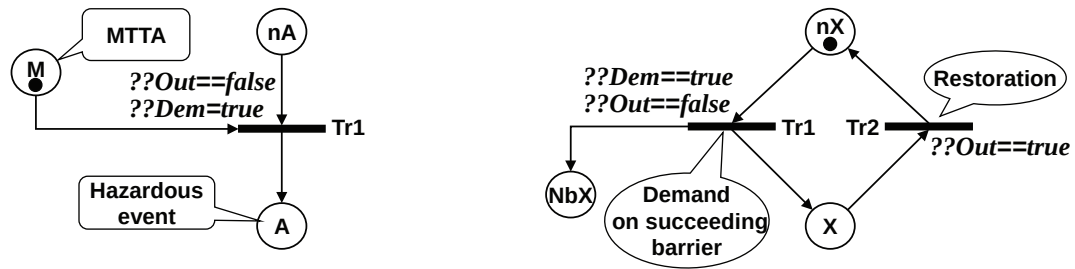


Figure M.7 — Hazardous event modelling in “on demand” mode of operation

M.4 Common cause failure modelling

Modelling the common cause failures by the β factor approach is rather simple:

- 1) If the CCF is repaired by itself, it can be modelled by one of the sub-PN presented on [Figure M.2](#).
- 2) If the CCF fails the components then each component are going to be repaired independently and this implies the implementation of the new mechanisms which are presented in [Figure M.8](#) with some other improvements.

In [Figure M.8](#), the improved blocks are represented on the right hand side where 2 of them have been organized into a redundant architecture (A, B). This constitutes the core of the model.

On the left hand side some auxiliary sub-PN are proposed:

- A simple model for CCF is drawn at the top. It communicates with the core of the model by the Boolean variable “CCF”. The components A and B reach immediately their respective FU states when “CCF” becomes “true”.
- From that point the behaviour is similar as above in [Figure M.3](#). Note that, in this model the CCF is detected by the periodic tests of components A and B.
- The simple model for the repair team sharing as been replaced by a more sophisticated sub-PN representing the mobilization of a support vessel to perform the repairs:
- When A or B is detected to be failed, the variable “Dem” becomes “true” and this message is received by the mobilization sub-PN.
- The mobilization starts immediately, except if it is winter.
- After a certain delay the support vessel is mobilized and able to go on the location of the failure.
- When it is on location the repair of the failed component begins. Only one repair is possible at the same time.
- When the repair is finished then the support vessel is freed and a token is added in the place “Mob”.
- Then the support vessel sails toward the next location.
- At this point it undertakes the next repair or it is demobilized if there is no more failure to repair.

Note that the place “Urg” in [Figure M.8](#) counts the number of failed items and the place “W” measures how much time the support vessel has been mobilized.

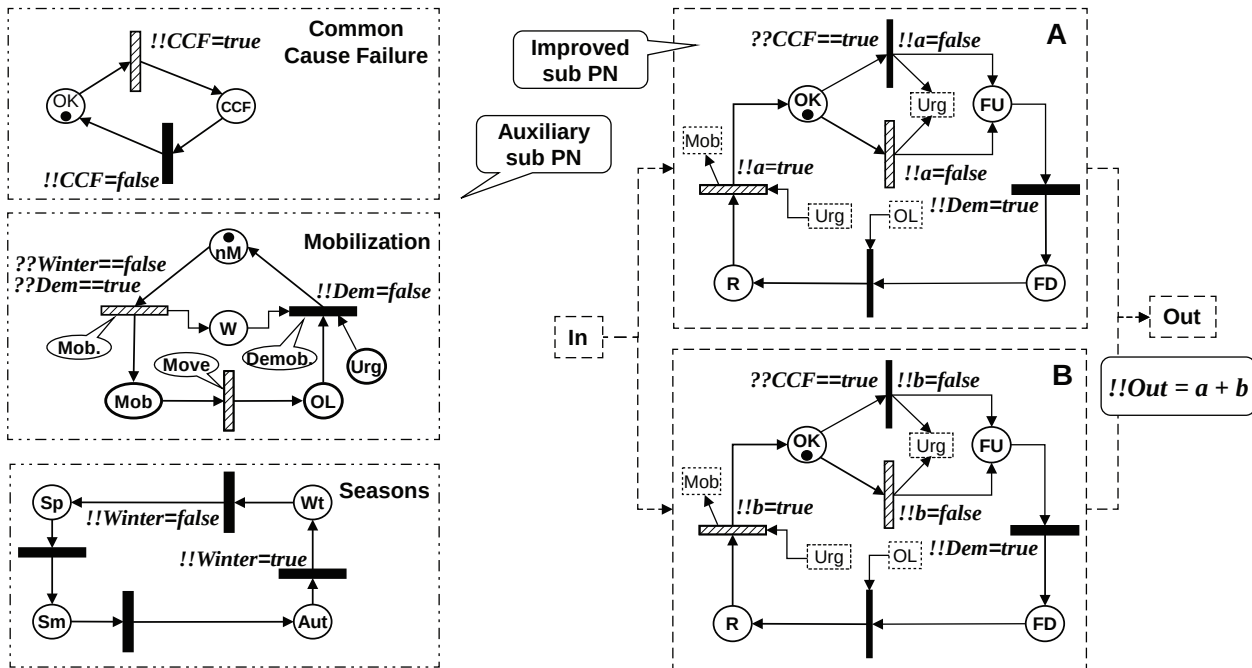


Figure M.8 — Improved sub-PN and auxiliary sub-PN

When there is more than 3 or 4 similar components the β factor approach becomes unrealistic and the shock method may be used instead. This consists in considering that the CCF is split in two parts:

- a lethal shock similar to the β factor.
- a non lethal shock which provokes the failure of the related components with a given probability γ .

