

Annex B

(normative)

Management of faults for safety-related functions

B.1 Introduction

This annex concerns the ability of the system, subsystem or equipment to continue to meet its specified safety requirements in the event of random hardware faults and, as far as reasonably practicable, systematic faults.

The annex is structured as follows:

- B.2 General concepts
- B.3 Effects of faults

Within B.3, subclauses (B.3.x) are structured in accordance with the headings of the Technical Safety Report Section 3 (as defined in 7.2):

- B.3.1 Effects of single faults
- B.3.2 Independence of items
- B.3.3 Detection of single faults
- B.3.4 Action following detection
- B.3.5 Effects of multiple faults
- B.3.6 Defence against systematic faults

B.2 General concepts

B.2.1 Detection and negation times

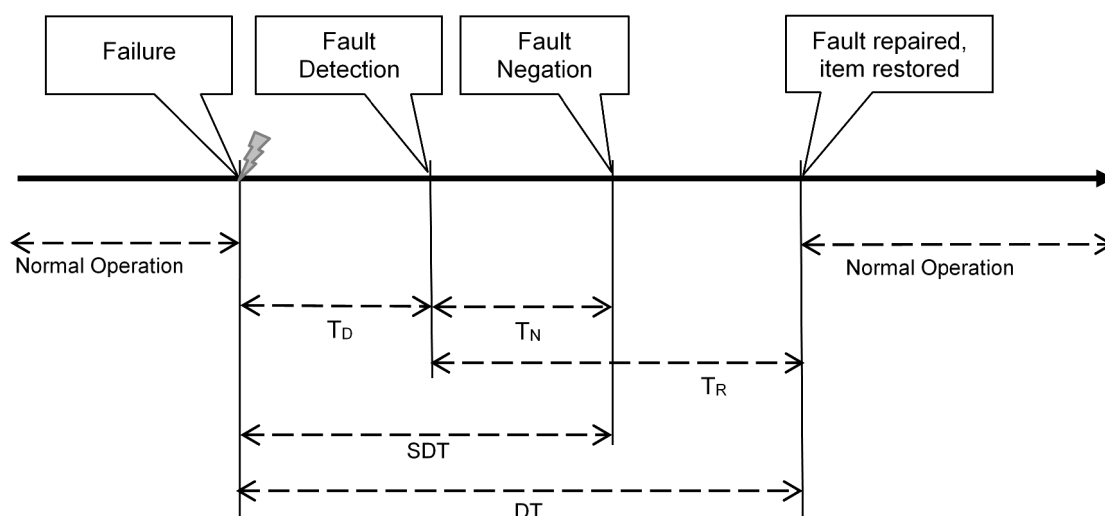


Figure B.1 — Detection and negation times

After a failure within an item has appeared, normal operation will be restored only when the following two events have taken place:

- the fault is detected and negated (this means a safe state is enforced);
- the fault is repaired and the item restored.

Detection time (T_D)

The fault detection time is the test interval in the case of detection by automatic and periodic testing, or the maintenance interval in the case of detection by staff. If no measures for detection are specified, implemented or planned, the fault detection time is the lifetime of the system.

In the case of equipment in storage, it can be linked to the interval between periodic testing by maintenance personnel.

Where fault detection is achieved through operational use, the fault detection time is mean time between uses.

While in a reliability context the detection time is usually not significant, this time becomes important in the safety context. Safety-critical applications might not rely on self-tests or similar measures, but the detection and negation might be performed independently of the item. Sufficient failure detection and negation mechanisms shall be demonstrated in the safety case.

Negation time (T_N)

The negation time is the time taken for the relevant part of the system to enforce a safe state, either automatically or by human action.

Detection plus negation times are defined here as Safe Down Time (SDT).

The Safe Down Time is the time required to detect and negate the last fault before the output becomes hazardous with respect to the defined TFFR. The SDT can have different meanings for different

architectures: in the case of composite fail-safety the system remains in a safe state, but an additional failure could be hazardous (if the first one is not negated); in the case of reactive fail-safety the system could result in a transient permissive output. See Figure B.4 for SDT in the context of composite fail-safety and Figure B.5 for SDT in the context of reactive fail-safety.

