

Chapter 2

Concepts and Requirements

This document is basically a summary of chapter 2 in the textbook *Reliability of Safety-Critical Systems* by Marvin Rausand.

2.1 Subsystems, Groups, Channels, and Elements

A SIS has at least three subsystems:

1. A sensor subsystem with one or more sensors that are installed to detect a possible undesired event in the EUC and send signals to the logic solver subsystem.
2. A logic solver subsystem with one or more logic solvers that receive the signals from the sensor subsystem, interpret these signals, and decide which actions should be taken.
3. A final element subsystem with one or more actuating elements (e.g., valves, circuit breakers, motors) that take a prescribed action to prevent harm.

The subsystems of a SIS are illustrated in Figure 2.1. Each subsystem may have one or more voted groups of channels. A *channel* is a structure of one or more elements and can independently perform a channel safety function. In Figure 2.1, the block denoted "pressure transmitter" is a channel that should (i) detect when the pressure goes beyond acceptable limits and (ii) send a signal to the logic solver sub-system. The block called "shutdown valve" in Figure 2.1 is also a channel, that upon signal from the logic solver subsystem shall shut and stop the flow in, for example, a pipeline. The channel "shutdown valve" has typically at least two elements: a pilot valve and the shutdown valve.

The main elements of a channel of a logic solver subsystem are illustrated in Figure 2.2. A logic solver subsystem is usually a complex system with several channels, a range of elements, and a lot of software. The term element is the lowest level of indenture that is considered in this course. A subsystem may have several types of channels, with different functions, as indicated

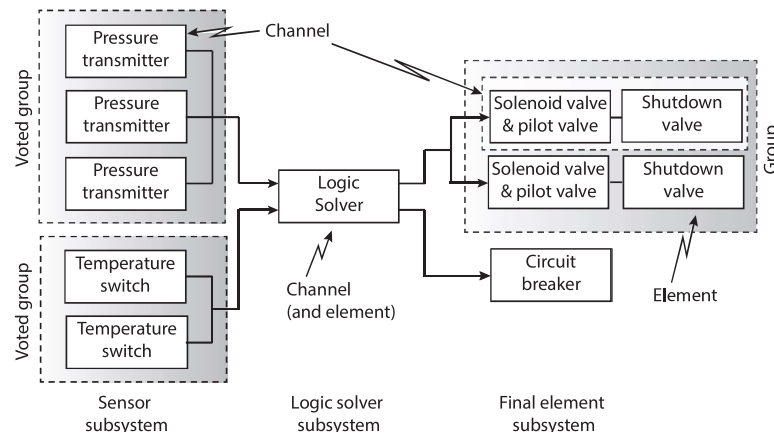


Figure 2.1: Subsystems, groups, and channels of a SIS (Rausand, 2014)

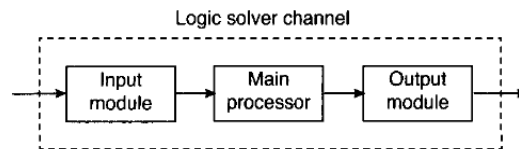


Figure 2.2: Elements of a channel (Rausand, 2014)

in Figure 2.1. Similar channels with the same function are often called a *voted group*. In Figure 2.1, the sensor subsystem has two voted groups: one voted group with three pressure sensor channels and one voted group with two temperature switch channels.

2.2 Redundancy

Redundancy means having two or more items, such that if one item fails, the system can continue to function by using the other item(s). This design principle is also referred to as fault tolerance. Redundancy can be implemented in many different ways. Two main categories are:

1. *Active redundancy.* All the redundant items are actively performing the duties. If the items carry a load, they share the load (e.g., pumps that should supply a given volume of a fluid).
2. *Standby redundancy.* One or more items perform the duties, while the rest of the items are in standby, waiting to be put into operation if one of the active items fails. While in standby, the items can be in cold standby or partly loaded standby. Items in cold standby are usually considered to be as-good-as-new when activated. The items may sometimes be activated and de-activated on a scheduled basis. Standby redundancy is also called dynamic redundancy.