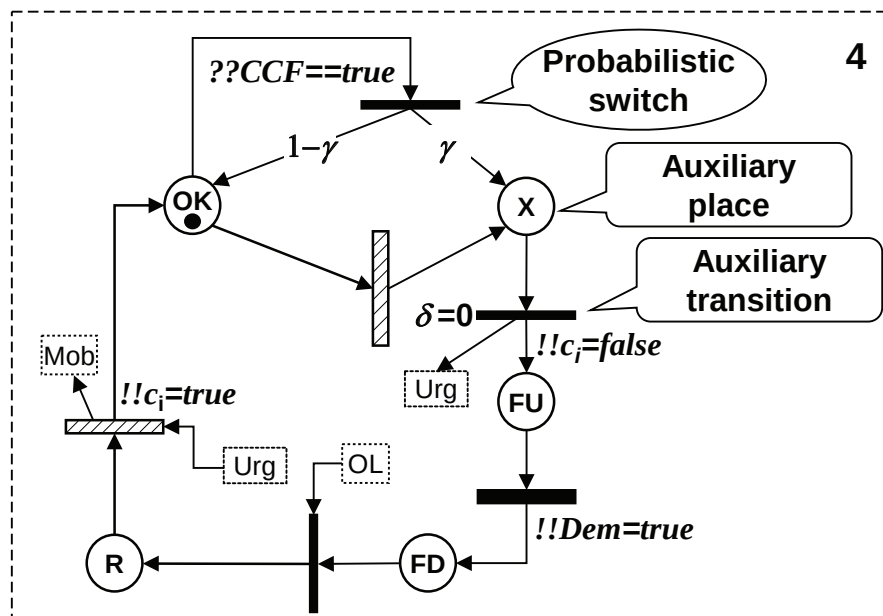


Figure M.9 — Sub-PN for modelling a non lethal shock

Figure M.9 models a non lethal shock. The sub-PN is derived from the previous one and can be used as a module within a Boolean driven Petri net modelling. A “probabilistic switch” is used to model the shock. When it occurs there is a probability γ that the component fails and $1-\gamma$ that it does not fail. As it is not possible to put the assertion $!!c_i = false$ on the probabilistic switch, one auxiliary place and one

auxiliary transition have been used for this purpose. This transition is fired immediately after it has been validated.

www.TeraStandard.com

Annex N (informative)

Monte Carlo simulation approach

The Monte Carlo simulation is not a modelling approach but a way of calculation which can be used with any model described in this Technical Report. It consists in building a random game and to play it a great number of times in order to obtain wanted results through statistical calculation. This is a powerful tool which still works when analytical approaches have reached their limits. Further development can be found in reference.^[25]

As, since the 18th century, Buffon has proposed to calculate π by throwing needled on the lathes of a wooden floor this is not really a new approach. Von Neumann was one of the main supporters of this approach in the 1940s. Its development follows the computer development.

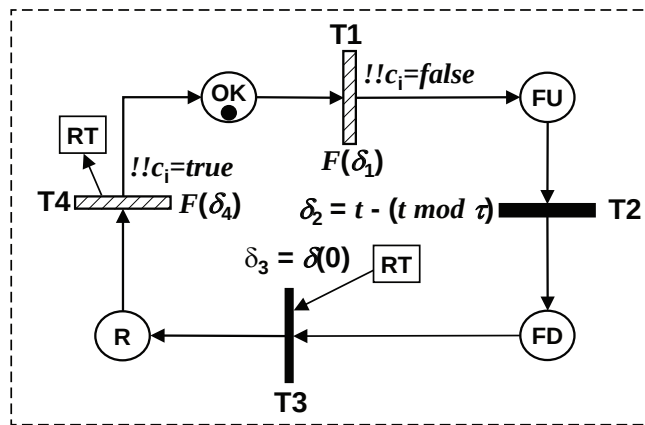


Figure N.1 — Simple PN for Monte Carlo simulation illustration

Performing Monte Carlo simulation needs a good behavioural model modelling both the functioning and dysfunctioning system behaviour. Any of the models described in this Technical Report can be used but the Petri nets are very effective to do that and the [Figure N.1](#) will help to show how to handle such a simulation. This Petri net is similar to the sub-PN 3 in [Figure M.2](#) and the names (T_1 , T_2 , T_3 and T_4) of the transitions and their attached probabilistic distributions have been added.