Repair time (T_R)

After fault detection, there is the time to repair (logistic time plus actual repair time: fault finding, repair, exchange, check up to restore equipment into operation).

Time between a failure and the next restore into operation is called Down Time (DT).

In a safety context generally the actual repair time can be neglected, as control measures are taken to enter and stay in an enforced safe state.

NOTE In some cases, (especially for components with low safety requirements), the negation occurs through the restoration of the item, therefore SDT = DT (no safe degraded operation exists, degraded operation is unsafe, or at least less safe than prior to the failure).

B.2.2 Composition of two independent items

For two items A and B implementing a function F, in the hypothesis that dangerous failures can be detected, and the two items are restorable, the following basic formula for the asymptotic hazardous functional failure rate (FFR) may be used:

$$FFR \approx \frac{FR_A}{SDR_A} \times \frac{FR_B}{SDR_B} \times (SDR_A + SDR_B) \quad (B.1)$$

NOTE The formula is an approximation that can be derived for a simple 2oo2 active redundant system with restorable elements (see e.g. IEC 61165:2006, Fig. B.7). In that case the down time (DT) would be used instead of the safe down time (SDT). By using SDR (inverse of SDT), it represents an upper bound (being therefore more conservative).

The formula becomes, when A and B are identical:

$$FFR \approx \frac{2FR^2}{SDR} = 2FR^2 \times SDT = 2FR^2 \times (T_D + T_N) \quad (B.2)$$

Validity of these formulas relies on the following assumptions which need to be justified for each application of the formulas:

- the rates are constant over time,
- any faults which could be hazardous if combined with a second fault are detected and negated,
- hazardous common cause failures (CCF) can be neglected,
- $FR \times SDT \ll 1$

EXAMPLES

Taking two identical items with a MTBF of 10 000 hours ($FR = 10^{-4}$) and a SDT of 1 hour, then the resulting failure rate for the parallel system (AND combination in failure logic) is 2×10^{-8} per hour.

If SDT is 1 000 hours (e.g. detection by maintenance), then the result is only 10^{-5} per hour, which is only a factor of 10 better than the MTBF of a single item. Note however that this result is affected by large approximation error: in fact in this case the condition $FR \times SDT \ll 1$ is no longer respected.

If SDT would be the lifetime, then the gain would become even more marginal.

B.3 Effects of faults

B.3.1 Effects of single faults

Any safety-related function shall meet its TFFR in the event of random fault.

It is necessary to ensure that SIL 3 and SIL 4 systems remain safe in the event of any kind of single hazardous random hardware fault which is recognized as possible. Faults whose effects have been demonstrated to be negligible may be ignored. This principle, which is known as fail-safety, can be achieved in three different ways:

1) composite fail-safety

With this technique, each safety-related function is performed by at least two items. Each of these items shall be independent from all others, to avoid common-cause failures. Non-restrictive activities can progress only if the necessary numbers of items agree. A hazardous fault in one item shall be detected and negated in sufficient time to avoid a co-incident fault in a second item, in relation with the TFFR target.

Hazardous random faults that could be dormant, causing errors only under particular conditions, shall be detected by online periodical monitoring.

2) reactive fail-safety

This technique allows a safety-related function to be performed by a single item, provided its safe operation is ensured by rapid detection and negation of any hazardous fault (for example, by encoding, by multiple computation and comparison, or by continual testing). Although only one item performs the actual safety-related function, the checking/testing/detection function can be regarded as a second item, which shall be independent to avoid common-cause failures.

3) inherent fail-safety

This technique allows a safety-related function to be performed by a single item, provided all the credible failure modes of the item are non-hazardous.

To use this technique the following requirements shall be applied:

- The effect of every credible failure mode (and relevant combinations of failure modes) of each component identified in Annex C shall be identified, analysed, demonstrated as non-hazardous through physical tests, technical justifications or simulation. Dedicated failure tests might be necessary to confirm the effect of credible failure modes.
- Any failure mode which is claimed to be incredible (for example, because of inherent physical properties) shall be justified using the procedure defined in Annex C. In this case, the hazard occurrence can be considered negligible and a hazardous failure rate of zero may be assumed.

