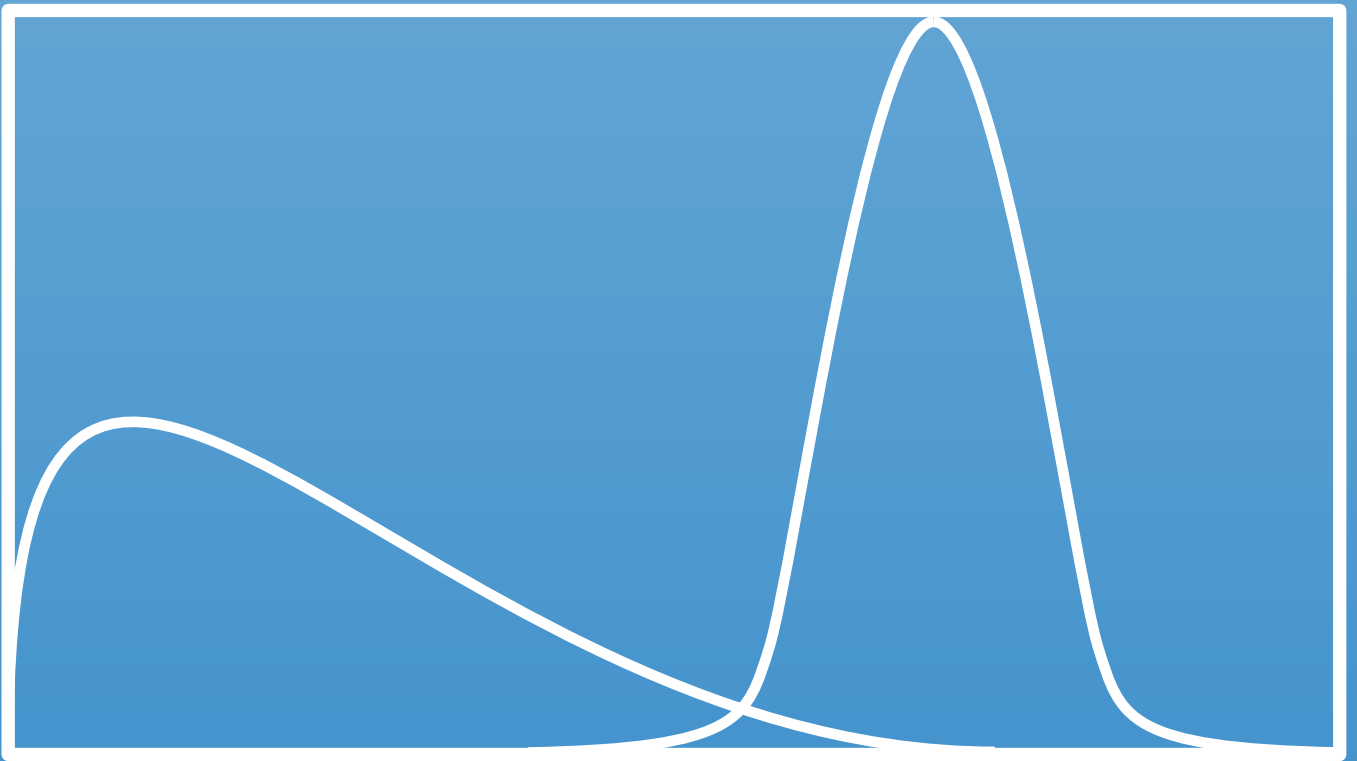


ASPECTS OF STRUCTURAL SAFETY



THE COMPENDIUM FOR
TKT4196 - Prosjektering Sikkerhetsforhold

Jochen Köhler
Høstsemester 2021

 **NTNU**



Norges teknisk-naturvitenskapelige universitet
Institutt for konstruksjonsteknikk

Compendium

TKT4196 - Prosjektering/Sikkerhetsforhold **Aspects of Structural Safety**

Prof. Dr.-Ing. Jochen Köhler

Synopsis:

This compendium is intended for the students of the Norwegian University of Science and Technology (NTNU) that take part in the course TKT4196 - Prosjektering Sikkerhetsforhold. It is still under continuous improvement and development. Any critical remarks and suggestions for improvement is welcome and highly appreciated.

Acknowledgement:

The present version of the compendium is produced in the period 2015-2020. The help and contribution of Jorge Mendoza, Michele Baravalle and Klodian Gradeci is highly acknowledged.

Table of Contents

	Page
Chapter 0. Introduction to TKT4196	1
1. Organization basics	1
2. Learning strategy	2
3. Literature & Standards	3
4. Rationale of the course	4
5. Course content	5
6. Safety - a requirement for structures	6
7. Summary	10
Chapter 1. Structural optimization	11
1. Theoretical aspects	11
1.1 Introduction	11
1.2 Optimization	12
1.3 Interest	16
1.4 Continuous renewal assumption	17
1.5 Risk acceptance	19
1.6 Conclusion of Chapter 1	21
2. Application examples	22
3. Practical exercises	28
Chapter 2. Structural reliability - Problem Statement and Monte Carlo Method . .	31
1. Theoretical aspects	31
1.1 Prelude	31
1.2 Introduction: The limit state principle	31

1.3	A simple case - the Fundamental Reliability Problem	33
1.4	Definitions	34
1.5	Monte Carlo Method	35
1.6	Sampling of correlated random variables	37
1.7	Generalization for multivariate distributions	40
1.8	Variance reducing methodologies	42
2.	Application examples	46
3.	Practical exercises	51
Chapter 3. First Order Reliability Method		55
1.	Theoretical aspects	55
1.1	Linear Limit State Functions and Normal Distributed Variables	55
1.2	Non-linear limit state function	56
1.3	Correlated and Dependent Random Variables	59
1.4	Non-Normal and Dependent Random Variables	62
2.	Application examples	65
3.	Practical exercises	71
Chapter 4. Introduction to Probabilistic Load and Material Modelling		73
1.	Recapturing basic principles of probability theory	73
2.	Theoretical aspects	73
2.1	Recapturing basic principles of probability theory	73
2.2	Interpretation of Probability	74
2.3	Uncertainties in Engineering Problems	76
2.4	Random Variables	78
2.5	Load bearing behaviour of materials	80
2.6	Scales of (variability) modelling	80
2.7	Brittle material	81
2.8	Probabilistic load modelling	82
2.9	Time variable load	84
2.10	Load combination	90
3.	Practical exercises	93
Chapter 5. Structural Design Standards		97
1.	Concepts	97
1.1	Levels of design	97

1.2	Risk informed decision making	98
1.3	Reliability based design	98
1.4	Semi-probabilistic design	99
1.5	Relevance and correspondence	99
2.	EUROCODE semi-probabilistic design.	100
2.1	General framework and design values	100
2.2	Partial factors	101
2.3	Reliability requirements in the Eurocodes	101
2.4	Direct correspondence between reliability and design values - the design value approach	102
2.5	Derivation of generalized α -values	106
2.6	The generalised α -values in the Eurocodes	107
3.	Application examples.	108
Chapter 6. System Reliability		117
1.	Theoretical aspects.	117
1.1	Introduction	117
1.2	Simple bounds for system reliability of failure	120
1.3	System Reliability using FORM	121
2.	Application examples.	125
3.	Practical exercises	129
Chapter 7. Engineering Decision Making		133
1.	Concept	133
2.	Bayes' Rule	135
2.1	Conditional Probability and Bayes' Rule	135
2.2	Example 2 - Using Bayes' Rule for Bridge Upgrading	137
2.3	Example 3 - Covid Test	138
3.	Engineering application.	139
3.1	A priori analysis	139
3.2	A posteriori analysis	140
3.3	A pre-posteriori analysis	141
3.4	Summary	144
Chapter 8. Engineering uncertainty quantification		147

1. Engineering uncertainty quantification	147
1.1 Cumulative Distribution and Probability Density Functions	147
1.2 Sampling and representation of data	155
1.3 Model building and parameter estimation	159
2. Application examples.....	166
Bibliography	179

Chapter 0

Introduction to TKT4196

In this first chapter of **TKT4196 - Aspects of Structural Safety**, a general overview of the course is given. After going through the practical and organisational aspects related to the course implementation, we will have a brief discussion on the role of safety aspects in structural engineering and in civil engineering in general. This chapter is meant to be a soft and “not too technical” start to the course and to the topic.

1. Organization basics

Contacts

Responsible lecturer:	Jochen Köhler	jochen.kohler@ntnu.no
Tutors:	Ramon Hingorani	ramon.hingorani@ntnu.no
Student assistants:	Jonas Nordbø Andreas H. Asheim	jonanord@stud.ntnu.no andreas.h.asheim@ntnu.no

Lectures and Exercises

The main language of the course is English. The lectures will be given in English and the written material and suggested literature is in English. The lecturers masters Norwegian to some degree and the students are very welcome to ask questions and start discussion in Norwegian.

Timing and location:

Mondays	10.15 - 12.00	R10
Mondays	12.15 - 14.00	Landmålerhallen
Fridays	08.15 - 10.00	Landmålerhallen

The lectures run from week 34 - 47. The first lecture is taking place on Monday, August 23, at 10:15 in room R10. It is recommended that the students are present, both, physically and mentally in all the announced classes. The lecture and exercises will be held in compliance with the official NTNU policies and procedures in regard to the COVID-19 situation.

Compulsory problem tasks

Three exercises must be delivered and accepted in order to be allowed to take the exam. The exercises will be graded and the cumulative grade counts 1/3 to the final character. Solutions will be published and presented during exercise classes.

The problem tasks will be solved and documented in groups by 3-5 students.

Exercises hand in: directly in the lecture or electronically following the lecture program and announcements.

Exam

The exam time and place will be announced in due time. It will be an oral exam with a duration of approximately 30 minutes. The exam is normally held in norwegian but english is also possible, of course. The grade counts 2/3 to the final character.

2. Learning strategy

One of the important premises of this course (and in general) is that *learning* will be achieved solely by the students themselves and this requires corresponding commitment. The successful participation of this course gives 7.5 credit points and it is expected that an average student has to work 10 hours per week for this course in order to reach an average grade. (Note that this rule holds generally for all courses with a corresponding number of credits).

Learning is *facilitated* by proper supervision and teaching. For this course the concept includes five major components. As a rule of thumb, 2 hours/week should be invested for each of the components.

1. **Compendium:** An accompanying compendium, which this chapter is a part of, is available to the students. It is divided into chapters for the different learning topics of the course. Each chapter is divided into three sections: 1. Theoretical aspects, 2. Application examples and 3. Practical exercises. 1. Theoretical aspects covers the content of the lectures. The material of some chapters is introduced during several lectures. 2. Application examples covers some applications of the theory, including the solution. These will be discussed during the exercise classes. 3. Practical exercises include some additional exercises, without the solutions, that the students should solve during the practical classes.
2. **Lectures:** All content of this course will be transmitted in the lectures. The theoretical concepts will be illustrated by examples. The lectures are oral and will be supported by blackboard demonstration and overhead slides. The Students are asked to follow actively the lectures and rise questions if something is unclear. An open dialog in the lecture will clearly increase the quality of the lecture. Keep in mind that no questions are “stupid” and that many other students may benefit from answering the same question. It is strongly recommended to make comprehensive notes during the lectures.
3. **Exercise classes:** In the exercise classes, the application examples will be demonstrated and discussed. Vital participation of the students (in terms of questions and comments) is absolutely required.

4. **Practical classes:** It is very important that the new obtained knowledge becomes an executable part of the students' skills. We will therefore spend some time on "hands-on" practical classes where students try to implement gained knowledge and understanding in practical problems. These exercises are largely based on the application examples. The main difference is that the solutions will not be provided beforehand. This brings an opportunity to the students to test the understanding of the course material avoiding any "narrative fallacy": exercises seem easier and more logical after the solution is given, but this cancels the student's creativity to think for herself on different ways how to answer a question. We will use some IT-tools for that (e.g. Python or Excel) and students should work in groups on their own computers. The classes will be tutored.
5. **Self-study:** Students are expected to constantly revise the current course material and read additional literature as it will be given in the lectures.

3. Literature & Standards

The following literature is recommended for further reading. Special chapters of some of the books will be defined as pensum during the lecture.

Reliability Basics	[1] Schneider J., 2006. Introduction to safety and reliability of structures. [Zürich] International Association for Bridge and Structural Engineering. [Available on Blackboard].
Reliability Further Reading	[2] Melchers R. E., Beck A.T., 2017. Structural Reliability Analysis and Prediction. Wiley.
Structural Codes	[3] Gulvanessian, H; Formichi, P; Calgaro, J. A 2009. Designers' guide to Eurocode 1: EN 1991-1-1 and -1-3 to -1-7, Actions on buildings / H. Gulvanessian, P. Formichi and J.-A. Calgaro [ebook at NTNU Library]. [4] Standard – EN 1990, Eurocode: Basis of structural design [via Standard Norge at NTNU Library]. [5] Standard – EN 1991, Eurocode 1: Actions on structures [via Standard Norge at NTNU Library].
Uncertainty Modelling	[6] Benjamin, J. R. and C. A. Cornell (1970). Probability, Statistics, and Decision for Civil Engineers.

4. Rationale of the course

Civil Engineers play and will play an important role in our society. Many current and future challenges, as related to the efficient management of (limited) financial and natural resources or the appropriate mitigation to the possible effects of climate change, are closely related to the build environment and require good civil engineering. As the challenges will become bigger, the demand for new and innovative solutions, i.e. well beyond traditional engineering practice, will increase. The role of future civil engineers is very nicely defined by the American Society of Civil Engineers (ASCE)¹ as a professional who is²:

“Entrusted by society to create a sustainable world and enhance the global quality of life, civil engineers serve competently, collaboratively, and ethically as master:

- *planners, designers, constructors, and operators of society’s economic and social engine, the built environment;*
- *stewards of the natural environment and its resources;*
- *innovators and integrators of ideas and technology across the public, private, and academic sectors;*
- *managers of risk and uncertainty caused by natural events, accidents, and other threats; and*
- *leaders in discussions and decisions shaping public environmental and infrastructure policy.*

As used in the vision, “master” means one who possesses widely recognized and valued knowledge, skills, and attitudes acquired as a result of education, experience, and achievement. Individuals within a profession who have these characteristics are willing and able to serve society by orchestrating solutions to society’s most pressing current needs while helping to create a more viable future.”

Many will agree that proper university education in the field of civil engineering develops a good understanding of the basic principles in science (being Mathematics, Physics, Applied Mechanics, and Chemistry) combined with good understanding of the practical engineering problems at hand. This combination of fundamentals and practical implications is very well reflected in the current study curriculum for “Bygg- og miljøteknikk” at NTNU, i.e. with fundamental classes and more practical classes for the individual study directions (that take up fundamentals and make them executable for practical engineering applications).

The understanding of uncertainty, risk and safety is an essential skill for civil and structural engineers. A basic and central reality of practical civil engineering is the fact that our world (especially the future side of it) is not exactly known but can be estimated or predicted by taking into account the uncertainties that are associated to this incomplete knowledge. Safety related

¹ASCE is the permanent organization representing the civil engineering profession in the United States. <http://www.asce.org/>

²Part of “ASCE - Civil Engineering Body of Knowledge for the 21st Century. Preparing the Civil Engineer for the Future” [LINK](#)

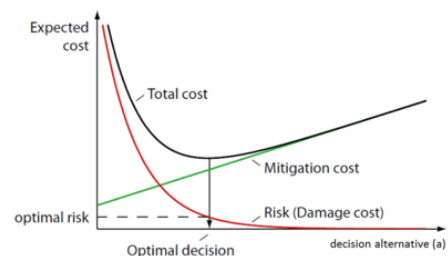
engineering decisions must take into account uncertainties and risks. Every responsible engineer should have at least basic knowledge on how to treat uncertainties in practical engineering problems and understand the background of the generally used regulations and design standards.

This course aims at the procurement of *basic* knowledge in the topics *uncertainty representation*, *risk* and *safety* and makes students develop a new mind-set expendable for interpreting all kinds of engineering problems. (Expert knowledge might be attained by specializing in master or PhD projects).

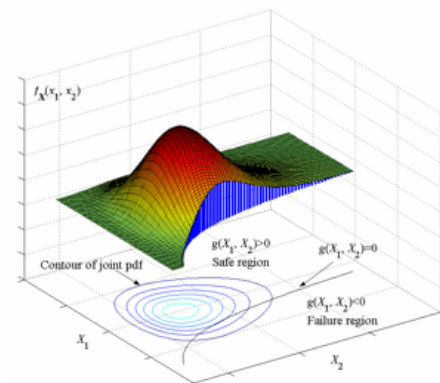
5. Course content

The course will develop basic knowledge, understanding and practical skills. The course content is organized in different thematic main blocks:

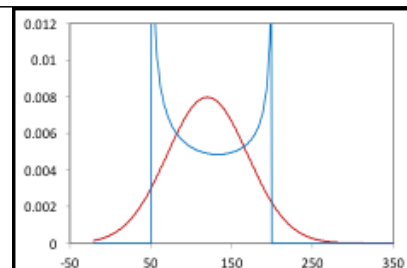
- A. **Optimization and Acceptance.** Evaluation of fundamental risk based performance criteria. An objective function for engineering structures is developed. It is discussed how human safety can be incorporated. It is seen that the structural failure probability is of particular interest.



- B. **Methods for structural reliability assessment.** Different methods for the estimation of the structural failure probability are introduced. The methods facilitate a realistic representation of structural engineering problems, however, we will learn that a proper engineering uncertainty representation is necessary.



- C. **Basic engineering uncertainty modelling.** It is demonstrated how uncertainties can be represented for different relevant phenomena. Practical methods are introduced to represent material resistance and time varying loads.



- D. **Standards and code calibration.** The methodical basis of structural codes is introduced and the general background of the overall safety concept is examined.



6. Safety - a requirement for structures

The built environment, i.e. infrastructural elements, industrial buildings and facilities as well as residential buildings, constitutes the basis for our economy and the continuous development of our society. In this respect structures play an important role, since the primary purpose of structures is to provide the functionality of the built environment. That structures are safe is an important premise for the success of our society, and structural safety is clearly a basic requirement of the society.

Correspondingly, structural *safety* is a requirement that is defined in legislation on several hierarchical levels – structures have to be safe by law. In Norway a general safety requirement can be found in the ”Plan og Byggningsloven” (§29.5); more detailed prescriptions for safety are found in ”Byggeteknisk forskrift (TEK17)”, Chapter 5 on ”Konstruksjonssikkerhet”. Here one can read:

”Byggverket skal plasseres, prosjekteres og utføres slik at det oppnås tilfredsstillende sikkerhet for personer og husdyr, og slik at det ikke oppstår sammenbrudd eller ulykke som fører til uakseptabelt store materielle eller samfunnsmessige skader.”

Several formal definitions of *safety* exist and accordingly, safety is used as absolute or relative characteristic in our regular language use. (“It is safe to use air traffic” is an absolute statement; “Air traffic is safer as car traffic” is a relative statement). In the context of structural engineering the following definition can be found in ISO 2394:2015³:

Structural safety: “*ability of a structure or structural member to avoid exceedance of ultimate limit states⁴, (...), with a specified level of reliability, during a specified period of time.*”

Remembering that structural reliability per time unit is defined as the complement of failure probability of structural failure per time unit; structural safety is an absolute description for situations where the probability of failure per time unit is sufficiently low.

QUESTIONS: *What further structural requirements does ”safety” include? What requirements are independent from safety?*

(space for your notes)

Attaining structural safety

Structural engineers are responsible that structures are safe. Structural engineers therefore make proper decisions during the entire life-cycle of the structure. The typical phases of the life-cycle of civil engineering infrastructure are illustrated in Figure 1.

For decision making, structural engineers take into account information e.g. about:

³ISO 2394:2015: General principles on reliability for structures.

⁴The exceedance of an ultimate limit states corresponds e.g. to structural failure.

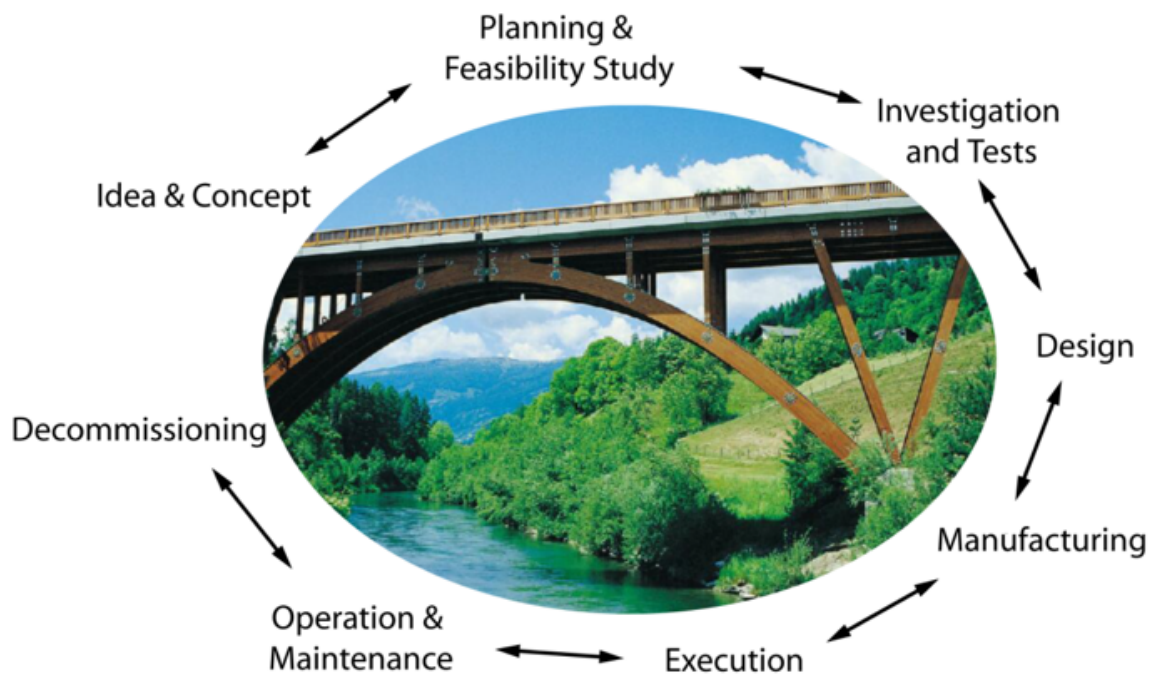


Figure 1: The life-cycle of engineering structures.

- Traffic volume
- Loads
- Resistances (material, soil, ...)
- Degradation processes
- Service life
- Manufacturing costs
- Execution costs
- Decommissioning costs

As a matter of fact, the information that has to be integrated is associated with uncertainties. Thus, the uncertainties have to be taken into account in the decisions that are made in the life-cycle of the structure.

The activities and the decision making of structural engineers are very well illustrated in the Swiss structural standard SIA 260 with the graph presented in figure 2. Here, a *structure* is defined by its *service criteria* and the *environment* it is embedded into. The structure then has the following phases during its life cycle:

- design

- execution
- use
- conservation
- decommissioning

QUESTION: What are examples for typical decisions that are made during the life-cycle of a structure? What is uncertain for these decisions?

(space for your notes)

In this course we will focus on the *design phase* of the structure. In figure 2 the design phase is further subdivided into “conceptual design”, “structural analysis” and “dimensioning”.

In **conceptual design** a structural concept is developed that is best to meet the requirements to the structure, i.e. in regard to the intended societal activity that should be supported by the structure these are requirements in regard to

- space and shelter created by the structure,
- requirements on limitation of vibrations and deflections (serviceability),
- requirement on the durability of the structure and the intended service-life,
- safety, i.e. preventing damage, loss of functionality, damage to the environment and to personnel,
- cost efficiency.

In the conceptional design phase, major decision are made in regard to principle structural layouts and material use. Build-ability is always a very important aspect here.

Structural analysis is intended to create a link between the external loads on the structure and the corresponding structural response in terms of stresses, deformations and deterioration. Mechanical models (analytical and numerical) are utilized and structural analysis is generally performed at a rather high level of modelling detail.

For structural **dimensioning** the structural dimensions are chosen such that the structure complies with some specified limit states regarding safety and serviceability. The compliance is demonstrated for several load scenarios (“design situations”).

In general it can be observed that the effort and level of detail is not consistently distributed among the different steps of structural design: Whereas a lot of attention is attributed to structural analysis, rather rudimentary criteria are used in dimensioning.

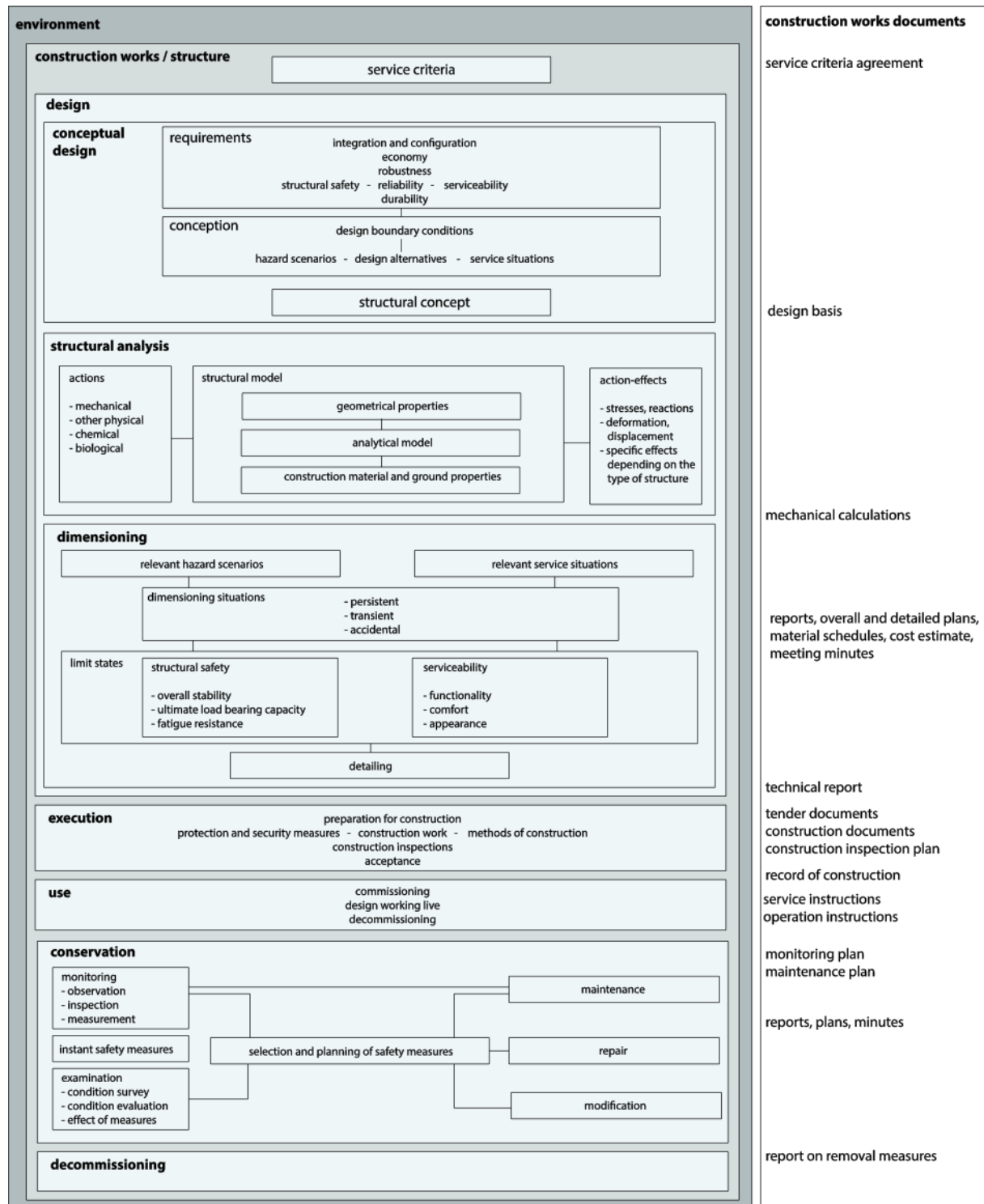


Figure 2: Schematic representation of structural engineering activities.

7. Summary

In this first lecture notes some organization information is shared. The contents and objectives of the course are defined and an initial approach to the term "*safety*" is given. It is demonstrated that the information that is necessary for structural engineering decisions is uncertain, and that this uncertainty has to be taken into account by the structural engineer when evaluating safety, i.e. the probability of failure has to be quantified.

But how safe is safe enough? This question will be followed up in the next lectures.

Chapter 1

Structural optimization

It is demonstrated how a simple structure can be designed in order to minimize the risk. The optimization could be the foundation for identifying sufficiently low failure probabilities. However, it should be critically assessed whether the optimality criterion is adequate for that.

The aspects of net present value are introduced. The simple optimization example is generalized and the infinite renewal assumption is introduced and discussed. Aspects that go beyond monetary optimization are argued and an acceptance criteria for life safety is introduced and discussed.

1. Theoretical aspects

1.1. Introduction

In the Chapter 0, a safe structure was defined as a structure with *sufficiently low failure probability* per time unit. We will now elaborate on the question “how safe is safe enough?”.

How safe is safe enough?

The question “how safe is safe enough?” is clearly related to the term *sufficiently low failure probability* in the definition above, i.e. to the meaning and quantification of *sufficiently low*.

Let us first look at the requirements for structures again¹. Safety is an overall requirement, but which further requirements does this include? Structures are used by people in many ways (residents in buildings, car drivers on bridges, workers on offshore platforms). The **risk to life** and limp induced by these structures should be sufficiently low. Furthermore, structures should not endanger the qualities of the natural environment excessively. And, it has to be considered that a large proportion of the societal wealth is invested into the development and maintenance of structures as a key part of the build environment.

In summary, structures have to fulfil the following requirements²:

¹Note that the functionality of the structure is THE basic requirement, since providing functionality for societal activities is the reason for creating structures. Functionality might include a multitude of further requirements like aesthetics, etc. However, in the present context we are primarily interested in the requirements related to the load bearing behaviour of structures.

²In regard to the load bearing behaviour.

1. Safety of personnel
2. Safety of the environment
3. Cost effectiveness

Let's look at the third requirement first. Cost effectiveness might be assessed with a cost-benefit analysis, i.e. where the expected benefit of the structure is maximized, or, the expected cost of the structure is minimized.

1.2. Optimization

We will approach this topic by introducing a practical example. Consider the following task that could have given to you in connection to a consulting project:

You have to design a beam that has to span $l = 10$ m and has to carry a load Q . The material that is available is glued laminated timber (Glulam) and the cross-section is specified to be rectangular with a width of $b = 300$ mm and height h .

For a project like this, you would work through the following steps:

1. *Conceptual Design*: this step includes the definition of the requirement for the structure and the setup of the structural concept. Requirements: after a careful discussion with the client it is specified that neither people nor environment are endangered by a possible failure of this structure. Thus the requirement is defined to be cost effectiveness only. The functionality does not require any permissions on deflection. Structural failure is the scenario that leads to consequences. The service period of the structure is specified to be 50 years.

Structural concept: You might choose simple support instead of fixed support due to economic reasons. (a moment resisting support is more expensive).

2. *Structural analysis*: in this phase, the loads are specified and the load effects and the load bearing capacity are estimated.

The load is given in this project to be a uniform distributed load that is represented by its 50 years maximum value Q . The material property of interest in this case is the bending strength of the Glulam F_m . The situation is illustrated in Figure 1.1.

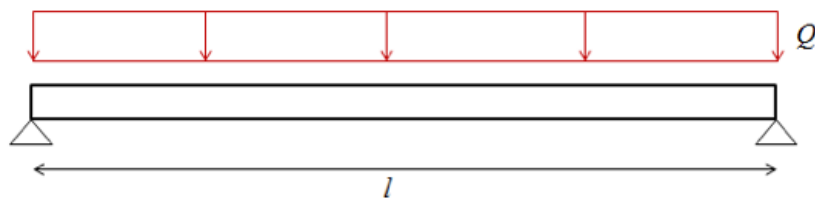


Figure 1.1: Beam geometry.

The maximum effect of the load S is the bending moment at mid-span that is given with $Ql^2/8$. The elastic bending load bearing capacity of the rectangular cross-section R is $F_m b h^2/6$.

3. *Dimensioning*: in this phase, limit states are formulated. Limit states are equations that describe events like, "failure" or "excessive deflection" that are relevant in regard for the description of requirements to the structure. Structural failure, for example, is defined as the event when the load on a structure is larger than its load bearing capacity, or the difference between the load bearing capacity and the load becomes negative. The corresponding limit state is general referred to as Ultimate Limit State (ULS) and for this example expressed as:

$$g(R, S) = R - S = \left(\frac{bh^2}{6}\right) F_m - \left(\frac{l^2}{8}\right) Q \leq 0 \quad (1.1)$$

The limit state equation contains different variables, some of them are uncertain or random, as Q, F_m , some of them are given or can be controlled, as l, b, h . We can now formulate the failure probability or the reliability index β as a function of the design choice h (we have chosen h as being the most effective design variable). If we assume that Q and F_m are represented as independent and both normal distributed random variables, the reliability index corresponding to a 50 years period³ can be assessed as a function of h using Eq. (1.2):

$$\beta(h) = \frac{\mu_R(h) - \mu_S}{\sqrt{\sigma_R^2(h) + \sigma_S^2}} = \frac{\mu_{F_m} \cdot bh^2/6 - \mu_Q \cdot \frac{l^2}{8}}{\sqrt{(\sigma_{F_m} \cdot bh^2/6)^2 + \left(\sigma_Q \frac{l^2}{8}\right)^2}} \quad (1.2)$$

Mean value and standard deviation for the two random variables are given as in Table 1.1. Based on this and remembering that the failure probability is related to the reliability index as $P_f = \Phi(-\beta)$, we can graphically represent the relation of h to the failure probability and to the reliability index correspondingly, cf. Figure 1.2. As it can be seen, a functional relationship between the design variable h and the failure probability can be established. However, the question remains **how it can be decided what failure probability is sufficiently low**.