In the following explanations T_i , δ_i and t_i are use to indicate the transition, its delay before firing and the time of its firing:

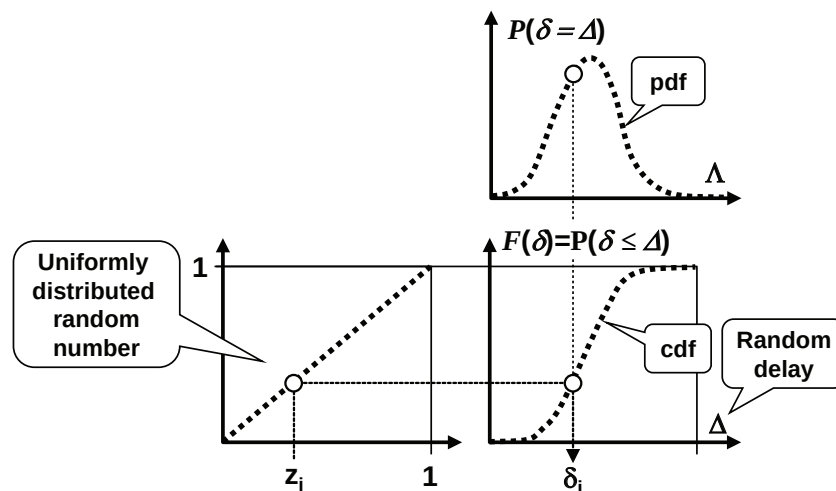
- 1) At the initial time, the system is OK (one token in the place OK).
- 2) This validates T_1 and the delay δ_1 can now be determined.
- 3) The use of a random number and of the distribution $F(\delta_1)$ give the delay δ_1 which is going to elapse before T_1 is fired. Therefore T_1 is going to be fired at $t_1 = \delta_1$.
- 4) When δ_1 is elapsed, the token is removed from OK, one token is added in FU and C_i is updated to the value "false".
- 5) This validates the transition T_2 and the delay δ_2 can now be determined.
- 6) The delay δ_2 is calculated thanks to a deterministic equation: this is the remaining time until the next periodic test. Therefore T_2 is going to be fired at $t_2 = t_1 + \delta_2$.
- 7) When δ_2 is elapsed then the token is removed from FU and one token is added in FD.

- 8) This is not sufficient to validate T_3 because as long as the place RT is not marked.
- 9) The marking of RT depends on other part of the overall model. Therefore T_3 is governed by external events. In fact RT is used to synchronize this sub-PN with the whole model.
- 10) Therefore T_3 is inhibited until one token is present in RT. When this occurs T_3 is fired immediately as δ_3 is governed by a Dirac distribution equal to 0. Then T_3 is fired t_3 .
- 11) The token is removed from FU and one token is added in R and this validates T_4 and the repair starts.
- 12) The delay of repair δ_4 is calculated by using one random number and the distribution $F(\delta_3)$.
- 13) When δ_4 is elapsed, the token is removed from R, one token is added in OK, one token is added in RT and the Boolean variable C_i becomes "true".
- 14) The situation is the same as at the beginning: system OK and repair team available. Then the steps 2 to 14 can be repeated again and again until the time of firing of one transition goes beyond T which in the upper boundary of the period under interest.
- 15) This provides one "history" of the system over $[0, T]$. All along the progress of one history the relevant results are recorded in order to constitute statistical samples.
- 16) The principle of the Monte Carlo simulation is to simulate plenty of such histories in order to gather statistical samples large enough to perform accurate statistical estimations. Therefore the steps 1 to 15 are, normally, repeated many times.

The results recorded during the simulation may be the marking of the places, the number of firing of the transitions, the occurrence of pre-defined states, etc.

Mixing deterministic and stochastic delays within the same model is very easy with the Monte Carlo when this may be very difficult with analytical calculations.

The remaining issue is the simulation of the delays according to their probabilistic distributions (i.e. their cumulative distribution function) by using random numbers. This is in fact rather simple because the random variable $Z = F(\delta)$ (where $F(\delta)$ is the cumulative distribution function of δ) is uniformly distributed over $[0, 1]$. Then a sample (δ_i) is obtained from a sample (z_i) uniformly distributed just by doing $\delta_i = F^{-1}(z_i)$.



NOTE The shape of the probabilistic density function (pdf) is similar to these of Weibull or log-Normal distributions.

Figure N.2 — Monte Carlo simulation of random variable

Figure N.2 illustrates the process used to generate a sample (δ_i) from uniformly distributed random numbers (z_i):

- 1) fire a random z_i number between 0 and 1 (good algorithms are now available to do that);
- 2) evaluate the corresponding value of the random variable, (e.g. δ_i), by using its inverse cdf - e.g. $F(\delta)$.

This is very easy when the analytical form of $F^{-1}(z)$ is known but effective algorithms exist for the main distributions used in reliability engineering. For example, a sample of delays of failure exponentially distributed is obtained by: $\delta_i = -\frac{1}{\lambda} \text{Log}(z_i)$ where λ is the failure rate.

