

1.1 Understanding Risk and Risk Analysis



Nothing in life is safe ...
[DMTP, 1994]

Risk unavoidably accompanies us in every act of our life. Where we live, what we live in, and what we do determine the risks we can be submitted to. When crossing the road there is a risk of being injured, or even killed, by a car. When staying at home, an accidental action can start a fire. The risk of natural disasters is something we all face, although for some of us this risk is higher than for others. Numerous examples show that technical/industrial risks cannot be totally avoided. Etc., etc. ...

Awareness of the risk by the public in general and perception of how it compares to other risks determine society's attitudes about reducing it: whether we do anything about it and how much we are ready to do. Understanding risks and their causes is therefore of paramount importance in this regard. The knowledge of what makes a person, a community or an asset more vulnerable than another to potential risks is necessary to decide what could or should be made to reduce the risks in question.

Society as a whole or people take collective or individual actions to protect themselves against risks thus identified, and have been relatively successful in doing so. Reducing the risk of disease for example has been one of the greatest achievements of the last 150 years. Average life expectancy for someone born in Europe in the middle of the 19th century was only 35 years (in great part because of the high rate of deaths shortly after birth or during the early years of childhood due to infectious diseases). Nowadays in most developed countries it is over 70 years, but not higher than 50 years in the 40 poorest countries. Generally speaking, societies appear to become safer and less tolerant to risks as they become more technologically advanced. But some technological or other human advances also bring with them new or increased risks: the automobile has cost many lives, many people are afraid of the possible consequences of a nuclear accident, or of the dispersion in nature of genetically modified organisms, cancer cases seem to increase, new mortal diseases like aids (HIV virus) and BSE ("mad cow disease") have appeared. Energy supplies and other industrial activities in particular are held responsible for the emergence of new or enhanced hazards. The benefits of new technologies or scientific advances appear however to generally outweigh the risks they bring.

As the risks diminish from common events like infectious diseases or traditional industrial activities, the risks posed by extraordinary events like natural hazards or exceptionally severe technical accidents assume a greater significance. This raises the question to know what risks we are really facing and how catastrophic events compare with other, more familiar and better accepted, risks, as well as what could be done against them. In many cases, it is far more effective to prevent disasters from occurring beforehand than to recover from them afterwards. In developing countries, the United Nations Development Program (UNDP) is promoting the goal of sustainable development, and it is argued that disaster awareness considerations should be incorporated into all development programming and planning, to protect the development process and reduce the risk of wasting scarce resources.

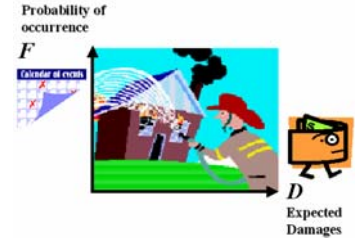
The aim of this learning unit is primarily to examine risk as a concept and present risk assessment techniques and their use in defining mitigation strategies: how can risks be assessed? and how can such results be used to make decisions about appropriate levels of safety?

But what do we mean exactly by “risk”? Within the context of risk analysis, “risk” refers to the possibility of injury, harm, or other adverse and unwanted effects. In other words, risk involves the occurrence, or potential occurrence, of some accident consisting of an event or sequence of events leading to undesirable consequences. For the Office of United Nations Disaster relief Coordinator (UNDRO), the term “risk” expresses the *expected losses* from a *given hazard* to a *given element at risk*, over a *specified time period* [UNDRO, 1979].

The definition of risk

From the above definitions, one can infer that “risk” is essentially a two-dimensional concept, these two dimensions being:

- the *probability of occurrence* or *frequency* (of the undesirable event), F , on one hand, and
- the *expected losses* or *damages* (resulting from the undesirable event, should this one occur), D , on the other hand.



Risk, a 2-D concept

This does not say anything however about how risk should be measured in practice. Depending on the circumstances, risk can actually be described and expressed in number of ways. A common method is to count all the people having – or susceptible to have - experienced the consequences (e.g. death, injuries) of a hazard over a defined time span and to divide this number by the total number of people exposed to this particular hazard. Thus, if the number of people who travel by train in any one year is ten million and then ten people are killed on average each year, then the annual risk of being killed in train travel is one in one million (10^{-6}). This implicitly means that a simple multiplicative law has in this case been used to combine the two dimensions of the risk R :

$$R = F \times D, \quad [1.1]$$

assuming for example that here on average $F = 2$ fatal train accidents/yr and $D = 5 \times 10^{-7}$ (expected relative life losses per train accident, i.e. 5 deaths per train accident on average for an exposed population of 10 millions people).

The division by the number of exposed people facilitates comparisons between different types and contexts of risks. This *average individual risk* is to be distinguished from the *societal risk*, which in the case of the above example would simply be 10 deaths per year, the expected total number of fatalities suffered by the considered human community in one year.

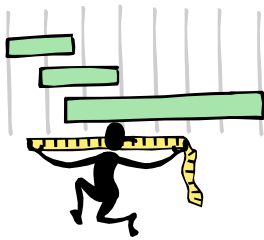
As another example, let us consider the case of a nuclear plant accident scenario that is expected to occur with a frequency of 1.1×10^{-5} per year. The consequences given the characteristics of this scenario would be 3.5×10^{-1} excess fatal cancers in an exposed population of 267'107 residents living in the plant Region of Influence (ROI). The societal risk from this potential accident would be 3.5×10^{-6} excess fatal cancers per year. The average individual risk in the ROI from this potential accident is defined as:

$$\frac{3.5 \times 10^{-6}}{267,107} = 1.3 \times 10^{-11} \text{ excess fatal cancers per year}$$

Societal and individual risks are well distinct notions that should not be mistaken one for the other. Let us assume that 300 people in a global population of 250 million are exposed to an average annual individual risk of 10^{-2} (1 death per 100 people); this represents a very serious individual risk for the exposed people, but a small societal risk for the whole population (3 deaths/yr). Now if we suppose that this same whole population can be affected by another kind of hazard resulting in 1 death per 100,000 people, the corresponding average individual risk (10^{-5}) is small but the societal risk (2,500 deaths/yr) is much greater than in the previous case.

The above examples show clearly that the definition of the population exposed greatly affects the assessment of risk. There is however no standard way of defining the population exposed to a risk. If in the example of the train travel the definition of the people exposed is relatively straightforward, in the case of the nuclear plant this task becomes already much more difficult. The risk is obviously higher for those who live nearest to the plant, and less for those who live further away. But where should the line be drawn to define the population exposed? Ten kilometers from the plant? A hundred? More? This means that any statistical expressions of risk need to be carefully defined and explained to be useful.

The measurement of risk



More generally, it should be acknowledged that risk remains somehow an abstract concept that cannot be defined in a univocal and indisputable manner. **Risk has no clear-cut unit.** It can as well be measured in lives lost over a given time period, as above, or in number of buildings experiencing heavy damages within a defined number of years, or in economic losses (in Swiss Francs, Euros, etc.) suffered during the same period, etc.

Moreover, even when a well-defined unit has been chosen, the accuracy of risk quantification still depends to a considerable extent on the amount of data available. The number of events on which information is known has to be large to be statistically significant. In addition, the quality or reliability of the data has to be adequate. These factors all pose problems for the risk assessor who has to identify “confidence limits” or range of uncertainty over any future risk figures offered. Some risks are obviously easier to quantify than others. The risks of damages from minor floods and small earthquakes, for example, are easier to predict than the ones resulting from catastrophic events of the same types, because they have been observed more often and there is more data on their occurrence and consequences. On the other hand, risks related to events that have not yet happened, such as the melt-down of a light water nuclear reactor with out-of-confinement radioactivity release for instance, have by definition no past statistics and so have to be estimated from probabilistic safety assessment methods of the type presented in the following chapters.

The presentation of risk figures

As there are different ways to measure risk, there is also no unique way to present risk figures. Tables of individual risks for example can be used to give approximate ranking of the respective probabilities of death from various causes, as shown in Table 1.1.