Table 1.1: Representation of load and resistance. The coefficient of variation V_X is defined as $V_X = \sigma_X/\mu_X$.

Variable	Distribution	Ch. value	μ_X	V_X
F_m	Normal	20 MPa	26.6 MPa	0.15
Q	Normal	39.4 N/mm	24.1 N/mm	0.30
$R = W_{el} f_m$	Normal		$W_{el} \mu_{f_m} = 1330h^2$ Nmm	0.15
$S = \left(\frac{l^2}{8}\right) Q$	Normal		$\left(\frac{l^2}{8}\right) \mu_Q = 301.25 \cdot 10^6$ Nmm	0.30

1.2.1. Identification of the optimal decision by minimizing expected costs

If cost effectiveness is taken as a criterion, as in this example, the optimal choice of h can be identified by minimizing the expected costs of the structure. The following costs are considered. The parameter values, units and symbols are summarized in Table 1.2.

³Note that in the entire example the reference period for the reliability index and the failure probability is 50 years.

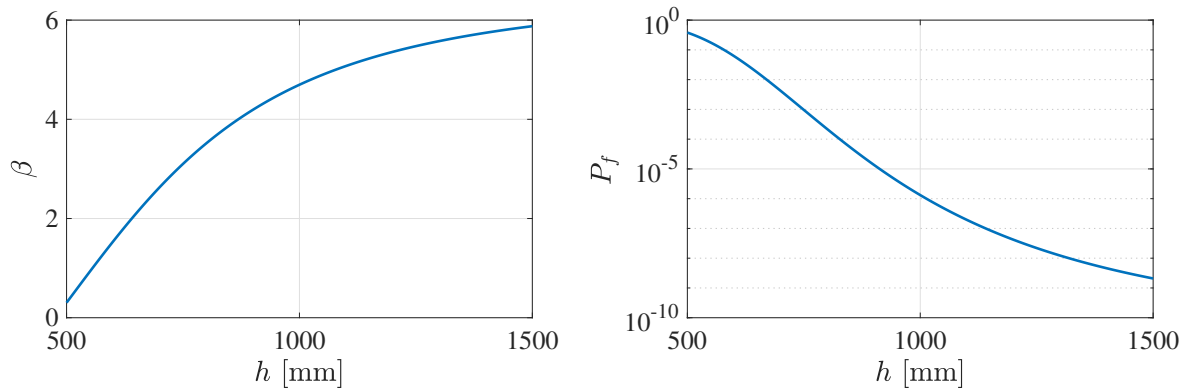


Figure 1.2: Reliability index β (left) and probability of failure P_f (right) as a function of h .

Table 1.2: Assumed values and relations.

Variable	Symbol	Value, relations and units
Cost of timber	c_{glulam}	4,000 NOK/m ³
Beam length	l	10,000 mm
Fixed construction costs	C_0	50,000 NOK
Variable construction costs coefficient	C_I	$b \cdot l \cdot c_{glulam}$
Indirect costs of failure	H	$r_{HC_0} \cdot C_0 = 20 \cdot C_0$

- **Construction costs:** These are the costs for implementing the structure. It covers transport and installation cost but also the material cost of the structural element. Correspondingly, the construction costs consist of a part depending on the chosen beam height and a fixed part C_0 (that is independent from the chosen height h).

The construction cost is a function of h , i.e. the possible decision alternatives for risk reduction. See Eq. (1.3).

$$C_{constr}(h) = C_0 + l \cdot c_{glulam} A_s = C_0 + l \cdot c_{glulam} (b \cdot h) = C_0 + C_I h \quad (1.3)$$

- **Expected failure costs:** The failure costs include the costs for re-installing the structure after failure, i.e. the construction costs. Additionally a cost term H is introduced. H contains all failure costs beyond the pure re-installation costs, e.g. damage to equipment, cleaning costs, downtime of services that are supported by the structure, etc. The failure cost is multiplied with the probability of failure in order to express the *expected* failure cost. Note that another term for the expected failure cost would be **Risk**.

$$E[C_{f,ULS}](h) = C_{f,ULS}(h) \cdot P_f(h) = [C_{constr}(h) + H] \cdot P_f(h) \quad (1.4)$$

- **Objective function:** The objective function to be minimized represents the total expected costs, cf. Eq. (1.5). The minimum is identified subject to h (the decision variable) and represents the optimal balance between investments into the structure and the expected adverse consequences.

$$\begin{aligned}
 E[C_{tot}](h) &= C_{constr}(h) + E[C_{f,ULS}](h) \\
 &= [C_0 + C_I h] + [C_{constr}(h) + H] P_f(h) \\
 &= [C_0 + C_I h] + [C_I h(h) + C_0(r_{HC_0} + 1)] P_f(h)
 \end{aligned} \tag{1.5}$$

$$h_{opt} = \arg \min_h E[C_{tot}](h) \tag{1.6}$$

The numerical minimization of the total expected costs can be performed with Matlab® or Python. The optimal height of the beam is found to be 781 mm.

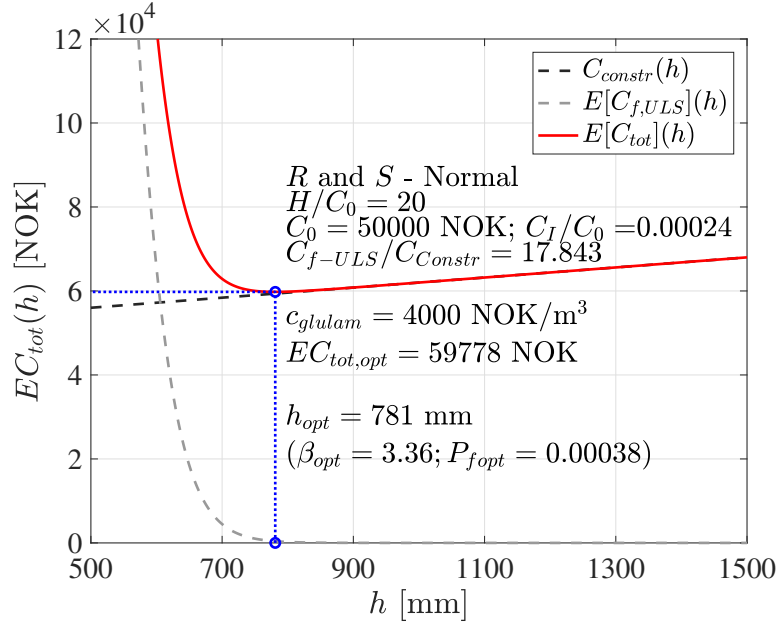


Figure 1.3: Expected monetary costs with economic optimum and input values.

The sensitivity of the optimal cross-section height is plotted in Figure 1.4 for different values of the input parameters.

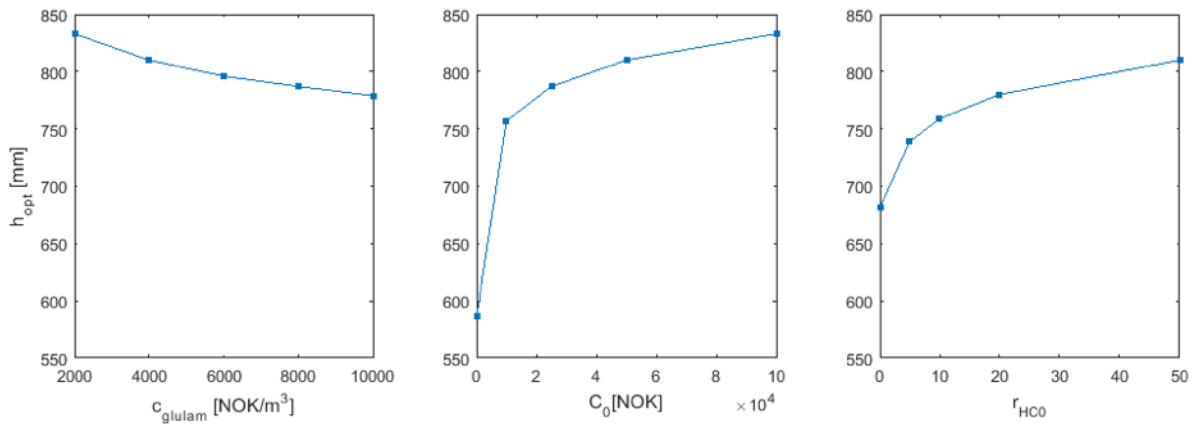


Figure 1.4: Sensitivity analysis for h_{opt} .

By formulating the objective function as in Eq. (1.5), it is assumed that the structure is replaced if it fails within the 50 years life cycle and therefore the construction costs have to be paid

(assuming we would then implement the same design variable h and the parameters of the cost function, C_0 and C_I are still valid) in addition to the damages to property beyond the construction itself, H .

1.3. Interest

What we did not take into account so far in this simple example is the fact that the expenses that are contained in the objective function take place at different times. I.e. when the structure is constructed and when the structure fails. In order to be comparable, the expenses have to be transformed to the value at the same point in time. The logical choice is the time when the decision takes place, i.e. when the structure is constructed that is now, presently. The failure costs take place in the future and have to be transferred back in time, given an assumed interest rate.

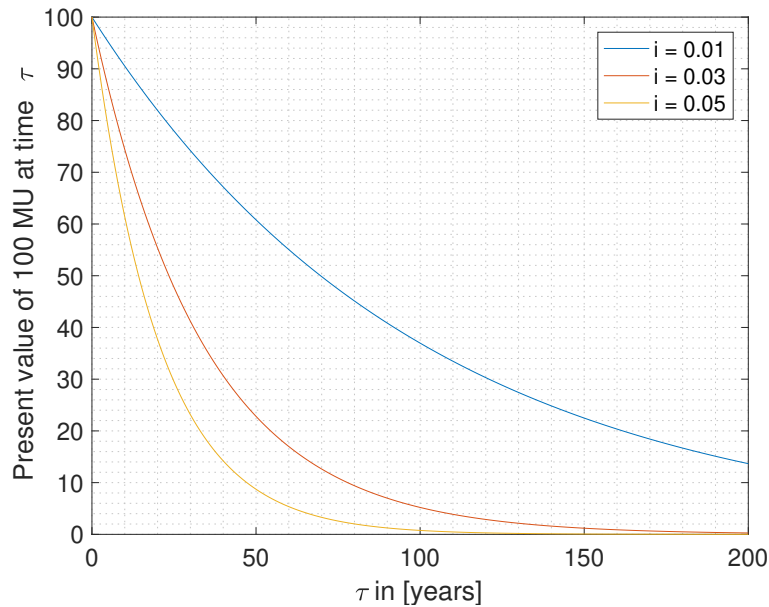


Figure 1.5: Illustration of the net present value: The graph shows the monetary units (MU) at time τ that correspond to 100 MU at present. The three lines correspond to the three different interest rates $i = 0.01$, $i = 0.03$, $i = 0.05$.

The net present value (NPV) of the cost $C_{t=\tau}$ that is generated in τ years can be generally transferred back as:

$$C_{t=0} = C_{t=\tau} \frac{1}{(1+i)^\tau} \quad (1.7)$$

where i is the expected average interest rate. See the illustration in Figure (1.5).

The objective function in Eq. (1.5) can be modified correspondingly, i.e. all cost terms that are paid in the future are scaled down by $(1+i)^{-\tau}$:

$$\begin{aligned} E[C_{tot}](h) &= C_{constr}(h) + E[C_{f,ULS}](h) \frac{1}{(1+i)^\tau} \\ &= [C_0 + C_I h] + [C_{constr}(h) + H] P_f(h) \frac{1}{(1+i)^\tau} \end{aligned} \quad (1.8)$$

Accounting for the interest has a significant effect on the result. Let us consider the beam example above. Assuming an interest rate of 3%, i.e. $i = 0.03$ and an average failure time after 25 years, the result presented in Figure 1.6 is obtained. Compared to the result without interest, both, the optimal height and the reliability are smaller. ($h_{opt} = 757$ mm instead of 781 mm and $\beta_{opt} = 3.16$ instead of 3.36).

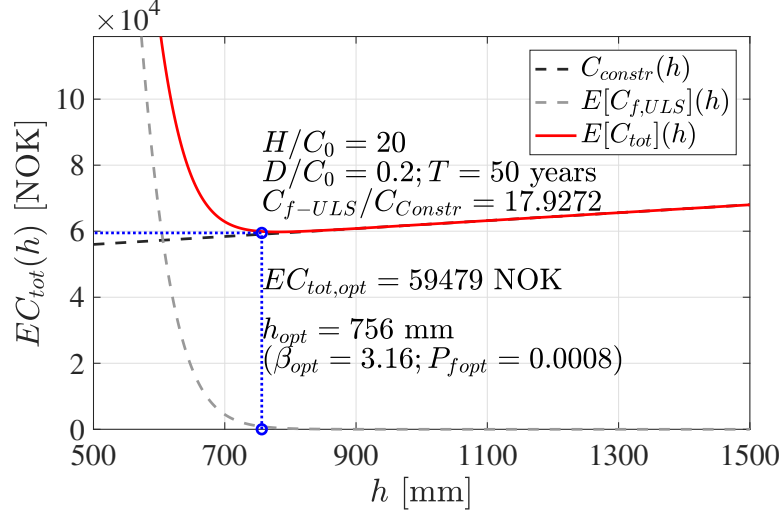


Figure 1.6: Expected monetary costs for net present value.

QUESTION: Why is h_{opt} and the optimal reliability index β_{opt} getting smaller when the interest is considered?

(space for your notes)

1.4. Continuous renewal assumption

In the example above, a simple structure was optimized for single use in a finite service period of 50 years. This perspective is valid if we want to optimize for the owner of the structure and for the case where the specific structure is demolished after the planned service period.

A more general perspective is the perspective from legislation and code writing. We are more generally interested in the reliability level that structures should have, from a societal perspective, in order to specify the target reliability level for structural codes and standards.

If we do so, the structures are no longer optimized individually for their finite service periods. What is optimized is the decision on structures that support a societal activity. It is assumed that the societal activities are continuously ongoing, whereas the structures that support these activities are continuously renewed.

Hereby, the act of renewal can become necessary for two reasons:

- a) the structure becomes obsolete and is replaced by a newer, more modern one;
- b) the structure failed.

The optimization is then performed on a yearly basis (with yearly probability of failure and yearly renewal costs). Regarding the simple previous beam example and representing the height of the beam h with a design parameter $h = p$, the interest rate with $\gamma = i$ the objective function would then be written as:

$$\begin{aligned} E[C_{tot}(p)] &= C_{const}(p) + E[C_f(p)] \frac{1}{\gamma} + E[C_{obs}(p)] \frac{1}{\gamma} \\ &= [C_0 + C_I p] + [C_0 + C_I p + H] \frac{\lambda P_f^{(1a)}(p)}{\gamma} + [C_0 + C_I p + D] \frac{\omega}{\gamma} \end{aligned} \quad (1.9)$$

Equation (1.9) consist of the following:

- Construction costs: $C_{const}(p) = C_0 + C_I p$ are the construction cost with a part $C_I p$ that is proportional to p , and a part C_0 that is independent of p .
- Expected annual failure cost: $E[C_f(p)] = [C_{const}(p) + H] \lambda P_f^{(1a)}(p)$, here H is introduced as a nonstructural failure cost, the construction costs are part of the failure cost as a similar decision (similar investment into p) is expected after failure. The failure cost are considered as annual expectations, so they are multiplied by the annual failure probability $P_f^{(1a)}(p)$. λ is introduced as the rate of a Poisson process with $\lambda = 1$ in order to interpret the annual failure probability as a rate of occurrence.
- γ is the annual interest rate that is selected as the societal interest rate $\gamma = \gamma_S$ if the preferences of the society are represented, or as the private interest rate $\gamma = \gamma_E$, if the optimisation is done relative to entrepreneurial preferences.
- Expected obsolescence cost: $E[C_{obs}] = [C_{const}(p) + D] \omega$, here D represents the demolition cost and ω is representing the expected rate at which the decision on p becomes obsolete.

The design parameter $p = p^*$ that is minimising the expected total cost is found by

$$\begin{aligned} \frac{d}{dp} \left\{ C_0 + C_I p + [C_0 + C_I p + H] \frac{\lambda P_f^{(1a)}(p)}{\gamma} + [C_0 + C_I p + D] \frac{\omega}{\gamma} \right\} \bigg|_{p=p^*} &\equiv 0 \\ \Rightarrow \frac{C_0 + C_I p^* + H}{C_I} &= \frac{1 + P_f^{(1a)}(p^*) \frac{1}{\gamma} + \frac{\omega}{\gamma}}{-\frac{dP_f^{(1a)}(p)}{dp} \bigg|_{p=p^*} \frac{1}{\gamma}} \end{aligned} \quad (1.10)$$

A simplified approximation can be formulated by considering that for typical structural engineering decision situations it is $P_f(p^*) \ll \omega + \gamma$:

$$\frac{C_I \cdot (\gamma + \omega)}{C_0 + C_I p^* + H} \approx -\frac{dP_f^{(1a)}(p)}{dp} \bigg|_{p=p^*} \quad (1.11)$$

Note that the expression for the ratio between the marginal safety costs (C_I) and the total failure costs ($C_0 + C_{Ip^*} + H$) can be further simplified. Usually, $C_{Ip^*} \ll C_0$ since the construction costs are dominated by the fixed costs. In these cases, the ratio on the left side of 1.11 is approximately equal to $C_I / (C_0 + H)$.

$$\frac{C_I \cdot (\gamma + \omega)}{C_0 + H} \approx - \left. \frac{dP_f^{(1a)}(p)}{dp} \right|_{p=p^*} \quad (1.12)$$

The optimal design p^* can be estimated from Eq. (1.11) and (1.12) when the functional relationship $P_f^{(1a)}(p)$ is known analytically in the vicinity of $p = p^*$. In other cases, the equation must be solved numerically. It is interesting to note that the optimal choice of p does only depend on $dP_f^{(1a)}(p)/dp$ and not on the absolute values of $P_f^{(1a)}(p)$, i.e. in the optimisation subject to p it is only necessary to represent the gradual change of $P_f^{(1a)}$ by changing p . This makes the results of the optimisation insensitive against the non-inclusion of failure scenarios those probability of occurrence cannot be influenced by gradually changing p . An example is the potential presence of gross human error. This definitely has a strong effect on the estimated absolute probability of failure, however, if it is assumed that the probability of failure conditional on gross human error is not (or very weakly) influenced by the particular choice of p , the optimum choice of $p = p^*$ derived by neglecting gross human error is still valid.

A reliability requirement can be determined from the above optimisation, i.e. if the choice of $p = p^*$ is minimising the expected total costs, then $P_f^{(1a)}(p^*)$ and $\beta(p^*) = -\Phi^{-1}(P_f^{(1a)}(p^*))$ are the corresponding failure probability and reliability index of the optimal choice. If for a decision problem at hand, it can be assumed that the attributes of the decision problem are sufficiently similar to the assumed attributes contained in Eq. (1.11), this probability / reliability should be chosen as a target. However, since such a target is expressed in terms of an absolute value it is of high importance that the reliability target is only applied to sufficiently similar scenarios.

The example will be worked further in the exercise class.

1.5. Risk acceptance

The economic optimisation that is obtained from Eq. (1.11) ensures that resources are expended optimally from a financial point of view. But it does not guarantee that the optimum decision $p = p^*$ is consistent with the societal preferences in regard to life safety. I.e. if human lives are at risks it has to be ensured that societal resources are allocated efficiently for preventing possible fatalities associated with structural failure. The marginal life-saving costs principle can be applied to derive risk acceptance limits (ISO 2015; Blomquist 1979). The principle ensures that the societal resources are allocated to efficient risk-reducing measures, i.e. measures that can save one additional life at a cost that the society is willing (or able) to pay. This is here referred to as the Societal Willingness To Pay (SWTP). The SWTP is multiplied with the expected number of fatalities given structural failure and inserted to the objective function in Eq. (1.8) which leads to the specification of the acceptable domain in Eq. (1.13), ISO 2015; Köhler et al. 2019.

$$- \frac{dP_f^{(1a)}(p)}{dp} \leq \frac{C_I (\gamma_S + \omega)}{SWTP \cdot N_F} = K_1 \quad (1.13)$$

The constant K_1 is introduced in (ISO 2015) as an indicator. Here, also typical values of K_1 are given. The minimum acceptable design (p_{acc}) just satisfies the inequality in Eq. (1.13) and depends on:

1. the uncertainty involved in the problem through the term $dP_f^{(1a)}(p)/dp$,
2. the marginal safety costs C_I measured in monetary units,
3. the societal ability to pay for saving one statistical life $SWTP$, measured in monetary units,
4. the number of expected fatalities given structural failure N_F , and
5. the obsolescence rate ω and the societal interest rate γ_S .

The condition for which the optimal design is within the acceptable domain (i.e. $p^* \geq p_{acc}$) is obtained inserting Eq. (1.13) into Eq. (1.12) :

$$\frac{C_I \cdot (\gamma_S + \omega)}{C_0 + H} \leq K_1 \quad (1.14)$$

How the $SWTP$ can be quantified is discussed in the following:

1.5.1. The societal willingness to pay for saving one additional life (SWTP)

The $SWTP$ corresponds to the amount of money that should be invested into saving one additional life for a given life risk reduction activity. The $SWTP$ is expressed in monetary terms and can be derived e.g. by means of the Life Quality Index (LQI) Nathwani et al. 1997. The LQI index informs about the revealed preferences of the society for investments into life safety by estimating the ability of a national economy to allocate monetary resources in risk reducing activities.

Values of the $SWTP$ that are derived from the LQI index are shown in Table 1.3 for some selected countries, a more complete Table is given in ISO 2015. The values in the Table are provided for different societal interest rates γ_S for considering the discounting of the future costs and benefits of the risk-reducing measures.

Table 1.3: $SWTP$ for five selected countries from ISO 2015 for γ_S equal to 2 %, 3 % and 4 % (all numbers in thousands purchasing power parity US dollars).

Country	2008 GDP	SWTP		
		$\gamma_S = 2\%$	$\gamma_S = 3\%$	$\gamma_S = 4\%$
Australia	35 624	4 840	4 298	3 843
Brazil	9 517	804	712	634
Norway	49 416	3 937	3 500	3 129
Mali	1 043	54	48	43
US	42 809	3 187	2 833	2 543

The use of monetary units in the context of risk acceptance should not be misunderstood. Indicators such as the cost of a statistical life or the money the society is willing (or better “able”)

to pay for saving one additional life (*SWTP*) are not to be considered the monetary value of a individual person's life, which has of course a value that cannot be expressed in monetary terms. The mentioned indicators are derived from small marginal changes in the probability of having a fatality caused by failure.

1.5.2. Expected values

The cost terms H , C_0 , C_I and N_F but also ω and γ_S are in general uncertain. A similar failure event might lead to different consequences depending, for example, on the occupancy of the building. All these terms need to be evaluated up to their expected value only since they affect the expected total cost linearly (compare Equation (1.8)). If in addition it can be assumed that they are independent of each other their probability density functions are irrelevant for the optimisation, see Raiffa et al. 1961. It follows that the values for the optimisation are not the highest or extreme values associated with the corresponding events, but the expected ones. However, the expectation operator is omitted in the text in order to ease the notation.

When evaluating the expected values of H and N_F , attention should be paid in cases where the assets and the lives lost in the event of failure might be themselves a significant cause of failure. A representative example might be the failure of a grandstand in a stadium due to static or dynamic load induced by the occupants. In this case, the likelihood of failure increases with the occupancy. This effect should be considered both in the formulation of the objective function and in the representation of the variables.

1.6. Conclusion of Chapter 1

The principles of structural optimisation and risk acceptance criteria for direct structural design and for code-making have been briefly discussed with the aim of providing a transparent background risk and reliability based design. The design requirements are derived by monetary optimisation minimising the expected costs over time. The analysis of the optimisation problem showed the four most important aspects that influence the optimal reliability levels. These aspects are

- the ratio between total failure costs and marginal safety costs;
- the uncertainty involved in the problem;
- the obsolescence rate; and
- the discounting rate.

The acceptability of the monetary optimum is not always satisfied. The acceptance criterion can be derived by the Marginal Life-Saving Cost Principle that was here implemented using the Life Quality Index (LQI). It was demonstrated that the acceptable threshold depends on:

- the marginal safety costs;
- the societal willingness to pay for saving one additional life;
- the number of expected fatalities given failure; and

- the obsolescence and the societal interest rates.

All considerations in this chapter depend on reasonable assumptions for the consequence and for the probability of failure as a function of the design decision variables. Whereas it is sufficient to represent the consequences as expected values, the accurate representation of the failure probability is of utmost importance. Accordingly, *more advanced reliability methods* and *probabilistic load and resistance modelling* will be treated further in the upcoming chapters of this lecture.

2. Application examples

E1.1.

Consider the structure in Figure 1.1. The rectangular cross-section dimensions are parametric on the cross-section height h . The width b and span l are known without uncertainty, being $b = 300$ mm and $l = 10$ m. The resistance f_m and load Q are random variables, whose distributions were given in Table 1.2.

In order to get started with a Python script you may import the necessary Python packages and do some presettings for the output:

```

1 import numpy as np
2 import scipy as sp
3 import matplotlib.pyplot as plt
4 import scipy.stats
5 fontsizes=18
6 plt.rcParams.update({'font.size': fontsizes})
7 plt.rcParams.update({"font.family": "serif"})
8 plt.rcParams.update({"mathtext.fontset" : "cm"})
9 plt.rcParams.update({'font.serif': 'Times New Roman'})
10 plt.close('all')
```

Then you specify some input parameters to the problem and compute the failure probability and the reliability index as functions of the height h :

```

1 # Geometry
2 l = 10000          # [mm] span
3 b = 300            # [mm] width
4 #=====
5 # Material properties
6 mu_fm = 26.6       # [MPa] mean material resistance
7 cov_fm = 0.15      # coeff. of variation
8 std_fm = mu_fm*cov_fm # [MPa] standard deviation
9
10 mu_R = mu_fm*b/6
11 std_R = std_fm*b/6
12 #=====
13 # Load
14 mu_q = 24.1        # [N/mm] mean load
```

```

15 cov_q = 0.3                # coeff. of variation
16 std_q = mu_q*cov_q        # [MPa] standard deviation
17
18 mu_S = (1**2/8)*mu_q
19 std_S = (1**2/8)*std_q
20 #=====
21 # Cost model
22 c_gl = 4000*1E-9           # [NOK/mm^3] Cost Glulam
23 C0 = 50000                 # [NOK] Fixed cost of construction
24 H = C0*20                  # Direct cost of failure
25 C1 = 1*b*c_gl              # [NOK/mm] Variable cost
26 i = 0.03                   # interest rate
27 T = 25                     # Service life
28 C_tau = 1/((1+i)**T)       # discount factor
29
30 #=====
31 # Reliability index as a function of the decision variable
32 BETA = lambda h,mu_S,std_S: (mu_R*h**2-mu_S)/(((std_R*h**2)**2+(std_S)**
33                                     2)**0.5)
34
33 h1 = np.linspace(500,1500, num=10000)
34 beta = BETA(h1,mu_S,std_S)
35 PF = sp.stats.norm.cdf(-beta)

```

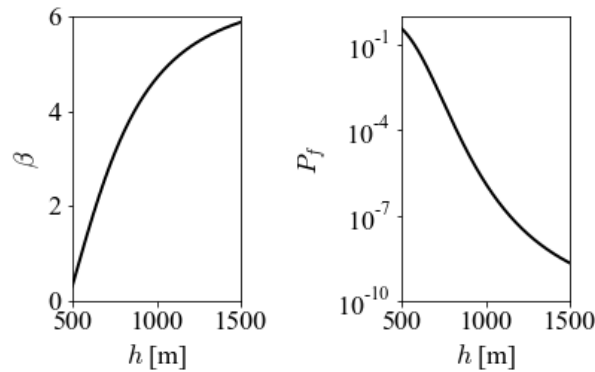
The results can now be plotted:

```

1 # Plot
2 plt.figure()
3
4 plt.subplot(121)
5 plt.plot(h1,beta, color='black',lw=2)
6 plt.xlabel('$h$ [m]',fontsize=fontsizes)
7 plt.ylabel('$\beta$',fontsize=fontsizes)
8 plt.xlim(500,1500)
9 plt.ylim(0,6)
10 plt.subplot(122)
11 plt.plot(h1,PF, color='black',lw=2)
12 plt.yscale('log')
13 plt.xlabel('$h$ [m]',fontsize=fontsizes)
14 plt.ylabel('$P_f$',fontsize=fontsizes)
15 plt.xlim(500,1500)
16 plt.ylim(1e-10,1e0)
17 plt.tight_layout()
18 plt.show()

```

Output:



E1.1.1. Structure with finite life (owner perspective) – neglecting interest rate

Make your own Matlab script to reproduce the results in Figure 1.3. The minimum expected costs are obtained by minimization:

```

1 # Not discounted costs
2 C_c = C0 + C1*h1                                # Construction cost
3 EC_f = (C_c + H)*PF                             # Expected failure cost
4 ECtot = C_c + EC_f                             # Objective function
5 # Compute optimum
6 idx_opt = ECtot.tolist().index(min(ECtot))
7 h_opt = h1[idx_opt]
8 beta_opt = beta[idx_opt]
9 Pf_opt = PF[idx_opt]
10 ECtot_min = ECtot[idx_opt]
11 # print("Optimum cross-section height: %.2f m" % (h_opt))
12 s1 = "Optimum cross-section height: {h:.0f} m\n".format(h=h_opt)
13 s2 = "Optimum probability of failure: {pf:.2e}\nOptimum beta: {b:.2f}\n"
    .format(pf=Pf_opt,b = beta_opt)
14 s3 = "Minimum expected total cost = {c:.0f} NOK\n".format(c = ECtot_min)
15 print(s1+s2+s3)

```

Optimum cross-section height: 781 m

Optimum probability of failure: 3.84e-04

Optimum beta: 3.36

Minimum expected total cost = 59778 NOK

And plot the objective function.

```

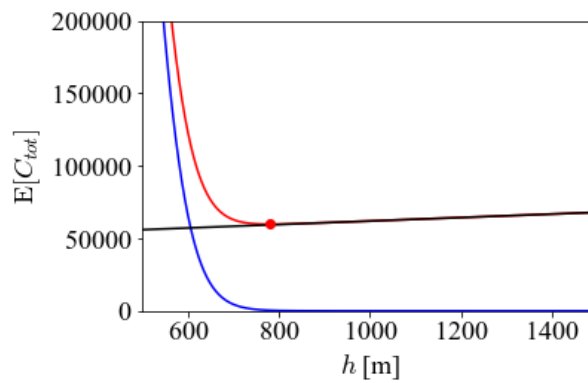
1 # Plot function
2 def pltECost(ECT_list,EC_f_list,EC_list):
3     # plt.figure(idx_plt)
4     plt.plot(h1,ECT_list[0],ECT_list[1],label = ECT_list[2])
5     plt.plot(h1,EC_f_list[0],EC_f_list[1],label = EC_f_list[2])
6     plt.plot(h1,EC_list[0],EC_list[1],label = EC_list[2])
7     plt.ylim(0,200000)
8     plt.xlim(500,1500)
9     plt.xlabel('$h$ [m]',fontsize=fontsizes)
10    plt.ylabel('$\mathrm{E}[C_{\text{tot}}]$',fontsize=fontsizes)
11    plt.tight_layout()

```

```

12     plt.ticklabel_format(style='sci',useMathText=True)
13     return
14
15 # Plot undiscounted costs
16 plt.figure()
17 plt.ECost([ECtot,'r','Undiscounted  $\mathrm{E}[C_T]$',[EC_f,'b','
                                Undiscounted  $\mathrm{E}[C_F]$',[
                                C_c,'k','$C_c$'])
18 # plt.plot(np.array([h_opt, h_opt, 0]),np.array([0, ECtot_min, ECtot_min
                                ]),':b')
19 plt.plot(h_opt,ECtot_min,'or')
20 plt.show()$$ 
```

Output:



E1.1.2. Structure with finite life (owner perspective) – present value of costs

Optimize the cross-section of the beam from the perspective of the owner, accounting for an interest rate of 3% and an average failure time τ of 25 years, i.e. reproduce the results in Figure 1.6.

E1.1.3. Structures with infinite renewal (societal perspective) – present value of costs

The optimization is now performed from the societal point of view. Under the infinite renewal assumption, structures are considered to be re-constructed after failure and demolished and re-constructed once they become obsolete.

Estimate the minimum expected cost and the optimal cross-section, probability of failure $P_{f,50yr}$ and reliability β_{50yr} (related to a 50 year period) by following this perspective. Some additional values to be used in the calculation are collected in Table 1.4. The 1 year maxima distributed load Q is assumed to be Normal with parameters $\mu_{Q,1yr} = 15.20$ N/mm and $\sigma_{Q,1yr} = 6.11$ N/mm.

Solutions The expected risk is plotted in Figure 1.7. The optimum cross-section height is $h_{opt} = 753$ mm. The reliability and failure probability associated with a 50 year period are computed assuming that failure between different years is independent. They result in $\beta_{50yr} = 3.12$ and $P_{f,50yr} = 9.09 \cdot 10^{-4}$, respectively.

Table 1.4: Assumed values and relations.