The remaining problem is to obtain uniformly distributed random numbers of good quality. This is not really a problem and there are several possibilities:

- Random number tables
- Physical methods: dice, roulette, thermal noise of a Zener diode, etc.
- Decimals of a transcendent number: several billons of decimals of $\pi = 3,142\ 592\ 653\ 589\ 793\ 23\dots$ are now established
- Built-in function of a programming language: RND, RANDOM, ...
- Algorithms in a computer: congruential methods like $x_{n+1} = (ax_n + b) \text{ modulo } m$

Using the powerful algorithms available in PC computers is the most practical solutions. Those algorithms are generally based on the congruential model mentioned just above. The limitation comes from the inherent cycles of the congruential generator: it exist a number x_n such as $x_n = x_0$. Therefore after x_n , the same random number set is generated as when starting from 0. The cycle depends of a , b and m and it is the job of the analyst to verify if this cycle is in adequacy with the intended Monte Carlo simulation to perform.

A criticism which is often done against the Monte Carlo simulation is that it wouldn't be accurate. This is an irrelevant criticism because the Monte Carlo simulation always provides the confidence interval of its estimations. This allows the analyst to verify straightforwardly if the results have converged toward accurate enough values.

NOTE In comparison, the accuracy of results obtained by analytical calculation would require to estimate the impact of all the approximations. This is rarely performed or even possible.

Let us consider a sample (x_i) of a given parameter x obtained through a Monte Carlo simulation. The statistical treatment provides the following result:

- Average: $\bar{x} = \frac{\sum x_i}{N}$,
- Variance: $\sigma^2 = \frac{\sum (x_i - \bar{x})^2}{N}$ and standard deviation: σ
- Confidence interval $\text{Conf}_{90\%} = 1.64 \frac{\sigma}{\sqrt{N}}$

This implies that $P[\bar{x} \in (\bar{x} - 1.64 \frac{\sigma}{\sqrt{N}}, \bar{x} + 1.64 \frac{\sigma}{\sqrt{N}})] = 90\%$ where x is the true value of the parameter x .

The confidence interval becomes narrow when the frequency of occurrence of X increases and when the number N of history increases. Nevertheless this is a slow process as this depends on the square root of N .

In addition it should be noted that the duration of the Monte Carlo calculation is linked with the number of occurring events and not at all with the size of the model. The negligible events just not occur during the simulations and, contrarily to the other approaches which often need simplified models to be handled, the Monte Carlo simulation is self approximating.

The current computation power of personal computers permits to handle safety system up to SIL4 (e.g. with average unavailabilities in the range 10^{-4} - 10^{-5} over one year) without big difficulties. Several safety barriers running in sequences may be also handled without too many difficulties. The probability of accident, if it is very low may cause some problems which can be solved by faster computers or a mix between analytical and Monte Carlo simulation as mentioned at the beginning of this annex.

.....

Annex O (informative)

Numerical uncertainties handling

NOTE This annex is not intended to deal with epistemic uncertainties (i.e. uncertainties due to a lack of knowledge) but only with uncertainties linked to data inputs.

The accuracy of the probabilistic calculations is directly linked to the uncertainties of the reliability parameters which are used as inputs. Therefore it is interesting to take those uncertainties into account in order to evaluate their impact on the results which are found.

Except in very simple case the analytical approaches do not work and the Monte Carlo simulation provides the only way to do that. This is illustrated in [Figure O.1](#) on a simple fault tree.

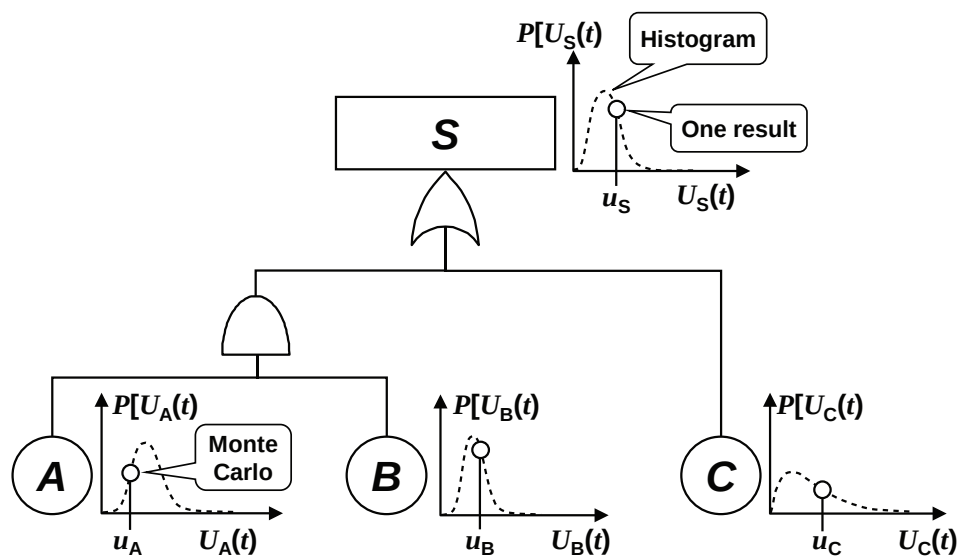


Figure O.1 — Uncertainty propagation

When handling the uncertainties, the input reliability parameters are no longer deterministic and they are replaced by random variables. Therefore the probabilities of failure of the leaves (e.g. unavailabilities) become also random variables in turn. The probability density functions (*pdf*) are more or less sharp or flat according to the degree of uncertainty which is involved. In [Figure O.1](#), and for a given time t , $U_B(t)$ is better known than $U_A(t)$ which is better known than $U_C(t)$.