Redundancy can further be categorized as:

- *Hardware redundancy* can be implemented by installing two or more items that can perform the same, or a similar, safety function. The redundancy can be implemented on element level, channel level, voted group level, subsystem level, and even on SIS level.
- *Software redundancy* is sometimes implemented by having two or more software routines, each written by independent coding teams and developed to give the same output for the same input. If there is no fault, the modules produce identical outputs at all times, but if one of the outputs differs, the software either has an undetected bug or the hardware it is running on has failed. In this case, the routine producing the deviating output is ignored and an error signal is sent.

✎ **kooon voted structure:** A structure of elements that is functioning when k -out-of- n channels are functioning, and which fails when $n - k + 1$ or more of its channels fail.

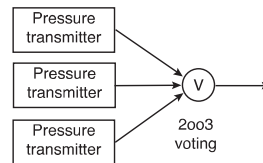


Figure 2.3: Voting of three channels (pressure transmitters) (Rausand, 2014)

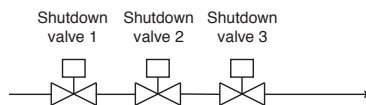


Figure 2.4: Three redundant shutdown valves (Rausand, 2014)

A 2oo3 voting of a group of three pressure transmitters is illustrated in Figure 2.3. In this figure, the voting is drawn as a circle with a V (for voting), but in most cases, the voting is done by the logic solver subsystem. The voted group of the three pressure transmitters is functioning when at least two of the transmitters are able to detect and transmit signal when the pressure goes beyond the acceptable limits. When the logic solver subsystem receives signals from at least two transmitters, the signals are treated and a decision about action is made. The channels may also be implemented to have a specific voting without any treatment by the logic solver subsystem. Such a system is illustrated in Figure 2.4 where three shutdown valves are installed in physical series in a pipeline. Each of the valves can stop the flow in the pipeline and we therefore have a 1oo3 voting with respect to the safety function “stop flow”.

2.3 Hardware Fault Tolerance

The concept hardware fault tolerance (HFT) is used to denote the ability of a hardware subsystem to continue to perform a required function in the presence of faults or errors. The HFT is given as a digit, where $HFT = 0$ means that if there is one fault, the function (e.g., to measure pressure) is lost. $HFT = 1$ means that if a channel fails, there is one other channel that is able to perform the same function, or that the subsystem can tolerate one failure and still be able to function. A subsystem of three channels that are voted 2oo3 is functioning as long as two of its three channels are functioning. This means that the subsystem can tolerate that one channel fails and still function as normal. The hardware fault tolerance of the 2oo3 voted group is, therefore, $HFT = 1$. A *koo* n voted group is functioning if at least k of its n channels are functioning and can tolerate up to $n - k$ channel failures without losing its ability to function. The HFT of a *koo* n voted group is therefore $n - k$. Note that hardware in this context covers the entire channel, included embedded software (e.g., for a “smart” transmitter).

2.4 Modes of Operation

Modes of operation were introduced in Chapter 1. IEC 61508 defines three modes of operation: low-demand mode, high-demand mode, and continuous mode but combines the two last modes and refers to these as the high-demand/continuous mode. IEC 61511, on the other hand, distinguishes between two modes of operation: demanded mode and continuous mode. There is a significant difference between a SIF that operates in demanded mode and one that operates in continuous mode. A SIF in demanded mode is passive in the sense that it does not perform any active function during normal operation but is an add-on to the EUC and is only called upon when something goes wrong, or starts to go wrong. A SIF that operates in continuous mode, on the other hand, plays an active role in the control of the EUC and a hazardous event will occur almost immediately when a dangerous failure of the SIF occurs.

For systems operating in demand mode IEC 61508 splits the demanded mode into two sub-modes:

- *Low-demand mode.* For this mode, the SIF is only performed when a demand occurs, in order to bring the EUC into a specified *safe state*, and where the demand rate (i.e., frequency of demands) is no more than once per year.¹ When the system is in continuous operation, this means that the demand rate is $\lambda_{de} < 1.15 \cdot 10^{-4}$ per hour. In low-demand mode, the EUC is usually kept in a safe state by an EUC control system and the SIF is called upon only when the EUC control fails, as illustrated by a simple event tree in Figure 2.5.
- *High-demand mode.* For this mode, the SIF is only performed on demand, in order to transfer the EUC to a specified safe state or to keep the EUC in a safe state, and where the

demand rate is greater than once per year. When the EUC is in continuous operation, this means that the demand rate is $\lambda_{de} > 1.15 \cdot 10^{-4}$ per hour. The mean number of demands per year will be one or more.