Variable	Symbol	Value, relations and units
Demolition costs	D	$r_{DC_0} \cdot C_0 = 0.2C_0$
Design life	T	50 years
Interest rate	i	3%
Rate of obsolescence	ω	$1/T$
Societal interest rate	γ_s	2%
Gross domestic product	g	(ISO 2398:)
Societal Willingness To Pay to save an additional life	$SWTP$	32,100,000 NOK (ISO 2394:2015 Table G.2, delta-regime)

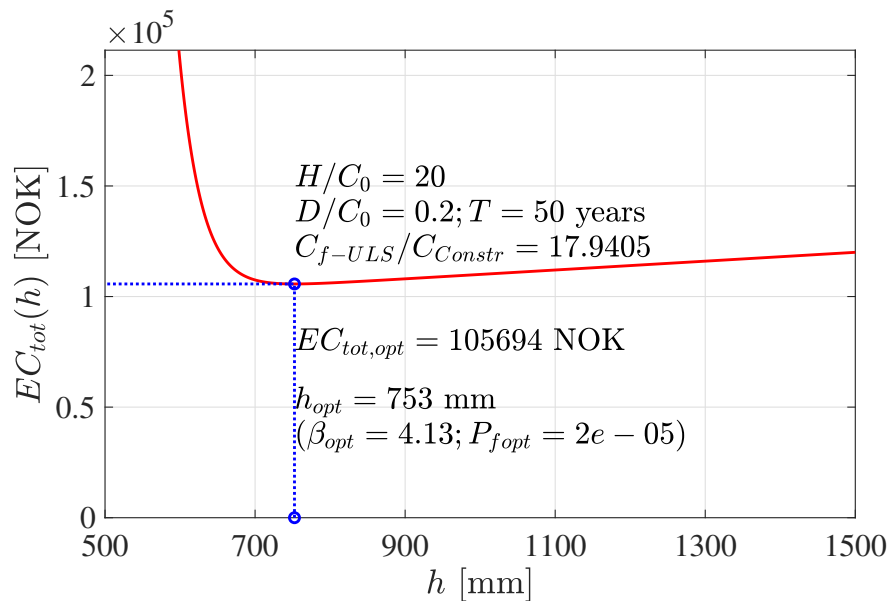


Figure 1.7: Expected cost assuming infinite renewal (E1.1.3).

E1.1.4. Risk acceptance, part 1

The principle of the marginal life-saving costs is applied to this example in order to evaluate the maximum accepted risk. The marginal life-saving cost principle is implemented through the use of the life quality index (LQI). The definition of acceptability is based on the requirement that all *efficient life-saving measures* have to be performed for societal risk acceptance. The *acceptable region* (i.e. acceptable values of h in our example) is equivalent to sufficient expenditure for life saving investments. The acceptable region is defined by the values of h satisfying the following inequality in Eq. (1.13). The criterion depends on the expected number of fatalities given structural failure N_F . Define the boundaries of the acceptable region for $N_F = 0.001$ and recalculate, h_{opt} , if necessary.

Solutions The solution is presented in Figure 1.8. Assuming $N_F = 0.001$, the acceptable region is $h > 656$ mm. This means that the optimal design found minimizing the expected costs is acceptable. Therefore, $h^* = h_{opt} = 753$ mm simultaneously minimizes monetary costs and ensures an acceptable level of risk.

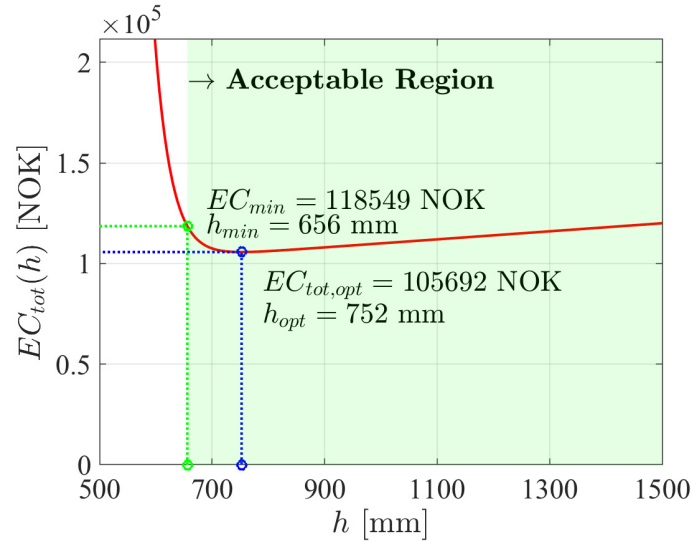


Figure 1.8: Monetary optimization with acceptance criteria $N_F = 0.001$ (E1.1.4).

E1.1.5. Risk acceptance, part 2

Re-do the calculations in the previous example for $N_F = 10$.

Solutions The solution is presented in Figure 1.9. Assuming $N_F = 10$ the acceptable region is $h > h_{min} = 971$ mm, meaning that the optimal design found minimizing the expected costs, i.e. $h_{opt} = 753$ mm is not acceptable. It follows that the design that simultaneously minimizes monetary costs and ensures an acceptable level of risk corresponds to $h_{opt} = h_{min} = 971$ mm.

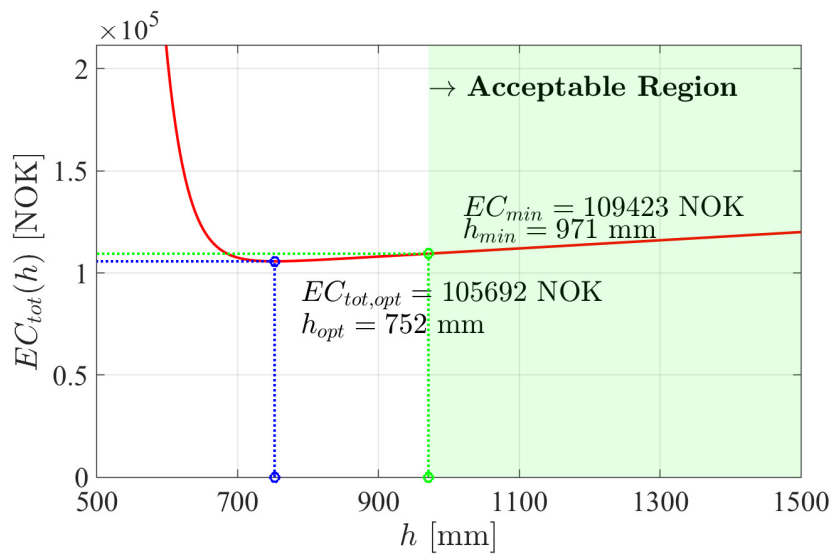


Figure 1.9: Monetary optimization with acceptance criteria $N_F = 10$ (E1.1.5).

3. Practical exercises

P1.1. Fixed-end beam optimization

Consider the beam in Figure 1.10. In this exercise, the beam is fixed at both ends. The span and the square cross-section are $l = 10$ m and $b = 300$ mm, as before. The basic random variables are also Normal distributed with mean and coefficient of variation given in Table 1.1. All other inputs are given in the example presented in the lecture.

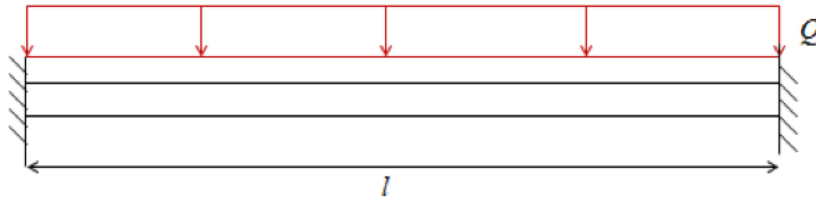


Figure 1.10: Boundary conditions of the beam in P1.1.

P1.1.1. Fixed-end beam optimization

Find the new optimum beam height h assuming finite life and using discounted costs.

Solutions The new optimum cross-section height is lower than for the simply supported beam: $h_{opt} = 623$ mm. However, it yields a higher optimum reliability index: $\beta_{opt} = 3.22$. This is due to the fact that the load effects of this boundary conditions are ca. 66% lower than for the simply supported beam.

P1.1.2. Additional efficient investment

Fixing the extremes of a beam is a more costly activity when compared to simply supporting them. The costs of producing a fix joint for beams of relatively similar cross-section dimensions can be assumed to be equal. Therefore, this cost can be considered part of the fixed construction costs C_0 .

Calculate the maximum additional investment that is sensible to make in fixing the extremes of the beam before it becomes a less optimum alternative.

Solutions The maximum additional investment in the fixed cost is found to be 2003 NOK. Note that this solution provides a higher reliability ($\beta = 3.22$) for the same expected optimum cost.

P1.2. Uncertainty in the modelling of the supports

A visual inspection of the beam that was designed in the previous practical exercise (see Figure 1.10) is conducted after construction due found defects in the joints. The uncertainties associated with the outcomes of the evaluation are given by the probability of both joints not

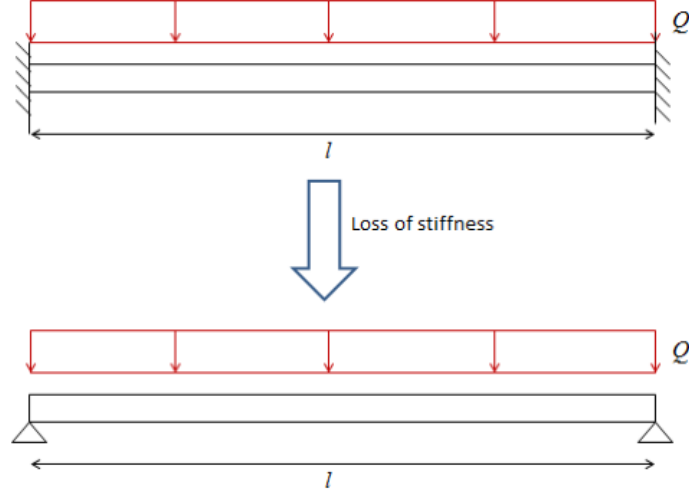


Figure 1.11: Lost of stiffness.

being stiff enough being $P_{joint} \approx 10^{-1}$. This is equivalent to having a simply supported static scheme instead of having the ends fixed, see Figure 1.11.

According to the the law of total probability, the probability of failure is now:

$$P_f^* = \Pr(\text{stiff joints} \cap (M_R(h_{opt}) \leq M_Q|\text{fix-fix})) + \Pr(\text{non-stiff joints} \cap (M_R(h_{opt}) \leq M_Q|\text{hinge-hinge})) \quad (1.15)$$

For independent events, it can be rewritten as:

$$P_f^* = \Pr(\text{stiff joints}) \cdot \Pr(M_R(h_{opt}) \leq M_Q|\text{fix-fix}) + \Pr(\text{non-stiff joints}) \cdot \Pr(M_R(h_{opt}) \leq M_Q|\text{hinge-hinge}) \quad (1.16)$$

1. Evaluate the probability of failure and the reliability index, and comment the obtained values.
2. Evaluate, the total expected costs and compare them to the ones obtained in P1.1.1. Comment the results. Evaluate and compare the expected failure costs only.
3. Imagine the reparation of the joints will cost ν NOK, what can be a possible decision criterion for deciding if repairing the joints or not?

Question::

How safe is safe enough?

Describe how a reliability criterion may be established based on optimization.

- Give an example for a possible *objective function*
- Represent the objective function graphically and explain the shape of the graph.
- What is the difference between the optimization of a single structure and the optimization of a large portfolio of (similar) structures? How are the objective functions different?
- How can the concept be adapted in order to represent societal preferences to prevent human fatalities?

Chapter 2

Structural reliability - Problem Statement and Monte Carlo Method

The fundamental structural reliability problem is revisited and extended to linear limit states of n normal distributed random variables.

The Monte Carlo simulation method is introduced as the simplest way to approach more realistic reliability problems.

It is demonstrated how samples of correlated random variables can be generated.

Latin Hypercube Sampling is introduced as a variance reduction method for Monte Carlo simulations.

1. Theoretical aspects

1.1. Prelude

Probability of failure is, as seen in the previous chapter, a crucial factor within the assessment of cost optimality and risk acceptability of civil engineering structures. In the introductory examples, we estimated the failure probability under strong assumptions. Namely, we represented failure by a linear limit state function ($R - S$) of Normal distributed variables and independent resistance and load.

However, the physical phenomena that generate structural resistance and structural loads are not very well represented by the Normal distribution and other distribution types should be used. Many relevant failure scenarios are better represented by non-linear limit state functions and the different variables of a structural problem might be correlated. In short, practical engineering situations are more complex and cannot be represented by the strong simplifications that we applied in the previous chapter. In this chapter, methods that allow for a more realistic representation of practical engineering situations are introduced.

1.2. Introduction: The limit state principle

The safety of an engineering structure depends on the type and magnitude of the applied load and the structural strength and stiffness. Whether the performance is considered satisfactory depends on the requirements which must be satisfied. Among others, these include reliability of

the structure against collapse, limitation of damages or of deflections, or other criteria.

In general, any state that may be associated with consequences in terms of costs, loss of lives and impact to the environment are of interest. In the following, it is not differentiated between these different types of states. For simplicity, it is referred to all of them as being failure events. It is convenient to describe failure events in terms of functional relations, which if fulfilled, define that the failure event $\mathbf{F}$ will occur:

$$\mathbf{F} = \{g(\mathbf{x}) \leq 0\} \quad (2.1)$$

where $g(\mathbf{x})$ is denominated as the limit state function. The components of the vector $\mathbf{x}$ are the realizations of the so-called basic random variables $\mathbf{X}$, representing all relevant uncertainties influencing the problem at hand. The failure event $\mathbf{F}$ is defined as the set of realizations of the limit state function $g(\mathbf{x})$, which are zero or negative. Some typical limit states are given in Table 2.1.

Table 2.1: Typical limit states for structures.

Limit State Type	Description	Examples
Ultimate	Collapse of the structure or part of it	Rupture, plastic mechanism, instability, progressive collapse, fatigue, deterioration, fire.
Damage	(often included in above)	Excessive permanent cracking, permanent irreversible deformation.
Serviceability	Disruption of normal use	Excessive deflection, vibration, local damage.

Due to the associated consequences, failure events linked with the most serious limit states, such as collapse and major damage, should be relatively rare events. The study of structural reliability is concerned with the assessment of the probability of failure of engineered structures related events at any stage during its service life.

The probability of occurrence of a failure event is a measure of the chance of its occurrence. This chance may be quantified by observing the long term frequency of the event for generally similar structures or components, or may be simply a subjective estimate of its numerical value. Engineering structures are mostly exclusive in regard to the structural assembly and their exposure to loads and environment, which in general precludes the assignment of relative frequencies of events within many similar structures. However, a more generic description can be given for structural components and materials. Thus, a combination of subjective estimates and frequency observations about structural components and structural assemblies is utilized in practice to assess the probability of limit state violation of a structure.

The probability of failure p_f may be generally determined by the following integral:

$$p_f = P(g(\mathbf{X}) \leq 0) = \int_{g(\mathbf{x}) \leq 0} f_{\mathbf{X}}(\mathbf{x}) d\mathbf{x} \quad (2.2)$$

where $g(\mathbf{X})$ is the limit state function, $\mathbf{X}$ is a vector of basic random variables and $f_{\mathbf{X}}(\cdot)$ is the joint probability function of the variables $\mathbf{X}$.

1.3. A simple case - the Fundamental Reliability Problem

Several methods can be found in the literature to calculate the probability of failure by solving the integral in Eq. (2.2). An overview and a discussion of different approaches is presented in e.g. J.Schneider 2006. The most straightforward method is that of Monte Carlo simulation, while probably the more efficient are the so-called approximate methods based on the calculation of the reliability index β , e.g. the first order reliability method (FORM) and the second order reliability method (SORM). There are also methods to increase the efficiency of the Monte Carlo simulation like Importance Sampling or Adaptive Sampling. In this chapter and the next one, only the main ideas of Monte Carlo Simulation and FORM are briefly outlined.

Linear Limit State Functions and Normal Distributed Variables

First, we consider the two dimensional case (i.e. two basic random variables), which is aligned to the simplification we applied in the previous chapter. A limit state function can be $g(r, s) = r - s$, and the corresponding safety margin is $M = R - S$. For R and S uncorrelated, the reliability index is:

$$\beta_C = \frac{E[M]}{D[M]} = \frac{E[R] - E[S]}{\sqrt{Var[R] + Var[S]}} = \frac{\mu_R - \mu_S}{\sqrt{\sigma_R^2 + \sigma_S^2}} \quad (2.3)$$

where $E[M]$ is the expected value of M and $D[M]$ is its standard deviation. In the multidimensional space (i.e. more than two basic random variables), the limit state function is defined as:

$$g(\mathbf{x}) = a_0 + \sum_{i=1}^n a_i x_i = a_0 + \mathbf{a}^T \mathbf{x} \quad (2.4)$$

where a vector notation has been introduced; $\mathbf{a}^T$ is a row vector with elements a_i , and $\mathbf{x}$ is a column vector with elements x_i . The so-called safety margin M is defined as:

$$M = a_0 + \sum_{i=1}^n a_i X_i = a_0 + \mathbf{a}^T \mathbf{X} \quad (2.5)$$

M is Normal distributed with expected value $E[M] = a_0 + \sum_{i=1}^n a_i E[X_i] = a_0 + \mathbf{a}^T E[\mathbf{X}]$ and standard deviation $D[M] = \sqrt{\mathbf{a}^T \mathbf{C}_X \mathbf{a}}$. Here $E[\mathbf{X}]$ is the vector of expected values (with $1 \times n$ elements):

$$E[\mathbf{X}] = [\mu_{X_1}, \mu_{X_2}, \dots, \mu_{X_n}] \quad (2.6)$$

and $\mathbf{C}_X$ is the matrix of covariance of $\mathbf{X}$ (symmetric, with $n \times n$ elements). For uncorrelated basic random variables it is:

$$\mathbf{C}_X = \begin{bmatrix} \sigma_{X_1}^2 & 0 & \cdots & 0 \\ 0 & \sigma_{X_2}^2 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \sigma_{X_n}^2 \end{bmatrix} \quad (2.7)$$

The reliability index can be written in a matrix form as:

$$\beta_C = \frac{E[M]}{D[M]} = \frac{a_0 + \mathbf{a}^T E[\mathbf{X}]}{\sqrt{\mathbf{a}^T \mathbf{C}_X \mathbf{a}}} \quad (2.8)$$

Defining the failure event as in Eq. (2.1), the probability of failure can be defined as:

$$p_f = P(g(\mathbf{X}) \leq 0) = P(M \leq 0) = \Phi(-\beta_C) \quad (2.9)$$

If the two-dimensional case of two independent Normal distributed random variables is considered, the reliability index β has a geometrical interpretation as illustrated in Figure 2.1.

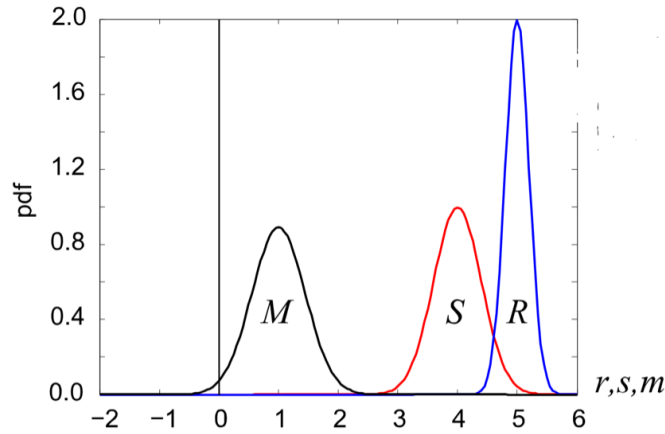


Figure 2.1: The safety margin $M = R - S$ for the case of independent and Normal distributed R and S . The reliability index β defines the mean value of M as a function of its standard deviation: $\mu_M = \beta\sigma_M$.

1.4. Definitions

Second-moment reliability theory: reliability theory based on the representation of uncertainties solely in terms of expected values (first moments) and covariance (second moments).

Basic variables: fundamental variables which define and characterize the behaviour and safety of a structure. They are denoted by X_i . The vector containing all the basic variables of a certain problem is denoted by $\mathbf{X} = [X_1, X_2, \dots, X_i, \dots, X_n]$. A realization of a basic variable is denoted by x_i while a realization of $\mathbf{X}$ is denoted by $\mathbf{x} = [x_1, x_2, \dots, x_i, \dots, x_n]$.

Failure surface/hyperplane or limit state surface/hyperplane: surface separating the failure domain (Ω_f) to the safe domain (Ω_s). The failure surface is denoted by L_X .

Limit state function: function $g(x_1, x_2, \dots, x_n) = g(\mathbf{x})$ that assumes the following values:

$$\begin{cases} g(\mathbf{x}) > 0, & \mathbf{x} \in \Omega_s \\ g(\mathbf{x}) = 0, & \mathbf{x} \in L_X \\ g(\mathbf{x}) < 0, & \mathbf{x} \in \Omega_f \end{cases} \quad (2.10)$$

Note that according to this definition a limit state function specifies the failure surface L_X and satisfies a sign convention outside L_X . The failure surface, on the other hand, does not define a unique failure function. The g -function usually results from mechanical analysis.

Safety margin: random variable obtained by replacing the parameters $x_1, x_2, \dots, x_n$ in the limit state function with the corresponding random variables $X_1, X_2, \dots, X_n$. It is usually denoted by M , being $M = g(\mathbf{X})$.

1.5. Monte Carlo Method

A very simple approach to estimate probability of failure, e.g. according to the general expression in Eq. (2.2), is the Monte Carlo Simulation (MCS) Method. Assuming that the basic random variables are represented through a set of independent random variables $\mathbf{X}$, the outcomes of the limit state function $g(\mathbf{X})$ can be processed by virtually sampling the basic random variables at random, according to their distribution functions. The outcomes might be in the failure domain ($g(\mathbf{x}) \leq 0$) or in the safe domain ($g(\mathbf{x}) > 0$). After an infinite number of “tests” the failure probability is:

$$p_f = P(g(\mathbf{x}) \leq 0) = \lim_{n \rightarrow \infty} \frac{n_f}{n} \quad (2.11)$$

where n is the total number of trials and n_f is the number of outcomes where $g(\mathbf{x}) \leq 0$.

Virtual random sampling or simulation is often consuming considerable computation time; therefore the number of simulations is limited to a certain extent. With a limited number of simulations, the failure probability can only be estimated and the uncertainty of these estimations is of interest. It can be demonstrated that the uncertainty associated with the estimate is proportional to $1/\sqrt{n_f}$.

Since in structural reliability analysis the failure probability of interest is small, i.e. in the order of 10^{-6} the number of simulated failures n_f is scarce. To estimate the failure probability of 10^{-6} with proper accuracy, e.g. a coefficient of variation of the estimate of 10%, the number of simulations should be 10^8 .

1.5.1. Generation of realizations x of random variables X .

Realizations $\mathbf{x}$ of random variables $\mathbf{X}$ are usually generated by computers. Nearly all relevant computer programs are able to generate samples of a random variable Y with uniform distribution in the interval $[0, 1]$. The corresponding cumulative distribution function is:

$$F_Y(y) = \begin{cases} 0 & y \leq 0 \\ y & 0 < y \leq 1 \\ 1 & 1 < y \end{cases} \quad (2.12)$$

Computer codes are inherently deterministic and are thus not able to generate true randomness. Nevertheless, so-called pseudo-random numbers can be generated. As an example in MS Excel the function `rand()` generates pseudo-random realization for y . As numbers generated by the computer are not truly random, the quality of the pseudo randomness should always be critically questioned. This is especially relevant for situations where the generation of large sets

of multiple random variables is necessary (as it is typical in structural engineering). E.g. the quality of the results of MCS is triggered by the quality of the pseudo-random variables.

True random numbers are nowadays provided in the internet (e.g. <https://www.random.org/>). Here, randomness is generated based on measures of true natural randomness, e.g. atmospheric noise.

1.5.2. Generation of a single random variable

For a reliability analysis with MCS we need to generate realizations of random variables with arbitrary distributions that are specified by the corresponding cumulative distribution functions.

In order to generate a random realization of an arbitrarily distributed variable, the so called transverse method can be used. Basis is taken in the cumulative distribution $F_X(x)$ of the random variable X . The cumulative distribution is defined as the probability that a realization x is smaller or equal than X , i.e.

$$F_X(x) = \Pr(x \leq X) = y \quad (2.13)$$

where y is defined in the interval $[0, 1]$.

The idea of the transverse method is to generate a realization of a uniform distributed random variable y and use the transposed cumulative distribution function of X to generate a realization x .

$$x = F_X^{-1}(y) \quad (2.14)$$

The principle is illustrated in Figure 2.2.

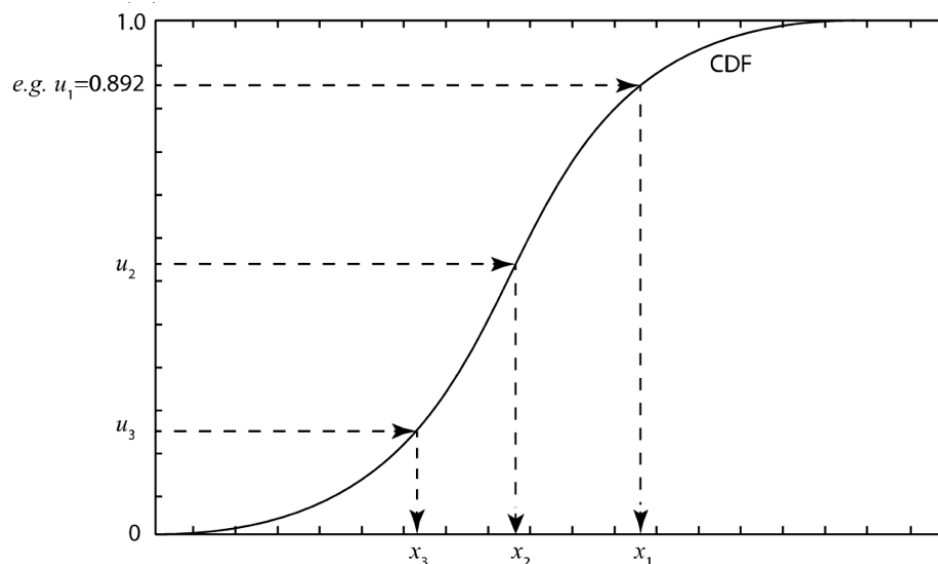


Figure 2.2: Generating random realizations x of a variable X based on realizations u of a uniform distributed random variable $U [0, 1]$.

Example

It is aimed to generate random numbers that follow an exponential distribution $F_X(x) = 1 - \exp(-\lambda x)$. The inverse of this function is computed as follows:

$$\begin{aligned}
1 - F_X(x) &= \exp(-\lambda x) \\
\ln(1 - F_X(x)) &= -\lambda x \\
-\frac{\ln(1 - F_X(x))}{\lambda} &= x
\end{aligned} \tag{2.15}$$

Finally, $F_X(x) = Y \sim U[0, 1]$ is set.

1.5.3. Random number generation in Python

There are several ways of generating random numbers from a given distribution function in Python. If one knows the inverse distribution function $F_X^{-1}(y)$, one can sample random numbers of $y \sim U[0, 1]$ first, and then proceed to evaluate those realizations using F_X^{-1} . You can sample uniform distributed numbers e.g. as follows:

```

1 import numpy as np
2 n_sim = int(1e5)
3 y = np.random.uniform(size=n_sim) # generate n_sim uniform distributed
                                     numbers

```

Alternatively, several Python packages offer built-in functions to generate random samples from the most common distributions. The NumPy package offers the followings:

<code>numpy.random.uniform</code>	$U(0, 1)$ distributed pseudo-random numbers
<code>numpy.random.normal</code>	Normal distributed pseudo-random numbers
<code>numpy.random.weibull</code>	Weibull distributed pseudo-random numbers
<code>numpy.random.lognormal</code>	Log-Normal distributed pseudo-random numbers
<code>numpy.random.exponential</code>	Exponential distributed pseudo-random numbers
<code>numpy.random.gamma</code>	Gamma distributed pseudo-random numbers
<code>numpy.random.gumbel</code>	Gumbel distributed pseudo-random numbers

Note that the inputs to these functions are usually the number of samples and the parameters of the distribution (not the moments). So you may need to compute the distribution parameters first, and then use the function of interest.

Other packages offer similar functions. For example, you can draw samples of the normal distribution using the SciPy package as `scipy.stats.norm.rvs`.

1.6. Sampling of correlated random variables

Many practical engineering problems require the consideration of correlated variables. Correlation has a large effect on structural reliability estimates.

Correlation between two random variables is generally expressed with the correlation coefficient $\rho \in [-1, 1]$. The effect of correlation between load and resistance on the realizations of a Monte Carlo simulation (and correspondingly on the estimate of the failure probability) can be seen in Figure 2.3.

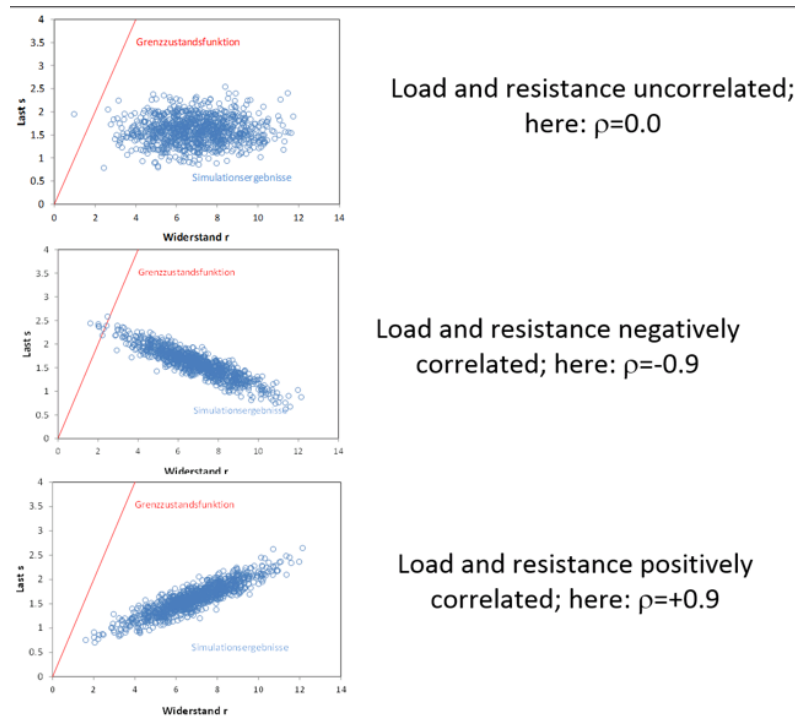


Figure 2.3: Realizations of load and resistance dependent on the correlation.

1.6.1. Generating correlated bivariate random variables

To generate pairs of realizations of Normal distributed correlated variables, first, realizations of independent variables are to be generated. Then, the coordinate system is to be rotated so that they become dependent.

The transformation of the coordinate system is performed as follows:

$$\begin{aligned} x' &= x \cdot \cos(\alpha) + y \cdot \sin(\alpha) \\ y' &= -x \cdot \sin(\alpha) + y \cdot \cos(\alpha) \end{aligned} \quad (2.16)$$

Setting the correlation coefficient to $\rho = \cos(\alpha)$ yields:

$$x' = x \cdot \rho + y \cdot \sqrt{1 - \rho^2} \quad (2.17)$$

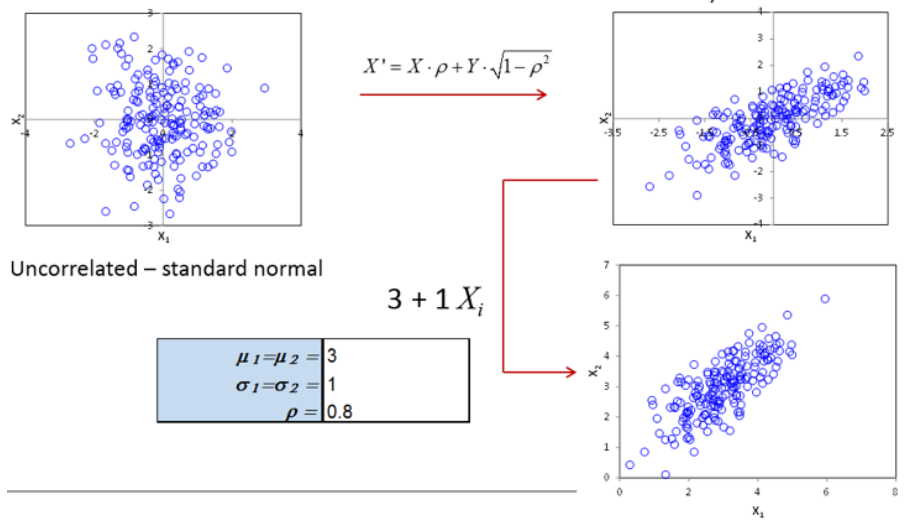
Algorithm for generating realizations of pairs of Normal distributed random variables:

1. Generate $X \sim N(0, 1)$, $Y \sim N(0, 1)$
2. Calculate $X' = X \cdot \rho + Y \cdot \sqrt{1 - \rho^2}$
3. Get the correlated pair of standard Normal distributed random variables $(X, X') = (X_1, X_2)$
4. Transform (X_1, X_2) to general bivariate Normal distribution $\mu_i + X_i \sigma_i$.

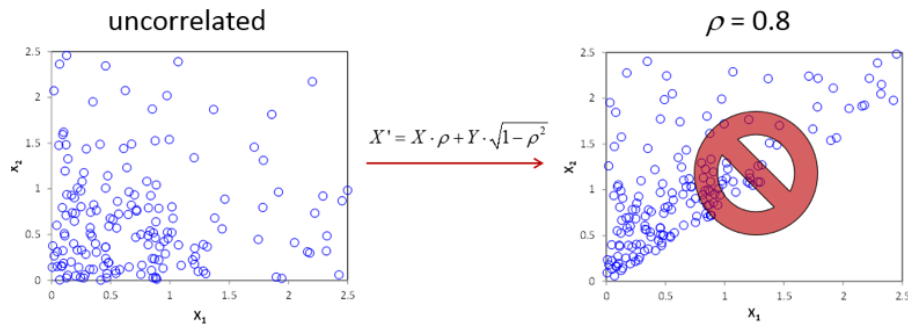
Caution: The direct application of this method is only valid for the standard Normal distribution.

For generation of general correlated random variables, the so-called “Normal to Anything” (NORTA) methods come into play.

Example: **Normal distribution**



Example: **Exponential distribution**



Transformation not valid for other distributions!