Table 1.1 Average individual risk figures (per year) [DMTP, 1994]

Hazard	Individual risk	One chance in
Smoking 10 cigarettes a day	5.00×10^{-3}	200
Heart disease	3.10×10^{-3}	320
Cancer	2.00×10^{-3}	500
All accidents	3.90×10^{-4}	2,560
Influenza	2.00×10^{-4}	5,000
Accident on the road (Europe)	1.25×10^{-4}	8,000
Leukaemia	8.00×10^{-5}	12,500
Earthquake (living in Iran)	4.35×10^{-5}	23,000
Accident at home	3.85×10^{-5}	26,000
Accident at work	2.30×10^{-5}	43,500
Drowning	2.10×10^{-5}	47,600
Radiation, working in radiation industry	1.75×10^{-5}	57,000
Homicide (in Europe; in USA: ~ 8 x more frequent)	1.00×10^{-5}	100,000
Civil aviation	5.20×10^{-6}	192,600
Water transport	3.90×10^{-6}	256,000
Accident on railway (in Europe)	2.00×10^{-6}	500,000
Earthquake (living in California)	5.00×10^{-7}	2,000,000
Hit by lightning	1.00×10^{-7}	10,000,000
Windstorm (northern Europe)	1.00×10^{-7}	10,000,000

Such tables give some idea of how disaster risks to an individual compares with other more common risks, and also how disaster risk may vary from place to place. The annual probability of being killed in an earthquake in Iran, for example, is obtained from the total number of people residing in this country killed by earthquakes between 1900 and 1990 (120,000) divided by 90 years. This gives an average of around 1,300 people killed annually. The population of Iran (in 2003: 68 million) averaged over the 90-years time period considered was less than 30 million, so the average probability of losing his life in an earthquake in Iran is given in Table 1.1 as one in 23,000. Of course, not everyone in Iran is equally at risk in this matter. Some parts of Iran are more seismic than others; those living in these zones are therefore more at risk. Moreover, people living in poorer quality houses have a greater probability of being killed than people who live in more seismic-resistant houses (*vulnerability* aspect).

Societal risks are for their part often described as $f:N$ curves. A $f:N$ curve plots the frequency f of events causing *greater than* a certain number N of fatalities. As an example, the $f:N$ curve for the risk resulting from the existence of a gas gathering pipeline system in a given region is presented in figure 1.1.

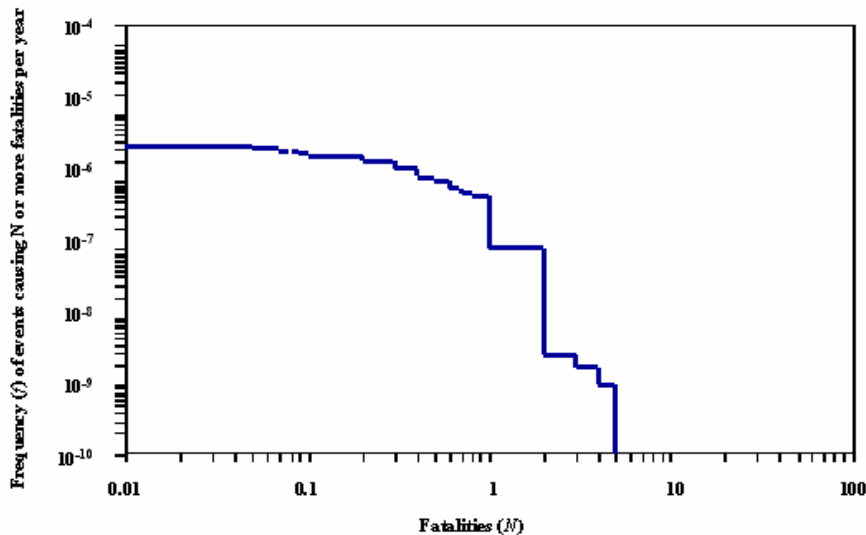


Figure 1.1 Societal risk ($f:N$ curve) related to a gas gathering pipeline system

It is interesting to compare with such graphical representation the societal risks posed by natural hazards with the ones resulting from technological accidents. As can be seen in figure 1.2, natural disasters greatly exceed technological accidents in their capability to cause massive loss of life (whereas it's the opposite for relatively small number of fatalities per event). This results from the fact that the scale of energy released during a natural event such as cyclone, flood, volcanic eruption or large earthquake, which may be equivalent to hundreds of atomic bombs, far outstrips any human-made source of energy.

The largest single-event life losses from a rapid-onset disaster to have occurred in the 20th century were from floods in China: an estimated 2 million people were killed in Northern China in flooding in 1956 and 1.4 million people were reportedly killed in a flood in 1931 on the Yangtze-Kiang river in China. The worst casualty rate from an earthquake in the same century was also in China, in Tangshan in 1976, when a quarter of a million people died [DMTP, 1994]. These exceptionally high death tolls result from extreme conjunction of very severe natural events, dense affected population and vulnerable communities. Fortunately, such disastrous conjunctions of circumstances are extremely rare; less severe situations resulting in lower numbers of lives lost happen more commonly.

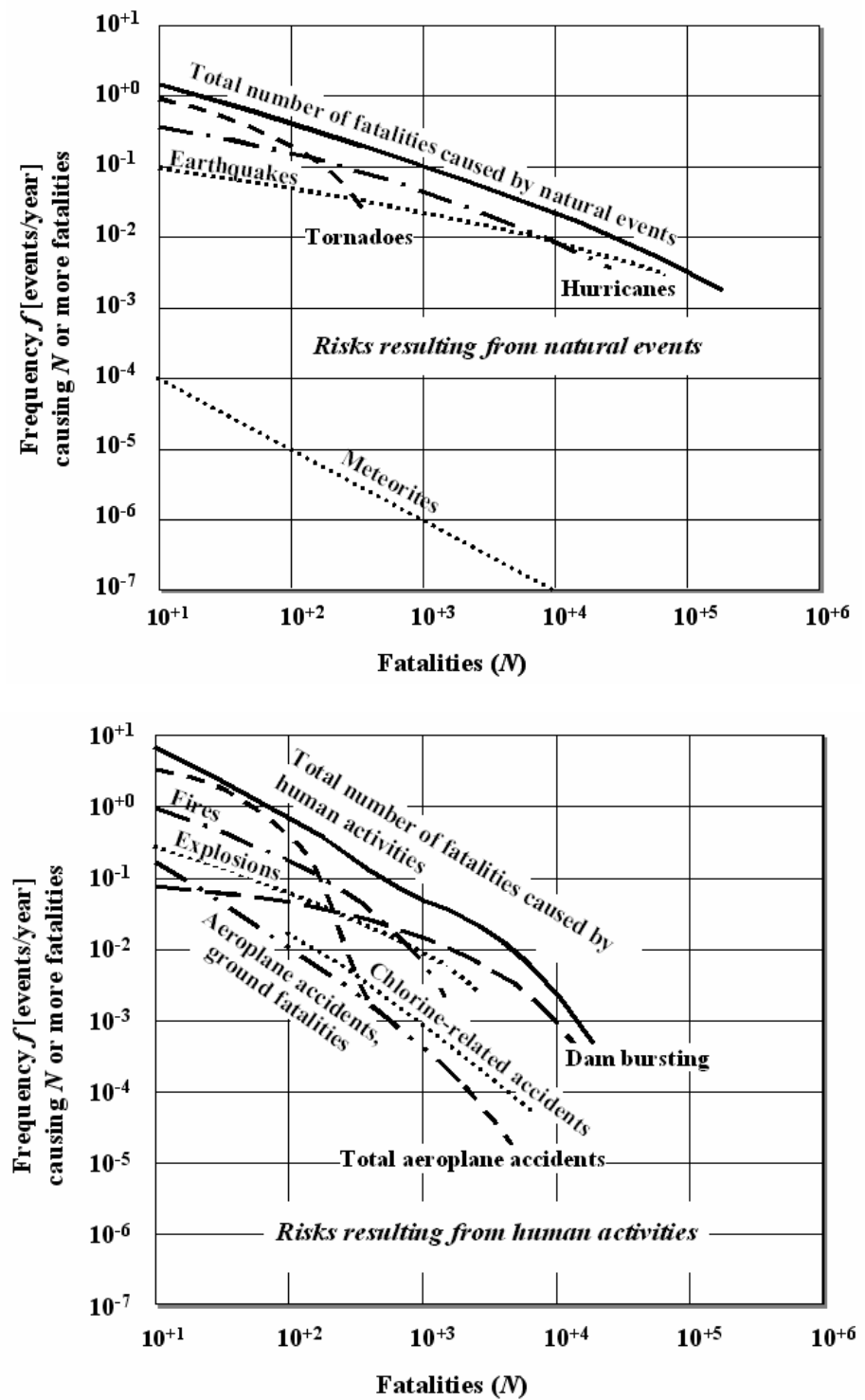


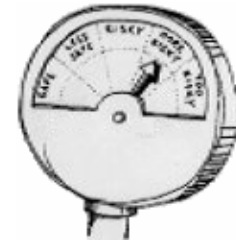
Figure 1.2 Societal risk curves related to natural and industrial hazards
(From WASH-1400, 1975)

According to the above graphs, an earthquake killing 10,000 people or more can for example be expected roughly once in a century on average.