As Monte Carlo simulation is implemented, the result at the top of the fault tree is no longer deterministic. It is obtained as histogram which is more or less sharp or flat according to the impact of the input uncertainties. [Figure O.1](#) illustrates the result obtained when the calculation is performed for a given time t .

Most of the time this Technical Report involves the calculation of averages values over a given duration $[0, T]$. Therefore the process described here for one value of the time has to be extended to several values covering the whole range $[0, T]$. The principle is the following:

- 1) Consider one time t and one leaf of the fault tree (e.g. A);
 - i) identify its random parameters - e.g. λ_a and μ_a ;
 - ii) chose a parameter - e.g. λ_a - and simulate its value by firing one uniformly distributed random numbers and apply the process described in [Figure O.2](#);

- iii) iterate ii) until all parameters have been processed;
 - iv) evaluate the unavailability of the leave at a given time t - e.g. $U_A(t)$.
- 2) Iterate step 1 until all the leaves have been processed. This provides a set of input probabilities.
 - 3) Perform the calculation of this set of input probabilities. This provides a single output result at time t - e.g. $U_S(t)$.
 - 4) Record this single output result.
 - 5) Iterate steps 1 to 4 until all the relevant number of time values have been processed to obtain a curve over $[0, T]$. This lead to a curve.
 - 6) Calculate the average value of this curve to obtain the relevant average probabilistic value over $[0, T]$ - e.g. the average system unavailability $\bar{U}_{S_i}(T)$. This provides 1 value for the statistical sample.
 - 7) Iterate steps 1 to 6 until a sufficient number of values $\bar{U}_{S_i}(T)$ is obtained (for example 100 or 1 000). This provides a statistical sample.
 - 8) Perform the statistical analysis of this statistical sample to obtain the average and the standard deviation of the average output result over $[0, T]$.

NOTE 1 It is essential that the average unavailability according to the uncertainties (punctual result at time t) and the average unavailability over a given period $[0, T]$ are not mixed-up.

NOTE 2 The probabilistic distributions of the individual reliability parameters should be selected according to the field feedback gathered about it or the engineering judgement when not enough field feedback is available. The distributions commonly used are the following: uniform distribution between two boundaries, triangular distribution with a most probable value, log-normal law with a given error factor, Chi square law, etc.

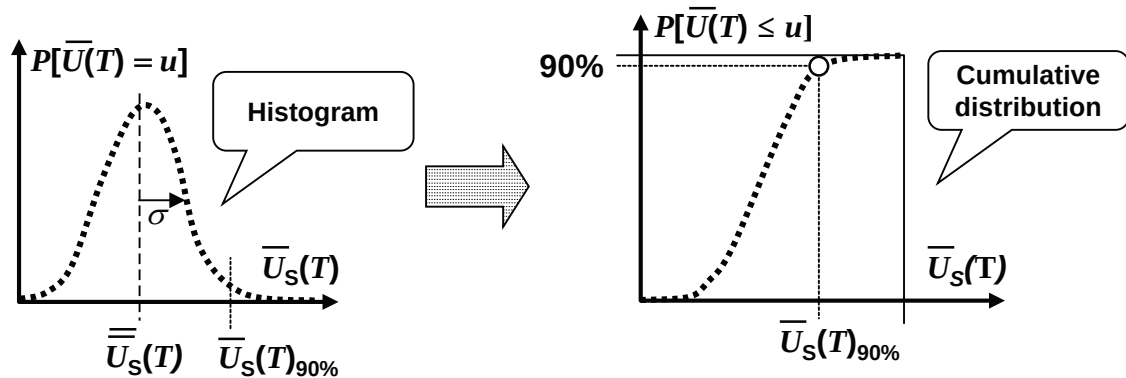


Figure 0.2 — Result at 90 % confidence level

The simplest statistical processing of the statistical sample obtained as output of step 8 is to build the corresponding histogram as this is shown on the left hand side of [Figure 0.2](#). The average of this histogram provides the estimation of the average unavailability, $\bar{U}_S(T)$. The double bar indicates the double averaging (over the time, and over the statistical sample). The sharpness of the histogram and/or the standard deviation σ gives also directly an idea of the uncertainty of the result.

Beyond the above classical results the cumulative distribution obtained from the statistical sample provides results with given confidence levels. This is shown on the right hand side of [Figure 0.2](#). For example this figure shows that the exact value of the average unavailability has 90 chances over 100 to be lower (i.e. better) than $\bar{U}_S(T)_{90\%}$ which is the upper boundary at 90 % confidence level.

$\bar{U}_S(T)_{90\%}$ is greater than $\bar{U}_S(T)$. Then this upper bound provides a conservative estimation of the average of the system unavailability. This is a good way to consider $\bar{U}_S(T)_{90\%}$ to take the uncertainties into account.