1.6.2. Normal to Anything

The basic idea of NORTA is to generate pairs of realizations of correlated bivariate standard Normal variables (as above). First, they are transformed into pairs of realizations of correlated Uniform distributed variables. These are then transformed using the distribution of choice.

Algorithm:

1. Generate a pair of correlated standard Normal random variables $(X, X') = (X_1, X_2)$ (see above)
2. Calculate $U_1 = \Phi(X_1)$ and $U_2 = \Phi(X_2)$. E.g.:
 Excel[®]: `Norm.dist(X,0,1)`
 Matlab[®]: `normcdf(X,0,1)`
 This gives pairs of realizations of correlated Uniform distributed variables (U_1, U_2)
3. Use U_i to generate correlated random variables by using the inverse transformation method.

NORTA: Example

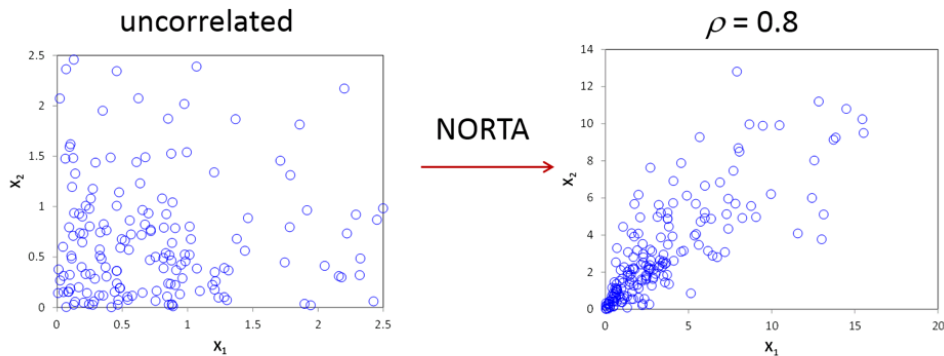
$X = X$
 $X' = X \cdot \rho + Y \cdot \sqrt{1 - \rho^2}$

$U_I = \Phi(X)$

$Z_i = -\frac{\ln(1 - U_i)}{\lambda}$

Nr. Sim	$N(\mu=0, \sigma=0)$	$N(\mu=0, \sigma=0)$	$MVN(0, 0, 1, 1, \rho = 0.8)$		$U(0,1)$	$U(0,1)$	$MVExp(\lambda_1, \lambda_2)$	
1	-1.721617252	0.040385677	-1.721617252	-1.353062396	0.04256944	0.08801785	0.13050624	0.27640458
2	1.352786938	-0.285493347	1.352786938	0.910933543	0.91193815	0.8188348	7.2891475	5.12503795
3	0.080014232	0.172940922	0.080014232	0.167775939	0.53188703	0.56662022	2.27713688	2.50842256
4	-0.166766125	0.319333788	-0.166766125	0.058187373	0.43377704	0.52320031	1.70630205	2.22197644
5	0.104898079	-1.739373857	0.104898079	-0.959705851	0.54177166	0.16860164	2.34116297	0.55393867
6	-0.536203406	0.200664025	-0.536203406	-0.30856431	0.29590899	0.37882649	1.05254297	1.42843448
7	1.097379345	0.59123515	1.097379345	1.232644566	0.8637622	0.8911458	5.98006018	6.65323769
8	-0.601325973	-0.045777285	-0.601325973	-0.508527149	0.27381145	0.30554185	0.95983675	1.09387014
9	-0.041260701	1.142412507	-0.041260701	0.652438943	0.48354403	0.74294098	1.98229573	4.07534866
10	-0.345012549	1.062599643	-0.345012549	0.361549747	0.36504247	0.64115574	1.36259148	3.07460039
11	-0.038455833	-1.428023401	-0.038455833	-0.887578707	0.48466212	0.18738371	1.98879756	0.62248873

Example: Exponential distribution



1.7. Generalization for multivariate distributions

The covariance structure of n random variables is described by the covariance matrix Σ . The covariance matrix of the random variables $X_1, X_2, \dots, X_n$ is written as:

$$\Sigma = \begin{bmatrix} \sigma_{X_1}^2 & \rho_{X_1 X_2} \sigma_{X_1} \sigma_{X_2} & \cdot & \rho_{X_1 X_n} \sigma_{X_1} \sigma_{X_n} \\ & \sigma_{X_2}^2 & \cdot & \rho_{X_2 X_n} \sigma_{X_2} \sigma_{X_n} \\ & & \cdot & \\ sym. & & & \sigma_{X_n}^2 \end{bmatrix} \quad (2.18)$$

The correlation matrix $\mathbf{C}$ is defined as:

$$\mathbf{C}[X_i, X_j] = \frac{\Sigma[X_i, X_j]}{\sqrt{\Sigma[X_i, X_i] \Sigma[X_j, X_j]}} \quad (2.19)$$

Note that for the multivariate standard Normal distribution $\Sigma = \mathbf{C}$ and can be decomposed

with $\mathbf{C} = \mathbf{R}^T \mathbf{R}$ in order to define correlated random variables $\mathbf{Y}$ out of uncorrelated $\mathbf{X}$:

$$\mathbf{Y} = \mathbf{X}\mathbf{R} \quad (2.20)$$

The matrix $\mathbf{R}$ is found by using e.g. a Cholesky Decomposition of the correlation matrix:

Example:

$$C = \begin{bmatrix} 1 & 0.8 & 0.2 \\ 0.8 & 1 & 0.3 \\ 0.2 & 0.3 & 1 \end{bmatrix} \xrightarrow{\text{Cholesky decomposition}} R = \begin{bmatrix} 1 & 0.800 & 0.200 \\ 0 & 0.600 & 0.233 \\ 0 & 0 & 0.952 \end{bmatrix}$$

Correlated standard normal distributed random variables Uncorrelated standard normal distributed random variables

$$\begin{bmatrix} 1.602 & 1.536 & 0.240 \end{bmatrix} = \begin{bmatrix} 1.602 & 0.424 & -0.188 \end{bmatrix} \begin{bmatrix} 1 & 0.800 & 0.200 \\ 0 & 0.600 & 0.233 \\ 0 & 0 & 0.952 \end{bmatrix}$$

When using Matlab[®], the function `mvnrnd` can be used for the generation of realizations of correlated multivariate Normal distributed random variables. NORTA can then be used to generate any distribution types based on that.

1.8. Variance reducing methodologies

Monte Carlo simulation becomes more and more common for solving reliability problems due to the computational power that is nowadays available even by using desktop computers. However, in many practical problems Monte Carlo is still too slow, e.g. when:

- the limit state function has no closed form and a single evaluation of the limit state requires considerable computation time,
- the failure probability has to be evaluated many times.

For these situations, the so-called variance reducing methodologies such as importance sampling, stratified sampling or Latin Hypercube Sampling (LHS) should be considered. In the following we will look closer to Latin Hypercube Sampling.

1.8.1. Latin Hypercube Sampling

The idea of latin Hypercube Sampling is to order random realizations into so-called "strats", in order to accelerate convergence.

In Figure 2.4, the direct comparison between crude sampling and ordered sampling is illustrated. For the ordered sampling the vertical axis is divided into (in this case 6) bins (or strats). Random realizations are generated for each of these individual bins. A faster convergence of the ordered sample towards the real character of the cumulative distribution function is observed.

If many variables are to be generated simultaneously, the same strategy can be used for ordering sampling. However, the order between the different variables is to be broken up carefully. This is done by the so-called hypercube.

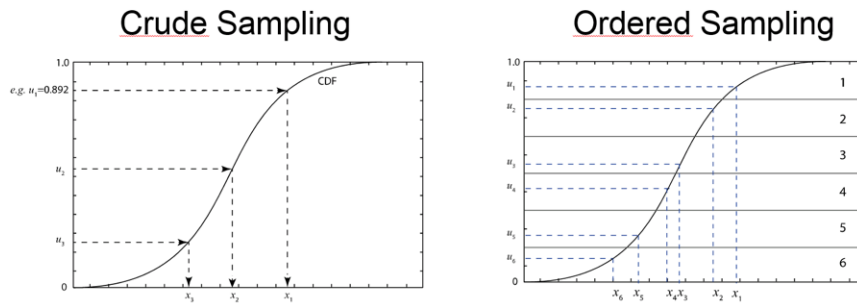


Figure 2.4: Difference between crude and ordered sampling.

The strategy is as follows:

1. Each variable is divided in n (same number) intervals.
2. A Latin Hypercube is constructed in order to combine the random variables.
3. The variables are combined and the outcome is computed.

An example for two random variables is illustrated in Figure 2.5.

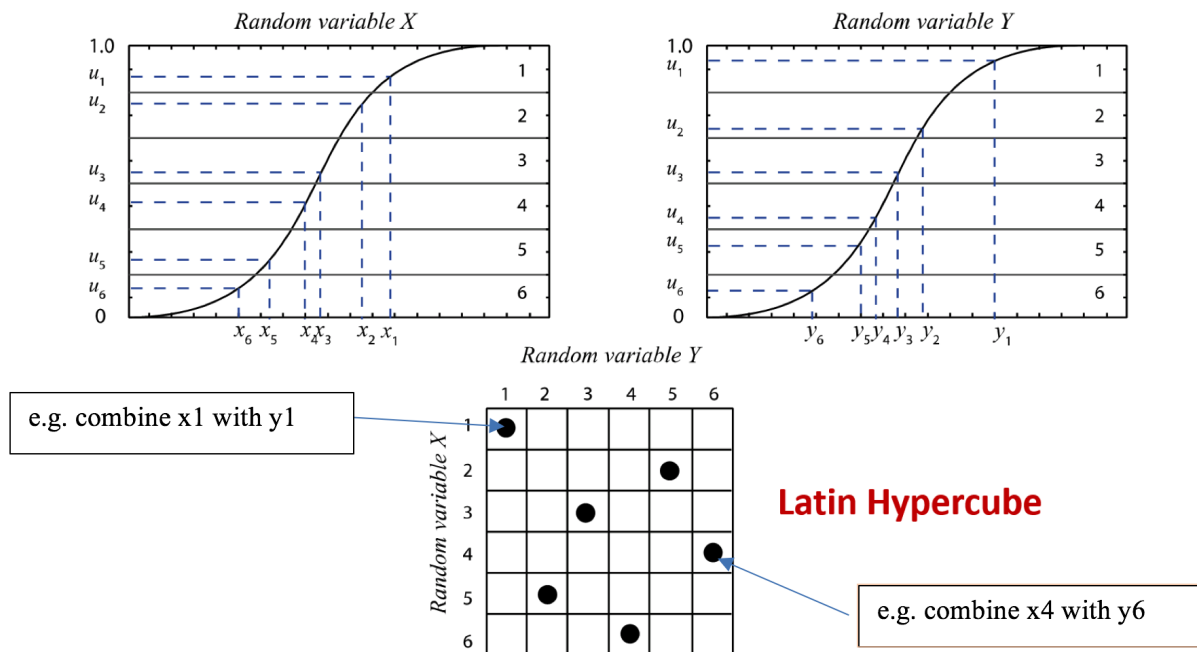


Figure 2.5: Example of the Latin Hypercube Sampling for two random variables.

1.8.2. Latin Hypercube Sampling with Python

Some Python packages include Latin Hypercube Sampling functions, e.g. [pyDOE](#). However, in order to increase understanding of the underlying sampling strategy it is worthwhile to implement LHS into a self-created function. There are several ways to do it; one possibility is shown

next, where the generation of realizations of standard Normal distributed variables is considered.

```

1  import numpy as np
2  import scipy as sp
3  import scipy.stats
4
5  def LHCNorms(NStrat, NVar):
6      """ Function for Latin Hypercube Sampling for realizations of
7          Standard
8          Normal distributed variables (JK, June 2016)
9
10         INPUT:
11             NStrat:: Size of the Strat as INT
12             Nvar::   Number of Variables as INT
13         OUTPUT:
14             Data::   Realizations as NSample x NVar array of RE
15
16         """
17         data = np.random.uniform(size=(NStrat, NVar)) # Random Numbers of U(
18                                                         0,1)
19         Data = np.zeros_like(data) # Definition of size
20                                     of output array
21         for i in range(NVar):
22             rank = np.random.permutation(NStrat) + 1 # Create a random
23                                                         sequence of ranks (Shuffle)
24             prob = (rank - data[:,i])/NStrat # Random realisation in
25                                               the interval below index/
26                                               NSample
27             Data[:,i] = sp.stats.norm.ppf(prob) # Output as Standard
28                                                 Normal (here any
29                                                 distribution could enter)
30
31         return Data

```

In line 21, the standard Normal distribution is specified. If other distribution types are of interest, the edit goes there. Alternatively, the pyDOE function lhs can be used to create a LHS experimental design and afterwards apply use that design for the distribution of interest. You can install the pyDOE package with the command `conda install pyDOE`.

1.8.3. Effect of LHS on convergence

The objective of LHS is to accelerate convergence to the “exact” result. Let us consider the simple example where realizations from a standard Normal distribution are sampled and the sample mean is to be calculated. The real mean value is known to be zero. However, the sample means will only converge to that value by increasing the sample size considerably.

We can create a little experiment where we evaluate the sample mean of m different samples of size n . We can compute m different sample means for each sample size. In Figure 2.6, the results of the experiment are displayed for $m = 50$ and $n = 50$. The green line represents the theoretical solution for the standard deviation of the sample means $\sigma_\mu = 1/\sqrt{n}$.

If sampling is now performed according to a Latin Hypercube design, e.g. according to the function introduced above, the convergence behavior can be enhanced considerably.

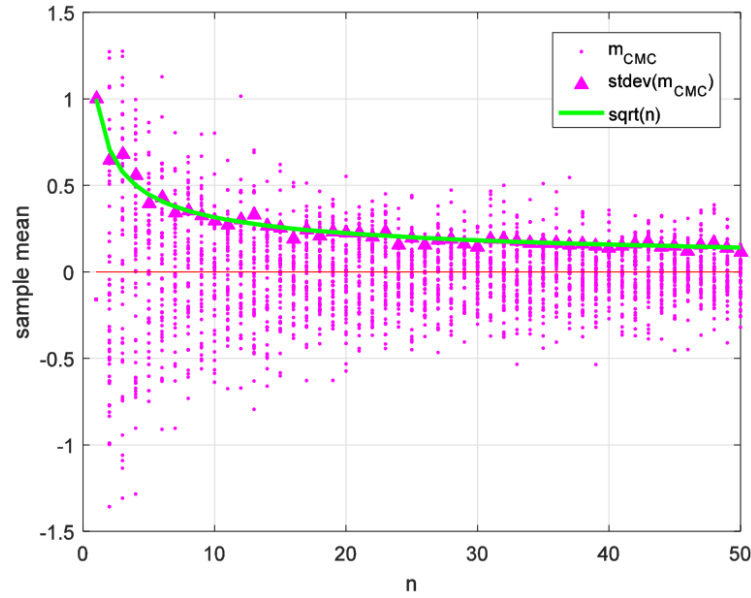


Figure 2.6: The mean value of different samples of random realizations of the standard Normal distribution of size n .

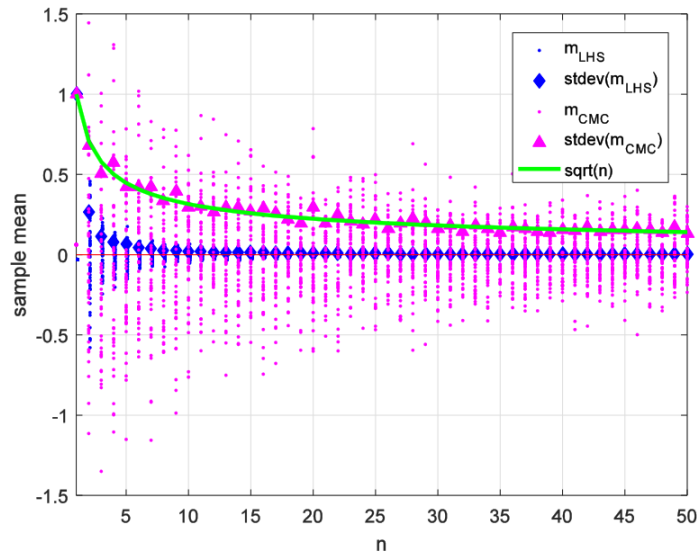


Figure 2.7: The blue dots represent the mean values of samples of size n as generated by LHS. The blue diamonds represent the standard deviations.

If LHS is applied for the solution of structural reliability problems, however, the increase of efficiency is less spectacular as will be shown in the following example.

Example

Crude Monte Carlo and Latin Hypercube Sampling are used for the estimation of the failure probability, that is defined as $P_f = \Pr(R - S < 0)$; whereas R and S are represented as independent Normal distributed random variables $R \sim N(\mu = 6.8, \sigma = 1.7)$ and $S \sim N(\mu = 1.6, \sigma = 0.3)$. The failure probability can be easily computed exactly as $P_f = \Phi\left(-\frac{(6.8 - 1.6)}{\sqrt{1.7^2 + 0.3^2}}\right) = 0.0013$.

In this example, however, we are interested in the convergence towards the true value. In

Figure 2.8, the convergence is illustrated by means of the standard deviation of the estimates of the failure probability as a function of the number of simulations. It can be seen that after $n_{CMC(10\%)} \approx 100/P_f = 7.7 \cdot 10^4$ simulations, the standard deviation is dropping below 10% of the failure probability. The same accuracy is reached after $n_{LHC(10\%)} \approx 4 \cdot 10^4$, which is about 45% less.

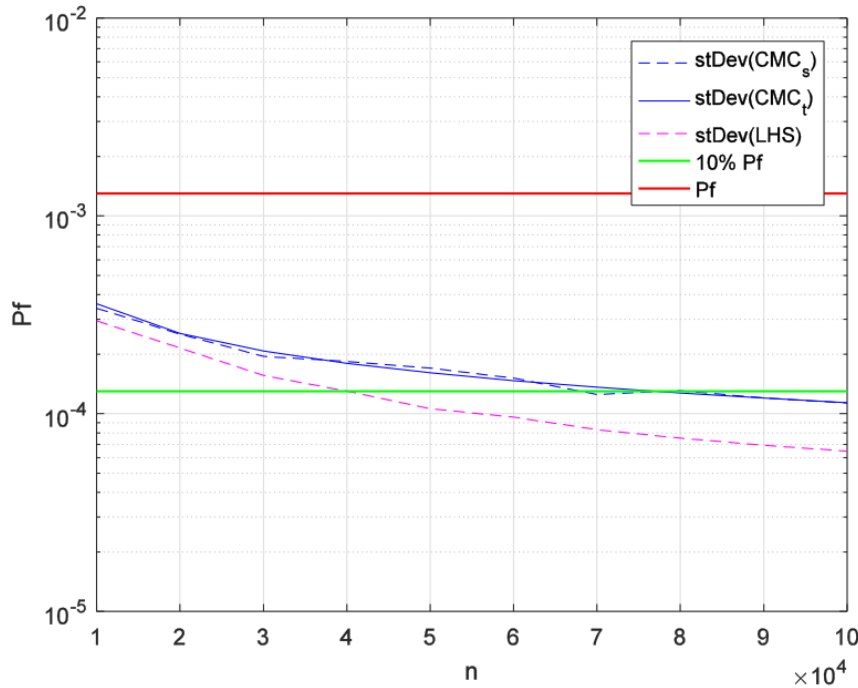


Figure 2.8: Development of the accuracy of estimates of the probability of failure with increasing number of simulations.

This sounds nice, but it should be considered that the generation of the LHC design also requires additional computational power, so, the gain in number of simulations might be well compensated by that.

2. Application examples

E2.1. Hyper-plane failure function with normal distributed variables

Consider the cantilever beam loaded by two concentrated forces in Figure 2.9 (from Madsen et al. 1985). The variables are uncorrelated and their properties are summarized in Table 2.2. Formulate the safety margin based on the cross-section capacity ultimate limit state, i.e. the beam fails if the moment capacity is exceeded. Write the second moment representation of F_m , P_1 and P_2 and calculate the reliability index.

Table 2.2: Basic variables properties.

X	μ_X	σ_X
F_m	167 MPa	20 MPa
W_y	1.5e6 mm ³	0
l	4 m	0
P_1	10 kN	3 kN
P_2	10 kN	3 kN

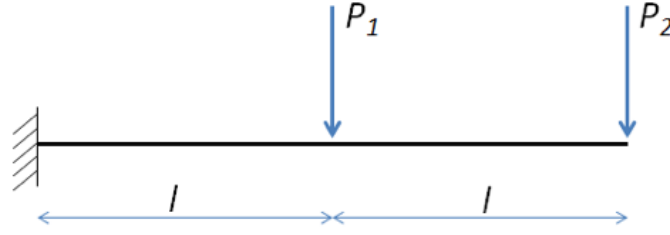


Figure 2.9: Cantilever beam.

Solution The linear safety margin is $M = W_y F_m - l P_1 - 2l P_2$. The second moment representation of the set of basic variables $\mathbf{X} = (F_m, P_1, P_2)$ is:

$$E[\mathbf{X}] = \begin{bmatrix} 167 \text{ MPa} \\ 10000 \text{ N} \\ 10000 \text{ N} \end{bmatrix}; \quad \mathbf{C}_X = \begin{bmatrix} 20^2 (\text{MPa})^2 & 0 & 0 \\ 0 & 3000^2 \text{ N}^2 & 0 \\ 0 & 0 & 3000 \text{ N}^2 \end{bmatrix}$$

$a_0 = 0$ and the vector of the coefficients is:

$$\mathbf{a}^T = [W_y \quad -l \quad -2l] = [1.5 \cdot 10^6 \text{ mm}^3 \quad -4000 \text{ mm} \quad -8000 \text{ mm}]$$

The reliability index is calculated by means of Eq. (2.8) yielding $\beta_C = 3.24$.

E2.2. Application of Crude Monte Carlo (CMC) simulation to the design of a timber column

This example is adapted from Köhler et al. (2008). The structural behaviour of columns subjected to axial compression is characterized by the non-linear (second order) increase of the deformation due to the axial load (P-delta effects). P-delta effects are due to imperfections of the structural members and strongly depended on the column slenderness and stiffness. Since columns subjected to axial compression become more slender and less stiff to deformation, the influence of P-delta effects increases and the structural behaviour is governed by column stability (buckling). Therefore, the stiffness properties (modulus of elasticity) of the structural members play an important role for a second order analysis, in particular for slender columns.

In the case of columns subjected to axial compression, the non-linear (second order) increase of the deformation due to the axial compression can be calculated considering the amplification

factor defined as:

$$\alpha = \frac{1}{1 - \frac{N}{N_{crit}}} \quad (2.21)$$

with N the load effect and $N_{crit} = \pi^2 EI / l^2$ the classic Euler's formula of buckling, i.e. with E the modulus of elasticity, I the moment of inertia and l the buckling length. The resulting moment of second order can be calculated as:

$$M_{II} = N\alpha\xi \quad (2.22)$$

The initial eccentricity ξ takes into account all geometric imperfections, as well as material imperfections.

In the case of a second order structural analysis, the ultimate limit state of columns subjected to axial compression can be verified with common interaction equations for combined bending and axial compression. Investigations by Buchanan showed that columns loaded by axial compression and bending tend to develop plastic deformations in the compression zone. Thus, a linear model for description of the interaction between bending and axial compression leads to conservative results. In EN 1995-1-1 and DIN 1052 this effect is taken into account by the following non-linear interaction curve for combined bending and axial compression:

$$\left(\frac{N}{N_R}\right)^2 + \frac{M}{M_R} \leq 1.0 \quad (2.23)$$

with N and M being the load effects. $N_R = k_{mod} F_c \cdot A$ and $M_R = k_{mod} \cdot F_m \cdot W$ are the resistances in compression and bending with, A the cross-sectional area and W the section modulus.

Eq. (2.23) can be used for the verification of the ultimate limit state of columns subjected to axial compression assuming the resulting moment M_{II} of second order:

$$\left(\frac{N}{N_R}\right)^2 + \frac{M}{M_R} = \left(\frac{N}{k_{mod} F_c A}\right)^2 + \frac{N\xi \left(\frac{1}{1 - \frac{N}{\frac{\pi^2 EI}{l^2}}} \right)}{k_{mod} \cdot F_m \cdot W} \leq 1.0 \quad (2.24)$$

Additionally, the simple check of axial compression without any second order moment and the check for buckling must be verified:

$$N \leq N_R \quad (2.25)$$

$$N \leq N_{crit} \quad (2.26)$$

A probabilistic model for the column load bearing capacity is derived from Eqs. (2.24) to (2.26). The failure domain is defined by:

$$\Omega_f = \{\mathbf{x} | g_1(\mathbf{x}) \leq 0 \text{ OR } g_2(\mathbf{x}) \leq 0 \text{ OR } g_3(\mathbf{x}) \leq 0\} \quad (2.27)$$

with:

$$\begin{aligned} g_1(\mathbf{x}) &= 1 - X_1 \left(\frac{g+q}{k_{mod} f_c h^2} \right)^2 - X_2 \frac{(g+q) \xi \left(\frac{1}{1 - \frac{g+q}{\pi^2 e h^4}} \right)}{\frac{k_{mod} f_m h^3}{6}} \\ g_2(\mathbf{x}) &= 1 - X_1 \frac{g+q}{k_{mod} f_c h^2} \\ g_3(\mathbf{x}) &= 1 - X_2 \frac{g+q}{\frac{\pi^2 e h^4}{12 l^2}} \end{aligned} \quad (2.28)$$

The following simplifications are introduced:

- $F = F_m = F_c$, i.e. same material strength for compression and bending load mode;
- the model uncertainties are not considered, i.e. $X_1 = X_2 = 1.0$ (they will be introduced later in the course);
- no duration of load effects, i.e. $k_{mod} = 1.0$.

The basic variables properties are given in Table 2.3. Calculate the probability of failure P_f and the reliability index β .

Table 2.3: Basic properties of the variables.

Basic Variables	PDF	COV	Mean	Standard deviation
F_m : bending strength [N/mm ²]	Normal	0.25	20	5
F_c : compression strength [N/mm ²]	Normal	0.25	20	5
E : MOE [N/mm ²]	Normal	0.13	11000	1430
G : Dead Load [kN]	Normal	0.1	10.15	1.02
Q : Live Load [kN]	Normal	0.4	28.23	11.29
ξ : Initial Bow [mm]	Normal	-	0	3
l [mm]	Deterministic	-	10000	-
h : cross-section dimension [mm]	Deterministic	-	250	-

Solution: For-cycle script The most intuitive way of programming is with the *for*-cycle. In each cycle, one value for basic variable is sampled; the cycle is repeated n_{sim} times. With 10^6 simulations, it is obtained $\beta = 3.83$; $P_f = 6.40 \cdot 10^{-5}$. The simulation time is 221 s.

Table 2.4: Development of the accuracy of estimates of the probability of failure with increasing number of simulations.

Simulation nr.	i	1	2	3	...	i	...	...	...	$n_{sim}=10^6$
sampled values	f	20.83	12.44	22.77	...	19.25	18.72	...	0.35	17.10
	e	9863.50	11738.04	10694.58	...	4442.14	11889.26	...	10516.73	9140.42
	g	11412.06	10789.79	9494.28	...	11142.47	11539.91	...	8840.47	11062.51
	q	60510.41	19574.54	6978.21	...	53780.87	25859.97	...	34953.30	64441.97
	xi	3.98	1.68	0.59	...	1.39	2.60	...	3.32	2.55
intermediate variables	n	71922.47	30364.33	16472.49	...	64923.33	37399.88	...	43793.77	75504.49
	nr	833061.68	497437.71	910862.72	...	769920.37	748791.20	...	14073.27	683993.57
	mr	27768722.67	16581257.05	30362090.54	...	25664012.47	24959706.77	...	469108.88	22799785.82
	ncrit	129798.48	154466.42	140735.02	...	58456.28	156456.41	...	138394.67	120283.14
	alpha	2.24	1.24	1.13	...	-9.04	1.31	...	1.46	2.69
evaluation of limit state function	g1	0.97	0.99	1.00	...	1.02	0.99	...	-9.14	0.97
	g2	0.91	0.94	0.98	...	0.92	0.95	...	-2.11	0.89
	g3	0.45	0.80	0.88	...	-0.11	0.76	...	0.68	0.37
Failure event counting	fails	0	0	0	...	1	1	...	64	64

cycle nr. i

Solution: Script without For-cycle The implementation of loops is not optimized in Matlab[®]. A better (faster) way of programming is achieved by using vectors. For each basic random variable, a $1 \times n_{sim}$ vector of sampled values is created with the command `x=norminv(rand(1,n_sim),mX,stdX)`. The limit state functions are evaluated using the element by element operation (`.*`, `./` and `.^`).

Set the elements of fails equal to 1 if failure occurs. The total number of failures is then equal to `sum(fails)`. The estimation of the probability of failure is $P_f = \text{sum(fails)} / n_{sim}$.

Table 2.5: Sampled vectors, limit state function evaluations and counting of failure events.

Simulation nr.	i	1	2	3	...	i	...	...	...	$n_{sim}=10^6$
sampled values	f	20.83	12.44	22.77	...	19.25	18.72	...	0.35	17.10
	e	9863.50	11738.04	10694.58	...	4442.14	11889.26	...	10516.73	9140.42
	g	11412.06	10789.79	9494.28	...	11142.47	11539.91	...	8840.47	11062.51
	q	60510.41	19574.54	6978.21	...	53780.87	25859.97	...	34953.30	64441.97
	xi	3.98	1.68	0.59	...	1.39	2.60	...	3.32	2.55
intermediate variables	n	71922.47	30364.33	16472.49	...	64923.33	37399.88	...	43793.77	75504.49
	nr	833061.68	497437.71	910862.72	...	769920.37	748791.20	...	14073.27	683993.57
	mr	27768722.67	16581257.05	30362090.54	...	25664012.47	24959706.77	...	469108.88	22799785.82
	ncrit	129798.48	154466.42	140735.02	...	58456.28	156456.41	...	138394.67	120283.14
	alpha	2.24	1.24	1.13	...	-9.04	1.31	...	1.46	2.69
evaluation of limit state function	g1	0.97	0.99	1.00	...	1.02	0.99	...	-9.14	0.97
	g2	0.91	0.94	0.98	...	0.92	0.95	...	-2.11	0.89
	g3	0.45	0.80	0.88	...	-0.11	0.76	...	0.68	0.37
Failure event counting	fails	0	0	0	...	1	0	...	1	0
										TOTAL
										784

1 x n_{sim} vector with the realizations of E

The results obtained from 10^7 simulations (i.e. 10 times more than before) are: $\beta = 3.80$; $P_f = 7.20 \cdot 10^{-5}$ and the time for simulation is 5.16 s. The simulation is much faster compared to the script with the for-cycle. The plots below show the convergence of the simulations.

It is important to keep in mind that the result for a finite number of simulations varies due to the randomness of the sampling. The scatter of the results decreases when the number of

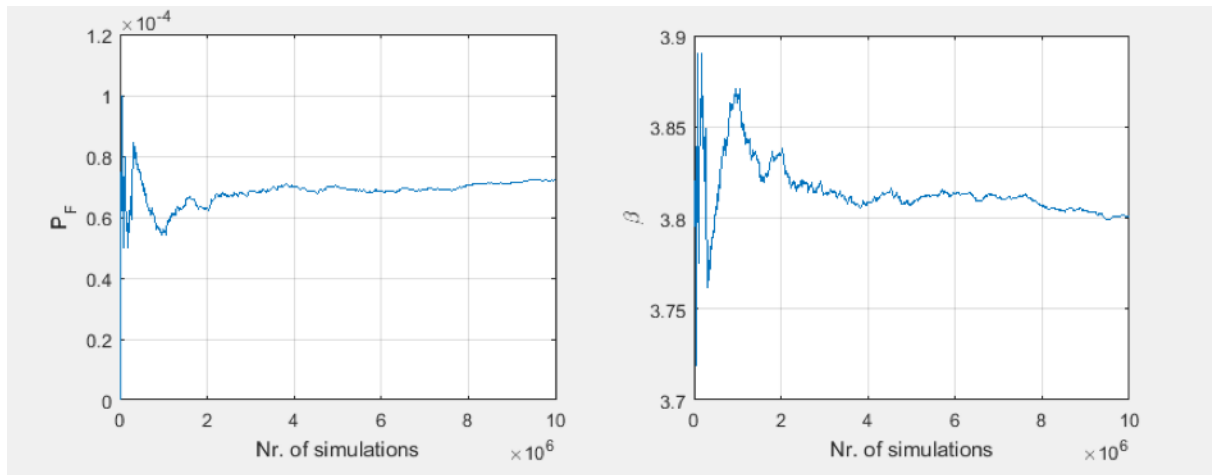
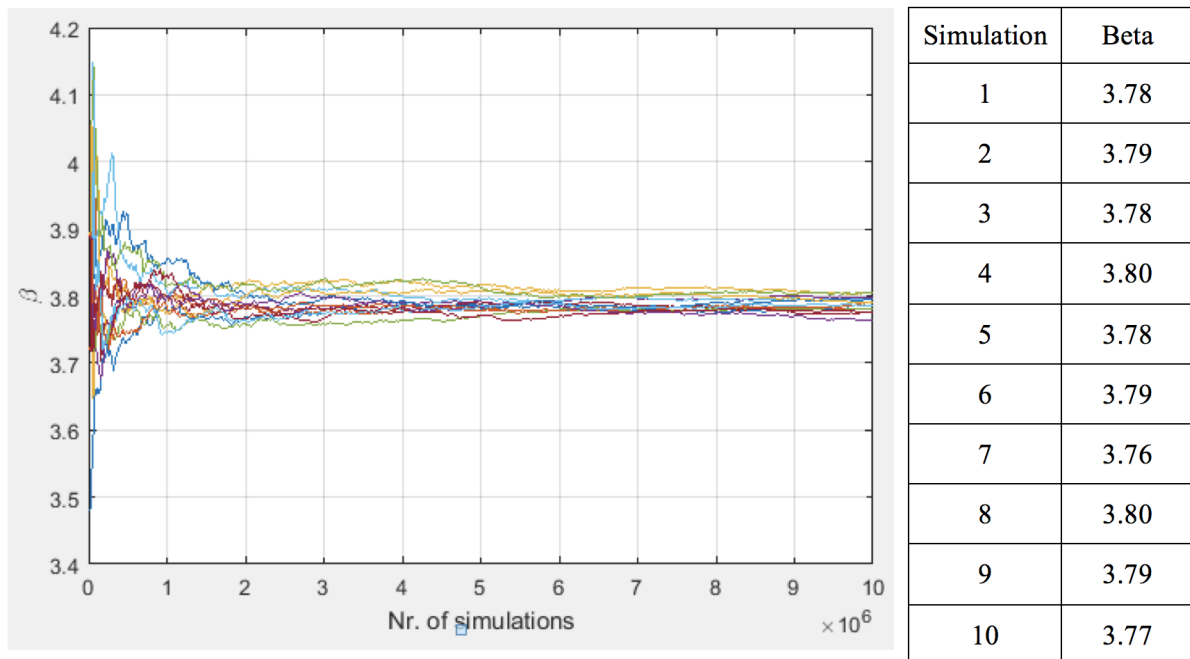


Figure 2.10: Convergence of the MCS.

simulations tends to infinite. The results of ten different MCS are reported in the table and figure below.