These graphs thus tell us something about the levels and scale of risks from different sources, based on statistical records. More sophisticated studies are however required to more precisely assess risk and vulnerability of people or assets in given situations.

Risk assessments (or risk analyses) are conducted to estimate the frequency and the level of damage or injury that can be expected from the exposure to given danger sources, and ultimately to assist in judging whether the resulting evaluations are great enough to require increased design effort, management or regulation.

The assessment of risk



Engineers and other scientific experts have a moral obligation to conduct risk analyses to design, construct and operate the safest possible systems (with the lowest probability of system failure and minimal consequences if the system does fail) within the given set of engineering, regulation and environmental constraints. By performing a risk analysis, appropriate information may be obtained about the system to redesign it, or operate it differently, and thus lower the probability of occurrence of an accident or mitigate the ensuing consequences. Alternatively, such a study may also serve to show that the probabilities of occurrence or the expected damages are negligibly small.

Depending on the kind of hazard and on the objectives of the study, the effects of primary concern might be human casualties in the public or among professionally exposed people, induced diseases such as cancer or other debilitating illnesses, damages to buildings or equipments, ecological impacts such as species extinction, loss of habitat and other kinds of ecosystem damages, etc. Cost is frequently used to measure risk effects because many types of losses can be converted into economic cost; it is thus a practical “currency” for considering and comparing a wide range of effects. Effects of this kind are known as *tangible losses*. *Intangible impacts* – like social disruption consequences (psychological, socio-cultural) for example – are in principle equally or even more important parameters to consider in many cases. Nevertheless, they are often neglected in risk analysis procedures because of the difficulty of quantifying such effects.

Risk assessment involves an analysis, processing and combination of both theoretical and empirical data concerning: the probabilities of occurrence of inventoried hazards of various types and intensities, and the losses (both physical and functional) expected to affect each element at risk from the exposition to the aforementioned hazards (*vulnerability analysis*, see below). It ranges widely in scope and complexity, depending on the application, from simple and low-cost screening analyses to major analytical efforts that can require years of effort and a substantial budget. Contemporary risk assessments ordinarily rely on many branches of science, on the methods and knowledge of disciplines such as systems engineering and related technical areas, toxicology, epidemiology, other health and environmental sciences, etc.

In the industrial domain, if the initiation of the undesirable event (accident, disaster) is non-deliberate, such as a dysfunction arising because of a natural catastrophe or from naturally occurring phenomena, wear-out due to abrasion for example, then the problem is one of *safety engineering* or *engineering systems risk assessment*. This includes the study of: (1) the causes of a system component failure and the probability of such a failure over a period of time, (2) the way in which the probability of failure of a component of a system influences the probability of failure of the system itself. In its broadest acceptance, safety engineering encompasses the analysis of *reliability* (probability for the system to function between the time $t=0$ and the time t), *availability* (probability for the system to be operational at time t) and *maintainability* (probability for a repairable system to be repaired at time t following a failure at time $t=0$).

If the initiation of the undesirable event is deliberate (“man-made” and on purpose), the problem is one of *safeguard engineering*. In such a case, threat assessments, involving social, economic, and political considerations, are required to study the initiating circumstances. Safeguards evaluations include both “overt” attacks on a component or system and “covert” attack. The distinction between these two notions being whether the initiating attack group is part of the system.

Although safeguards evaluations have become a burning question and big concern with the recent surge of large-scale international terrorism, this aspect of the question will not be tackled in the present document.

Another important risk assessment category nowadays is that of *ecological risk assessment*. Even though there is a long history of evaluating environmental and ecosystem impacts, the concept of ecological risk assessment has only recently emerged as a distinct field of risk assessment. The concern about ecological risk follows from the understanding that healthy ecosystems are necessary to provide renewable resources and food, water storage and flood control, biodegradation and removal of contaminants from air and water, pest and disease control, moderation of climatic extremes, recreational opportunities, and scenic beauty. The aim of an ecological risk assessment is to estimate the possibility of adverse impacts on one or more of these ecosystem dimensions from exposures to ecosystem/environmental stressors such as technology (buildings, roads, and other types of infrastructures or installations) and pollutants. Ecological risk assessment's greatest challenge is to work out practical concepts and measures about what constitute adverse (and, thereby, undesirable) ecological changes and what constitute beneficial (and thus desired) ecological changes. The scope of an ecological risk assessment may be fairly narrowly defined, such as the adverse effects of development in a wetland area, or widely encompassing, such as the worldwide issues of global climate change [American Chemical Society, 1998].

Comparative risk assessment



Risk assessments are often carried out in decision-making contexts, in view of selecting the options (technological, operational, organizational) ensuring the lowest level of risk for individuals or for the whole society (assuming, of course, that this constitutes the only, or at least main, criterion to consider, which is not necessarily the case). The level of a particular risk compared to other potential risks is moreover important in determining whether a community or an individual takes action to reduce it. Natural disaster risks are for example unlikely to be considered important in a community that faces much more immediate everyday threats of diseases, food shortage, or child mortality because of insufficient primary health care. By contrast, communities for which diseases and other “primary” risks are today very low and kept under control may initiate natural disaster mitigation program, even living in much less hazardous environment and much less vulnerable houses, because such risks appear *comparatively* important. The amount of resources available to invest in mitigation actions and the value of assets to be protected also determines how readily a community will invest in such actions (see *risk management* below)

Used in such contexts, risk assessment is called *comparative risk assessment*. In essence, comparative risk assessment is directed at developing risk rankings and priorities that would put various kinds of hazards on an ordered scale (small to large). There are two principal forms of comparative risk assessment [American Chemical Society, 1998]. *Specific risk comparison* refers to side-by-side evaluation of the risk (on an absolute or relative basis) associated with exposures to a few substances, products, or activities. Such comparisons may involve similar agents (e.g. the comparative cancer risks of two chemically similar pesticides) or widely different agents (the cancer risk from a nuclear installation compared with the risk of death or injury from automobile travel). The second form is *programmatic comparative risk assessment*, which seeks to make macro-level comparisons among many widely differing types of risks, usually to provide information for setting regulatory and budgetary priorities for hazard reduction. In this kind of comparison, risk rankings are based on the relative magnitude of risk (which hazards pose the greatest threat) or on relative risk reduction opportunities (i.e. the amount of risk than can be avoided with available technologies and resources).

The methods for conducting comparative risk assessment are still developing and remain controversial, as is the concept of using relative comparisons to establish priorities for hazard reduction.

The determination of risk requires assessing the hazards occurrence probability, the elements at risk, and also the *vulnerability* of these last ones. Vulnerability is the propensity of things to be damaged by a hazard. It measures the extent to which the element at risk is likely to be damaged or hurt by the impact of a particular hazard, on account of its nature, resistance characteristics and proximity to the hazard source. People's lives and health, buildings and infrastructure integrities, ecosystem health, etc., are at risk directly or indirectly from the destructive or damaging effects of the hazard. More generally, vulnerability involves also socio-economic aspects, including physical, social and economic considerations (both short and long term), and the extent to which essential services (health, energy, transport, etc.) are able to continue or resume functioning after a major disruption (notion of *resiliency*). Each type of hazard puts a somewhat different set of elements at risk. It is important for risk analysts to make appropriate effort to quantify at least the tangible aspects of vulnerability and loss to assist mitigation and preparedness planning and actions. For engineering purposes, vulnerability is a mathematical function defined as the degree of loss to a given element at risk, or system made of such elements, expected to result from the impact of a hazard of a given magnitude. It is specific to a given type of "target", and expressed on a scale from 0 (no damage) to 1 (total destruction).