The above principles are very general and may be applied on any kind of models: analytical formulae, Markov processes and even Petri nets. A double step Monte Carlo simulation is needed when the main calculation already implement a Monte Carlo simulation (e.g. Petri nets).

www.TeraStandard.com

Bibliography

- [1] INNAL. F., *Contribution to modelling safety instrumented systems and to assessing their performance Critical analysis of IEC 61508 standard*. Thesis of the doctoral school of physical and engineering sciences. University of Bordeaux, 2008
- [2] IEC 61508 ed. 2, *Functional safety of electrical/electronic/programmable electronic (E/E/PE) safety related systems*, part 1-7
- [3] IEC 61511 ed. 2, *Functional safety — Safety instrumented systems for the process industry sector*, part 1-3. (to be published)
- [4] IEC 61078, *Analysis techniques for dependability — Reliability Block Diagram and Boolean methods*
- [5] IEC 61025, *Fault tree Analysis (FTA)*
- [6] IEC 61165, *Application of Markov Techniques*
- [7] IEC 62551 ed. 1, *Analysis techniques for dependability — Petri nets techniques*
- [8] IEC 62502 ed. 1, *Analysis techniques for dependability — Event Tree Analysis*
- [9] CENTER FOR CHEMICAL PROCESS SAFETY / AICHE (AMERICAN INSTITUTE OF CHEMICAL ENGINEERS). *Layer of Protection Analysis: Simplified Process Risk Assessment*. Wiley, 2001
- [10] RISO-M-1374. *Cause consequence diagrams*. 1971
- [11] POINT. G., *AltaRica: Contribution à l'unification des méthodes formelles et de la sûreté de fonctionnement*. These de docteur, Université Bordeaux 1. 2000
- [12] ARNOLD. A, POINT, G. and als, *The AltaRica formalism for describing concurrent systems, Fundamenta informaticae 34*. IOS Press, 2000
- [13] *Reliability Prediction Method for Safety Instrumented Systems - PDS Method Handbook*, ISBN 978-82-14-05524-5, SINTEF A23881. 2013
- [14] IEC 60500-191 ed 2, *International electrotechnical vocabulary — Part 191: dependability*. (to be published)
- [15] ISO 14224 ed.2, *Petroleum, Petrochemical and gas industries — Collection and exchange of reliability and maintenance data for equipment*
- [16] ISO 20815, *Petroleum, petrochemical and natural gas industries — Production assurance and reliability management*
- [17] IEC 60300-3-2, *Dependability management — Part 3-2: Application guide — Collection of dependability data from the field*
- [18] *Reliability Data for Safety Instrumented Systems - PDS Data Handbook*, ISBN 978-82-14-05523-8, SINTEF A23880. 2013
- [19] OREDA, ed. 5, *Offshore Reliability Data Handbook, Volume 1 - Topsides Equipment, Volume 2 - Subsea Equipment*. 2009
- [20] KUMAMOTO. H & HENLEY, E. J., *Probabilistic risk assessment and management for engineers and scientist*. IEEE Press, 1996
- [21] ATWOOD. C. L., *The binomial failure rate common cause model. Technometrics*. 1986, **28** (2) pp. 139-148

- [22] RAUSAND. M. & HOYLAND, A., *System Reliability Theory / models, Statistical Methods and Applications*. Wiley Series in Probability and Statistics, 2004
- [23] COCOZZA-THIVENT. C., *Processus stochastiques et fiabilité des systèmes. Collection Mathématiques & Applications*. Springer, 1997
- [24] SIGNORET. j.-P., *Analyse des risques des systèmes dynamiques: approche markovienne*. Techniques de l'ingénieur SE 4071. 2005
- [25] SIGNORET. j.-P., *Analyse des risques des systèmes dynamiques: Réseaux de Petri*. Techniques de l'ingénieur SE 4073. 2008
- [26] VAN NOORTWIJK. J. M., & als, *Expert judgment in maintenance optimization*. *IEEE Trans. Reliab.* 1992, **41** (3)
- [27] COOKE. R. M., *Experts in Uncertainty: Expert Opinion and Subjective Probability in Science*. Oxford University Press, 1992
- [28] ISO 31000, *Risk management — Principles and guidelines*
- [29] ISO 17776, *Petroleum and natural gas industries — Offshore production installations — Guidelines on tools and techniques for hazard identification and risk assessment*
- [30] ISO 13702, *Petroleum and natural gas industries — Control and mitigation of fires and explosions on offshore production installations — Requirements and guidelines*
- [31] ISO 23251, *Petroleum, petrochemical and natural gas industries — Pressure-relieving and depressuring systems*
- [32] API STD 521 ed5, *Guide for Pressure-relieving and Depressuring Systems: Petroleum petrochemical and natural gas industries-Pressure relieving and depressuring systems*, 2008
- [33] ISO 10418, *Petroleum and natural gas industries — Offshore production installations — Basic surface process safety systems*
- [34] ISO 14723, *Petroleum and natural gas industries — Pipeline transportation systems — Subsea pipeline valves*
- [35] API RP 580 ed1, *API Recommended Practice 580, Risk-Based Inspection*, First Edition. 2002
- [36] HOLLNAGEL. E., *Cognitive Reliability and Error Analysis Method*. Elsevier Science Ltd, Oxford, 1998
- [37] KIRWAN. B., *A Guide to Practical Human Reliability Assessment*. Taylor & Francis, London, UK, 1994
- [38] LEES. F.P., *Loss Prevention in the Process Industries Vol.1 Hazard Identification and Control*, edited by Sam Mannan. 3rd Edition. Ch 14. Human Factors and Human Error. Elsevier. 2005
- [39] McNICHOL. D., *A primer of signal detection theory*. Lawrence Erlbaum Associates Inc, USA, 2005
- [40] REASON, J., *Human Error*. Cambridge University Press, UK, 1990
- [41] WICKENS. T. D., *Elementary Signal Detection Theory*. Oxford University Press, 2002
- [42] WOODWORTH. R.S., SCHLOSBERG, H., *Experimental Psychology. Ch 9 Psychophysics II Scaling methods*. Methuen & Co Ltd, London. 1954
- [43] *Alarm Systems — A Guide to Design, Management and Procurement*, EEMUA Publ. 191
- [44] *Process Plant Control Desks Utilising Human-Computer Interfaces — A Guide to Design, Operational and Human Interface Issues*, EEMUA Publ. 201
- [45] API 770, *A Manager's Guide to Reducing Human Errors Improving Human Performance in the Process Industries*. 2001