Figure 2.11: Obtained values of the reliability index for 10^7 simulations.



3. Practical exercises

P2.1. Calculation of the reliability index with correlated random variables

Consider the cantilever beam loaded by two concentrated forces presented in E2.1. If we assume that P_1 and P_2 are originated from the same physical phenomenon (e.g. both are induced by the snow) then P_1 and P_2 will be positively correlated, with correlation coefficient $\rho_{P_1 P_2}$.

Formula the covariance matrix as a function of $\rho_{P_1 P_2}$ and plot the reliability index for correlation coefficients between 0 and 1. Does the probability of failure increases or decreases for increasing correlation? why?

P2.2. Crude Monte Carlo (CMC) simulation for the column example

Consider the column capacity problem in E2.2. For the timber material the bending strength F_m and the compression strength F_c are different but positively correlated; moreover, both properties are positively correlated with the elastic modulus E . The correlation of the three random variables is given in the Table 2.6. Compute the reliability index taking into account the correlation among the random variables and compare the results with the ones in E2.2. Use a cross-section height $h = 200$ mm in this case.

Table 2.6: Correlation ρ of the random variables.

	F_c	E
F_m	0.8	0.8
F_c	1	0.6

P2.3. Latin Hypercube sampling for the column example

Solve again the exercise P2.2 using the Latin Hypercube sampling (LHS) technique and compare the computation time.

P2.4. Non-normal distributed variables

Some random variables are better described by other distributions than the Normal distribution. Distributions that are often used in structural reliability are the Weibull, Exponential, Pareto etc. The appendix of Jordaan (2005) provides a good summary of the most used ones. The material properties (F_m , F_c and E) are better represented by lognormal distributions while the live load yearly maxima (Q) are represented by a Gumbel distribution. Adapt the example E2.2 accordingly and compute the reliability index. Assume that the same moments of the random variables as were given in Table 2.3. Use a cross-section height $h = 200$ mm in this case.

Gumbel distribution

The Gumbel distribution is representing maxima of random variables with parent distributions having exponential tail. In civil engineering, it is used for representing the yearly maxima of

environmental actions. The cumulative density function is:

$$\begin{aligned}
 F_X(x|a, b) &= \exp[-\exp[-a(x - b)]] \\
 \text{mean: } \mu_X &= b + \frac{a}{\gamma}, \quad (\gamma \approx 0.577216 \text{ is the Euler's constant}) \\
 \text{standard deviation: } \sigma_X &= \frac{\pi}{a\sqrt{6}} \\
 \text{parameters: } a &= \frac{\pi}{\sigma_X\sqrt{6}}; \quad b = \mu_X - \frac{\gamma}{a}
 \end{aligned} \tag{2.29}$$

Log-Normal distribution

A variable X is said to be Log-Normal distributed if $Y = \ln(X)$ is Normal distributed. In civil engineering, it is used to represent material properties and many variables that are strictly positive; in fact, the lognormal distribution is defined only on the positive real axis. The cumulative density function is derived from the normal distribution:

$$\begin{aligned}
 F_X(x|\mu_{\ln(x)}, \sigma_{\ln(x)}) &= \Phi\left(\frac{\ln(x) - \mu_{\ln(x)}}{\sigma_{\ln(x)}}\right) \\
 \text{mean: } \mu_X &= \exp\left(\mu_{\ln(x)} + \frac{\sigma_{\ln(x)}^2}{2}\right) \\
 \text{standard deviation: } \sigma_X &= \mu_X \sqrt{\exp(\sigma_{\ln(x)}^2) - 1} \\
 \text{parameters: } \mu_{\ln(x)} &= \ln(\mu_X) - \frac{\sigma_{\ln(x)}^2}{2}; \quad \sigma_{\ln(x)} = \sqrt{\ln\left(1 + \left(\frac{\sigma_X}{\mu_X}\right)^2\right)}
 \end{aligned} \tag{2.30}$$

Question::

Limit state and structural reliability.

- Give some examples for relevant events that can be expressed by a limit state.
- How is “failure” computed? (Hint: The reliability integral.)
- How would you solve a reliability problem that can be formulated as a linear combination of 5 normal distributed correlated random variables?
- Consider the simple case $g(r, s) = r - s$. Represent the problem graphically:
 - In a 2 dimensional diagram.
 - In a 3 dimensional diagram. (indicate the limit state $g(r, s) = 0$)
 - In the standard normal space.

Question::

Monte Carlo Simulation Method.

- Explain the main principle of the Monte Carlo Simulation Method.
- How many simulations do you have to perform in order to get a reliable result.
- How can you accelerate convergence to the true failure probability? Explain the method and the challenges.
- How do you consider correlated random variables?
- What are sources of correlation in structural reliability? Give some examples.
- What is the effect of correlation on the reliability? Give an illustrative example.

Chapter 3

First Order Reliability Method

The first order reliability method (FORM) is introduced as a very efficient approximation method to compute the probability of failure.

1. Theoretical aspects

1.1. Linear Limit State Functions and Normal Distributed Variables

Let's revisit the simple case, where the limit state function $g(\mathbf{x})$ is a linear function of Normal distributed random variables $\mathbf{X}$. Then, the limit state function can be written as:

$$g(\mathbf{x}) = a_0 + \sum_{i=1}^n a_i x_i \quad (3.1)$$

The so-called safety margin M is defined as:

$$M = a_0 + \sum_{i=1}^n a_i X_i \quad (3.2)$$

and is Normal distributed with mean value μ_M and standard deviation σ_M as:

$$\mu_M = a_0 + \sum_{i=1}^n a_i \mu_{X_i} \quad (3.3)$$

$$\sigma_M^2 = a_0^2 + \sum_{i=1}^n a_i^2 \sigma_{X_i}^2 + \sum_{i=1}^n \sum_{j=1, j \neq i}^n \rho_{ij} a_i a_j \sigma_{X_i} \sigma_{X_j} \quad (3.4)$$

where ρ_{ij} are the correlation coefficients between the variables X_i and X_j . The probability of failure can be defined as:

$$p_f = P(g(\mathbf{X}) \leq 0) = P(M \leq 0) \quad (3.5)$$

which in this case, reduces to the evaluation of the standard Normal distribution function as:

$$p_f = \Phi(-\beta) \quad (3.6)$$

where β is the so-called reliability index:

$$\beta = \frac{\mu_M}{\sigma_M} \quad (3.7)$$

For the two-dimensional case of two independent Normal distributed random variables, the reliability index β has the geometrical interpretation that is illustrated in Figure 3.1. The limit state function $g(x_1, x_2)$ is transformed into the limit state function $g(u_1, u_2)$ by normalization of the independent Normal distributed random variables X_1, X_2 into standardized Normal distributed random variables U_1, U_2 :

$$U_i = \frac{X_i - \mu_{X_i}}{\sigma_{X_i}} \quad (3.8)$$

Note that the operation results in the random variables U_i having zero means and unit standard deviations.

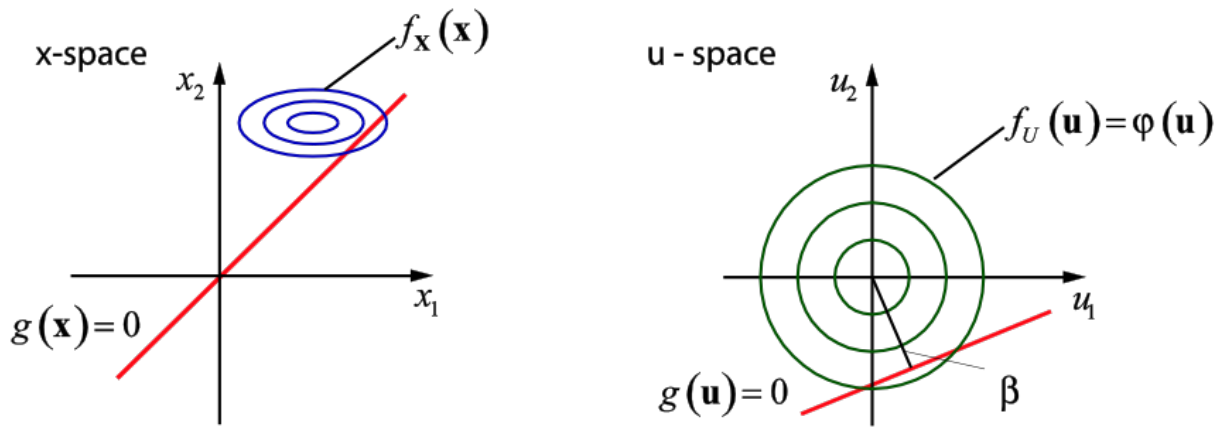


Figure 3.1: Illustration of the two-dimensional case of a linear limit state function and standardized Normal distributed variables U_1, U_2 .

The reliability index β has the geometrical interpretation of the distance to the origin of the line $g(u) = 0$, which is the boundary between the safe domain and the failure domain. For k -dimensions, the line can be generalized as the $k - 1$ dimensional hyper-plane. The point contained in $g(u) = 0$ that is closest to the origin is referred to as the design point u^* .

1.2. Non-linear limit state function

When the limit state function is non-linear with respect to the basic random variables $\mathbf{X}$, the situation is not as simple as outlined in the previous. One feasible approximation is to represent the failure domain in terms of a linearization of the boundary between the safe domain and the failure domain, i.e. the failure surface. However, the question remains on how to choose the linearization point.

Hasofer and Lind (1974) suggests performing this linearization at the design point of the failure surface represented in the normalized space. The situation is illustrated in the two-dimensional space in Figure 3.2.

The failure surface is linearized at the design point $\mathbf{u}^*$ by the line $g'(\mathbf{u}) = 0$. The α -vector is the outward directed normal vector to the failure surface in the design point $\mathbf{u}^*$, i.e. the point of

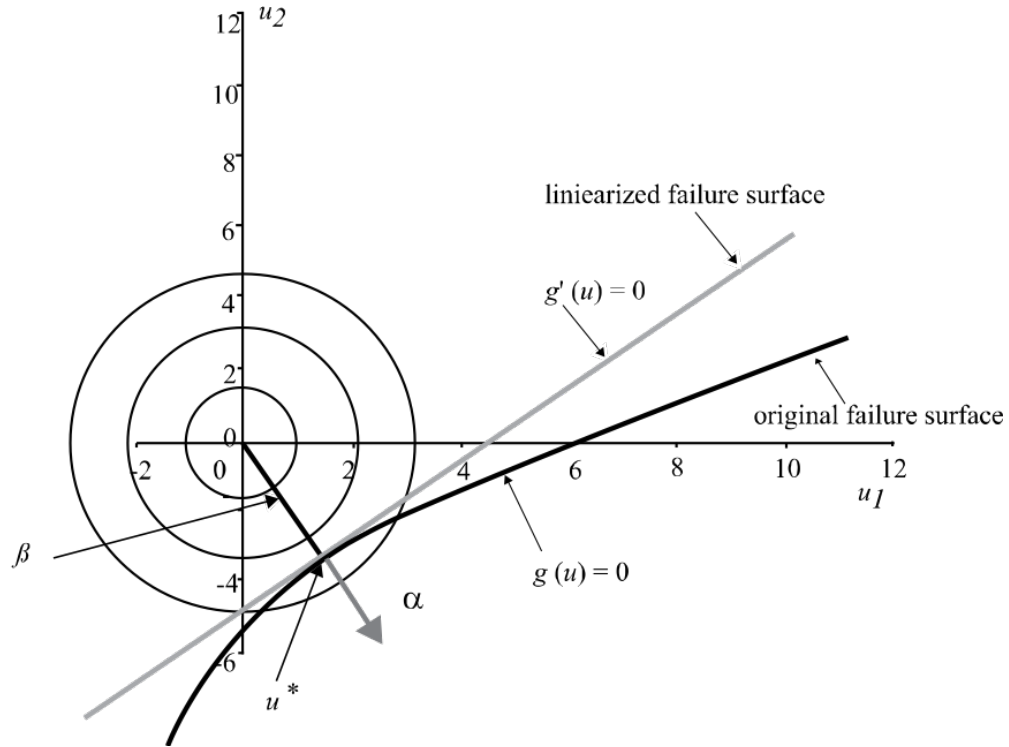


Figure 3.2: Illustration of the linearization proposed by Hasofer and Lind¹ in the standard Normal space.

the linearized failure surface with the shortest distance to the origin β .

As the limit state function is in general non-linear, one does not know the design point in advance. This has to be found iteratively, e.g. by solving the optimization problem in Eq. (3.9).

$$\beta = \min_{u \in \{g(u)=0\}} \sqrt{\sum_{i=1}^n u_i^2} \quad (3.9)$$

This problem may be solved in a number of different ways. Provided that the limit state function is differentiable, the following simple iteration scheme may be followed:

$$\alpha_i = \frac{-\left. \frac{\partial g}{\partial u_i} \right|_{(\beta \cdot \alpha)}}{\sqrt{\sum_{i=1}^n \left(\left. \frac{\partial g}{\partial u_i} \right|_{(\beta \cdot \alpha)} \right)^2}}, \quad i = 1, 2, \dots, n \quad (3.10)$$

$$g(\beta \cdot \alpha_1, \beta \cdot \alpha_2, \dots, \beta \cdot \alpha_n) = 0 \quad (3.11)$$

First, a design point is guessed ($\mathbf{u}^* = \beta \cdot \alpha$) and inserted into Eq. (3.10), whereby a new normal vector α to the failure surface is achieved. This α -vector is then inserted into Eq. (3.11), from which a new β -value is calculated.

¹Hasofer, A. M., & Lind, N. C. (1973). An exact and invariant first-order reliability format. Solid Mechanics Division, University of Waterloo.

The iteration scheme will converge in a few iterations (say normally 6-10) and provides the design point $\mathbf{u}^*$, as well as the reliability index β and the outward normal vector to the failure surface at the design point α . As already mentioned, the reliability index β may be related directly to the probability of failure. The components of the α -vector may be interpreted as sensitivity factors, which give the relative importance of the individual random variables for the reliability index β .

The Second Order Reliability Methods (SORM) follow the same principles as FORM. However, as a logical extension of FORM, the failure surface is expanded to the second order at the design point. The result of a SORM analysis may be given as the FORM β multiplied with a correction factor, evaluated on the basis of the second order partial derivatives of the failure surface at the design point. Notice that the SORM analysis becomes exact for failure surfaces that are defined by second order polynomials of the random variables. However, generally speaking, the result of a SORM analysis can be shown to be asymptotically exact for any shape of the failure surface as β approaches infinity. The interested reader is referred to the literature² for the details of SORM analyses.

Example: Non-linear Safety Margin

Consider a steel rod that has resistance R and is loaded in tension by the force S . The cross-section area of the steel rod A is assumed to be uncertain. The steel yield stress R is Normal distributed with mean values and standard deviation $\mu_R = 350$ MPa, $\sigma_R = 35$ MPa. The loading S is Normal distributed with mean value and standard deviation $\mu_S = 1500$ N, $\sigma_S = 300$ N. Finally, the cross-section area A is assumed Normal distributed with mean value and standard deviation $\mu_A = 10$ mm², $\sigma_A = 1$ mm².

The limit state function may be written as:

$$g(x) = r \cdot a - s$$

The first step is to transform the Normal distributed random variables R , A and S into standardized Normal distributed random variables, i.e.:

$$U_R = \frac{R - \mu_R}{\sigma_R}$$

$$U_A = \frac{A - \mu_A}{\sigma_A}$$

$$U_S = \frac{S - \mu_S}{\sigma_S}$$

The limit state function may now be written in the space of the standardized Normal distributed random variables as:

$$\begin{aligned} g(u) &= (u_R \sigma_R + \mu_R)(u_A \sigma_A + \mu_A) - (u_S \sigma_S + \mu_S) \\ &= (35u_R + 350)(u_A + 10) - (300u_S + 1500) \\ &= 350u_R + 350u_A - 300u_S + 35u_R u_A + 2000 \end{aligned}$$

²Hasofer, A. M., & Lind, N. C. (1974). Exact and invariant second-moment code format. *Journal of the Engineering Mechanics division*, 100(1), 111-121.

The reliability index and the design point may be determined in accordance with Equations (3.10) and (3.11) as:

$$\begin{aligned}\beta &= \frac{-2000}{350\alpha_R + 350\alpha_A - 300\alpha_S + 35\beta\alpha_R\alpha_A} \\ \alpha_R &= -\frac{1}{k}(350 + 35\beta\alpha_A) \\ \alpha_A &= -\frac{1}{k}(350 + 35\beta\alpha_R) \\ \alpha_S &= \frac{300}{k}\end{aligned}$$

with

$$k = \sqrt{\sum_{i=1}^n \left(\frac{\partial g}{\partial u_i} \Big|_{(\beta, \alpha)} \right)^2} = \sqrt{(350 + 35\beta\alpha_A)^2 + (350 + 35\beta\alpha_R)^2 + (300)^2}$$

which by calculation gives the iteration history shown in the following Table:

Table 3.1: Iteration history for the non-linear limit state example.

Iteration	Start	1	2	3	4	5
β	3.0000	3.6719	3.7399	3.7444	3.7448	3.7448
α_R	-0.5800	-0.5701	-0.5612	-0.5611	-0.5610	-0.5610
α_A	-0.5800	-0.5701	-0.5612	-0.5611	-0.5610	-0.5610
α_S	0.5800	0.5916	0.6084	0.6086	0.6087	0.6087

From Table 3.1, it is seen that the basic random variable S modelling the load on the steel rod is slightly dominating with an α -value equal to 0.6087. Furthermore it is seen that both the variables R and A are acting as resistance variables as their α -values are negative. The failure probability for the steel rod is determined as $P_F = \Phi(-3.7448) = 9.02 \cdot 10^{-5}$.

1.3. Correlated and Dependent Random Variables

The situation where basic random variables $\mathbf{X}$ are stochastically dependent is often encountered in practical problems. For Normal distributed random variables, the joint probability distribution function may be described in terms of the first two moments, i.e. the mean value vector and the covariance matrix. This is, however, only the case for Normal or Log-Normal distributed random variables.

Correlation and dependency may be treated completely along the same lines as described in the context of Monte Carlo Simulation and the NORTA method. Note that, in addition to the transformation from a limit state function expressed in variables $\mathbf{X}$ to a limit state function expressed in variables $\mathbf{U}$, a transformation in between is introduced where the considered random

variables are first normalized before they are made uncorrelated. I.e. the row of transformations yields:

$$X \rightarrow Y \rightarrow U$$

In the following, the case of Normal distributed random variables is considered. It will be seen how this transformation may be implemented in the iterative procedure outlined previously. Let us assume that the basic random variables $\mathbf{X}$ are correlated with covariance matrix given as:

$$\mathbf{C}_\mathbf{X} = \begin{bmatrix} Var[X_1] & Cov[X_1, X_2] \dots & Cov[X_1, X_n] \\ \vdots & \vdots & \vdots \\ Cov[X_n, X_1] & \dots & Var[X_n] \end{bmatrix} \quad (3.12)$$

and correlation coefficient matrix $\rho_\mathbf{X}$:

$$\rho_\mathbf{X} = \begin{bmatrix} 1 & \dots & \rho_{1n} \\ \vdots & 1 & \vdots \\ \rho_{1n} & \dots & 1 \end{bmatrix} \quad (3.13)$$

Note that the random variables are uncorrelated if $\rho_\mathbf{X}$ is the identity matrix.

As before, the first step is to transform the n -vector of basic random variables $\mathbf{X}$ into a vector of standardized random variables $\mathbf{Y}$ with zero mean values and unit variances. This operation may be performed by:

$$Y_i = \frac{X_i - \mu_{X_i}}{\sigma_{X_i}}, \quad i = 1, 2, \dots, n \quad (3.14)$$

whereby the covariance matrix of $\mathbf{Y}$, denoted $\mathbf{C}_\mathbf{Y}$, is equal to the correlation coefficient matrix of $\mathbf{X}$, i.e. $\rho_\mathbf{X}$.

The second step is to transform the vector of standardized basic random variables $\mathbf{Y}$, into a vector of uncorrelated basic random variables $\mathbf{U}$. This last transformation may be performed in several ways. The approach described in the following utilizes the Cholesky factorization from matrix algebra and is efficient for both hand calculations and for implementation in computer programs. The transformation that is sought after may be written as:

$$\mathbf{Y} = \mathbf{T}\mathbf{U} \quad (3.15)$$

where $\mathbf{T}$ is a lower triangular matrix, i.e. $T_{ij} = 0$ for $j > i$.

It is then seen that the covariance matrix $\mathbf{C}_\mathbf{Y}$ can be written as:

$$\mathbf{C}_\mathbf{Y} = E[\mathbf{Y} \cdot \mathbf{Y}^T] = E[\mathbf{T} \cdot \mathbf{U} \cdot \mathbf{U}^T \cdot \mathbf{T}^T] = \mathbf{T} \cdot E[\mathbf{U} \cdot \mathbf{U}^T] \cdot \mathbf{T}^T = \mathbf{T} \times \mathbf{T}^T = \rho_\mathbf{X} \quad (3.16)$$

where the superscript T denotes the transpose of a matrix. The components of $\mathbf{T}$ may be

determined from Eq. (3.16) as:

$$\begin{aligned}
 T_{11} &= 1 \\
 T_{21} &= \rho_{12} \\
 T_{31} &= \rho_{13} \\
 T_{22} &= \sqrt{1 - T_{21}^2} \\
 T_{32} &= \frac{\rho_{23} - T_{31} \cdot T_{21}}{T_{22}} \\
 T_{33} &= \sqrt{1 - T_{31}^2 - T_{32}^2} \\
 &\vdots
 \end{aligned} \tag{3.17}$$

Example (Cont.)

The previous example of the rod is again considered. Additionally, the random variables A and R are assumed to be correlated, with correlation coefficient matrix given as:

$$\boldsymbol{\rho}_{\mathbf{X}} = \begin{bmatrix} 1 & \rho_{AR} & \rho_{AS} \\ \rho_{RA} & 1 & \rho_{RS} \\ \rho_{SA} & \rho_{SR} & 1 \end{bmatrix} = \begin{bmatrix} 1 & 0.1 & 0 \\ 0.1 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

with $\rho_{AR} = \rho_{RA} = 0.1$ and all other correlation coefficients equal to zero. The transformation matrix T can now be calculated as:

$$\mathbf{T} = \begin{bmatrix} 1 & 0 & 0 \\ 0.1 & 0.995 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

The components of the vector $\mathbf{Y}$ may then be calculated as:

$$\begin{aligned}
 Y_A &= U_A \\
 Y_R &= 0.1 \cdot U_A + 0.995 \cdot U_R \\
 Y_S &= U_S
 \end{aligned}$$

Finally, the components of the vector $\mathbf{X}$ are determined as:

$$\begin{aligned}
 X_A &= U_A \cdot \sigma_A + \mu_A \\
 X_R &= (0.1 \cdot U_A + 0.995 \cdot U_R) \cdot \sigma_R + \mu_R \\
 X_S &= U_S \cdot \sigma_S + \mu_S
 \end{aligned}$$

whereby the limit state function can be written in terms of the uncorrelated and normalized random variables $\mathbf{U}$ as follows:

$$g(u) = ((0.1 \cdot u_A + 0.995 \cdot u_R) \cdot \sigma_R + \mu_R)(u_A \sigma_A + \mu_A) - (u_S \sigma_S + \mu_S)$$

from which the reliability index can be calculated as in the previous example.

In case that the stochastically dependent basic random variables are not Normal or Log-

Normal distributed, the dependency can no longer be described completely in terms of correlation coefficients. Thus, the above-described transformation is thus not appropriate. In such cases, other transformations must be applied, as described in the next section.

1.4. Non-Normal and Dependent Random Variables

As stated in the previous, the joint probability distribution function of a random vector $\mathbf{X}$ can only be completely described in terms of the marginal probability distribution functions for the individual components of the vector $\mathbf{X}$ and the correlation coefficient matrix, when all the components of $\mathbf{X}$ are either Normal or Log-Normal distributed.

First of all, the simpler case of the components of $\mathbf{X}$ being independent but non-Normal distributed is considered. Thereafter, it shall be seen how, in some cases, the situation of jointly dependent and non-Normal distributed random variables may be treated.

1.4.1. The full transformation

Let g be a limit state function, dependent on n random variables $\mathbf{X} = \{X_1, \dots, X_n\}$. The full transformation transforms a random realization x_i into the u -space with the following the mapping:

$$x_i = F_X^{-1}[\Phi(u_i)] \quad (3.18)$$

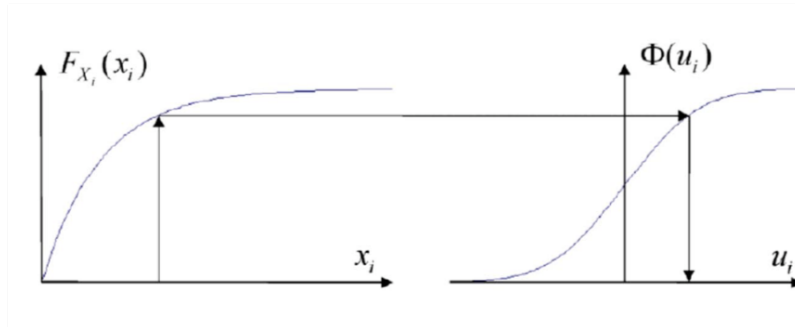


Figure 3.3: Full-transformation scheme.

The limit state function g can be transformed to the standard Normal space as:

$$g_U(\mathbf{u}) = g(F_{X_1}^{-1}[\Phi(u_1)], \dots, F_{X_n}^{-1}[\Phi(u_n)]) \quad (3.19)$$

The derivative of g_U with respect to u_i is then computed by using the chain rule, i.e.:

$$\frac{\partial g}{\partial u_i} = \frac{\partial g}{\partial x_i} \frac{\partial x_i}{\partial u_i} \quad (3.20)$$

Explicit formulations for x_i exist for some particular distribution families. In particular, two cases are here regarded:

X is Log-Normal distributed

The cumulative distribution function of a Log-Normal distributed random variable X with

expected value μ and standard deviation σ can be expressed as:

$$F_X(x) = \Phi\left(\frac{\ln(x) - \mu_L}{\sigma_L}\right) \quad (3.21)$$

Where μ_L and σ_L are the distribution parameters given by:

$$\begin{aligned} \sigma_L &= \sqrt{\ln\left(\frac{\sigma^2}{\mu^2} + 1\right)} \\ \mu_L &= \ln \mu - \frac{1}{2}\sigma_L^2 \end{aligned} \quad (3.22)$$

The transformation in Eq. (3.18) can be then be rewritten as:

$$x = F_X^{-1}[\Phi(u)] = \exp(\sigma_L \cdot u + \mu_L) \quad (3.23)$$

The derivative of x with respect to u :

$$\frac{\partial x}{\partial u} = \frac{\varphi(u)}{f_X(F_X^{-1}(\Phi(u)))} = \sigma_L \exp(\mu_L + \sigma_L u) \quad (3.24)$$

X is Gumbel distributed

The cumulative distribution function of a Gumbel distributed random variable X with expected value μ and standard deviation σ can be expressed as:

$$F_X(x) = \exp[-\exp[-a(x - b)]] \quad (3.25)$$

Where a and b are the distribution parameters given by:

$$\begin{aligned} a &= \frac{\pi}{\sqrt{6}\sigma} \\ b &= \mu - \frac{0.5772}{a} \end{aligned} \quad (3.26)$$

The transformation in Eq. (3.18) can be then be rewritten as:

$$x = F_X^{-1}[\Phi(u)] = b - \frac{1}{a} \ln[-\ln \Phi(u)] \quad (3.27)$$

The derivative of x with respect to u :

$$\frac{\partial x}{\partial u} = \frac{\varphi(u)}{f_X(F_X^{-1}(\Phi(u)))} = -\frac{\varphi(u)}{a\Phi(u) \ln(\Phi(u))} \quad (3.28)$$

1.4.2. The Normal-tail Approximation

Another approach to consider the problem of non-Normal distributed random variables, within the context of the iterative scheme given in Eqs. (3.10) and (3.11) for the calculation of the reliability index β , is to approximate the actual probability distribution by a Normal probability distribution at the design point. Due to the fact that the design point is usually located in the

tails of the distribution functions of the basic random variables, the scheme is often referred to as the “Normal tail approximation”.

Denoting the design point by x^* , the approximation is introduced as:

$$F_{X_i}(x_i^*) = \Phi\left(\frac{x_i^* - \mu'_{X_i}}{\sigma'_{X_i}}\right) \quad (3.29)$$

$$f_{X_i}(x_i^*) = \frac{1}{\sigma'_{X_i}} \varphi\left(\frac{x_i^* - \mu'_{X_i}}{\sigma'_{X_i}}\right) \quad (3.30)$$

where μ'_{X_i} and σ'_{X_i} are, respectively, the unknown mean value and standard deviation of the approximating Normal distribution.

Solving Eqs. (3.29) and (3.30) with respect to μ'_{X_i} and σ'_{X_i} results in the following expressions:

$$\sigma'_{X_i} = \frac{\varphi(\Phi^{-1}[F_{X_i}(x_i^*)])}{f_{X_i}(x_i^*)} \quad (3.31)$$

$$\mu'_{X_i} = x_i^* - \Phi^{-1}[F_{X_i}(x_i^*)] \cdot \sigma'_{X_i}$$

This transformation may easily be introduced in the iterative evaluation of the reliability index β as a final step before the basic random variables are normalized.

1.4.3. The Rosenblatt Transformation

If the joint probability distribution function of the random vector $\mathbf{X}$ can be obtained in terms of a sequence of conditional probability distribution functions such as:

$$F_X(x) = F_{X_n}(x_n | x_1, x_2, \dots, x_{n-1}) \cdot F_{X_{n-1}}(x_{n-1} | x_1, x_2, \dots, x_{n-2}) \cdot \dots \cdot F_{X_1}(x_1) \quad (3.32)$$

then the transformation from the $\mathbf{X}$ -space to the $\mathbf{U}$ -space may be performed using the so-called *Rosenblatt transformation*:

$$\begin{aligned} \Phi(u_1) &= F_{X_1}(x_1) \\ \Phi(u_2) &= F_{X_2}(x_2 | x_1) \\ &\vdots \\ \Phi(u_n) &= F_{X_n}(x_n | x_1, x_2, \dots, x_{n-1}) \end{aligned} \quad (3.33)$$

where n is the number of random variables, $F_{X_i}(x_i | x_1, x_2, \dots, x_{i-1})$ is the conditional probability distribution function for the i^{th} random variable, given the realizations of $x_1, x_2, \dots, x_{i-1}$ and $\Phi(\cdot)$ is the standard Normal probability distribution function. From the transformation given by Eq. (3.33), the basic random variables $\mathbf{X}$ may be expressed in terms of standardized Normal distributed random variables $\mathbf{U}$ by:

$$\begin{aligned} x_1 &= F_{X_1}^{-1}[\Phi(u_1)] \\ x_2 &= F_{X_2}^{-1}[\Phi(u_2) | x_1] \\ &\vdots \\ x_n &= F_{X_n}^{-1}[\Phi(u_n) | x_1, x_2, \dots, x_{n-1}] \end{aligned} \quad (3.34)$$

In some cases the Rosenblatt transformation cannot be applied because the required conditional probability distribution functions cannot be provided. In such cases, other transformations may be useful such as, e.g. the Nataf transformation. Standard commercial software for FORM analysis usually include a selection of possibilities for the representation of dependent non-Normal distributed random variables.

2. Application examples

E3.1. Reliability assessment of the SLS of a bridge deck using FORM

The limitation of beam deflections is usually a governing requirement in design. Deflections larger than the maximum allowed violate a limit state which is categorized as Serviceability Limit State (SLS), since no failure is involved.

The probability of having a mid-span deflection larger than 8 mm needs to be calculated for the beam illustrated in Figure 3.4 given the distributed load Q . The basic variables are reported in the Table 3.2.

Table 3.2: Basic Variables.