The vulnerability aspect

The vulnerability of a set of buildings to a hurricane of 130 km/hr may for example be defined as:

"20% of buildings (vulnerability = 0.2) suffering heavy damage or worse, when experiencing 130 km/hr winds"

The distinction between this definition of "vulnerability" and that of "risk" is important to note. Risk combines the expected losses from all levels of hazard severity and takes moreover into account their occurrence probability.

As scientifically sound as the risk assessment studies could be, there is notwithstanding room for large interpretations of their results. What still complicates things in this matter is that different people or groups of people very often do not perceive given risks, even "objectively" assessed, the same way. Assessment implies evaluation, and any valuation requires a value system and a metric for measurement; both the value system and the metric used to assess risk can widely vary from one person to another (and, what is even worse, even for the same person in different circumstances!).

The perception of risk



If there were a unique value system and metric, according to the definition given in page 2 the two risks represented in the figure 1.3 below should be considered as equivalent. The reality is obviously quite different.

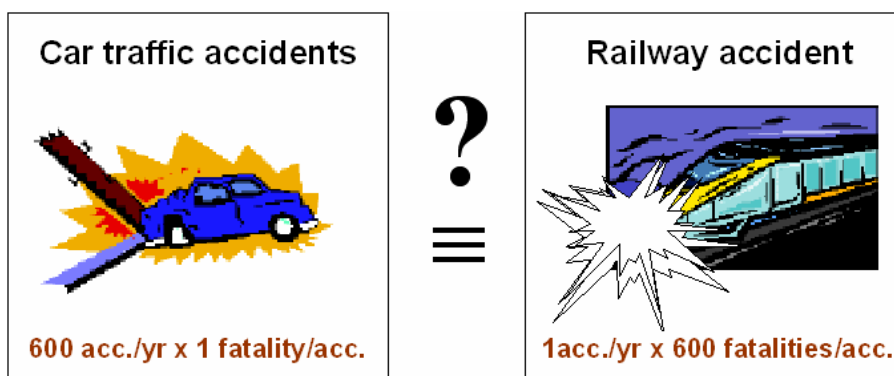


Figure 1.3 *Two equivalent risks ?*

There are many reasons why the way particular risks are perceived often does not fit the "objective" measure of these risks as defined by equation [1.1].



There is evidence that risk perception is considerably influenced by availability of information [DMTP, 1994]

Perception of risk has been an important field of psychological research in the USA and many other industrialized countries for quite a long time. These studies confirm that the mental process of evaluating risk – making sense out of a complex collection of different types of information – tends to differ significantly between individuals and groups. Independently of these societal differences, an important element in the psychology of risk perception seems in other respects to be the “availability” of information. The more “available” the information on a given hazard, the more “dreadful” it appears to the public. Thus, if the perception of risk, which determines society’s motivation for reducing it, appears to be dependent to some extent on exposure to the risk - i.e. to its probability of occurrence - the “dread factor”, which is related to the scale of the potential catastrophe, clearly has a greater impact. Dramatic information rich in deaths and disastrous effects tends to be highly memorable. The way media report catastrophic events is therefore extremely influential on risk perception. For most people, personal contact with hazards is (fortunately) fairly rare and so knowledge of them is acquired more through the news media than from first-hand experience. The fact that the media tend to concentrate on the more unusual and dramatic events has thus for result to make them often perceived to be more frequent than they actually are. Disaster killing numbers of people at once – an airplane crash for example – has for this reason a far greater impact on the public than an equivalent number of casualties from the more frequent but less deadly car crashes (which, in great part, explains the “paradox” of the example given in figure 1.3).

An illustration of that is given by experiments made in the state of Oregon in the USA, in which people were asked to judge the frequency of various causes of death, like diseases, accidents and natural hazards. The results of these experiments (see figure 1.4) show that the, fairly well-informed, interviewees tend to know in general which are the most common and least frequent lethal events but have a general tendency to overestimate the incidence of rare causes of deaths and underestimate the frequency of the more common ones.

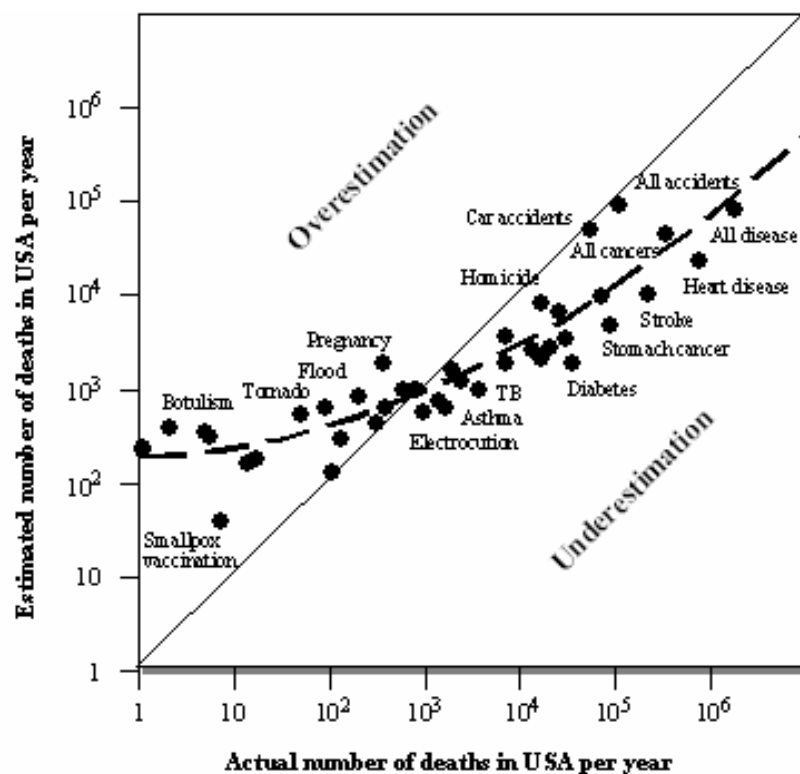


Figure 1.4 Perception of risk in USA (well-informed group, Oregon, 1978)

It is interesting to note that the overestimated risks correspond with favorite American newspaper and media topics, which confirms that, in a society with strong media exposure, perception of risk is highly influenced by media treatment. Experiments of the same kind conducted at EPFL over many years with several groups of students, although less significant from a statistical viewpoint (20-30 students per group), had given very similar results.

Research has also shown that the frequent reiteration of the fact that certain events (like an aircraft crashing) are rare may have the opposite affect on the audience for who it is meant. People may perceive only the fact of the event (the possibility of air crashes) and not the intended message (that air crashes are rare), thus reinforcing the psychological “availability” of information on air travel risk. The same effect has been observed with the reiterated affirmation that nuclear power plants present a very low level of residual risk. The general public obviously retains much more the repeated information that nuclear power plants may experience serious accidents, than the fact that these facilities are designed to make such events highly improbable.



There have been no psychological studies of perception of risk among groups much less exposed to media coverage, but there is evidence that such less-informed communities generally underestimate the risks they face. Their perception of risk is likely to be shaped more by personal experience, local and recent events and verbal information transmission than by media presentation of risks. Information horizons – the distance from which they are brought news and the length of history available to them – may not encompass the rarer events that represent their major threat.

Another important finding of the research in risk perception is that the abstraction of risk is more easily accepted than the personalization of risk. “*It will never happen to me*”, is a common attitude in both richer and poorer societies. The risk of death or injury to a group of people, even one that includes the individual himself, is from this standpoint more readily accepted than the risk to the individual personally.

In general, the research in risk perception shows that the quantitative aspects of risk are often less important than some of its qualitative attribute, the most important of these factors being:

Familiarity – Personal experience, or knowledge from other information sources, of the hazardous event.

Preventability – Degree to which the hazard is perceived as controllable or its effects preventable.

Dread – Mental image of horror associated to the hazard, its scale and its consequences.

Voluntary exposure – Risks voluntarily assumed are ranked differently from those imposed by others or the society.

Unfairness – Substantial outrage is a more likely result if people feel they are being wrongly exposed.