Inherent fail-safety can also be used for certain parts within Composite and Reactive fail-safe systems, for example to ensure independence between items, or to enforce shut-down if a hazardous fault is detected.

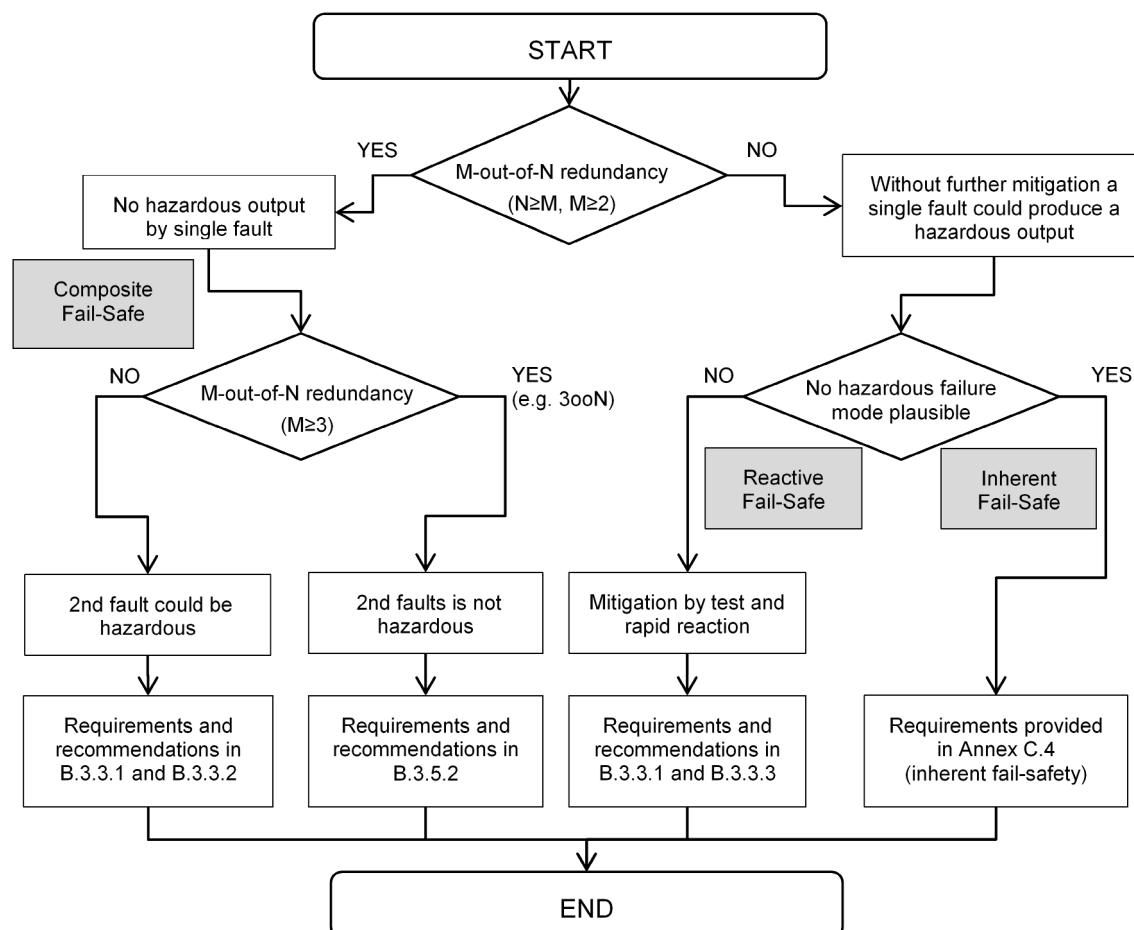
Whichever technique or combination of techniques is used, assurance that no single random hardware component failure mode is hazardous shall be demonstrated using appropriate structured analysis methods.