	Distribution	Mean	Standard deviation
b	Deterministic	50 mm	-
E	Deterministic	205000 N/mm ²	-
l	Deterministic	3000 mm	-
H	Normal	100 mm	5 mm
Q	Normal	5 N/mm	1 N/mm

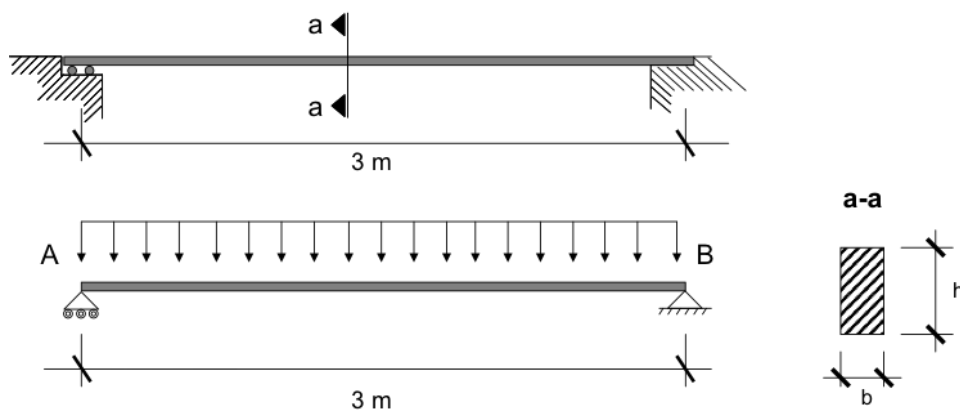


Figure 3.4: Beam geometry and cross-section.

Approach 1

The deflection at mid-span is calculated as:

$$w = \frac{5ql^4}{384E \frac{bh^3}{12}} \quad (3.35)$$

The limit state function which bounds the failure domain is:

$$g(q, h) = 8 \text{ mm} - \left(1.235 \cdot 10^6 \frac{\text{mm}^5}{\text{N}} \right) \frac{q}{h^3} = 0 \quad (3.36)$$

Realizations of Q and H which gives $g(q, h) \leq 0$ are associated with violation of the SLS, i.e. deflections larger than 8 mm. Inserting the deterministic values in Eq. (3.36) yields:

$$g(q, h) = 8 \text{ mm} - \left(1.235 \cdot 10^6 \frac{\text{mm}^5}{\text{N}} \right) \frac{q}{h^3} = 0 \quad (3.37)$$

Or equivalently:

$$g(q, h) = (8 \text{ mm})h^3 - \left(1.235 \cdot 10^6 \frac{\text{mm}^5}{\text{N}} \right) q = 0 \quad (3.38)$$

As it can be observed, $g(q, h)$ is non-linear.

The transformation of the basic random variables (H, Q) into normal standard variables (U_H, U_Q) reads:

$$U_H = \frac{H - \mu_H}{\sigma_H} ; U_Q = \frac{Q - \mu_Q}{\sigma_Q} \quad (3.39)$$

This corresponds to:

$$H = U_H \sigma_H + \mu_H ; Q = U_Q \sigma_Q + \mu_Q \quad (3.40)$$

The limit state function in the normal standard space is obtained by substituting Eq. (3.40) into Eq. (3.38):

$$\begin{aligned} g'(u_Q, u_H) &= 8(u_H \sigma_H + \mu_H)^3 - 1.235 \cdot 10^6(u_Q \sigma_Q + \mu_Q) = 0 \\ g'(u_Q, u_H) &= u_H^3 + 60u_H^2 + 1200u_H - 1235u_Q + 1825 = 0 \end{aligned} \quad (3.41)$$

The reliability index is the distance between the origin and the design point **in the U-space**. The **design point** has coordinates (u_Q^*, u_H^*) in the U -space and is defined as the point on the line $g'(u_Q, u_H) = 0$ that is closest to the origin, see Figure 3.6. The corresponding point in the X -space has coordinates $q^* = u_Q^* \sigma_Q + \mu_Q, h^* = u_H^* \sigma_H + \mu_H$. The design point can therefore be expressed as $(u_Q^* = \beta \alpha_Q, u_H^* = \beta \alpha_H)$. Substituting this expression in Eq. (3.41) leads to:

$$g'(u_Q, u_H) = (\beta \alpha_H)^3 + 60(\beta \alpha_H)^2 + 1200(\beta \alpha_H) - 1235(\beta \alpha_Q) + 1825 = 0 \quad (3.42)$$

The FORM algorithm for finding the reliability index is reported in the following pages. The algorithm is very efficient even when the starting vector $\alpha^{(1)}$ is far from the final one. Setting $\alpha^{(1)} = (1/\sqrt{7}, -6/\sqrt{7})$ the algorithm converges in 3 – 4 iterations, see Figure 3.7.

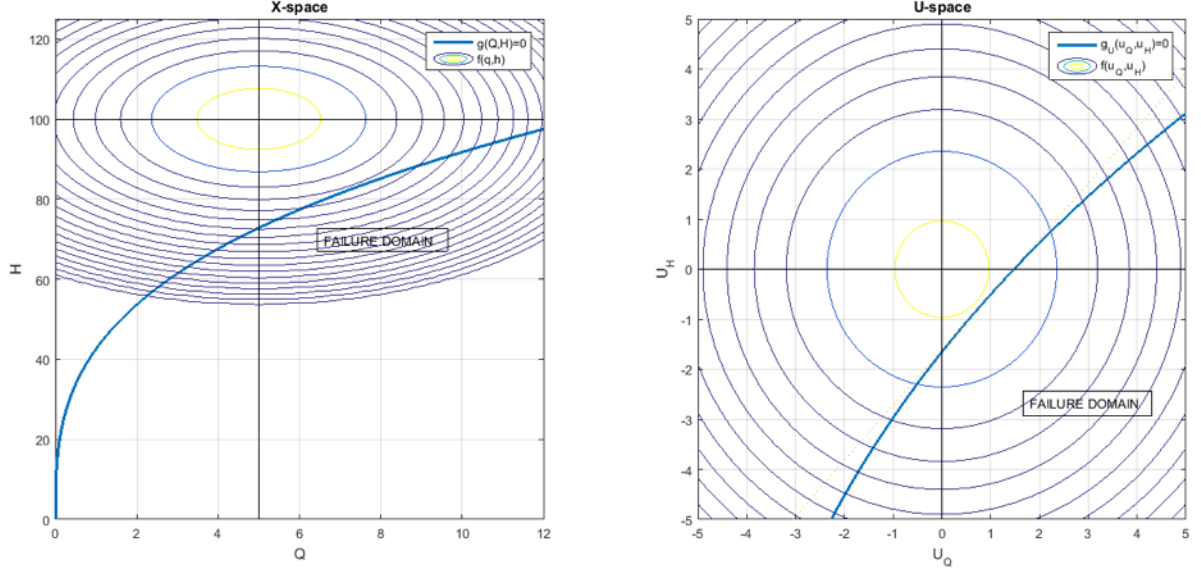


Figure 3.5: Limit state function in the X and U -space with joint probability density function.

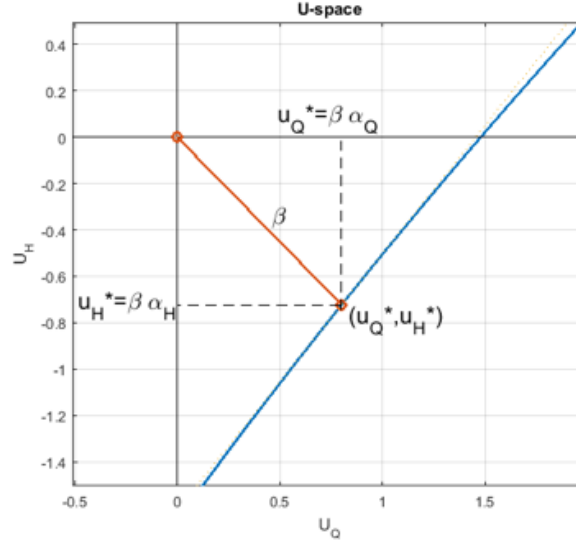


Figure 3.6: Geometrical interpretation of β and design point. The vector with coordinates (α_Q, α_H) has unitary length and points towards the design point.

Approach 2

An alternative approach is proposed in order to avoid solving the limit state function for finding the reliability index $\beta^{(i)}$. Its convergence is slightly slower but still satisfactory. The key point is to collect beta in Eq. (3.42):

$$g'(u_Q, u_H) = \beta [\alpha_H(\beta\alpha_H)^2 + 60\beta\alpha_H^2 + 1200\alpha_H - 1235\alpha_Q] + 1825 = 0 \quad (3.43)$$

Isolating β results in:

$$\beta = \frac{-1825}{[\alpha_H(\beta\alpha_H)^2 + 60\beta\alpha_H^2 + 1200\alpha_H - 1235\alpha_Q]} \quad (3.44)$$

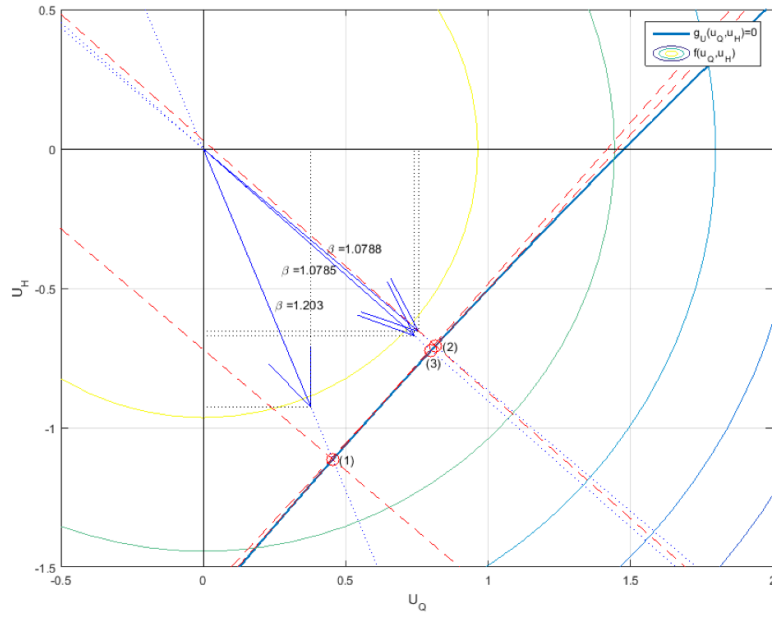
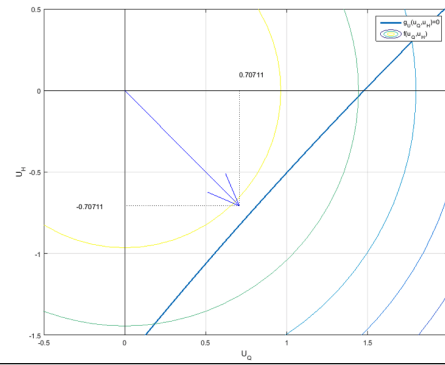


Figure 3.7: Convergence for $\alpha^{(1)} = (1/\sqrt{7}, -6/\sqrt{7})$.

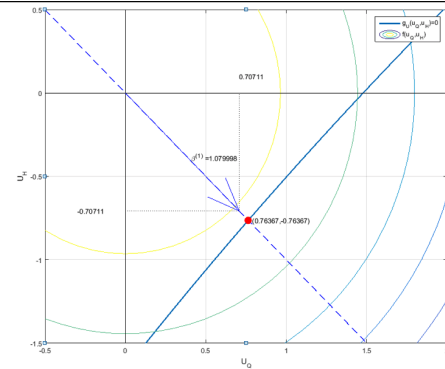
Iteration 1:

Step 0 (only in first iteration): guess an α -vector. The vector must have unitary length; negative coordinates for “resistance-variables” and positive coordinates for “load-variables”. Having two random variables a possible start is: $\alpha^{(1)} = (1/\sqrt{2}; -1/\sqrt{2})$. The norm of $\alpha^{(1)}$ is $\|\alpha^{(1)}\| = \sqrt{(1/\sqrt{2})^2 + (-1/\sqrt{2})^2} = 1$. The design point is $(u_Q^{(1)} = \beta^{(1)} \cdot \alpha_Q^{(1)}; u_H^{(1)} = \beta^{(1)} \cdot \alpha_H^{(1)})$ with $\beta^{(1)}$ unknown calculated in step 2.



Step 1: only from 2nd iteration

Step 2: calculate $\beta^{(1)}$ by solving Eq. (8): $g'(\beta^{(1)} \cdot \alpha_Q^{(1)}, \beta^{(1)} \cdot \alpha_H^{(1)}) = 0$. The equation is difficult to solve for $\beta^{(1)}$ therefore a numerical solver is used, it gives $\beta^{(1)} = 1.079998$. (check that the design point is on the limits state function: $g'(1.080 \cdot (1/\sqrt{2}), 1.080 \cdot (-1/\sqrt{2})) \approx 0$, OK)



We define $\beta^{(i)}$ as the reliability index at the iteration i . We can approximate $\beta^{(i)}$ by calculating it from the reliability index $\beta^{(i-1)}$ and $\alpha^{(i-1)}$ estimated in the previous iteration. A good guess of the starting values is:

- $\beta^{(i)} = 1.00$ for SLS and $\beta^{(i)} = 3.00$ or 4.00 for typical ULS.
- $\alpha^{(1)} = (1/\sqrt{n}, 1/\sqrt{n}, \dots, 1/\sqrt{n}, -1/\sqrt{n})$, with n representing the number of basic random variables.

Iteration 2

Step 1: evaluate the new alpha vector $\alpha^{(2)}$ i.e. the unitary vector orthogonal to $g'(\cdot)$ at the design point:

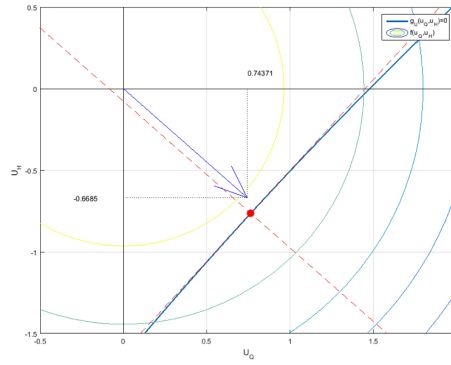
$$\alpha^{(2)} = \left(\frac{-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(1)}, u_H^{(1)} \right)}{k}, \frac{-\frac{\partial g'}{\partial u_H} \left(u_Q^{(1)}, u_H^{(1)} \right)}{k} \right) \text{ where}$$

$$-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(1)}, u_H^{(1)} \right) = -(-1235.00)$$

$$-\frac{\partial g'}{\partial u_H} \left(u_Q^{(1)}, u_H^{(1)} \right) = -\left(3(u_H^{(1)})^2 + 120(u_H^{(1)}) + 1200 \right) = -1110.11$$

$$k = \sqrt{\sum_{i=1}^2 \left(\frac{\partial g'}{\partial u_i} \left(u_Q^{(1)}, u_H^{(1)} \right) \right)^2} = \sqrt{1235.00^2 + (-1110.11)^2} = 1660.59$$

$$\alpha^{(2)} = (0.7437; -0.6685)$$



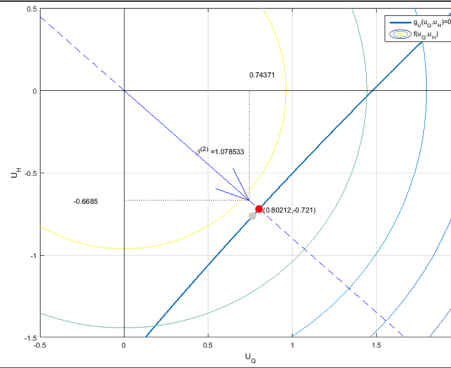
Step 2: calculate $\beta^{(2)}$ by solving Eq. (8)

$$g'(\beta^{(2)} \cdot \alpha_Q^{(2)}, \beta^{(2)} \cdot \alpha_H^{(2)}) = 0.$$

Numerical solution is $\beta^{(2)} = 1.078533$.

(check:

$$g'(\beta^{(2)} \cdot \alpha_Q^{(2)}, \beta^{(2)} \cdot \alpha_H^{(2)}) = 0 = 9.04 \cdot 10^{-8} \approx 0)$$



Iteration 3

Step 1: evaluate the new alpha vector $\alpha^{(3)}$

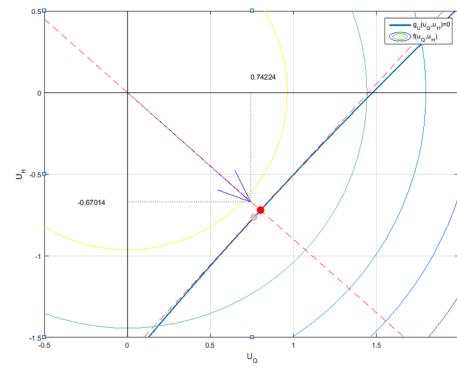
$$\alpha^{(3)} = \left(\frac{-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(2)}, u_H^{(2)} \right)}{k}, \frac{-\frac{\partial g'}{\partial u_H} \left(u_Q^{(2)}, u_H^{(2)} \right)}{k} \right) \text{ where}$$

$$-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(2)}, u_H^{(2)} \right) = -(-1235.00)$$

$$-\frac{\partial g'}{\partial u_H} \left(u_Q^{(2)}, u_H^{(2)} \right) = -\left(3(u_H^{(2)})^2 + 120(u_H^{(2)}) + 1200 \right) = -1115.04$$

$$k = \sqrt{\sum_{i=1}^2 \left(\frac{\partial g'}{\partial u_i} \left(u_Q^{(2)}, u_H^{(2)} \right) \right)^2} = \sqrt{1235.00^2 + (-1115.04)^2} = 1663.89$$

$$\alpha^{(3)} = (0.7422; -0.6701)$$



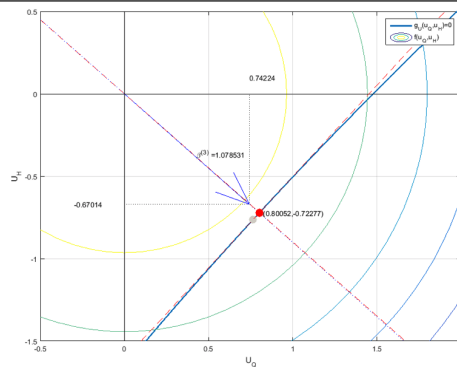
Step 2: calculate $\beta^{(3)}$ by solving Eq. (8)

$$g'(\beta^{(3)} \cdot \alpha_Q^{(3)}, \beta^{(3)} \cdot \alpha_H^{(3)}) = 0.$$

Numerical solution is $\beta^{(3)} = 1.078531$.

(check:

$$g'(\beta^{(3)} \cdot \alpha_Q^{(3)}, \beta^{(3)} \cdot \alpha_H^{(3)}) \approx 0)$$



The iterations are similar to the previous method:

Iteration 4**Step 1:** evaluate the new alpha vector $\alpha^{(4)}$

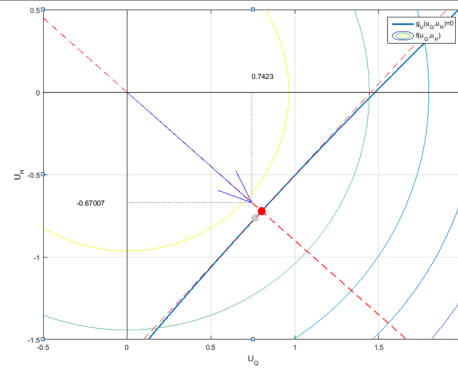
$$\alpha^{(4)} = \left(\frac{-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(3)}, u_H^{(3)} \right)}{k}, \frac{-\frac{\partial g'}{\partial u_H} \left(u_Q^{(3)}, u_H^{(3)} \right)}{k} \right) \text{ where}$$

$$-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(3)}, u_H^{(3)} \right) = -(-1235.00)$$

$$-\frac{\partial g'}{\partial u_H} \left(u_Q^{(3)}, u_H^{(3)} \right) = -\left(3(u_H^{(3)})^2 + 120(u_H^{(3)}) + 1200 \right) = -1114.84$$

$$k = \sqrt{\sum_{i=1}^2 \left(\frac{\partial g'}{\partial u_i} \left(u_Q^{(3)}, u_H^{(3)} \right) \right)^2} = \sqrt{1235.00^2 + (-1114.84)^2} = 1663.76$$

$$\alpha^{(4)} = (0.7423; -0.6701)$$

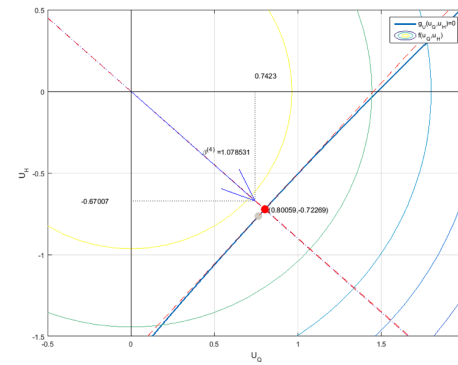
**Step 2:** calculate $\beta^{(4)}$ by solving Eq. (8)

$$g'(\beta^{(4)} \cdot \alpha_Q^{(4)}, \beta^{(4)} \cdot \alpha_H^{(4)}) = 0.$$

Numerical solution is $\beta^{(4)} = 1.078531$.

(check:

$$g'(\beta^{(4)} \cdot \alpha_Q^{(4)}, \beta^{(4)} \cdot \alpha_H^{(4)}) \approx 0)$$

Stop since $|\beta^{(3)} - \beta^{(4)}| < tol$ **Step 1:** Evaluate the new vector alpha

$$\alpha^{(i)} = \left(\frac{-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(i-1)}, u_H^{(i-1)} \right)}{k}, \frac{-\frac{\partial g'}{\partial u_H} \left(u_Q^{(i-1)}, u_H^{(i-1)} \right)}{k} \right) \text{ where}$$

$$-\frac{\partial g'}{\partial u_Q} \left(u_Q^{(i-1)}, u_H^{(i-1)} \right) = -(-1235.00)$$

$$-\frac{\partial g'}{\partial u_H} \left(u_Q^{(i-1)}, u_H^{(i-1)} \right) = -\left(3(u_H^{(i-1)})^2 + 120(u_H^{(i-1)}) + 1200 \right)$$

$$k = \sqrt{\sum_{i=1}^2 \left(\frac{\partial g'}{\partial u_i} \left(u_Q^{(i-1)}, u_H^{(i-1)} \right) \right)^2}$$

Step 2: Evaluate $\beta^{(i)}$ with:

$$\beta^{(i)} = \frac{-1825}{\left[\alpha_H^{(i)} (\beta^{(i-1)} \cdot \alpha_H^{(i)})^2 + 60 \cdot \beta^{(i-1)} \cdot (\alpha_H^{(i)})^2 + 1200 \cdot \alpha_H^{(i)} - 1235 \cdot \alpha_Q^{(i)} \right]}$$

In general, we obtain $g'(\beta^{(i)} \cdot \alpha_Q^{(i)}, \beta^{(i)} \cdot \alpha_H^{(i)}) \neq 0$ so the design point is NOT on the curve $g'(\cdot) = 0$, but it will be quite close to it.

The second method converges to the same result and has the advantage of avoiding the analytical calculation of β from the limit state function. β is calculated numerically. The analytical solution is reported in Figure 3.8. Note that in many situations is not possible to find the analytical solution.

$$\beta = \frac{1}{6} - \frac{\left(\frac{14820\sqrt{15}}{\alpha_h^2} \sqrt{\frac{135\alpha_h^3 - 1080\alpha_h^2\alpha_g + 2160\alpha_h\alpha_g^2 - 988\alpha_g^3}{\alpha_h^2} + 666900\alpha_g - 2667600\alpha_g^2} \right)^{1/3}}{\alpha_h \left(\frac{14820\sqrt{15}}{\alpha_h^2} \sqrt{\frac{135\alpha_h^3 - 1080\alpha_h^2\alpha_g + 2160\alpha_h\alpha_g^2 - 988\alpha_g^3}{\alpha_h^2} + 666900\alpha_g - 2667600\alpha_g^2} \right)^{1/3}} + \frac{2470\alpha_g}{\alpha_h \left(\frac{14820\sqrt{15}}{\alpha_h^2} \sqrt{\frac{135\alpha_h^3 - 1080\alpha_h^2\alpha_g + 2160\alpha_h\alpha_g^2 - 988\alpha_g^3}{\alpha_h^2} + 666900\alpha_g - 2667600\alpha_g^2} \right)^{1/3}} - \frac{20}{\alpha_h^2}$$

Figure 3.8: Analytical solution for the reliability index β .

3. Practical exercises

P3.1. Extension to non-Normal random variables

Consider the variable Q to be Gumbel distributed with same expected value and standard deviation. Compute again the reliability index by using FORM. Address the extension of the FORM algorithm to non-Normal distributed variables with both the *full transformation* and the *normal-tail approximation*. Comment on the solution and the difference with the previous case of Q being Normal distributed with the same moments.

Short solution The iteration steps are reported below; the starting values are $\sigma_Q^{(1)} = \sigma_Q$; $\mu_Q^{(1)} = \mu_Q$ and $\alpha^{(1)} = (1/\sqrt{n}, 1/\sqrt{n}, \dots, 1/\sqrt{n}, -1/\sqrt{n})$, with n representing the number of basic random variables. After 4 iterations, it is obtained $\beta = 1.1585$.

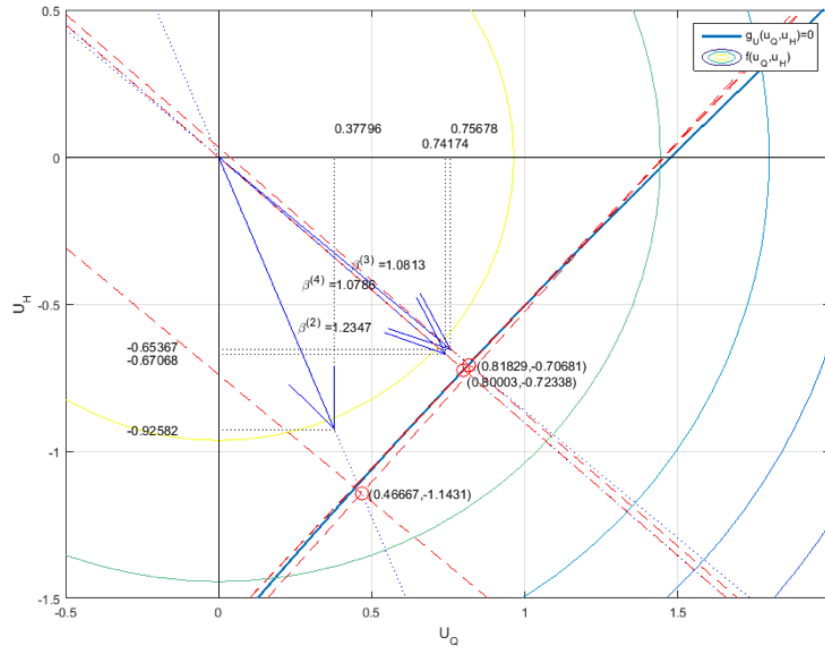
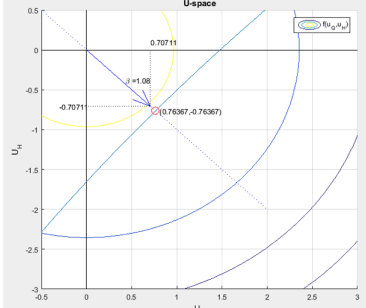
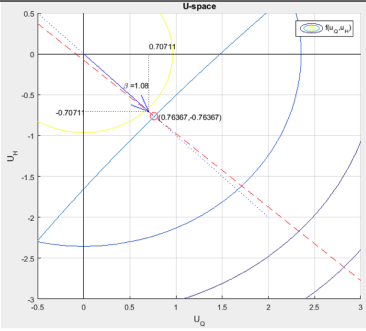
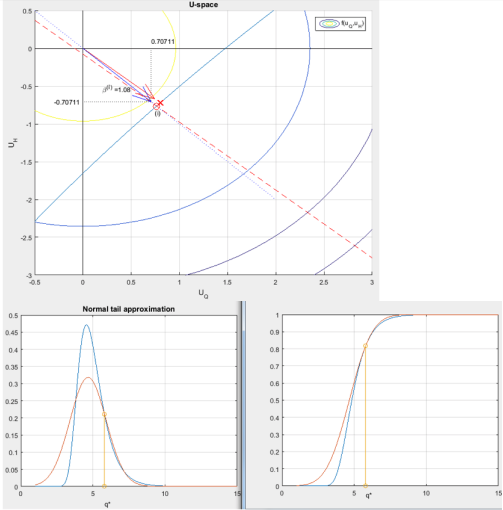


Figure 3.9: FORM convergence following Approach 2 ($\beta^{(1)} = 2.00$ and $\alpha^{(1)} = (1/\sqrt{7}, -6/\sqrt{7})$).

Iteration i Step 1: find the value of $\beta^{(i)}$ by solving $g'(\beta^{(i)}, \alpha_{\varrho}^{(i-1)}, \beta^{(i)}, \alpha_H^{(i-1)}) =$ $= 8(\beta^{(i)} \cdot \alpha_H^{(i-1)} \sigma_H + \mu_H)^3 - 1.235 \cdot 10^6 (\beta^{(i)} \cdot \alpha_{\varrho}^{(i-1)} \sigma_{\varrho}^{(i-1)} + \mu_{\varrho}^{(i-1)}) = 0$	
Step 2: evaluate the new alpha values (i.e. direction of the red dashed line) $\alpha^{(i)} = \left(\frac{\frac{\partial g'}{\partial u_{\varrho}} \Big _{(u_{\varrho}^{(i-1)}, u_H^{(i-1)})}}{k}, \frac{\frac{\partial g'}{\partial u_H} \Big _{(u_{\varrho}^{(i-1)}, u_H^{(i-1)})}}{k} \right) = \left(\frac{1235000 \sigma_{\varrho}^{(0)}}{k}, -\frac{(24(\beta^{(i)} \alpha_H^{(i-1)} \sigma_H + \mu_H)^2 \sigma_H)}{k} \right)$ $k = \sqrt{(1235000 \sigma_{\varrho}^{(0)})^2 + \left(-\left(24(\beta^{(i)} \alpha_H^{(i-1)} \sigma_H + \mu_H)^2 \sigma_H \right) \right)^2}$	
Step 3: Normal tail approximation for Q in the approximated design point (red cross) $\sigma_{\varrho}^{(i)} = \frac{\phi[\Phi^{-1}(F_{\varrho}(q^{*(i)}))]}{f_{\varrho}(q^{*(i)})}$ $\mu_{\varrho}^{(i)} = q^{*(i)} - \Phi^{-1}(F_{\varrho}(q^{*(i)})) \cdot \sigma_{\varrho}^{(i)}$ with $q^{*(i)} \approx \beta^{(i)} \alpha_{\varrho}^{(i)} \sigma_{\varrho} + \mu_{\varrho}$	
Run steps 1-2-3 until convergence.	

Question::

First order reliability method

- Explain the principle steps of the method alongside an example. Make sketches and give the relevant formulas.
- Why is the method called “first order”?
- How are non-normal random variables handled? Mention two possibilities.
- Is the method exact?

Chapter 4

Introduction to Probabilistic Load and Material Modelling

Any reliability analysis depends on the choice of an appropriate probabilistic representation of loads effects and resistances.

Some basic principles are introduced in this chapter.

It is understood that modelling is an art of reasonable simplification of reality such that the outcome is sufficiently explanatory and predictive in an engineering sense. An important aspect of an engineering model is its operationability, i.e. the ease in handling it in applications.

1. Recapturing basic principles of probability theory

2. Theoretical aspects

2.1. Recapturing basic principles of probability theory

The uncertainty about the occurrence of events is generally expressed in terms of probability. Probability of an event A can take numbers in the closed interval between zero and one, $P(A) \in [0, 1]$; loosely speaking $P(A) = 0$ indicated that A is impossible, $P(A) = 1$ means that A is certain. Probability theory is the branch of mathematics concerned with probability calculus. Formally, probability is defined via the 3 Kolmogorov axioms, named after the Russian mathematician Andrey Kolmogorov (published in 1933):

- **First axiom** (basic measure):

The probability of an event is a non-negative real number:

$$P(A) \in \mathbb{R}, P(A) \geq 0 \quad \forall A \in \Omega$$

where Ω is the event space.

- **Second axiom** (unit measure):

The probability that some (at least one) event will occur is one, i.e.

$$P(\Omega) = 1.$$

The second axiom is also referred to as axiom of unitarity.

- **Third axiom** (combination):

Any countable sequence of mutually exclusive events $A_1, A_2, \dots$ (synonymous with disjoint sets) satisfies

$$P(\bigcup_{i=1}^{\infty} A_i) = \sum_{i=1}^{\infty} P(A_i)$$

The third axiom is also referred to as axiom of σ -additivity.

Statistics is a branch of mathematics dealing with the collection, organization, analysis, interpretation and presentation of data.

Based on the probability theory and statistical assessments it is possible to represent the uncertainties associated with a given engineering problem in the decision making process. These topics are addressed with some detail in this chapter aiming to provide a fundamental understanding of the notion of uncertainty. An appropriate representation of uncertainties is available through probabilistic models such as random variables and random processes. The characterization of probabilistic models utilizes statistical information. The general principles for this are shortly outlined in this chapter as well.

2.2. Interpretation of Probability

The purpose of the theory of probability is to enable quantitative assessment of probabilities. However, the real meaning and interpretation of probabilities and probabilistic calculations as such is not a part of the theory. Consequently, two people may have completely different interpretations of the probability concept, but still use the same calculus. In the following, three different interpretations of probability are introduced and discussed based on simple cases. A formal presentation of a basic set of theory and the axioms of probability theory may be found in basic standard literature.

2.2.1. Frequentistic Definition

The *frequentistic* definition of probability is the typical interpretation of probability of the *experimentalist*. In this interpretation the probability is simply the relative frequency of occurrence of the event as observed in an experiment with trials; i.e. the probability of an event is defined as the number of times that the event occurs divided by the number of experiments that is carried out:

$$P(A) = \lim_{n_{\text{exp}} \rightarrow \infty} \frac{N_A}{n_{\text{exp}}} \quad (4.1)$$

where N_A is the number of experiments where A occurred and n_{exp} is the total number of experiments.

If a frequentist is asked what is the probability for achieving a “head” when flipping a coin, she would principally not know what to answer until she would have performed a large number of experiments. If say after 1000 experiments (flips with the coin) it is observed that “head” has occurred 563 times the answer would be that the probability for “head” is 0.563. However, as the number of experiments is increased the probability would converge towards 0.5. In the mind of a frequentist, **probability is a characteristic of nature**.