A safe state is a state of the EUC, whether the system is operating or shut down, such that an undesired event cannot occur. The safe state must be achieved in a timely manner. In the process industry, the time allowed to bring the process to a safe state is called the *process safety time*.

A demand is defined as follows:

■ *Demand*: An event or a condition that requires a SIF to be activated (i) to prevent an undesired event from occurring or (ii) to mitigate the consequences of an undesired event. In the process industry, a demand is also called a *process upset* or a *process deviation*.

Demands are most often regarded as shocks that occur without any significant duration, but in some cases it is also relevant to assume a certain *demand duration*, either as a specific time interval or as a random variable. An example of a SIF with a prolonged duration is an automatic fire extinguishing system. To perform its safety function, the final element subsystem consisting of fire pumps must start and survive as long as the fire lasts.

2.5 Testing of Safety-Instrumented Functions

A SIS is often a passive system that is activated only when a demand occurs. Failures may therefore occur and remain hidden until the system is demanded or tested. There are two main categories of testing:

- *Proof-Testing*. To verify that a SIS is able to perform its SIFs, the system is usually proof-tested at regular intervals of length τ . The time interval between two consecutive proof tests is often called the proof test interval. Dangerous failures detected by proof-testing are called dangerous undetected (DU) failures.
- *Diagnostic testing*. A diagnostic test is an automatic partial test that uses built-in self-test features to detect failures. Dangerous failures detected by a diagnostic test are called dangerous detected (DD) failures. The identified faults are announced as alarms, locally at the equipment and in the control room.

2.6 Safety Integrity Levels (SILs)

The term safety integrity is used as a performance measure for a SIF.

■ *Safety integrity*: Probability of a SIS satisfactorily performing the specified SIFs under all the stated conditions within a stated period of time.

IEC 61508 does not specify detailed probability values but divides the requirements into four safety integrity levels, SIL1, SIL2, SIL3, and SIL4, with SIL4 being the most reliable and SIL 1 being the least reliable. To demonstrate that the requirements for a specific SIL are fulfilled, we must verify that the requirements are met for: (i) Hardware safety integrity, (ii) Software safety integrity and (iii) Systematic safety integrity.

The hardware safety integrity requirements concern the hardware reliability of the SIS and is the main focus of this course. The hardware safety integrity is split into two categories of requirements:

- *Quantitative Reliability Requirements*. This part requires that the reliability of the SIF is analyzed and quantified. Two different reliability measures are used:
 - The average probability of (dangerous) failure on demand (PFDavg)
 - The average frequency of dangerous failures per hour (PFH)
- *Architectural Constraints*. The reliability quantification is connected with uncertainty of several types, such as completeness uncertainty, model uncertainty, and parameter uncertainty. In addition to the requirements for the quantified reliability, IEC 61508 also gives requirements for the robustness of the structure of the SIS. These requirements are given as architectural constraints and set restrictions to the designer's freedom to choose hardware architecture on the basis of PFDavg and PFH calculations alone

There are two routes for how to determine the architectural constraints in IEC 61508:

1. Route 1_H
2. Route 2_H

where the following parameters are of interest:

- Minimum hardware fault tolerance (minimum HFT)
- Category A and B
- Safe failure fraction (SFF)

Architectural constraints are applied at the sub-system level (sensor system, logic solver and final elements), one by one, so that in the end the whole SIF is covered.

2.6.1 Minimum hardware fault tolerance

Hardware fault tolerance was defined as the number of dangerous failures tolerated before the sub-system loses its safety function. Depending on the required SIL we can define a minimum HFT:

■ *Minimum HFT*: The HFT mandated by the architectural constraints in light of the SIL requirement

2.6.2 Safe failure fraction (SFF)

Safe failure fraction (SFF) is a measure of how safe the component responds in the presence of faults. SFF is computed by the following formula:

$$\text{SFF} = \frac{\lambda_S + \lambda_{DD}}{\lambda_S + \lambda_{DD} + \lambda_{DU}} \quad (2.1)$$

where λ is the failure rates and the subscripts give the failure category, S for safe, DU for dangerous undetected, and DD for dangerous detected.