- [46] ISA TR84. 00.002, *Safety Instrumented Functions (SIF) — Safety Integrity Level (SIL)*. Evaluation Techniques, 2002
- [47] NAMUR NE 130, *Proven-in-use-devices for safety instrumented systems and simplified SIL calculation*. 2010
- [48] EN 13306, *Maintenance — Maintenance terminology*
- [49] IEC 61810-2, *Electromechanical elementary relays — Part2: Reliability*
- [50] IEC 61164 ed2, *Reliability growth — Statistical test and estimation methods*
- [51] RAUZY A. *BDD for Reliability Studies*. Pages 381–396, Handbook of Performability Engineering. K.B. Misra ed., Elsevier. 2008
- [52] ISO 16708 ed2, *New Petroleum and natural gas industries — Pipeline transportation systems — Reliability-based limit state methods*
- [53] ISO 13623 ed2, *Petroleum and natural gas industries — Pipeline transportation systems*
- [54] ISO 26000, *Guidance on social responsibility*
- [55] ISO 13628-1 ed2, *Petroleum and natural gas industries — Design and operation of subsea production systems — Part 1: General requirements and recommendations*
- [56] ISO 13628-6 ed2, *Petroleum and natural gas industries — Design and operation of subsea production systems – Part 6: Subsea production control systems*
- [57] ISO 13628-7, *Petroleum and natural gas industries — Design and operation of subsea production systems- Part 7: Completion/workover riser systems*
- [58] API RP 14C, *Recommended Practice for Analysis, Design, Installation, and Testing of Basic Surface Safety Systems for Offshore Production Platforms*, 2001
- [59] API RP 170 ed1, *Recommended Practice for Subsea High Integrity Pressure Protection Systems (HIPPS)*, 2009
- [60] API RP 17G ed2, *Recommended Practice for Completion / Workover Risers*, (Identical to ISO 13628-7:2005), 2006

British Standards Institution (BSI)

BSI is the national body responsible for preparing British Standards and other standards-related publications, information and services.

BSI is incorporated by Royal Charter. British Standards and other standardization products are published by BSI Standards Limited.

About us

We bring together business, industry, government, consumers, innovators and others to shape their combined experience and expertise into standards-based solutions.

The knowledge embodied in our standards has been carefully assembled in a dependable format and refined through our open consultation process. Organizations of all sizes and across all sectors choose standards to help them achieve their goals.

Information on standards

We can provide you with the knowledge that your organization needs to succeed. Find out more about British Standards by visiting our website at bsigroup.com/standards or contacting our Customer Services team or Knowledge Centre.

Buying standards

You can buy and download PDF versions of BSI publications, including British and adopted European and international standards, through our website at bsigroup.com/shop, where hard copies can also be purchased.

If you need international and foreign standards from other Standards Development Organizations, hard copies can be ordered from our Customer Services team.

Subscriptions

Our range of subscription services are designed to make using standards easier for you. For further information on our subscription products go to bsigroup.com/subscriptions.

With **British Standards Online (BSOL)** you'll have instant access to over 55,000 British and adopted European and international standards from your desktop. It's available 24/7 and is refreshed daily so you'll always be up to date.

You can keep in touch with standards developments and receive substantial discounts on the purchase price of standards, both in single copy and subscription format, by becoming a **BSI Subscribing Member**.

PLUS is an updating service exclusive to BSI Subscribing Members. You will automatically receive the latest hard copy of your standards when they're revised or replaced.

To find out more about becoming a BSI Subscribing Member and the benefits of membership, please visit bsigroup.com/shop.