2.2.2. Classical Definition

The *classical probability* definition originates from the days when the probability calculus was founded by Pascal and Fermat¹. The inspiration for this theory was found in the games of cards and dice. The classical definition of the probability of the event A can be formulated as:

$$P(A) = \frac{n_A}{n_{tot}} \quad (4.2)$$

where n_A represents the number of equally likely ways by which an experiment may lead to A and n_{tot} is the total number of equally likely ways in the experiment.

According to the classical definition of probability, the probability of achieving a “head” when flipping a coin would be 0.5 as there is only one possible way to achieve a “head” and there are two equally likely outcomes of the experiment.

In fact, there is no real contradiction to the frequentistic definition, but the following differences may be observed:

1. The experiment does not need to be carried out as the answer is known in advance.
2. The classical theory gives no solution unless all equally possible ways can be derived analytically.

2.2.3. Bayesian Definition

In the *Bayesian interpretation* the probability $P(A)$ of the event A is formulated as a *degree of belief* that A will occur:

$$P(A) = \text{degree of belief that } A \text{ will occur} \quad (4.3)$$

Coming back to the coin-flipping problem, an engineer following the Bayesian interpretation would argue that there are two possibilities. As she has no preferences as to “head” or “tail” she would judge the probability of achieving a “head” to be 0.5.

The degree of belief is a reflection of the state of mind of the individual person in terms of experience, expertise and preferences. In this respect the Bayesian interpretation of probability is *subjective* or more precisely person-dependent. This opens up the possibility that two different people may assign different probabilities to a given event and thereby, contradicts the frequentistic interpretation that probabilities are a characteristic of nature.

The Bayesian statistical interpretation of probability includes the frequentistic and the classical interpretation in the sense that the subjectively assigned probabilities may be based on experience from previous experiments (frequentistic) as well as considerations of e.g. symmetry (classical).

The degree of belief is also referred to as a *prior belief* or *prior probability*, i.e. the belief, which may be assigned prior to obtaining any further knowledge. It is interesting to note that Immanuel Kant² developed the purely philosophical basis for the treatment of subjectivity at

¹Blaise Pascal, philosopher and mathematician, 1623-1662 ; Pierre de Fermat, mathematician, 1601-1665. The probability in definition in Eq. (4.2) is often called “Laplace probability” (Pierre-Simon de Laplace, mathematician, 1749-1827).

²Immanuel Kant, philosopher, 1724-1804

the same time as Thomas Bayes³ developed the mathematical framework, later known as the *Bayesian statistics*.

Modern structural reliability and risk analysis is based on the Bayesian interpretation of probability. However, the degree of freedom in the assignment of probabilities is in reality not as large as indicated in the above. In a formal Bayesian framework, the subjective element should be formulated before the relevant data are observed. Arguments of objective symmetrical reasoning and physical constraints, of course, should be taken into account.

2.2.4. Practical Implications of the Different Interpretations of Probability

In some cases probabilities may adequately be assessed by means of *frequentistic information*. This is e.g. the case of the probability of failure of massively produced components, such as pumps, light bulbs and valves. However, in order to utilize reported failures for the assessment of the probability of failure for such components, it is a prerequisite that the components are in principle identical, that they have been subject to the same operational and/or loading conditions, and that the failures can be assumed to be independent.

In other cases when the considered components are e.g. bridges, high-rise buildings, ship structures or unique configurations of pipelines and pressure vessels, these conditions are not fulfilled. In these cases, the number of identical structures may be very small (or even just one) and the conditions in terms of operational and loading conditions are normally significantly different from structure to structure. In such cases the Bayesian interpretation of probability is far more appropriate.

The basic idea behind the Bayesian statistics is that lack of knowledge should be treated by probabilistic reasoning, similarly to other types of uncertainty. In reality, decisions have to be made despite the lack of knowledge and probabilistic tools are a great help in that process.

2.3. Uncertainties in Engineering Problems

For the purpose of discussing the phenomenon uncertainty in more detail, let us initially assume that the universe is deterministic and that our knowledge about the universe is perfect. This implies that it is possible to achieve perfect knowledge about any state, quantity or characteristic, which otherwise cannot be directly observed or has yet not taken place, by means of e.g. a set of exact equation systems and known boundary conditions by performing an analysis. Thus, the future, as well as the past, would be known or accessible with certainty. Considering the beam design optimization problem, it would thus be possible to assess the load bearing resistance of the beam and the exact magnitude of the maximum load which would occur in a given reference period. Therefore, an optimal decision can be achieved by a deterministic *cost benefit analysis*.

Whether the universe is deterministic or not was a rather deep philosophical and physical question that cumulated to a kind of showdown in the 1920's. Albert Einstein was, e.g. quoted (freely translated from [German](#)) as "(God) is not throwing dices" with the corresponding implication in 1924. Only 3 years later Werner von Heisenberg formulated his "[Uncertainty Principle](#)" in 1927 and since then the assumption of a fundamental random character of the world has become common knowledge.

³Thomas Bayes, reverend and mathematician, 1702-1761

In engineering decision analysis, subjected to uncertainties such as *Quantitative Risk Analysis* (QRA) and *Structural Reliability Analysis* (SRA), a commonly accepted view angle is that *uncertainties* should be interpreted and differentiated in regard to their type and origin. In this way, it has become standard to differentiate between uncertainties due to *inherent natural variability*, *model uncertainties* and *statistical uncertainties*. Whereas the first mentioned type of uncertainty is often denoted *aleatory* (or Type 1) uncertainty, the two latter are referred to as *epistemic* (or Type 2) uncertainties.

Considering again the beam example, it can be imagined that an *engineering model* might be formulated where future extreme load levels are predicted in terms of a regression of previously observed annual extremes. In this case, the uncertainty due to inherent natural variability would be the uncertainty associated with the annual extreme load level. The model chosen for the annual extreme load level events would introduce model uncertainties by itself. The parameters of the model would introduce statistical uncertainties as their estimation would be based on a limited number of observed annual extremes. Finally, the extrapolation of the annual extreme model, to extremes over longer periods of time, would introduce additional model uncertainties.

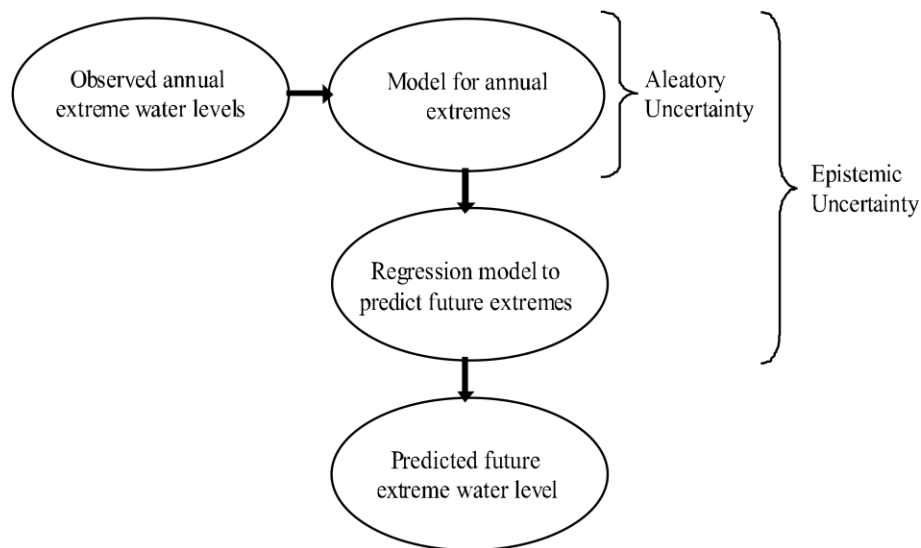


Figure 4.1: Illustration of uncertainty composition in a typical engineering problem.

The uncertainty associated with the future extreme load level is thus composed as illustrated in Figure 4.1. Whereas the so-called inherent natural variability is often understood as the (aleatoric) uncertainty caused by the fact that the universe is not deterministic. It may also be interpreted simply as the uncertainty that cannot be reduced by means of collection of additional information. It is seen that this definition implies that the amount of uncertainty due to inherent natural variability depends on the models applied in the formulation of the engineering problem. Presuming that a refinement of models corresponds to looking in more detailed at the problem at hand one could say that the uncertainty is scale dependent.

The probability of flooding within a given reference period can be assessed, having formulated a model for the prediction of future extreme load levels and the load bearing capacity of the beam, taking into account the various prevailing types of uncertainties.

It is interesting to notice that the type of uncertainty associated with the state of knowledge is time dependent. Following Figure 4.2, it is possible to observe an uncertain phenomenon when

it has occurred. In principle, if the observation is perfect without any errors, the knowledge about the phenomenon is perfect. However, the modelling of the same phenomenon in the future is uncertain, as this involves models subjected to natural variability, model uncertainty and statistical uncertainty. Often but not always, the available models tend to lose their precision rather fast, so that phenomena lying just a few days or weeks ahead can be predicted only with significant uncertainty. An extreme example of this concerns the prediction of the weather.

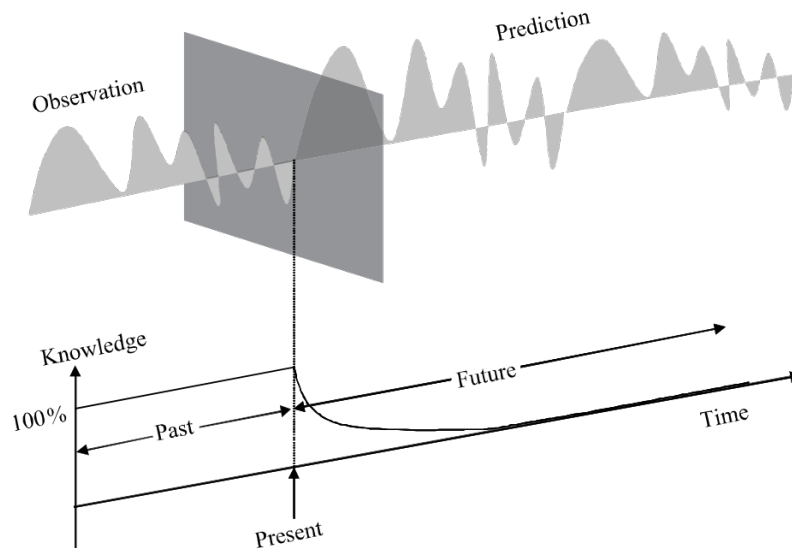


Figure 4.2: Illustration of the time dependence of knowledge.

The above discussion shows another interesting effect, namely that the uncertainty associated with a model concerning the future transforms, from a mixture of aleatory and epistemic uncertainty to a purely epistemic uncertainty, when the modelled phenomenon is observed. This transition of the type of uncertainty has a significant importance, because it facilitates that the uncertainty is reduced by utilization of observations, i.e. updating.

2.4. Random Variables

The performance of an engineering system, facility or installation (in the following referred to as system) may usually be modelled in mathematical physical terms, in conjunction with empirical relations. For a given set of model parameters, the performance of the considered system can be determined on the basis of this model. The basic *random variables* are defined as the parameters that carry the entire uncertain input to the considered model.

The basic random variables must be able to represent all types of uncertainties that are included in the analysis. The uncertainties that must be considered are, as previously mentioned, the physical uncertainty, the statistical uncertainty and the model uncertainty. The *physical uncertainties* are typically uncertainties associated with the loading environment, the geometry of the structure, the material properties and the repair qualities. The *statistical uncertainties* arise due to incomplete statistical information, e.g. due to a small number of materials tests. Finally, the model uncertainties must be considered to account for the uncertainty associated with the idealised mathematical descriptions used to approximate the actual physical behaviour of the structure.

Modern methods of reliability and risk analysis allow for a very general representation of these uncertainties, ranging from non-stationary stochastic processes and fields to time-invariant random variables, see e.g. Melchers (1987). In most cases, it is sufficient to model the uncertain quantities by random variables, with given cumulative distribution functions and distribution parameters estimated on the basis of statistical and/or subjective information. Therefore, the following is concerned with a basic description of the characteristics of random variables.

2.4.1. Cumulative Distribution and Probability Density Functions

A random variable, which can take on any value, is called a *continuous random variable*. The probability that such a random variable takes on a specific value is zero. The probability that a continuous random variable, X , is less than or equal to a value, x , is given by the *cumulative distribution function*:

$$F_X(x) = P(X \leq x) \quad (4.4)$$

Generally speaking, capital letters denote a random variable and lowercase denote an outcome or a realization of a random variable. Figure 4.3 illustrates an example of a continuous cumulative distribution function. For continuous random variables, the *probability density function* is given by:

$$f_X(x) = \frac{dF(x)}{dx} \quad (4.5)$$

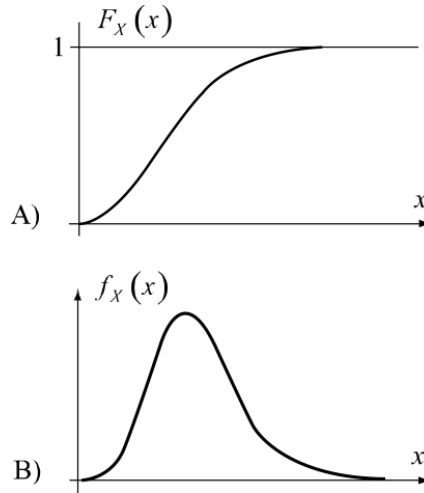


Figure 4.3: Illustration for a continuous random variable of A) a cumulative distribution function and B) a probability density function.

The probability of an outcome in the interval $[x; x + dx]$ where dx is small, is given by $P(X \in [x; x + dx]) = f_X(x) dx$.

Random variables with a finite or infinite countable sample space are called *discrete random variables*. For discrete random variables, the cumulative distribution function is given as:

$$P_X(x) = \sum_{x_i \leq x} p_X(x_i) \quad (4.6)$$

where $p_X(x_i)$ is the probability density function given as:

$$p_X(x_i) = P(X = x_i) \quad (4.7)$$

A discrete cumulative distribution function and probability density function is illustrated in Figure 4.4.

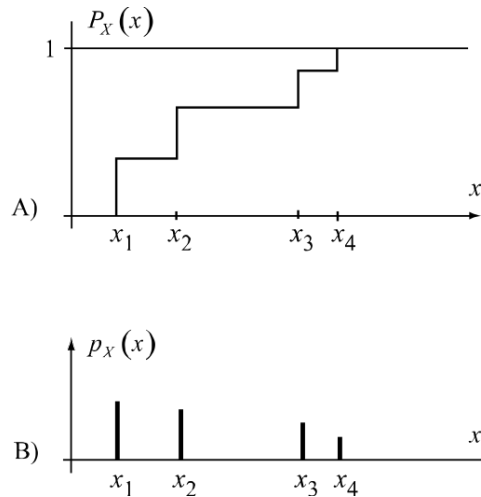


Figure 4.4: Illustration for a discrete random variable of A) a cumulative distribution function and B) a probability density function.

2.5. Load bearing behaviour of materials

In general, it is distinguished between stiffness related material properties, as the Modulus of Elasticity or the Shear Modulus, and the strength related material properties as ultimate strength or yield strength. Probabilistic models for strength and stiffness related properties should take into account the cause of variation and also the mechanical mechanism that leads to the corresponding material behaviour.

2.6. Scales of (variability) modelling

Material properties vary locally in space and, possibly, in time. As far as the spacial variations are concerned, it is useful to distinguish between three hierarchical levels of variation: macro (global), meso (local) and micro, cf. Table 4.1.

Table 4.1: Hierarchical levels of variation.

Scale	Population	Reference name
macro (global)	set of structures	gross supply
meso	set of elements	lot
meso (local)	one element	unit
micro	aggregate level	reference volume

A so-called hierarchical model for the material strength can be formulated, based on the different scales of modelling. Accordingly, the variability of the material strength is subdivided. As an example, the steel yield capacity might be represented by the following model:

$$X = \exp(Y_1 + Y_2)$$

where Y_1 is representing the logarithm of the variable mean value of a lot of steel samples. Y_2 has a zero mean value and is representing the logarithm of the difference between the yield capacity of one specimen and the logarithm of the mean value. This hierarchical representation is convenient as it makes it easy to take into account the correlation of strength within one lot.

2.7. Brittle material

Some materials are brittle and rupture is initiated locally by the weakest material particle. For these material the spatial variability of material strength is of importance. Therefore, the model for ideal brittle materials is also called the weakest link theory, proposed by Weibull (1939). If the strength X of the individual particles are assumed to be mutually independent and identically distributed with distribution function $F_X(x)$, the distribution function of the strength R of a body with nv_{vol} identical loaded particles⁴ is:

$$F_R(x) = 1 - (1 - F_X(x))^{nv_{vol}} = 1 - \exp(nv_{vol} \ln(1 - F_X(x))) \quad (4.8)$$

The behaviour of $\ln(1 - F_X(x))$ for small arguments can be represented as:

$$\ln(1 - F_X(x)) \cong -c(x - x_0)^k; \quad x \geq x_0 \quad (4.9)$$

where x_0 is the smallest possible value of x , and c and k are positive constants. Thus, Eq. (4.8) can be written as:

$$F_R(x) = 1 - \exp\left(-cnv_{vol} \left((x - x_0)^k\right)\right) \quad (4.10)$$

An isotropic solid under a homogeneous stress distribution is considered. A reference volume $v_{vol,0}$, e.g. the volume of a standard test specimen, is introduced. Substituting cn by $1/(v_{vol,0}w^k)$ in Eq. (4.10) leads to:

$$F_R(x) = 1 - \exp\left(-\frac{v_{vol}}{v_{vol,0}} \left(\frac{x - x_0}{w}\right)^k\right) \quad (4.11)$$

which is the expression of a 3-parameter Weibull distribution. w is called the shape parameter and k the scale parameter. The expected value $E[R]$ and the variance $Var[R]$ of R can be given as:

$$E[R] = x_0 + w\Gamma\left(1 + \frac{1}{k}\right) \left(\frac{v_{vol}}{v_{vol,0}}\right)^{-1/k} \quad (4.12)$$

$$Var[R] = w^2 \left(\Gamma\left(1 + \frac{2}{k}\right) - \Gamma^2\left(1 + \frac{1}{k}\right) \right) \left(\frac{v_{vol}}{v_{vol,0}}\right)^{-2/k} \quad (4.13)$$

⁴ n particles of finite volume v_{vol} .

where $\Gamma(\cdot)$ is the gamma distribution.

The expected value function decreases with the considered volume v_{vol} . For $x_0 = 0$, the coefficient of variation is independent of v_{vol} . The size effect is observed for many brittle materials such as hardened steel, concrete and stiff clay. The coefficient of variation and the correlation of the strength of the reference volumes $v_{vol,0}$ determine the importance of the volume effect.

Assuming unique failure probability for two material bodies with different volumes $v_{vol,1}$ and $v_{vol,2}$ gives the following expression:

$$\begin{aligned} F_R(x_1, v_{vol,1}) &= 1 - \exp \left(-\frac{v_{vol,1}}{v_{vol,0}} \left(\frac{x_1 - x_0}{w} \right)^k \right) = \\ &= 1 - \exp \left(-\frac{v_{vol,2}}{v_{vol,0}} \left(\frac{x_2 - x_0}{w} \right)^k \right) = F_R(x_2, v_{vol,2}) \end{aligned} \quad (4.14)$$

which results in:

$$\frac{x_2 - x_0}{x_1 - x_0} = \left(\frac{V_1}{V_2} \right)^{1/k} \xrightarrow{x_0=0} \frac{x_2}{x_1} = \left(\frac{V_1}{V_2} \right)^{1/k} \quad (4.15)$$

Eq. (4.15) can directly be used to compare the load bearing capacity of different volumes, for loading modes resulting in constant stress fields and brittle failure modes.

If an in-homogeneous stress distribution can be considered by introducing the stress $s(\xi_1, \xi_2, \xi_3)$ as a product of a reference stress s , (e.g. the maximum stress in the body) and a non-dimensional function $h(\xi_1, \xi_2, \xi_3)$, the probability distribution function of the strength r of the body can be written as:

$$F_R(x) = 1 - \exp \left\{ -\frac{1}{v_{vol,0}} \int_{xh(\xi_1, \xi_2, \xi_3) > x_0} \left(\frac{xh(\xi_1, \xi_2, \xi_3) - x_0}{w} \right)^k dv_{vol} \right\} \quad (4.16)$$

Eq. (4.16) is utilized in Johnson (1953) to compare the mean strength of beams, for various support and loading conditions, with the standard deviation of the reference strength. A graphical representation of the result of these calculations is presented in Figure 4.5. The ordinate in the figure is the mean strength of the beam, normalized with the strength of a beam loaded by a constant bending moment. It is seen that the mean strength of beams exposed to other load configurations is larger. Also, it increases for increasing coefficient of variation of the material strength.

2.8. Probabilistic load modelling

In the following, the term load will be related to forces acting on structural components and systems. Nevertheless, the notions and concepts introduced will, to a large extent, be valid for other types of influences, such as temperature, aggressive chemicals and radiation "acting from the outside" of the engineering system of consideration.

Loads and/or load effects are uncertain due to:

1. Random variations in space and time
2. Model uncertainties

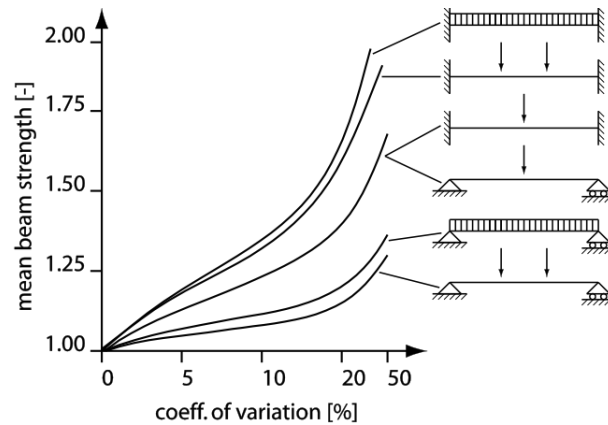


Figure 4.5: Comparison of mean strength of rectangular beams with various loading conditions and standard deviation of reference volume strength. From Johnson (1953).

3. Statistical uncertainties.

Whereas the model uncertainties associated with the physical model, used to represent the loads and/or load effects in the reliability analysis, may be represented by random variables, the loads themselves are usually time and space varying quantities. Therefore, they are best modelled by stochastic processes.

The following features of the loads may be characterized: their nature in regard to the variability of their magnitude in time; the variability in regard to their position in time; and their nature in regard to their effect on the engineering system. Knowing these characteristics is a prerequisite for the probabilistic modelling. It is often helpful to categorize loads according to the following descriptors:

1. permanent or variable
2. fixed or free
3. static or dynamic.

By way of example, consider the case where the reliability in regard to ultimate collapse of a reservoir structure is analyzed. The considered load acting on the structure is the hydrostatic pressure due to the water contained in the reservoir. As the hydrostatic pressure varies with the water level in the reservoir, which is dependent on the amount of rainwater flowing into the reservoir, the loading must be considered to be variable in time. Furthermore, due to the characteristics of the hydrostatic pressure, the loading is assumed fixed in space. Finally, as the reliability of the reservoir structure is considered in regard to an ultimate collapse failure mode and is dynamically insensitive, the loading may be idealized as acting static.

Having identified the characteristics of the considered load, the probabilistic modelling may proceed by:

- specifying the definition of the random variables used to represent the uncertainties in the loading;
- selecting a suitable distribution type to represent the random variable;

- assigning the distribution parameters of the selected distribution.

In practical applications, the first point is often more important. It requires a clear definition of the reliability problem to be considered. For the considered example, the random variable can e.g. only be concerned about the modelling of the random variations in time of the reservoir water level. The second point cannot be performed on the basis of data alone. It requires the simultaneous consideration of data, experience and physical understanding.

2.9. Time variable load

Many loads are variable in time and might be appropriately represented by random processes. However, if the problem at hand can be assumed time invariant, it is sufficient to represent the variable loads with extreme value distributions.

Statistical Assessment of Extreme Values

In risk and reliability assessments, *extreme values* (small and large) of random processes in a specified reference period are often of special interest. This is e.g. the case when considering the maximum sea water level, maximum wave heights, minimum ground water reservoir level, maximum wind pressures, strength of weakest link systems, maximum snow loads, etc.

For continuous time-varying loads, which can be described by a scalar, e.g. the water level or the wind pressure, one can define a number of related probability distributions. Often, the simplest option, namely the *arbitrary point in time* distribution, is considered.

If $x(t^*)$ is a realization of a single time-varying load at time t^* , then $F_X(x)$ is the arbitrary point in time cumulative distribution function of $X(t)$ defined by:

$$F_X = P(X(t^*) \leq x(t^*)) \quad (4.17)$$

In Figure 4.6, time series of measurements of the maximum wind speed are given per six months, per year and per five years. In addition to that, the histograms shows the sampling frequency distribution. These are the arbitrary point in time frequency distribution and the frequency distribution.

In the following, some results are given concerning the extreme events of trials of random variables and random processes, see also Madsen et al. (1985) and Benjamin et al. (1970). Asymptotic results are given taking basis in the tail behaviour of cumulative distribution functions, leading to the so-called *extreme value distributions*.

2.9.1. Extreme Value Distributions

When *extreme events* are of interest, the arbitrary point in time distribution of the load variable is not of immediate relevance. The distribution of the maximal values of the considered quantity over a given reference period is then regarded.

If the random process $X(t)$ may be assumed to be ergodic⁵, the distribution of the largest extreme in a *reference period* T $F_{X,T}^{\max}(x)$ can be thought of as being generated by sampling values of the maximal realization $x_{\max}$ from successive reference periods T .

⁵A process is denominated ergodic if its time average converges to its mean. Thus, its statistical properties can be inferred from a sufficiently long random sample.

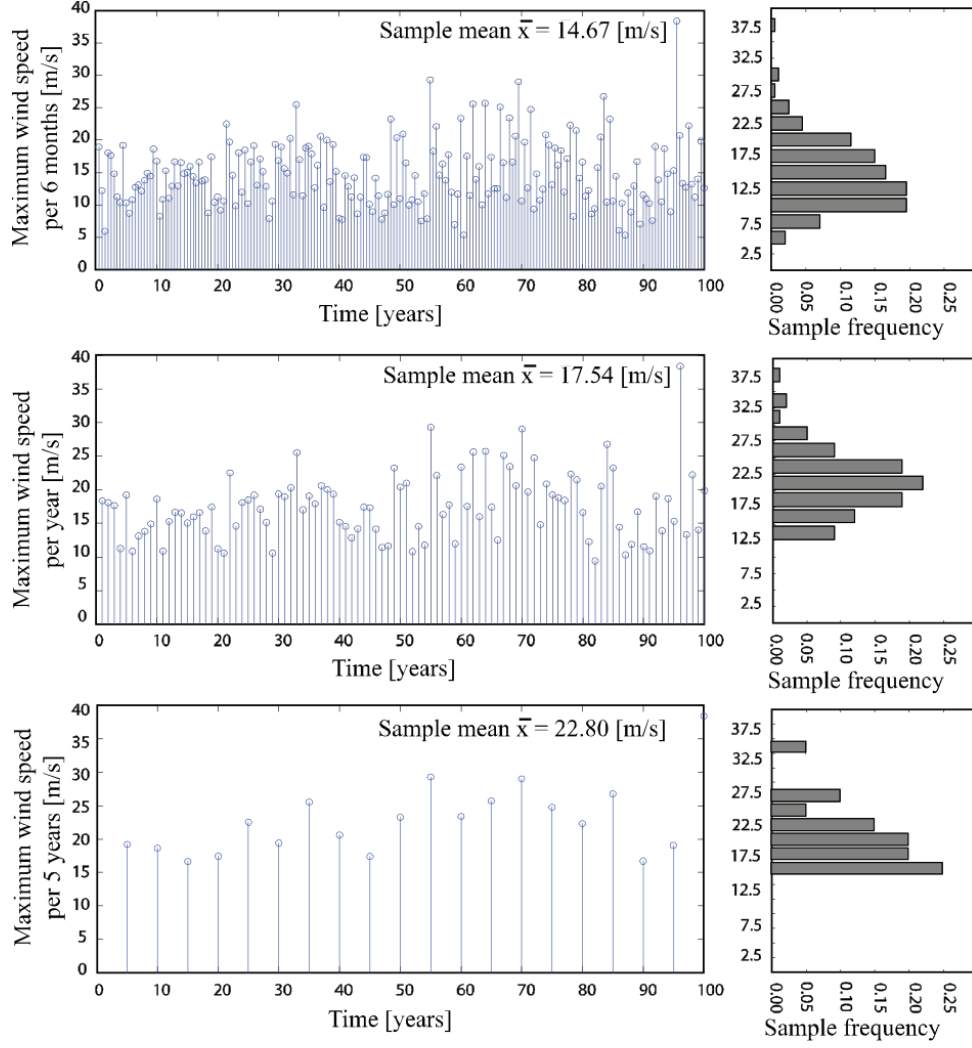


Figure 4.6: Time series and corresponding sample frequency histograms of recorded half yearly, annual and five year maximum observed wind velocities.

The cumulative distribution function of the largest extreme in a period of nT , $F_{X,nT}^{\max}(x)$ may be determined from the cumulative distribution function of the largest extreme in the period T , $F_{X,T}^{\max}(x)$, by:

$$F_{X,nT}^{\max}(x) = F_{X,T}^{\max}(x)^n \quad (4.18)$$

Eq. (4.18) follows from the multiplication law for independent events. The corresponding probability density function may be established by its differentiation, yielding:

$$f_{X,nT}^{\max}(x) = n F_{X,T}^{\max}(x)^{n-1} f_{X,T}^{\max}(x) \quad (4.19)$$

In Figure 4.7, the case of a Normal distribution with mean value equal to 10 and standard deviation equal to 3 is illustrated for increasing n .

Similarly to the derivation of Eq. (4.18), the cumulative distribution function for the extreme minimum value in a considered reference period T , $F_{X,nT}^{\min}(x)$ may be found as:

$$F_{X,nT}^{\min}(x) = 1 - (1 - F_{X,T}^{\min}(x))^n \quad (4.20)$$

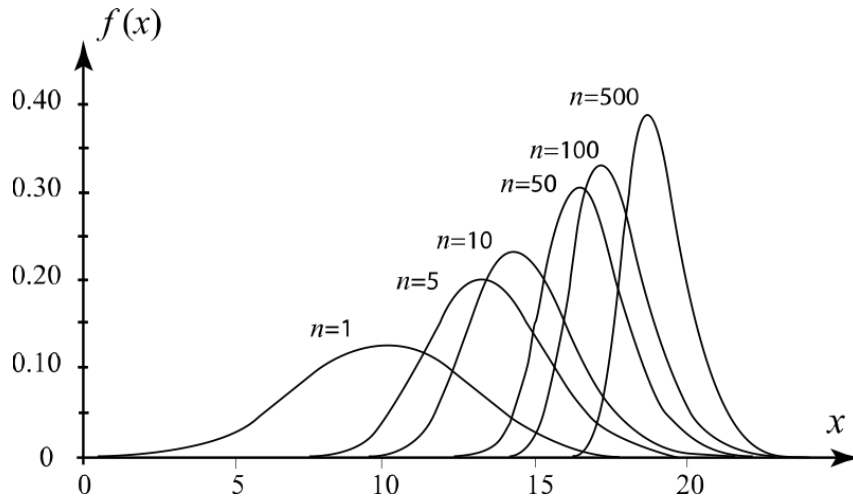


Figure 4.7: Normal extreme value probability density functions.

Subject to the assumption that the considered process is ergodic, it can be shown that the cumulative function for an extreme event $F_{X,nT}^{\max}(x)$ converges asymptotically (as the reference period nT increases) to one of three types of *extreme value distributions*, i.e. type I, type II or type III. To which type the distribution converges depends only on the tail behaviour (upper or lower) of the considered random variable generating the extremes, i.e. $F_{X,T}^{\max}(x)$.

In the following, the three types of extreme value distributions are introduced and the conditions under which they may be assumed are discussed. In Table 4.2, the definition of the extreme value probability distributions and their parameters and moments is summarized.

Table 4.2: Probability distributions, Schneider (1994).