Uncertainty – Scientific uncertainty about the effect, severity, or prevalence of a hazard tend to escalate unease.

Untrustworthiness – The level of outrage is higher if the source of the risk is not trusted.

Such considerations can be extremely influential in shaping the public’s reaction to a given hazard – even to the point of overpowering scientific findings about the magnitude of the risk. This has direct consequences on the acceptability of risk, as explained below.

The acceptance of risk

Studies of what people actually do about risks have been carried out to try to derive an understanding about the acceptability of risk. The figure 1.5 shows for example results obtained in the U.S.A. that give some insight into the factors influencing the acceptance of risk. It indicates that the accepted level of risk increases (up to a certain point) with the benefit expected from the concerned activity or action. The risk associated with driving to work, for example, is generally considered acceptable because the benefit of this action for the person taking the risk is obvious. The figure also makes clear that the tolerance about risk can be as much as 1,000 times higher for risks taken voluntarily than for risks undergone involuntarily. Finally, it suggests that the background risk of death from diseases in society as a whole may provide a yardstick from which acceptability of involuntary risks may be judged.

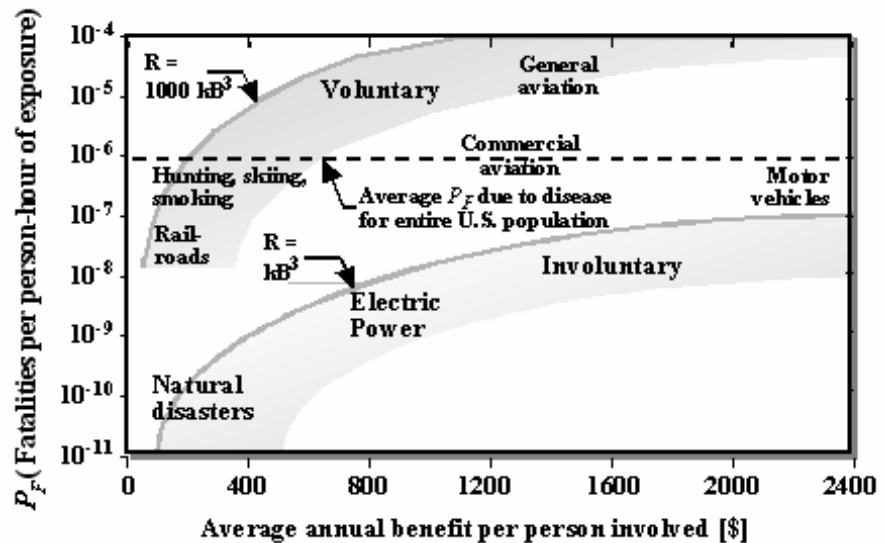


Figure 1.5 Studies of accepted risk levels (U.S.A.)

The concept of risk tolerance and the thresholds of unacceptability are what determine, ultimately, whether public or private money is invested to prevent given risks or mitigate their potential consequences. Because of all the above-mentioned complex psycho-sociologic aspects, the decision about what should be done is something that cannot be based exclusively on scientific arguments but depends as much on subjective determination using value judgments (Fig. 1.6).

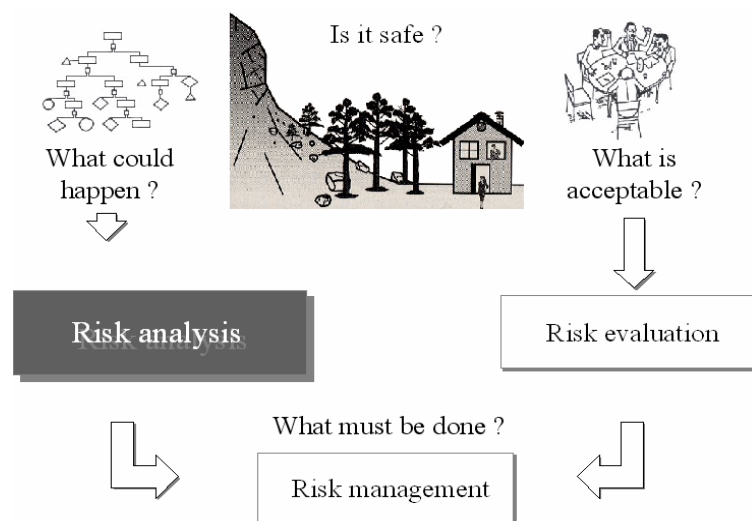


Figure 1.6 Risk analysis, risk evaluation and risk management [Buwal, 1999]

As shown in figure 1.6, risk analysis (or risk assessment) aims at providing an answer to the question “what could happen?”, whereas the question answered by risk evaluation (or risk acceptance) is “what is acceptable?”. Finally, both answers are required to be able to decide about “what must be done?”, the object of *risk management*.

The essential tasks of risk management are:

- to determine what hazards present higher level of danger than society (as represented by governments, public authorities, political parties, pressure groups, voters, etc.) is willing to accept;
- to consider what options are available to master and control these risks and at what cost;
- to decide on appropriate measures and/or actions to reduce (or eliminate when feasible) unacceptable risks, taking into account the material and human resources at disposal as well as other pertinent practical constraints.

The management of risk



At the broadest level, risk management includes a range of management and policy-making activities: agenda setting, risk reduction decision-making, program implementation, and outcome evaluation.

There are several policy approaches to hazard reduction. One avenue, which has a long-standing history, is that of command and control measures. This includes regulations, permits and enforcement actions. Other options rely on market-based economic incentives that prompt desired changes in industrial production decisions and consumer behaviors, voluntary reductions of risk-producing activities, promotion of pollution prevention, and information and education programs to modify behaviors by alerting technology, services and other consumer goods users of the risk involved in their choices and life habits.

To be effective, risk management implies *risk communication*. Risk communication covers a range of activities directed at increasing the public's knowledge about risk issues and participation in the decision-making process. It only emerged as a recognized part of risk management in the 1980s. At this time it was realized that a large fraction of the public was not familiar with the nature of risk and that risk management decisions could not simply be made by technical experts and public authorities to be finally imposed to the persons concerned with justifications after the fact. The problem of risk communication is challenging because the public's response to risk issues is complex, multidimensional, and diverse due to the fact that there are actually many publics with differing values and stakes in these issues.

Ultimately, those with the responsibility of managing risk in society's interest have to make appropriate decisions about risk control and hazard reduction, and work through issues that are not easy to resolve, including legal obligations, uncertainties in the risk assessment evidence, and trade-offs among competing interests to protect the public's health and welfare [American Chemical Society, 1998].

Entering into these considerations is beyond the scope of this course, which is focused on risk analysis (or “risk assessment” when the emphasis is less on quantitative results) approaches and methodology. After a brief review of the history of the development of risk analysis methods in the following section, the necessary mathematical foundations and tools are briefly recalled in chapter 2. The question of the nature of failure data is also tackled in this chapter. Chapter 3 is devoted to methods for the risk analysis of elementary systems, whereas the following chapters present available approaches for the risk analysis of more complex systems, distinguishing between qualitative (chapter 4), semi-quantitative (chapter 5) and quantitative (chapter 6) methods. When required, appendixes are added at the end of the concerned chapters. A general bibliography is provided at the end of the document.

The course content



Auguste Comte
1798-1857

1.2 A Brief Historical Account of Risk Analysis

“On ne connaît pas vraiment une science tant qu’on n’en sait pas l’histoire”

Auguste Comte, French Philosopher, Founder of the Positivism School of Thought

The concern of preserving man from the negative consequences of his own creations or of natural events is probably as old as the emergence of human consciousness. Our ancestors of the prehistoric times certainly became aware of these questions from the very moment they started fabricating their first primitive tools or arms. For thousands of years however, controlling the reliability or the safety of human inventions had relied only on empirical approaches, proceeding through successive iterations supported by the experience grown out of the utilization of these artefacts.

Success and failure, risk and benefit, are indeed inseparable in the long history of the technical progress, to the extent that this one can as well be described as a succession of successes or as a succession of failures. It is primarily by learning from observed failures, or even catastrophic events, that man was able to correct initial mistakes and progress towards new successes.

If *risk analysis*, *risk assessment* and *risk management* are relatively new terms in public debate, they are nevertheless practices with lengthy histories. According to historians, the first professional risk assessors were from ancient Babylon (3200 B.C.); these were a special set of people who served as consultants offering advices on risky, uncertain, or difficult decisions in life – such as marriage proposals or selecting building sites [American Chemical Society, 1998].