2.6.3 Type A and Type B

Type A or type B are two categories used to distinguish proven/low-complexity components from unproven/more complex components. A component is classified as type A if ALL the following criteria are fulfilled:

- Failure modes of the element (and all its constituent components) are well defined
- The behaviour of the element under fault conditions can be completely determined
- There is sufficient dependable failure data to show that the claimed rates of failure for DD and DU failures are met

An element is type B if one or more of the above criteria are not met.

2.6.4 Route 1_H and 2_H

There are two routes for how to determine the architectural constraints in IEC 61508:

1. Route 1_H: Determining minimum HFT with SFF
2. Route 2_H: Determining minimum HFT without SFF. This route can only be applied when extensive field data is available. The approach does not allow only point-values for PFD or PFH, but requires also that the confidence levels for the λ values are determined.

Table 2.1: Minimum HFT-SFF-SIL relationship

	minimum HFT with type A			minimum HFT with type B		
	0	1	2	0	1	2
<60%	SIL1	SIL2	SIL3	-	SIL1	SIL2
≥60%, <90%	SIL2	SIL3	SIL4	SIL1	SIL2	SIL3
≥90%, <99%	SIL3	SIL4	SIL4	SIL2	SIL3	SIL4
≥99%	SIL3	SIL4	SIL4	SIL3	SIL4	SIL4

Route 1_H is discussed in the following where the minimum HFT is defined with basis in

- The SIL requirement
- The safe failure fraction (SFF) of subsystem component(s)/elements.
- The component category, type A or type B, defined on the basis of system complexity and maturity

Table 2.1 shows the minimum HFT-SFF-SIL relationship is proposed for subsystems of identical components:

Example: A smart sensor is often classified as type B. Assume that the SFF has been calculated to 85%. If a SIL 2 requirement is specified, it is necessary to select an architecture with minimum HFT of 1. This could be a 1oo2 architecture or a 2oo3.

2.7 Systematic Safety Integrity

The systematic safety integrity is specified by qualitative requirements.

■ *Systematic failure*: Failure, related in a deterministic way to a certain cause, which can only be eliminated by a modification of the design or of the manufacturing process, operational procedures, documentation or other relevant factors

Systematic safety integrity needs an extended examination of the design, production, and test procedures of both hardware and software. The higher the SIL claimed, the more detailed the examination has to be, and suppliers have to provide the required evidence.

2.8 Reliability metrics

2.8.1 Probability of Failure on Demand

The probability of (dangerous) failure on demand, $PFD(t)$ is the probability that the SIS has a dangerous fault and that it is not able to perform its SIFs at time t . The notion probability of failure on demand may indicate that we are dealing with a conditional probability, given that a demand has occurred. This is not correct and $PFD(t)$ may be expressed as

$$PFD(t) = \Pr(\text{The SIS is not able to perform its SIF at time } t) \quad (2.2)$$

irrespective of whether a demand occurs or not. If a demand should occur at time t , $PFD(t)$ is the probability that the SIS fails to perform its SIF. In many cases, it is not necessary to determine the PFD as a function of time and we can suffice with an average value. If the SIF is proof tested after regular intervals of length τ and the system is considered to be as-good-as-new after each proof test, the long-term average probability of failure on demand can be expressed as

$$PFD = \frac{1}{\tau} \int_0^{\tau} PFD(t) dt \quad (2.3)$$

2.8.2 Average Frequency of Dangerous Failures per Hour

For SIFs that are operated in high-demand or continuous mode, IEC 61508 requires that the reliability is specified by the average frequency of dangerous failures (PFH) where the frequency is given as number of dangerous failures per hour. The abbreviation PFH is retained from the previous version of IEC 61508 where the metric was called “average probability of (dangerous) failure per hour.” The idea behind using the PFH as a reliability metric is that demands will occur so often that when a dangerous failure of the SIF occurs, it is most likely that a demand will occur and a hazardous event will be manifested before we can bring the EUC to a safe state.

2.8.3 Reliability Metrics and SIL

To fulfil the requirements of a safety integrity level, a SIF in low-demand mode must have a PFD in the corresponding interval specified in Table 2.2. Similarly, a SIF in high-demand or continuous mode must have a PFH in the corresponding interval specified in Table 2.2.