The component failure modes to be considered in the analysis shall be identified using the procedures defined in Annex C.

A top-down failure analysis method should be used, such as Fault Tree Analysis (FTA). This should be supported, if necessary, by a bottom-up method such as Failure Modes and Effects Analysis (FMEA). See also guidance given in Table E.6.

Failure analyses shall be qualitative and also quantitative where credible data are available. Random hardware failure rates, or probabilities of component failure, should be based on field data if possible. Apportionment of an overall component failure rate between its failure modes shall be justified in the analysis.

A flow-chart showing the application of the three different kinds of fail-safety, and related applicable provisions, is given in Figure B.2.



NOTE M-out-of-N redundancy means redundancy wherein at least m of the total n items function and therefore meet the requirements.

Figure B.2 — Single and multiple fault control

B.3.2 Influences between items

B.3.2.1 General requirements

In systems containing items whose simultaneous malfunction could be hazardous, independence between items is a mandatory precondition for safety concerning single random faults.

This is applicable to Composite Fail-Safety (independence between the two items performing the safety-related function) as well as to Reactive Fail-Safety (independence between the item performing the function and the checking/testing/detection function).

Appropriate rules or guidelines shall be fulfilled to ensure this independence. The measures taken shall be effective for the whole life cycle of the system. In addition, the system design shall be arranged to minimize potentially hazardous consequences of loss of independence caused, for instance, by a systematic design fault.

As an example, the various types of influence in a system consisting of two operating items are represented in Figure B.3. This figure can be extended to systems consisting of more than two operating items.

In accordance with Figure B.3, several types of influences can be identified:

— Internal influences

Type A: physical internal influences

Type B: functional internal influences

— External influences

Type C: physical external influence

Type D: functional external influences

Functional influences (types B and D) are influences via transmitted information (e.g. data or signals): this is the case when faulty information in one item influences another item in a hazardous manner. Protection against those influences requires specific measures to be defined on a case by case basis.

Types A and C are influences resulting from other physical causes, as further detailed in the following subclauses.

B.3.2.2 Type A for SIL 3 and SIL 4

If no physical connection exists between internal items of a system, there are no physical influences. Therefore, internal independence is achieved.

NOTE A physical connection is any medium between items, for example:

- galvanic connection;
- electromagnetic coupling.

This is a condition hardly reachable in reality, so measures shall be taken to avoid non-intentional physical internal influences.

These measures may be scaled in case of "primary" or "secondary" independence.

Independence of two items whose simultaneous malfunction could be hazardous is called "primary independence". Related items are referred to as "main items".

A N-out-of-M system consists of M independent main items. Each main item can have one or more so-called "additional items" checking the main item. The degree of independence between a main and an additional item is called "secondary independence".

Secondary independence may be less stringent than primary independence, since simultaneous malfunction between a main item and an additional item would not be hazardous, provided that any fault of a main item is detected before it can become hazardous through further faults of another main item.

The following recommendations are given for primary independence.

- 1) Measures should be taken to avoid non-intentional galvanic connections (protection of internal galvanic insulation). Insulation distances (creepage distances and clearances) should be dimensioned at least in accordance with the requirements for reinforced insulation of EN 50124-1.

- a) Insulation between lines on the same layer of a printed-circuit board.
 - b) Insulation between lines on different layers of a multilayer printed-circuit board.
 - c) Insulation between insulated wires in the same cable.
 - d) Insulation between insulated windings in the same transformer. Maximum temperature inside transformers should be limited (including fault conditions), to avoid carbonization.
 - e) Insulation between insulated items inside an opto-coupler. Maximum temperature inside opto-couplers should be limited (including fault conditions), to avoid carbonization.
- 2) Measures should be taken to avoid non-intentional effects via intentional connections (protection of internal interfaces).

Interfaces should be protected by means of devices with inherent fail-safe properties.

- 3) Measures should be taken to avoid non-intentional effects via electromagnetic coupling (protection against internal cross-talk).

Cross-talk between electronic networks should be prevented as follows.