The *science* of risk analysis is however much younger. It can be dated back to the infancy of the probabilistic “reliability” assessment of the German V-1 flying-bomb during World War II, followed immediately after the end of this one, by the progressive development of a whole arsenal of methods and tools for reliability and safety assessment purposes in the aircraft (later in the spacecraft) and nuclear industries in particular.

Let us recall here that a science is characterised by the fact that it relies on *models* (inevitably an approximation of the reality), developed on the basis of the *experience*, able to deliver qualitative and quantitative *forecasts* when feed with appropriate *data*. The results thus obtained can in their turn be compared with measurements of real parameters, leading in case of discrepancies to possible improvements of the model or input data (iterative process, see Fig. 1.7).



Risk or benefit ?
Success or failure ?

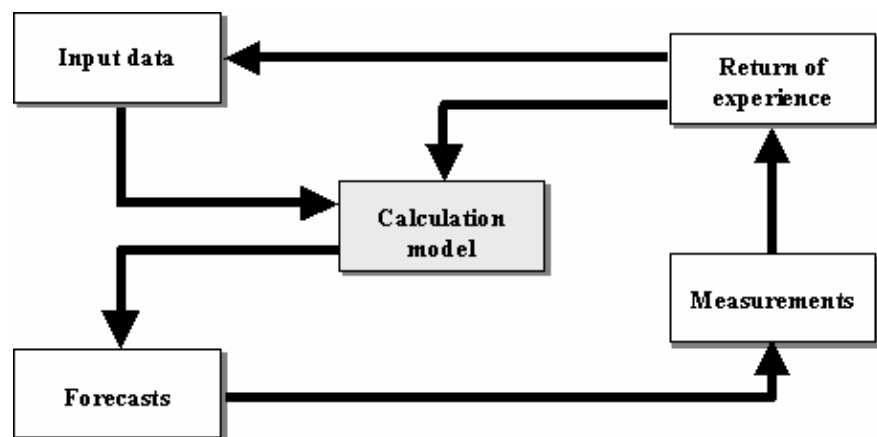


Figure 1.7 The iterative scheme of any scientific approach

Till the 1940s, the aspects of safety and reliability were generally taken into account essentially in a qualitative way, without solid methodological bases. The approach of these questions by the engineers of the time was largely empirical on the scientific and technical level, mostly based on experts' knowledge, itself based solely on experience and the consideration of fundamental laws of nature. A striking illustration of the limits of such an approach has been given by the tragic fate of the *Titanic*. This transatlantic liner, the biggest and most powerful of its time, sank on April 15, 1912, with 2,228 passengers on board, after the starboard side of its bow scraped against an iceberg, although its designers had conceived the ship to be in theory "unsinkable". This tragedy cost the lives of over 1,500 passengers and crew members. The reasons of this catastrophic death toll were many. There were the lack of sufficient lifeboats (what's the use of lifeboats on an "unsinkable" ship?!), the steaming ahead at full-speed despite various warnings about the icebergs-field, the lack of binoculars for the lookout, the poor procedures with the newly invented wireless and use of SOS signals. But above all there was the overconfidence of the *Titanic*'s designers in the safety ensured by the fact that the ship's was divided by 15 transverse bulkheads, a layout that at the time was supposed to correspond to the state of the art as regards maritime navigation safety. *Titanic* could float with any two of her 16 compartments flooded and only the worst possible accident, a collision right at a bulkhead, could even flood two. Unfortunately, the transverse bulkheads were far to be watertight, for passengers' attention and ease reasons: stewards could serve them more easily if doors were cut in the bulkheads, moreover a grand staircase required a spacious opening at every ship's level, making a watertight deck impossible (this explains also why the bulkheads extended only 10 feet above the waterline). We can now estimate that the collision with the iceberg resulted in a gash in the ship's hull of almost 100 meters length (it was probably more an irregular series of holes and rips). Five compartments were flooding rapidly and a sixth was leaking. As the front compartments filled, and the bow sank, the transverse hull between the fifth and the sixth compartments dropped over 10 feet below the waterline. The water was thus able to spill into the next compartment. So the ship sank further, and water spilled over into the next compartment, and the next, and so on. The pumps could only slightly delay this fatal "chained reaction".

Engineers working today in such safety-conscious designs as nuclear power plants use the concept of "defence in depth". Behind the first safety system lies another, and behind that, still another ... each with its own backups. Nothing of the kind was into force in the case of the *Titanic*. The ship's designer had contented themselves with but a "single layer" of defence, the all-too-short transverse bulkheads. Soon after the disaster, the sister ship *Olympic*, as well as many other liners of comparable designs, were being expensively retrofitted with an inner, second hull (it had initially been claimed that this would have uselessly eaten up valuable space for passengers and cargo in the *Titanic* case). This type of layouts had since then become safety standards for passengers' ships, independently of the cost of such measures. This could be compared, all things considered, with the case of the seat belts in cars; a few decades ago, such devices were thought a little-needed luxury, whereas today vastly more expensive airbags have become a standard feature. The public has either resigned itself to the added cost, or embraced the measure with enthusiasm, once the statistics of death and crippling were considered (Brander, 1995).

In spite of the progress of the science and technology, the *Titanic*'s disaster had its counterpart in the air, almost exactly a quarter of a century after the former catastrophe. The fire that destroyed the *LZ 129 Hindenburg* dirigible, in Lakehurst (USA) in 1937, killed 22 of 61 crew members, 13 of 36 passengers and one member of the ground crew. Many survivors were badly burnt. Here again, overconfidence in the skill of the designers of this type of airships and underestimation of the danger of the presence on board of great quantities of hydrogen for buoyancy were at the origin of the disaster. The interesting fact, is that the impact of this spectacular (but not that life-costly) accident on the public resulted in the almost complete disappearing of the dirigibles as passengers' transport means.

Before World War II



*The Titanic, the Hindenburg,
... and a Boeing 747*

This is quite unique in the history of technologies, because, as said before, usually an accident of this kind is the motive for rethinking the safety of the concerned system and achieving new progresses that finally will make the system even more successful after some time.

The World War II years

Until the Second World War, analysis methods of the safety or reliability of technological systems had more things in common with an art than with a science. It was generally thought sufficient to amply strengthen the pieces or components that were supposed to be critical, based on the idea that the strength of a chain was determined by the strength of its weakest link.

The deficiencies of this way of looking at the safety or reliability of systems was made apparent by the initial failures encountered in the development of the German V-1 flying bomb. The first series of V-1 - ten or so missiles – displayed a failure rate of 100%; the engines of war exploded on the launching ramp or fell into the English Channel. In spite of the fact that this flying bomb was a relatively simple system (see figure 1.8), it revealed itself much less reliable than its weakest component.

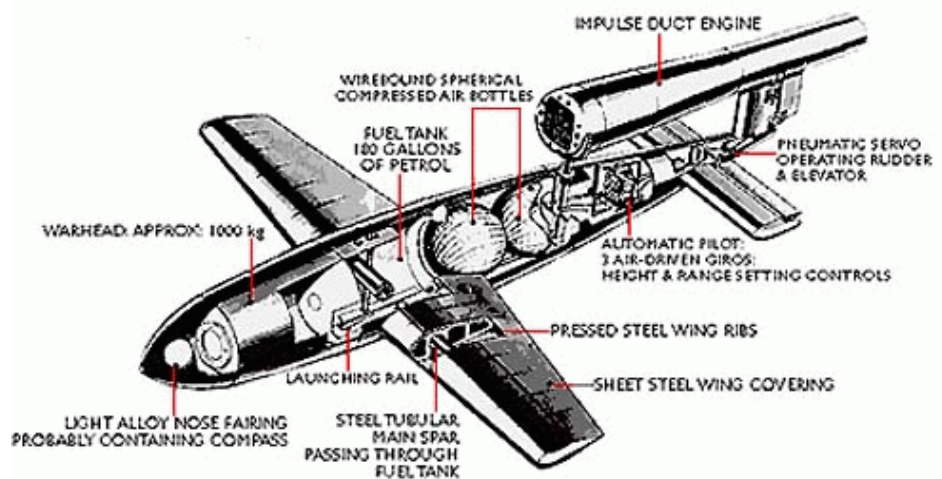


Figure 1.8 *Cutaway view of the German V-1 flying bomb*

Improving after each launch the reliability of the particular component that had been found responsible of the previous failure did however nothing to prevent the next one, because each time a new component “gave way”. This led to the abandon of the idea of the “weakest link” for the concept of some mean reliability value that should be achieved by the whole set of system components. Again, this did not solve the problem, because new tests revealed that the engine of war was clearly less reliable than this mean reliability value. It was finally Eric Pieruschka, a mathematician working with Von Braun’s team, who provided the key answer to the problem; he explained that if the reliability (probability of survival) of a component is $1-x$ (with $x \ll 1$ in principle), the reliability of a set made up of n identical components of this kind linked in series will only be $(1-x)^n \cong 1-nx$. This means that the reliabilities of the components have to be much higher than the reliability demanded from the system. It seems that the V-1 was the first “industrial” system for which, in the latest phases of its development, the reliability of the whole system was successfully achieved (unfortunately for the people of London and some other main European cities!) on the basis of strict reliability requirements set - and experimentally verified prior to assembly – for its components.