With a **Multi-User Network Licence (MUNL)** you are able to host standards publications on your intranet. Licences can cover as few or as many users as you wish. With updates supplied as soon as they're available, you can be sure your documentation is current. For further information, email bsmusales@bsigroup.com.

Revisions

Our British Standards and other publications are updated by amendment or revision.

We continually improve the quality of our products and services to benefit your business. If you find an inaccuracy or ambiguity within a British Standard or other BSI publication please inform the Knowledge Centre.

Copyright

All the data, software and documentation set out in all British Standards and other BSI publications are the property of and copyrighted by BSI, or some person or entity that owns copyright in the information used (such as the international standardization bodies) and has formally licensed such information to BSI for commercial publication and use. Except as permitted under the Copyright, Designs and Patents Act 1988 no extract may be reproduced, stored in a retrieval system or transmitted in any form or by any means – electronic, photocopying, recording or otherwise – without prior written permission from BSI. Details and advice can be obtained from the Copyright & Licensing Department.

Useful Contacts:

Customer Services

Tel: +44 845 086 9001

Email (orders): orders@bsigroup.com

Email (enquiries): cservices@bsigroup.com

Subscriptions

Tel: +44 845 086 9001

Email: subscriptions@bsigroup.com

Knowledge Centre

Tel: +44 20 8996 7004

Email: knowledgecentre@bsigroup.com

Copyright & Licensing

Tel: +44 20 8996 7070

Email: copyright@bsigroup.com

BSI Group Headquarters

389 Chiswick High Road London W4 4AL UK



EAM

Enterprise Asset Management

مدیریت دارایی های سازمان

CMMS

Computerized Maintenance Management System

سیستم مکانیزه مدیریت نگهداری و تعمیرات

MKMS

Maintenance Knowledge Management System

سیستم مدیریت دانش نگهداری و تعمیرات

نرم افزار مدیریت دارایی های سازمان EAM مبتنی بر مدیریت دانش نگهداری و تعمیرات تجهیزات MKMS
براساس استاندارد ISO 55001 و ISO 14224 و استقرار آن مبتنی بر مدل Step Modules Model

بزرگ ترین خطر در هنگام تغییر، خود تغییر نیست، بلکه عمل کردن با منطق دیروز است! "پترو دراکر"

ارزیابی (ممیزی) مدیریت دارایی های فیزیکی و سیستم نگهداری و تعمیرات

مشاوره و طرح ریزی فرآیندها و سیستم مدیریت نگهداری و تعمیرات

برگزاری کلاس های آموزش سیستم مدیریت نگهداری و تعمیرات

طرح ریزی شاخص های کلیدی عملکرد فرآیندهای نگهداری و تعمیرات



مدل تعالی مدیریت دانش نگهداری و تعمیرات صنایع ایرانی iMKMS >>>

مدیریت اطلاعات و داده های قابلیت اطمینان
(EAM - MKMS - CMMS)

مهندسی نگهداری و تعمیرات

قابلیت اطمینان بهینه سازی تحلیل خرابی ها مدیریت ریسک هزینه های نت

استانداردسازی نگهداری و تعمیرات

مدیریت ابزار کالیبراسیون پایش وضعیت پیش تعمیرات اساسی بازرسی فنی نگهداری پیشگیرانه نت اپراتوری HSE پیمانکاران نت شاخص های نت

قطعات یدکی ساخت قطعات تامین قطعات مدیریت موجود گردش کارهای نت کارگاه های تعمیراتی خطوط تولید / عملیات انبار قطعات یدکی نظام آراستگی

آموزش، شایستگی، مهارت، تخصص، ساختار سازمانی، تامین نیروی انسانی، شرح وظایف

نیروی انسانی

مدیریت دارایی های فیزیکی ISO 55001
(استراتژی، اهداف، برنامه ها)

* Ali Moradi / 09125908046

مشاوران تدبیر پرداز آیوژه (سهامی خاص)
عضو رسمی انجمن مشاوران مدیریت ایران (IMCA)
مشهد - بلوار معلم بین معلم ۱۲ و ۱۴ پلاک ۲۹۰ کدپستی ۹۱۸۸۸۶۶۶۱۳
تلفن ۰۵۱-۳۶۰۹۷۲۴۴ فاکس ۰۵۱-۳۶۰۹۷۲۴۳
www.mtpa.ir info@mtpa.ir
نت بهره ور فراگیر www.irantpm.ir
نرم افزار نگهداری و تعمیرات www.cmms.ir
مدیریت دانش نگهداری و تعمیرات www.mkms.ir

مدیرعامل - ناصر جلالی ۳۳۹۳ ۳۱۶ ۰۹۱۵
مدیر فناوری اطلاعات - مهدی محمدی ۴۸۴۰ ۵۲۱ ۰۹۱۵
مدیر مهندسی - بهزاد غلامزاده ۵۵۷۲ ۱۰۲ ۰۹۱۵





نقشه راه استقرار نرم افزار مدیریت دارایی های فیزیکی سازمان

Road Map of Implementation EAM – CMMS – MKMS



ترتیب مراحل پیاده سازی ماژول های نرم افزار بر اساس مدل Step Modules