- a) If different items are on the same printed-circuit board, they should be supplied by different power-supply networks. If not, then the impedance of the ground network should be sufficiently low to avoid cross-talk, even in the event of faults.
- b) If different lines on the same board need to be protected against cross-talk occurring between them, the necessary separation distance depends on the technology used, the coupling length and the coupling mechanism. This distance should be demonstrated for the normal operational mode by theoretical calculations and/or by practical measurements.
- c) If necessary to avoid coupling in the event of faults, additional measures (for example, shielding or doubling of distance) should be taken. Effectiveness should be demonstrated by theoretical calculations and/or by practical measurements.

The following simplifications to primary independence are given for secondary independence.

- 1) Insulation distances (creepage distances and clearances) should be dimensioned at least in accordance with the requirements for basic insulation of EN 50124-1.
- 2) Protecting devices do not require inherent properties. (Only a second fault may be able to inhibit the independence between a main item and an additional item).
- 3) Voltage monitoring is recommended. At least the voltage-monitoring (additional item) should not be influenced by the faults of the power-supply (i.e. main item).

B.3.2.3 Type C for SIL 3 and SIL 4

A physical external influence could cause a loss of independence between items.

NOTE These could be due to, for example,

- environmental stresses such as EMI, ESD, climatic, mechanical and chemical,
- the power supply, and
- the external inputs and outputs.

Measures shall be taken to avoid non-intentional physical external influences. Subclause 7.2 (Section 4 of the Technical Safety Report) contains requirements for external influences which shall be considered.

The following recommendations are given to provide physical external independence.

- 1) Measures should be taken to avoid non-intentional effects of EMI/ESD disturbing correct operation, in accordance with EN 50121-4.
- 2) The specified climatic conditions should normally be complied with. Measures should be taken to minimize the risk of the system being operated outside its specified climatic conditions.
- 3) Measures should be taken to avoid non-intentional effects by mechanical stresses disturbing the correct operation:
 - a) measures should be taken to ensure reliable correct operation in spite of mechanical stress-conditions;
 - b) protection should be compliant with EN 50125-1 and/or EN 50125-3 as appropriate.
- 4) Measures should be taken to ensure reliable correct operation under chemical stress-conditions.
- 5) Measures should be taken to avoid non-intentional operation under non-permitted power-supply voltages (protection of supply-voltages):
 - a) non-permitted supply voltages (outside data-sheet values for supplied systems, subsystems or equipment) should be disclosed by voltage-monitoring triggering an enforced safe state before hazardous situations are possible;
 - b) voltage-monitoring should operate correctly for the whole life cycle. Voltage-monitoring redundancy should be applied if disclosure of voltage-monitoring faults is not possible.
- 6) Measures should be taken to avoid non-intentional hazardous effects caused by external voltages across input and output ports disturbing the correct operation (protection of external interfaces):
 - a) worst-case external voltages should be assumed (process-voltages and all possible EMI-induced voltages on cables and lines);
 - b) clearances between live parts and exposed conductive parts/earth/circuits whose correct operation needs to be protected should be dimensioned in accordance with surge voltages specified in EN 50124-1;
 - c) creepage distances between live parts and exposed conductive parts/earth/circuits whose correct operation needs to be protected should be dimensioned at least in accordance with EN 50124-1 and considering maximum rated r.m.s. voltages during operation;
 - d) for dimensioning insulation, the larger distance (clearance or creepage distance) is decisive.

B.3.2.4 Type A and C for SIL 1 and SIL 2

For SIL 1/SIL 2 functions, the same recommendations as given for SIL 3/SIL 4 functions (see B.3.2.2 and B.3.2.3) apply, with the following exceptions:

- 1) the European Standard EN 50124-1 requirements for basic insulation may be applied, provided that failures related to rupture of insulation do not exceed the specified TFFR;
- 2) protecting devices may not require inherent fail-safety properties, provided that failures related to loss of protection do not exceed the specified TFFR;

- 3) voltage monitoring may not be present, provided hazardous faults related to power supply do not exceed the specified TFFR.