It is however in the U.S.A., immediately after the end of World War II, that important theoretical and methodological developments actually laid the foundation for this new branch of engineering science: the science of risk analysis.

One important incentive for these developments was the poor initial performance of the newborn electronic industry, characterized by the chronic unavailability of the devices using the technology in question. According to a study of the U.S. Department of Defense in the early 1950s, the equipments of this kind were in working order only 30% of the time, and the maintenance of the equivalent of 1 US \$ of equipment cost up to 2 US \$ annually. This led the U.S. Air Force captain and engineer Edward A. Murphy to make the statement that since then has become famous under the name “Murphy’s law”: “*If anything can go wrong, it will!*”.

To break the spiral of increasing maintenance costs, the U.S. Department of Defense and the whole electronic industry founded in 1952 the *Advisory Group on Reliability of Electronic Equipment* (AGREE). This commission recommended making reliability studies henceforward an integral part of the development cycle of all electronic equipments. It advocated testing any new electronic system in the most extreme environmental conditions it could possibly have to face in operation and correcting the detected weaknesses before to start mass production. It defined also appropriate indicators for quantifying the reliability of any new equipment, e.g. the *Mean Time between failure* (MTBF). This was the outset of the general enforcement of *reliability clauses* to mass production components.

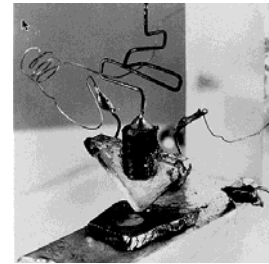
A greater importance is also given in the same decade to safety problems, more particularly in the civil aviation domain, due to the important development of this means of transport and, in the late 1950s, in the nuclear industry with the start of operation of the first civil nuclear power plants in the U.S.A. (Shippingport in 1957, Yankee Row in 1960).

The 1960s saw the emergence of new reliability techniques and a broader range of industrial applications. It is the beginning of detailed analyses about the failures of components and their effects on the operation of systems, or about the safety of persons as well as of public or private goods. Many methods – *Success Path Diagram, Failure Mode and Effect Analysis, Fault Tree, Event Tree* – are perfected during this period and applied to the *probability assessments of system failures*.

At the end of the 1960s, the use of these methods became general in the aviation - *CONCORDE* project in France and Great Britain, jumbo-jet development in the USA - and space industries. The National Aeronautics and Space Administration (NASA), in particular, had largely recourse to such methodological approaches in the framework of the *APOLLO* program and, more specifically, for the risk and reliability analyses of the *Lunar Module* spacecraft (LM, see Fig. 1.9). It is considered that the thorough and systematic use of probabilistic risk and reliability assessment methods was a key factor of the success of the *Apollo* program, which allowed American astronauts to walk on the Moon at six different occasions and carry out numerous scientific experiments – with only one accident (*Apollo XIII*) that had moreover nothing to do with a design problem according to the post-accident inquiry. This was far to be obvious; some experts had cast doubts in the beginning of the 1960s on the feasibility of mastering this way – at the prototype stage - risks by definition in great part totally unknown, without moreover any possibility to fully test beforehand the equipments in their future operating environment.

Unfortunately, this did not prevent later on the American space shuttle program from knowing two dramatic accidents: *Challenger* in 1986 and *Columbia* in 2003. As in the case of the *Apollo XIII* accident however, the reasons of these tragic events were more of a managerial than of a purely technical nature, at least for the former. The problem with the O-ring seals of the re-usable booster rockets, which was directly at the origin of the oxygen-tank explosion, had been known for a rather long time before the accident, and an investigation on this subject was in progress, but NASA officials did not think necessary to postpone further launches until its results were known due to lack of proper communication between different levels of NASA management, economic considerations, political pressures, and scheduling backlogs.

After WWII, the 1950s

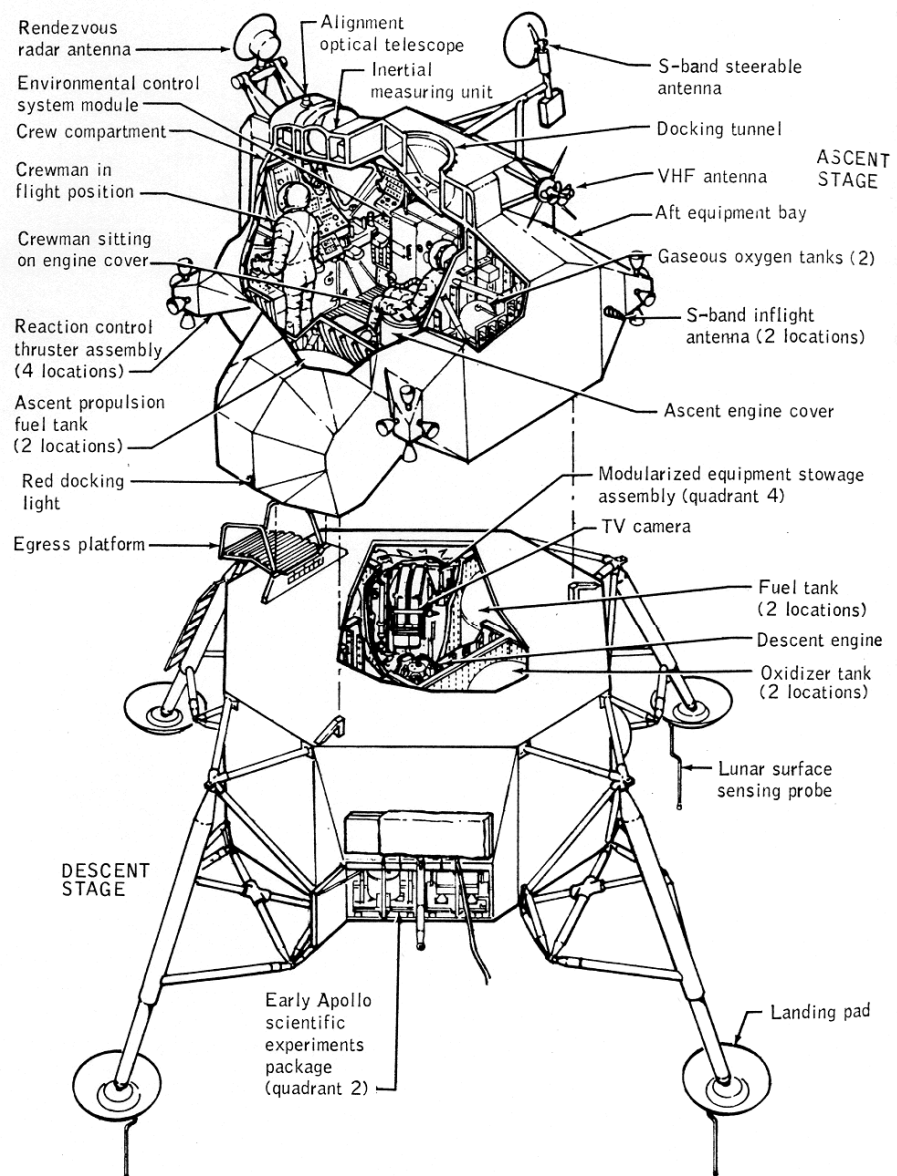


The first transistor
(Bell Laboratories, 1947)

The 1960s, a turning point
decade in the development
of risk analysis methods



Challenger accident (1986)



LUNAR MODULE CONFIGURATION FOR INITIAL LUNAR LANDING

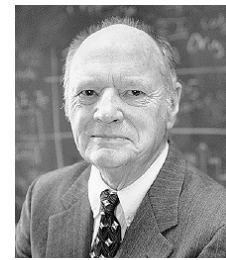
Figure 1.9 *Cutaway view of the Apollo Lunar Module* (doc. NASA)

Coming back to the 1960s, it is in the middle of this decade that emerged the concept of *maintainability* – scientific planning of the maintenance operations – to curtail the increasing costs resulting from the unavailability of more and more complex equipments representing important investments. It is also worth mentioning that the first reliability data banks are created at the end of this period.