2.9 Hazardous Event

A hazardous event occurs when a SIF fails when a demand for the SIF occurs. If several barriers are implemented with corresponding SIFs, there will be a sequence of hazardous events with

Table 2.2: SIL requirements vs PFD/PFH

SIL	Low demand mode of operation (Average probability of failure to perform its design function on demand)	High demand mode of operation (Average probability of failure per hour to perform its design function)
4	$10^{-5} \leq \text{PFD} < 10^{-4}$	$10^{-9} \leq \text{PFH} < 10^{-8}$
3	$10^{-4} \leq \text{PFD} < 10^{-3}$	$10^{-8} \leq \text{PFH} < 10^{-7}$
2	$10^{-3} \leq \text{PFD} < 10^{-2}$	$10^{-7} \leq \text{PFH} < 10^{-6}$
1	$10^{-2} \leq \text{PFD} < 10^{-1}$	$10^{-6} \leq \text{PFH} < 10^{-5}$

higher and higher severity, and in the ultimate situation an accident will occur.

When several barriers are in place to prevent an accident, the last barrier is called an ultimate safety barrier because demands that pass this safety barrier will affect assets and lead to an accident.

For a SIF that is operated in demanded mode, a hazardous event can occur in two different ways:

1. A demand occurs while the SIF has a dangerous fault (i.e., either a DD or a DU fault)
2. A dangerous failure of the SIS occurs while a demand situation is present.

For a SIF that is operated in continuous mode, a hazardous event occurs more or less immediately when a dangerous failure of the SIF occurs.

2.9.1 Hazardous Event Frequency (HEF)

A SIF that operates in demanded mode should ideally stop all demands that it has been installed to stop, but because it is not 100% reliable, some demands will not be stopped and will create hazardous events. The frequency of hazardous events will hence depend on the frequency of the demands, λ_{de} and the reliability of the SIF, and is given by

$$\text{HEF} = \text{PFD}_{\text{avg}} \cdot \lambda_{de} \quad (2.4)$$

Formula (2.4) is only correct when the demand is a shock with a negligible duration. When the demand duration is not negligible and has a mean demand duration (MDD), the HEF is approximately given by:

$$\text{HEF} = (\text{PFD}_{\text{avg}} + \bar{\lambda}_{\text{SF}}^* \text{MDD}) \lambda_{de} \quad (2.5)$$

where $\bar{\lambda}_{\text{SF}}^*$ is the average dangerous failure rate of the SIF when the demand is present.

Example

Consider a fire pump. If a fire occurs there is a probability, say $\text{PFD}_{\text{avg}} = 2\%$ that the fire pump will not start upon the demand, i.e., the fire. Then assume that the mean duration before the fire has been extinguished is $\text{MDD} = 3$ hours and $\bar{\lambda}_{\text{SF}}^* = 0.01$ per hour. Further a fire occurs every 20 years. Thus:

$$\text{HEF} = (\text{PFD}_{\text{avg}} + \bar{\lambda}_{\text{SF}}^* \text{MDD}) \lambda_{\text{de}} = (0.02 + 0.01 \cdot 3) 1/20 = 0.05 \cdot 0.05 = 2.5 \cdot 10^{-4}$$

which means that a fire where which is not extinguished occurs every 400 years. Note that the time unit for PFD_{avg} and MDD is hours, whereas the time unit for the fire frequency is years. In the example this is no problem because $\bar{\lambda}_{\text{SF}}^* \text{MDD}$ has no unit, but in general all parameters should have the same unit.

2.10 Allocation of Safety Functions to Protection Layers

The fire-pump example calls for an approach to determine if this is good enough. That is we need:

1. To allocate safety functions to protection layers
2. To determine required SIFs
3. To determine the required SIL for each SIF

The safety function for the fire pump is to extinguish the fire and in more explicit terms the required SIFs are:

1. Start upon demand
2. Continue to pump water as long as required during the fire
3. Do not pump water when not asked for

To determine the required SIL for each SIF we need to consider the required risk reduction. Without the fire-pump the risk is essential a function of λ_{de} and the severity of the fire, say $\mathbf{S}$. Often this function is a product of the factors, i.e.,:

$$R = \lambda_{\text{de}} \mathbf{S} \quad (2.6)$$

Risk acceptance criteria for a system or an activity may be defined by using a risk metric. One such criterion is $\text{FAR} < 10$ where FAR is the fatal accident rate which is the expected number of

fatalities in a defined population per 100 million hours of exposure. (If 1 000 persons are working 2000 hours per year during 50 years, they are exposed in 100 million hours).