The 1970s, a decade of decisive developments in nuclear safety analyses

At the same time, the taking into account of potential accidents already at the design stage reach in the nuclear industry a level never achieved before in any industrial activity. All potential failures are identified and classified by categories of importance as well as estimated occurrence frequency, taking systematically into account the worst possible situations. For each of these potential accidents, it is then conclusively proven that the considered design arrangements will insure a safety level judged acceptable. This led in 1975 to the publication in the U.S.A. of an important report – *WASH-1400 report* - on the safety of nuclear power plants, which marked an epoch in the development of probabilistic safety assessment methods.

The 9-vol. WASH_1400 technical report is perhaps better known under the appellation *Rasmussen report*, from the name of the MIT Professor who directed the study. The task for Dr. Rasmussen was to replace the "maximum credible accident" approach to safety analysis (which, precisely, was not sufficiently ... credible). The solution that Dr. Rasmussen developed to address safety analysis was *Probabilistic Risk Assessment* (PRA). PRA revolutionized not only safety analysis, but also quantified, for the first time, accident risk. The pressurized water reactor (PWR) plant selected for the study was the 778 MWe Surrey Power Station, Unit I; the boiling water reactor (BWR) was the 1065 MWe Peach Bottom Atomic Power Station, Unit II. These plants were the largest of each type that were about to start operation in 1972 when the study begun. The combination of fault trees and event trees used to determine the probability of accidents and the amount of radioactivity potentially released was applied to these two types of reactors. In order to make the results of the study also applicable to other PWR or BWR power stations (i.e. to perform a "generic" risk analysis rather than a "site-specific" one), the consequences of the potential accidents were analyzed by using weather and population data for a set of generic sites by averaging such data over the 68 sites of the first 100 nuclear plants scheduled for operation in the U.S.A. The Rasmussen Study Executive Summary has a table asserting that, for 100 nuclear reactors, the frequency of an accident causing more than 100 fatalities is once in 100,000 years of operation, and the frequency of an accident causing more than 1000 fatalities is once in a million years of operation. The table gives exactly the same frequencies for fatalities due to meteor strikes. Although the Rasmussen Report is flawed in some respects, and tends to understate the dangers from nuclear reactors, nevertheless there is much in the report that is innovative, illuminating and important.



Dr. Norm Rasmussen,
MIT Professor Emeritus

After the U.S.A., many countries adopted the approach developed by the MIT team for the safety analyses of their own nuclear power plants park. In particular, a comprehensive study designed to evaluate the risk from nuclear power plants was conducted in 1979 in the Federal Republic of Germany. The study was aimed at determining the risk involved in incidents and accidents in nuclear power plants, taking account of German conditions. Initial risk investigations (phase A) mainly aimed at assessing the risk involved in accidents in nuclear power plants and comparing it with other risks of civilization and nature. In phase B of the German risk study, extensive research into the incident behavior was carried out. In this context, the time sequence of the incidents, the impact involved and the intervention of the safety systems provided to control the accident or incident were thoroughly analyzed. These investigations showed the importance of in-plant accident management. The analyses demonstrated that in many cases nuclear power plants still have safety reserves when the safety systems do not intervene as intended and safety-related design limits are exceeded.

The confidence in the conclusions of such PRA analyses in the nuclear field was however shaken by the accident at the Three Mile Island Unit 2 (TMI-2) nuclear power plant in Pennsylvania, U.S.A., on March 28, 1979. This accident, which led to a loss of coolant and partial core meltdown, was the most serious in commercial nuclear power plant operating history¹, even though it involved no deaths or injuries to plant workers or members of the nearby community. Based on a series of investigations, the main factors appear to have been a combination of personnel error, design deficiencies, and component failures. For some antinuclear militants, the TMI accident was the demonstration that the conclusions of nuclear PRA studies were not credible. Experts, who went back into the Rasmussen Report and found on the contrary that the possibility of an accident of this kind had been duly identified and its risk reasonably estimated, refuted this assertion.

¹ There had been a partial core meltdown accident before (1969, January 21) at the experimental underground reactor of Lucens (VD), Switzerland, due to a coolant malfunction. This resulted in a large amount of radiation released, but solely into the reactor cavern, which was then sealed and later cleaned.

One of the striking aspects of the Rasmussen Study is indeed that it suggested that accidents other than the large break loss of coolant accident that had been used as the design basis for reactors build up to that time were greater contributors to the risk of core melt. One is called "TMLQ" in the code of the study; it means loss of feedwater plus failure of a safety valve. It is exactly what happened at Three Mile Island. As the WASH-740 study contemplated serious reactor accidents and reported consequences and estimated frequencies, it is not correct to claim that this kind of accident had been said to be "impossible".

The TMI accident brought about sweeping changes involving emergency response planning, reactor operator training, human factors engineering, radiation protection, and many other areas of nuclear power plant operations. It also caused the U.S. Nuclear Regulatory Commission and similar organizations around the world to tighten and heighten their regulatory oversight, all measures that had the effect of enhancing nuclear safety in Western countries.

Risk analysis studies are also initiated in other industrial sectors during the 1970s decade. An example of a comprehensive risk assessment related to petrochemical installations is provided by the Canvey study carried out in the United Kingdom by the "Health and Safety Executive" between 1976 and 1978 at the request of the British government. Canvey is a small island situated on the north bank of the Thames River, downstream of London. It is a highly industrial area, which employed more than 3'200 people at the time working in various petrochemical facilities occupying a surface of 50 km². Public investigation leading to a report was to investigate the overall risks to health and safety of the nearby population (around 33'000 people) arising from any possible interaction between existing and proposed installations. A similar investigation was conducted in the beginning of the 1980s for the Rijnmond region, another important petrochemical complex and harbor situated between Rotterdam and the North Sea.



Bhopal Memorial

Still in the 1970s, an important effort is made to take into account the *human factor* in risk analyses, from a qualitative as well as from a quantitative viewpoint. The TMI accident, and later the Chernobyl disaster - steam explosion and fire having released at least five percent of the radioactive reactor core into the atmosphere as a result of a flawed reactor design moreover operated with inadequately trained personnel and without proper regard for safety (1986) - in the nuclear sector, the Bhopal tragedy - a runaway chemical reaction caused by the inadvertent addition of water in a methylisocyanate tank, resulting in the death of at least 10'000 people and in the serious intoxication of some 500'000 more (1984) - in the chemical industry, the Challenger explosion in the space domain, etc., have amply demonstrated the importance of this factor in the triggering and/or development of many serious accidents.

The 1980s and beyond

The 1980's witnessed an unprecedented number of media-fed disasters, some of them already mentioned above - core breeches in nuclear reactors (Chernobyl), sinking ships, oil spills, chemical leaks (Bhopal), etc. With a nearly continuous spectacle of large-scale technological calamity ... the mass media declared the 1980's, the "age of limits". Either because, or in spite, of these events, one observes in the 1980s a generalization of the use of probabilistic safety, reliability and availability assessment studies in a growing number of sectors such as: oil industry, chemistry, railways, car industry, treatment and purification of waste water, etc., covering that way very different activities and technological structures.

All the developments shortly evoked in this section have finally given birth to a genuine and new engineering science: the science of *reliability and safety engineering*. It will be the aim of the following chapters to present the essential foundations and applications of this science in more details.