If R in equation (2.7) is the expected number of fatalities per 100 million hours of exposure, and $R = \lambda_{de} \mathbf{S} > 10$ risk reduction is required. And more explicitly we need to allocate SIL to the SIFs such that

$$R = \lambda_{de} \mathbf{S} (\text{PFD}_{\text{avg}} + \bar{\lambda}_{\text{SF}}^* \text{MDD}) < 10 \quad (2.7)$$

Other relevant risk metrics could be:

- *Individual risk per annum (IRPA)*. The probability that an individual will be killed due to a specific hazard or by performing a certain activity during one year's exposure.
- *Potential loss of lives (PLL)*. The expected number of fatalities per year for an activity, for example number of fatalities in Norwegian rail transportation per year.

2.11 Safety lifecycle

The safety life cycle outlines a sequential pathway from initiation of a new system till it is installed and eventually removed.

■ *Safety lifecycle*: An engineering process designed to manage the design and operation of safety-critical systems, in light of requirements in standards like IEC 61508 [Slightly modified version of definition from the textbook].

2.12 Safety lifecycle phases

IEC 61508 suggests the following life cycle phases to achieve the desired level of functional safety:

- Preparation (familiarization)
- Analysis (risk assessment and overall design specification)
- Implementation (realization of SIS, based on design specification)
- Operation (follow-up of SIS, including maintenance, modifications and eventually de-commissioning)

Parallel phases:

- Planning, auditing, verification, and validation,++

Figure 2.5 illustrates the safety life cycle.

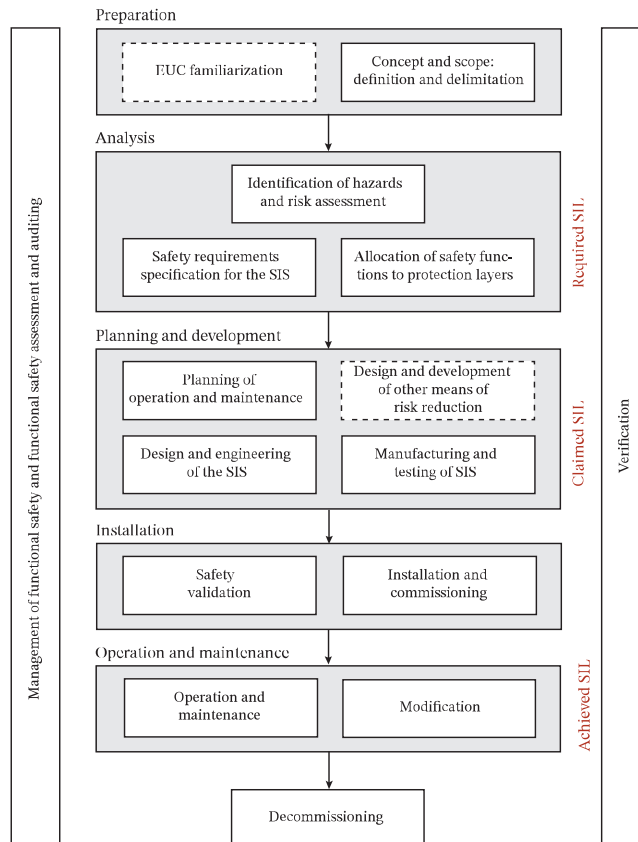


Figure 2.5: Safety life cycle (Rausand, 2014)

Figure 2.5 also emphasise the different meaning of SIL, i.e., we distinguish between:

- Required SIL (“SIL requirement”): The SIL level that is required for a SIF on the basis of the risk analysis
- Claimed SIL: The SIL that can be claimed or predicted for a specific SIF on the basis of design and results from analyses, checks, and tests before the system is installed
- Achieved (or actual/experienced) SIL: The SIL that is claimed based on operational experience and failure reporting.

In relation to the phases in the safety life cycle we have:

- Analysis phase:

- Focuses on the specification of functional safety and the *required SIL*
- Planning and development phase:
 - Focuses on the demonstration of functional safety requirements and the *claimed SIL*
- Operation and maintenance phase:
 - Focuses on the demonstration of of functional safety requirements and the *achieved SIL*

Bibliography

Rausand, M. (2014). *Reliability of Safety-Critical Systems*. Wiley Series in Probability and Statistics. Wiley.