Distribution type	Param.	Mean and Standard deviation
Extreme Type I <i>Gumbel max</i> $-\infty \leq x \leq \infty$ $f_X(x) =$ $\alpha \exp(-\alpha(x - u) - \exp(-\alpha(x - u)))$ $F_X(x) = \exp(-\exp(-\alpha(x - u)))$	u $\alpha > 0$	$\mu = u + \frac{0.577216}{\alpha}$ $\sigma = \frac{\pi}{\alpha\sqrt{6}}$
Extreme Type I <i>Gumbel min</i> $-\infty \leq x \leq \infty$ $f_X(x) = \alpha \exp(\alpha(x - u) - \exp(\alpha(x - u)))$ $F_X(x) = 1 - \exp(-\exp(\alpha(x - u)))$	u $\alpha > 0$	$\mu = u - \frac{0.577216}{\alpha}$ $\sigma = \frac{\pi}{\alpha\sqrt{6}}$

Extreme Type II <i>Frechet max</i> $\varepsilon \leq x \leq \infty, \quad u, k > 0$ $f_X(x)$ $\frac{k}{u - \varepsilon} \left(\frac{u - \varepsilon}{x - \varepsilon} \right)^{k+1} \exp \left(- \left(\frac{u - \varepsilon}{x - \varepsilon} \right)^k \right)$ $F_X(x) = \exp \left(- \left(\frac{u - \varepsilon}{x - \varepsilon} \right)^k \right)$	$= \begin{array}{c} u > 0 \\ k > 0 \\ \varepsilon \end{array}$	$\begin{array}{l} \mu = \varepsilon + (u - \varepsilon) \Gamma \left(1 - \frac{1}{k} \right), \quad k > 1 \\ \sigma = (u - \varepsilon) \sqrt{\Gamma \left(1 - \frac{2}{k} \right) - \Gamma^2 \left(1 - \frac{1}{k} \right)}, \quad k > 2 \end{array}$
Extreme Type III <i>Weibull min</i> $\varepsilon \leq x \leq \infty, \quad u, k > 0$ $f_X(x)$ $\frac{k}{u - \varepsilon} \left(\frac{x - \varepsilon}{u - \varepsilon} \right)^{k-1} \exp \left(- \left(\frac{x - \varepsilon}{u - \varepsilon} \right)^k \right)$ $F_X(x) = 1 - \exp \left(- \left(\frac{x - \varepsilon}{u - \varepsilon} \right)^k \right)$	$= \begin{array}{c} u > 0 \\ k > 0 \\ \varepsilon \end{array}$	$\begin{array}{l} \mu = \varepsilon + (u - \varepsilon) \Gamma \left(1 + \frac{1}{k} \right) \\ \sigma = (u - \varepsilon) \sqrt{\Gamma \left(1 + \frac{2}{k} \right) - \Gamma^2 \left(1 + \frac{1}{k} \right)} \\ k \approx \left(\frac{\sigma}{\mu} \right)^{-1.09} \end{array}$

Type I Extreme Maximum Value Distribution - Gumbel max

For upwards unbounded distribution functions $F_X(x)$, where the upper tail falls off in an exponential manner, the cumulative distribution of extremes in the reference period T $F_{X,T}^{\max}(x)$ has the form shown in Eq. (4.21). This is the case for the exponential function, the Normal distribution and the Gamma distribution.

$$F_{X,T}^{\max}(x) = \exp(-\exp(-\alpha(x - u))) \quad (4.21)$$

with corresponding probability density function:

$$f_{X,T}^{\max}(x) = \alpha \exp(-\alpha(x - u) - \exp(-\alpha(x - u))) \quad (4.22)$$

This is also denoted the *Gumbel distribution* for extreme maxima. The mean value and the standard deviation of the Gumbel distribution may be related to the parameters u and α as:

$$\begin{aligned} \mu_{X_T}^{\max} &= u + \frac{\gamma}{\alpha} = u + \frac{0.577216}{\alpha} \\ \sigma_{X_T}^{\max} &= \frac{1}{\alpha\sqrt{6}} \end{aligned} \quad (4.23)$$

where γ is the Euler's constant.

The Gumbel distribution has the useful property that the standard deviation is independent on the considered reference period, i.e. $\sigma_{X_{nT}}^{\max} = \sigma_{X_T}^{\max}$, and that the mean value $\mu_{X_{nT}}^{\max}$ depends on n in the following simple way:

$$\mu_{X_{nT}}^{\max} = \mu_{X_T}^{\max} + \frac{\sqrt{6}}{\pi} \sigma_{X_T}^{\max} \ln(n) \quad (4.24)$$

Finally, manipulating Eq. (4.21) by utilizing a Taylor expansion to the first order of $\ln(p)$ in

$p = 1$, it can be shown that the characteristic value x_c , corresponding to an annual exceedance probability of p and corresponding return period $T_R = 1/p$ can be written as:

$$x_c \approx u + \frac{1}{\alpha} \ln(T_R) \quad (4.25)$$

This shows that the characteristic value, i.e. a typical engineering decision parameter, increases with the logarithm of the considered return period. This expression is valid for a Gumbel max distribution for large return periods.

Type I Extreme Minimum Value Distribution - Gumbel min

In case the cumulative distribution function $F_X(x)$ is downwards unbounded and the lower tail falls off in an exponential manner, symmetry considerations leads to a cumulative distribution function for the extreme minimum $F_{X,T}^{\min}(x)$ within the reference period T of the following form:

$$F_{X,T}^{\min}(x) = 1 - \exp(-\exp(\alpha(x - u))) \quad (4.26)$$

with corresponding probability density function:

$$f_{X,T}^{\min}(x) = \alpha \exp(\alpha(x - u) - \exp(\alpha(x - u))) \quad (4.27)$$

This is also denoted the *Gumbel distribution* for extreme minima. The mean value and the variance of the Gumbel distribution can be related to the parameters u and α as:

$$\begin{aligned} \mu_{X_T^{\min}} &= u - \frac{\gamma}{\alpha} = u - \frac{0.577216}{\alpha} \\ \sigma_{X_T^{\min}} &= \frac{\pi}{\alpha\sqrt{6}} \end{aligned} \quad (4.28)$$

Type II Extreme Maximum Value Distribution - Frechet max

For cumulative distribution functions downwards limited at zero and upwards unlimited with a tail falling off in the form:

$$F_X(x) = 1 - \beta \left(\frac{1}{x} \right)^k \quad (4.29)$$

the cumulative distribution function of extreme maxima in the reference period T $F_{X,T}^{\max}(x)$ has the following form:

$$F_{X,T}^{\max}(x) = \exp\left(-\left(\frac{u}{x}\right)^k\right) \quad (4.30)$$

with corresponding probability density function:

$$f_{X,T}^{\max}(x) = \frac{k}{u} \left(\frac{u}{x} \right)^{k+1} \exp\left(-\left(\frac{u}{x}\right)^k\right) \quad (4.31)$$

which is also denoted the *Frechet distribution* for extreme maxima. The mean value and the

variance of the Frechet distribution can be related to the parameters u and k as:

$$\begin{aligned}\mu_{X_T^{\max}} &= u\Gamma\left(1 - \frac{1}{k}\right) \\ \sigma_{X_T^{\max}}^2 &= u^2 \left[\Gamma\left(1 - \frac{2}{k}\right) - \Gamma^2\left(1 - \frac{1}{k}\right) \right]\end{aligned}\quad (4.32)$$

where it is noticed that the mean value only exists for $k > 1$ and the standard deviation only exist for $k > 2$. In general, it can be shown that the i^{th} moment of the Frechet distribution exists only when $k > i$.

Type III Extreme Minimum Value Distribution - Weibull min

In the case where the cumulative distribution function $F_X(x)$ is downwards limited at ε and the lower tail falls of towards ε in the form:

$$F(x) = c(x - \varepsilon)^k \quad (4.33)$$

leads to a cumulative distribution function for the extreme minimum $F_{X,T}^{\min}(x)$ within the reference period T of the following form:

$$F_{X,T}^{\min}(x) = 1 - \exp\left(-\left(\frac{x - \varepsilon}{u - \varepsilon}\right)^k\right) \quad (4.34)$$

with corresponding probability density function:

$$f_{X,T}^{\min}(x) = \frac{k}{u - \varepsilon} \left(\frac{x - \varepsilon}{u - \varepsilon}\right)^{k-1} \exp\left(-\left(\frac{x - \varepsilon}{u - \varepsilon}\right)^k\right) \quad (4.35)$$

which is also denoted the *Weibull distribution* for extreme minima. The mean value and the variance of the Weibull distribution can be related to the parameters u , k and ε as:

$$\begin{aligned}\mu_{X_T^{\min}} &= \varepsilon + (u - \varepsilon)\Gamma\left(1 + \frac{1}{k}\right) \\ \sigma_{X_T^{\min}}^2 &= (u - \varepsilon)^2 \left[\Gamma\left(1 + \frac{2}{k}\right) - \Gamma^2\left(1 + \frac{1}{k}\right) \right]\end{aligned}\quad (4.36)$$

Return Period for Extreme Events

The *return period* T_R for an extreme event may be defined by:

$$T_R = nT = \frac{1}{(1 - F_{X,T}^{\max}(x))}T \quad (4.37)$$

where T is the reference period for the cumulative distribution function of the extreme events $F_{X,T}^{\max}(x)$. If as an example the annual probability of an extreme load event is 0.02, the return period for this load event is 50 years.

2.10. Load combination

Load variables and other actions are typically time-varying quantities and are usually modelled as stochastic processes $\{Q(t)\}$.

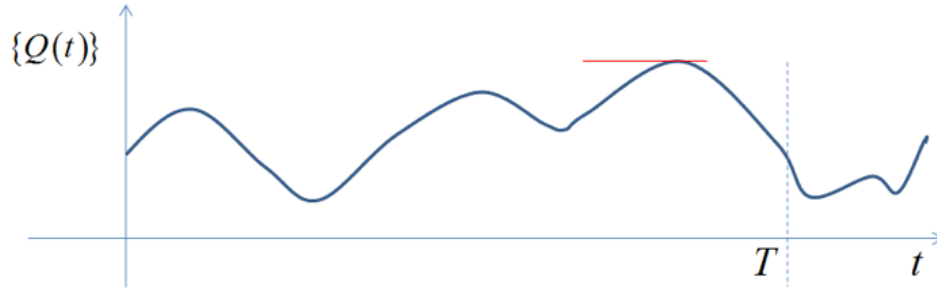


Figure 4.8: Stochastic process representing the time-varying load Q and maximum value in period T .

When the process is continuous, the probability distribution representing the maximum value in a reference period T is very closely approximated by one of the asymptotic extreme value distributions. The most commonly applied in structural reliability is the Gumbel distribution. In such a way, the load is modelled by a stochastic variable Q_T^{max} instead of modelling a stochastic process. The random variable Q_T^{max} is independent on point in time t but strictly related to the reference period T , which is usually one year.

The above simplification cannot be applied when more than one time-varying load act simultaneously on the structure then. The reason is that the determination of the distribution of the combined load effect requires knowledge of the detailed variation in time of the individual loading processes. This is clear in Figure 4.9, where the maxima of $\{Q_1(t)\}$, $\{Q_2(t)\}$ and $\{Q_1(t) + Q_2(t)\}$ in the reference period T do not appear at the same time. Moreover, the maxima of the sum-process is smaller than the sum of the maxima of Q_1 and Q_2 . It is clear that detailed knowledge of the variation in time of the two processes is required and that the knowledge of the maxima of the individual processes gives insufficient information to evaluate the combined effect.

The characteristic of the sum-process can be obtained from theory of stochastic processes. We will focus on a simpler solution presented in the following.

2.10.1. Ferry-Borges and Castanheta load model

A simple but rather accurate method for combining loads is the so-called Ferry-Borges and Castanheta load model (FBC). In this model real processes are largely simplified in such a way that the mathematical problem connected with estimating the distribution function of the maxima is greatly simplified.

For each load process $\{Q_i(t)\}$ it is assumed that the load changes after equal so-called *elementary intervals* of time τ_i , see Figure 4.10. The reference period T (e.g. 1 year) is divided into n_i time intervals of equal duration $\tau_i = T/n_i$ (n_i is called *repetition number*). Moreover, the load is constant in each elementary interval. In other words, the process is modelled as a *rectangular pulse* process. The loads in each elementary interval are mutually independent and

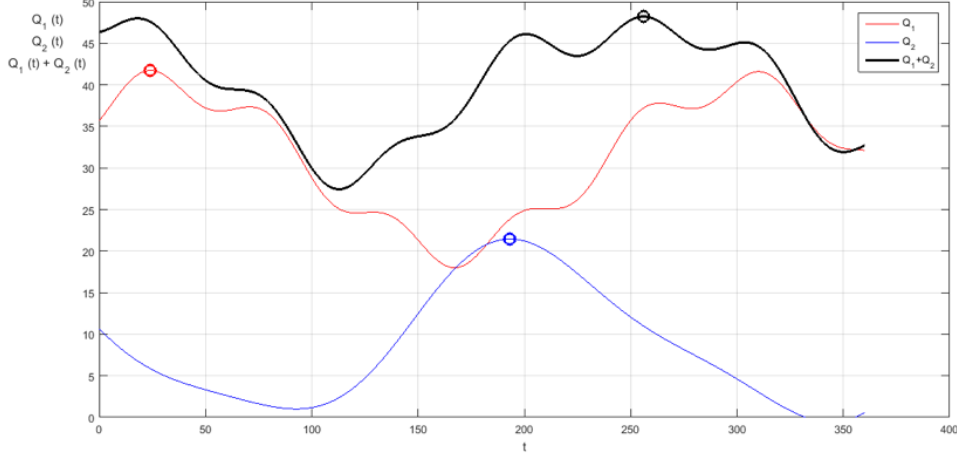


Figure 4.9: Stochastic processes Q_1 and Q_2 and their sum.

identically distributed random variables, the distribution function F_{Q_i} is called *point-in-time* distribution.

The distribution representing the maxima in the reference period T is then $F_{Q_{i,T}^{max}} = (F_{Q_i})^{1/n_i}$.

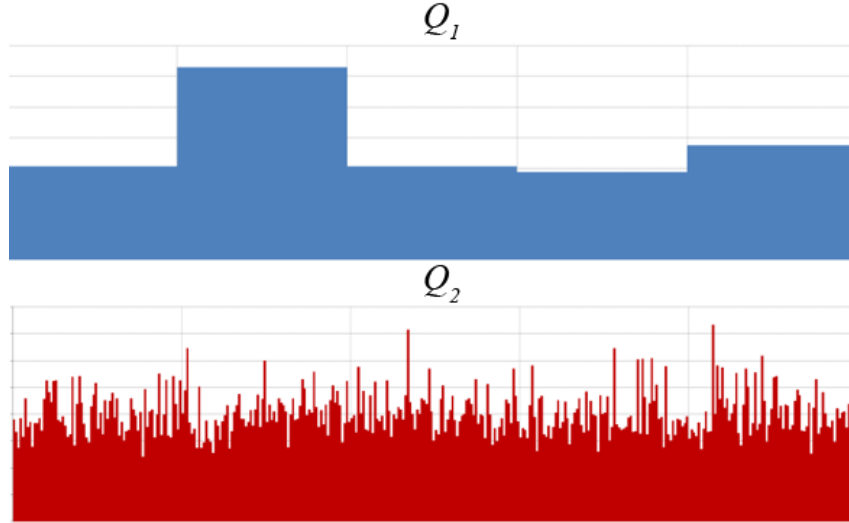


Figure 4.10: FBC processes for Q_1 ($n_1 = 5$) and Q_2 ($n_2 = 360$) in a reference period $T = 360$.

When two or more load processes $\{Q_1(t)\}$, $\{Q_2(t)\}$, ..., $\{Q_r(t)\}$ are considered, it is assumed in the FBC that they are mutually independent with integer repetition numbers $n_1 \leq n_2 \leq \dots \leq n_r$ and the ratios n_i/n_{i-1} to be positive integer numbers for all i .

FBC for two loads

The combination rule for two variable loads is then:

- Load combination 1: Q_1 leading + Q_2 accompanying:
 $\max\{n_1 \text{ reps. of } Q_1\} + \max\{n_2/n_1 \text{ reps. of } Q_2\}$
- Load combination 2: Q_2 leading and Q_1 accompanying:
 $\max\{n_2 \text{ reps. of } Q_2\} + \max\{1 \text{ rep. of } Q_1\}$

The procedure is illustrated in Figure 4.11 and Figure 4.12.

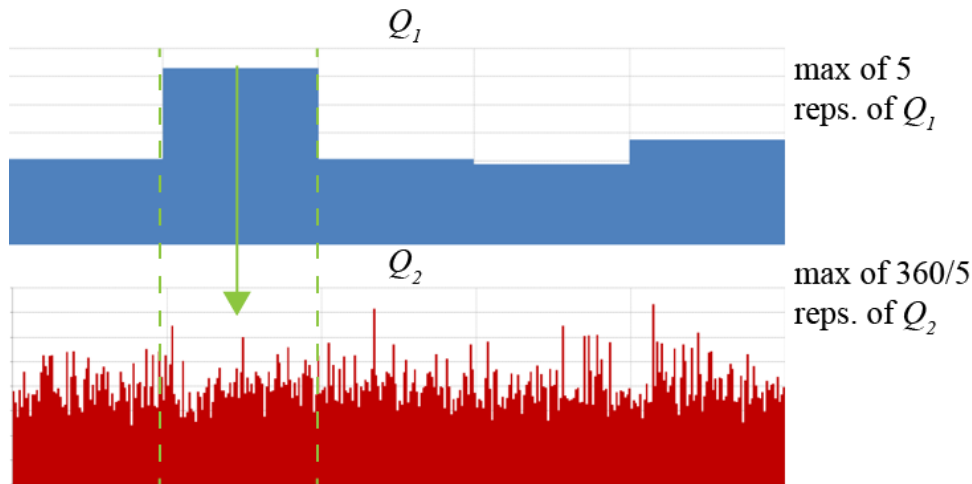


Figure 4.11: LC1 for Q_1 ($n_1 = 5$) and Q_2 ($n_2 = 360$).

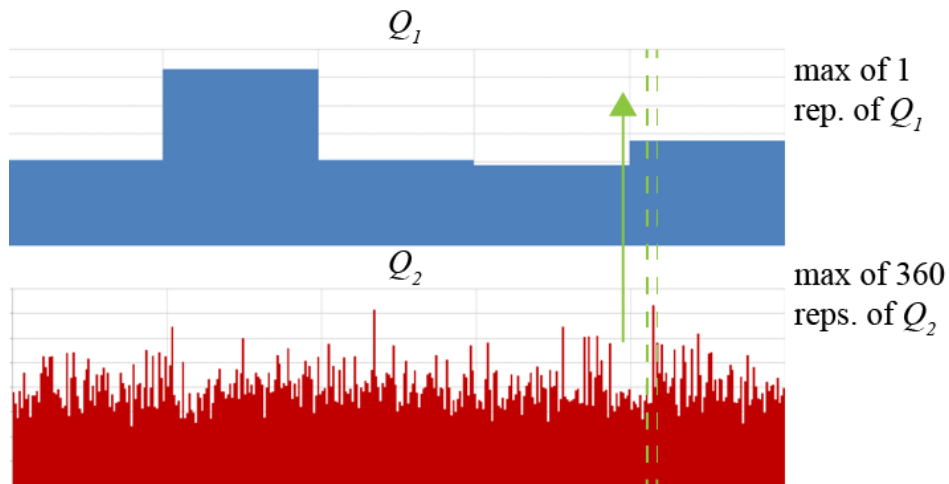


Figure 4.12: LC2 for Q_1 ($n_1 = 5$) and Q_2 ($n_2 = 360$).

FBC for three loads

The combination rule for three variable loads is then:

- Load combination 1: Q_1 leading + Q_2, Q_3 accompanying:
 $\max\{n_1 \text{ reps. of } Q_1\} + \max\{n_2/n_1 \text{ reps. of } Q_2\} + \max\{n_3/n_1 \text{ reps. of } Q_3\}$
- Load combination 2: Q_2 leading and Q_1, Q_3 accompanying:
 $\max\{n_2 \text{ reps. of } Q_2\} + \max\{1 \text{ rep. of } Q_1\} + \max\{n_3/n_2 \text{ reps. of } Q_3\}$
- Load combination 3: Q_1 leading + Q_3, Q_2 accompanying:
 $\max\{n_1 \text{ reps. of } Q_1\} + \max\{n_3/n_1 \text{ reps. of } Q_3\} + \max\{1 \text{ rep. of } Q_2\}$
- Load combination 4: Q_3 leading and Q_1, Q_2 accompanying:
 $\max\{n_3 \text{ reps. of } Q_3\} + \max\{1 \text{ rep. of } Q_1\} + \max\{1 \text{ rep. of } Q_2\}$

FBC for r loads

In general, for r loads there are 2^{r-1} load combinations to be considered.

3. Practical exercises

P4.1. Load combination

Consider the simply supported beam in Figure 4.13, with span $l = 5000$ mm, rectangular cross-section with area $A = bh$, with $b = 200$ mm and $h = 300$ mm. The material strength is designated by F_m . The beam is subject to two time-variant loads Q_1 and Q_2 , which are modelled as a FBC process. The distributions representing the material strength and the load **yearly maxima** are given in the Table 4.3. The rectangular pulse process is illustrated in Figure 4.10.

Table 4.3: Basic random variables, reference period $T = 1$ yr.

	mean	C.o.V.	n_i
F_m (Normal)	20 MPa	0.1	/
$Q_{1,T}^{max}$ (Gumbel)	2 N/mm	0.2	5
$Q_{2,T}^{max}$ (Gumbel)	4 N/mm	0.4	360

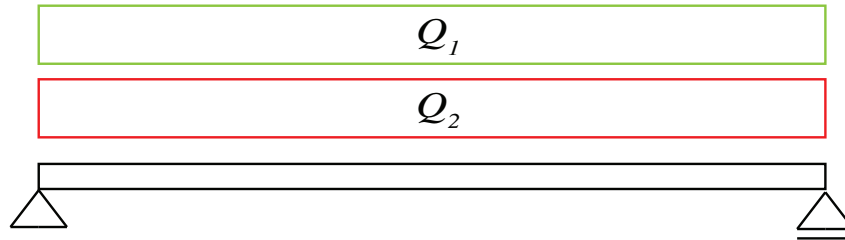


Figure 4.13: Simply supported beam with two time-variant loads.

P4.1.1. LC formulation

Formulate the load combinations to be considered and their associated limit state function.

P4.1.2. Computation of failure probability

Compute the yearly probability of failure of the beam using FORM. Note that the limit state function is highly non-linear and a numerical solver need to be used. The approximated method used in the course is not applicable since it is not possible to collect beta.

P4.1.3. Additional exercises

1. Re-derive the load combination for 2 and 3 variable loads according to FBC method.
2. Derive the combination rule for 4 variable loads.

3. Evaluate the reliability index of the beam presented in the exercise considering yearly maxima of Q_1 and Q_2 acting simultaneously. Do you expect a lower or larger reliability index? Why?
4. Derive the mean and standard deviation of the Gumbel distribution representing the maxima of Q_1 and Q_2 over a reference period of 5, 10 and 100 years.
5. Evaluate the reliability index of the beam for a 50-year period assuming that failure among years is independent. Hint: the structure fails in the period $T = 50$ yr if it fails in any one year. This is a series system with 50 components, each of them has a failure probability of $P_{f,comp_i} = P_{f,1yr} = \Phi(-\beta_{1yr})$, then we have that $P_{f,50yr} = P_{f,system}$.

Question::

Probability in Engineering

- What different interpretations of probability do you know? Which one is most suited for engineers and why?
- Mention some examples for epistemic and aleatory uncertainties. Why is the differentiation relevant for engineering decision making?
- Mention 3 different probability distributions and explain which relevant phenomena can be represented by them and why.

Question::

Extreme Value Distribution

- What phenomena can be represented with an extreme value distribution?
- How can an extreme value distribution be derived?
- How are scale or time effects represented by extreme value distributions? Make an example.
- Explain a method for representing the combined effect of two independent variable loads.

Chapter 5

Structural Design Standards

Structural design standards are used for the daily design of structures. They comprise of simple rules that represent the current best practice and are generally applied to regular structures. Structural design standards do not account for risk, reliability and uncertainty in an explicit manner. However, the safety concept that is implemented is taking reference to uncertainty and reliability implicitly in terms of partial safety factors and characteristic values.

In this chapter, it is demonstrated that structural design standards correspond to the societal preference for safety and cost effectiveness. A direct correspondence between reliability and the choice of partial safety factors is developed for particular cases. However, it is also shown that the set up of structural design codes is merely a problem of proper generalisation and simplification.

1. Concepts

1.1. Levels of design

The design of a structure is in principle a decision problem with the objective to identify a structural design with a performance that maximizes the expected utility (or equivalently minimizes the expected cost) for the corresponding decision maker, see the first chapter of this compendium and Rackwitz (2000). A common attribute of these decisions is that they have to be performed subject to uncertainties. In general, different levels of detail for the assessment of the structural performance are distinguished (ISO 2015); risk informed decision making, reliability based design, and semi-probabilistic design (see overview in Figure 5.1).

The three different approaches are often associated to numbered levels: The semi-probabilistic approach corresponds to Level 1, the reliability based approach to Level 2 or 3 (dependent on the level of detail in uncertainty representation), and the risk informed approach corresponds to Level 4. The approaches on the different levels closely correspond to each other. The ability to account for the particular conditions of specific design situations and therefore identify a more optimal design solution is also increasing with increasing level. That is the reason why higher level approaches should be used to verify or calibrate lower level approaches. However, the complexity and difficulty is increasing with increasing level and from a regulative perspective, i.e. where a unified set of rules and assumptions for broad application shall be standardized, low level approaches are advantageous since their inherent low level of detail goes along with the allowance for broad generalization.

	Commonly applied when:
Risk-informed decision making: - decisions are taken with due consideration of the total risks (e.g. loss of lives, injuries, environmental and monetary losses.	Exceptional design situations in regard to uncertainties and consequences. Derivation of reliability requirements.
Reliability-based design and assessment: - reliability requirement to fulfil.	Unusual design situations in regard to uncertainties. Code Calibration.
Semi-probabilistic approach: - safety format prescribing the design equations and the analysis procedures to be used.	Usual design situations in regard to consequences and uncertainties. Default method of most design codes.

Figure 5.1: Different levels of detail for the assessment of the structural performance as arranged in ISO (2015)

1.2. Risk informed decision making

Risk informed decision making allows for the highest level of detail and follows broadly the scheme of Bayesian Decision Analysis as outlined in the introduction to the course, see also JCSS (2008). In the structural engineering context, decision alternatives may include structural measures, i.e. the dimension of cross sections, the choice of material grades, the design of a superstructure, repair and strengthening etc., and non-structural measures as e.g. the implementation of quality control and checking, inspection or the installation of structural health monitoring. Simplification or enhancement in the representation of the decision problem by models could also be considered as a non-structural measure. It is assumed that executive decision making is not mechanically following the results of the formal decision analysis but that the well documented assessment informs the decision maker, such that he or she is able to identify his or her final decision.

Risk informed decision making is very flexible and can be applied for systems of different scale in space and time. Depending on the definition of the system boundaries, structural performance attributes as e.g. sustainability, resilience, robustness and reliability can be addressed within an analysis (Faber et al. 2018). The risk informed decision framework can also be applied for the calibration of lower level decision making methods.

Guidance and standardisation for risk based decision making can be found in ISO (2015). However, the flexibility and generality of the method requires a large amount of expertise and experience from the person or group of persons elaborating on a risk based decision analysis.

1.3. Reliability based design

In reliability based design, the design decision is chosen such that it complies to a predefined reliability requirement. The reliability requirement is defined based on past experience, i.e. specified as the inherent reliability of traditionally accepted design solutions (Baravalle et al. 2017),

or it is based on formal calibration using risk informed methods. Following the latter, it may be differentiated between different types of decisions on structures in regard to consequences and the relative cost of implementing the decision as discussed in the later sections of this paper. In reliability based design, it is possible to assess the effect of a decision on the structural performance on a component level, i.e. only one possible failure mode is considered, or on a system level, i.e. the interaction of different failure modes in a structure are considered. However, as consequences are not represented explicitly in a reliability analysis, subordinate structural performance attributes as robustness, sustainability and resilience can not be considered explicitly by reliability based design.

Reliability based design is addressed in the international standard ISO (2015), more detailed guidance on structural reliability methods and uncertainty representation in regard to models, structural resistance and structural demands is found in the Probabilistic Model Code of the Joint Committee on Structural Safety, (JCSS 2001).

1.4. Semi-probabilistic design

The semi-probabilistic approach corresponds to the lowest level of detail. Here, a design decision is chosen such that it complies with the criterion that a design value of a resistance is larger than a design value of a corresponding load effect. Design values for the load bearing capacity r_d are chosen to have a sufficiently low non-exceedance probability and design values for loads e_d are chosen to have a sufficiently low exceedance probability such that the design criterion in the limit ($r_d = e_d$) corresponds to the required level of reliability.

In the so-called load and resistance factor design (LRFD) format (Ravindra et al. 1978) design values are estimated based on characteristic values and partial safety factors γ as e.g. $r_d = r_k / \gamma_M$ for resistance variables and $e_d = \gamma_E \cdot e_k$ for the effects of applied loads. Both, the definition of the characteristic value and the choice of partial safety factor, is made in order to meet the reliability requirements. However, the correspondence to reliability requirements is generally made for domains of structures for that generalised assumptions in regard to consequences and uncertainties are made. With semi-probabilistic design, structural performance on a component/failure mode level can be assessed. The explicit consideration of the interaction of failure modes in a structure, i.e. system effects, is not accommodated.

The principles of semi-probabilistic design are outlined in ISO (2015). It is the method of choice for most structural design decision problems and executive guidance and standardisation is found in several national and international design standards as, e.g. the Eurocodes (CEN 2002).

1.5. Relevance and correspondence

The different levels of engineering decision making are all relevant and support the super-ordinate objective of the safe and optimal development and maintenance of structures in the build environment. And the approaches on the different levels closely correspond to each other. The ability to account for the particular conditions of specific design situations and therefore identify a more optimal design solution is increasing with increasing level. That is the reason why higher level approaches are used to verify or calibrate lower level approaches. However, the com-

plexity and difficulty is increasing with increasing level and from a regulative perspective, i.e. where a unified set of rules and assumptions for broad application shall be standardised, semi-probabilistic approaches are advantageous since their inherent low level of detail goes along with the allowance for broad generalisation. The design, calibration and layout of simplified decision making approaches, has to be based, directly or indirectly, on risk informed decision making, as this is the only level of detail that allows for the explicit consideration of the superordinate objectives in engineering decision making.

2. EUROCODE semi-probabilistic design

2.1. General framework and design values

In the Eurocodes a semi-probabilistic design method is introduced through partial factor design that is generally applied as default design method for common design situations. The partial factor design format is formulated such that it facilitates the identification of acceptable and feasible design solutions concerning new structures. The partial factor design format comprises:

- consequence class categorization (see EN1990 Annex B);
- design situations (see EN1990 and EN1991-EN1999);
- design equations (see EN1990 and EN1991-EN1999);
- design values (as determined according EN1990, section 6).

Design equations can in general be formulated for failure modes involving in principle both failure of individual cross sections of the structures, as well as for failure modes involving the failures of several cross sections of the structures. The principle form of the design equation is given as (compare EN1990, equation 6.1):

$$r_d - e_d \geq 0 \quad (5.1)$$

where the design values for the load bearing capacity r_d and the action effect e_d are obtained from:

$$r_d = r_d(\mathbf{x}_d; \mathbf{a}_d; \boldsymbol{\theta}_d) \quad (5.2)$$

$$e_d = e_d(\mathbf{f}_d; \mathbf{a}_d; \boldsymbol{\theta}_d) \quad (5.3)$$

where:

$\mathbf{x}_d$ is a vector of design values of material properties;

$\mathbf{f}_d$ is a vector of design values of actions;

$\mathbf{a}_d$ is a vector of design values of geometrical properties;

$\boldsymbol{\theta}_d$ is a vector of design values of model uncertainties.

2.2. Partial factors

The design value of a basic variable Y is defined as the multiplication or division of a corresponding partial safety factor γ_Y and the characteristic value y_k . Accordingly, Eq. (5.1) can be rewritten in an extended form as:

$$\frac{r_k}{\gamma_R} = r_d \geq e_d = \gamma_E e_k \quad (5.4)$$

A characteristic value y_k is taken as a specified p – fractile value from the statistical distribution $F_Y(y)$ that is chosen to represent the basic variable, as:

$$y_k = F_Y^{-1}(p) \quad (5.5)$$

where $F_Y(y)$ is the cumulative probability distribution function of the basic variable. Note: Typical values for p are:

- resistance related variables: $p = 0.05$;
- permanent actions: $p = 0.5$;
- time-variable actions: $p = 0.98$.

The partial safety factors for the various actions and materials characteristics entering the design equations are determined such as to account for the uncertainties associated with the loads and resistances which are of relevance for the given design situation in such a way that design solutions that fulfill the inequality in Eq. (5.4) are consistent with the reliability requirements of the Eurocodes.

2.3. Reliability requirements in the Eurocodes

A basic requirement stated in the Eurocode is that structures shall sustain its anticipated loads with appropriate level of reliability and in an economical way, see section 2.1(1) of EN 1990:2002. The term “appropriate level of reliability” is specified further in 2.2(3) where its choice is related to relevant factors, including:

- (a) the possible cause and /or mode of attaining a limit state;
- (b) the possible consequences of failure in terms of risk to life, injury, potential economical losses;
- (c) public aversion to failure;
- (d) the expense and procedures necessary to reduce the risk of failure.

A formal differentiation according to (a), (c) and (d) is not further considered in the present version of the Eurocode. According to (b), structures and components are classified, e.g., regarding the consequence of failure, whereas the reliability levels might be given either for classified structures as a whole or for classified components (see 2.2(4) of EN 1990:2002).

In EN 1990:2002, Annex B3.1, a classification of buildings and structures in terms of consequences is defined (3 classes). It is mentioned that importance of a failure mode for the consequences should be considered (see B3.1(2,3) of EN 1990:2002). This might lead to different classification of different failure modes in one structure. In EN 1990:2002, B3.2 minimum reliability levels for 3 consequence classes are recommended.

Table B2 - Recommended minimum values for reliability index β (ultimate limit states)

Reliability Class	Minimum values for β	
	1 year reference period	50 years reference period
RC3	5,2	4,3
RC2	4,7	3,8
RC1	4,2	3,3

Figure 5.2: Reliability requirements as stated in the EUROCODES.

2.4. Direct correspondence between reliability and design values - the design value approach

For specific cases, a direct correspondence between the design value and the reliability requirements can be established by the so called design value method. In order to demonstrate this, the simple reliability problem is revisited, where one normal distributed resistance variable R is compared with a normal distributed load variable S by a safety margin M that is defined correspondingly as normal distributed (for independent R and S), i.e. $M = R - S$. Failure is then defined as $M = R - S < 0$ and the failure probability as

$$\begin{aligned}
 p_F &= \Pr(M = R - S < 0) = \Phi\left(\frac{0 - \mu_M}{\sigma_M}\right) \\
 &= \Phi\left(\frac{-\mu_R + \mu_S}{\sqrt{\sigma_R^2 + \sigma_S^2}}\right) \\
 &= \Phi(-\beta)
 \end{aligned} \tag{5.6}$$

Thus, for this simple case, the reliability index β is defined as

$$\beta = \frac{\mu_R - \mu_S}{\sqrt{\sigma_R^2 + \sigma_S^2}} \tag{5.7}$$

In structural design e.g. according to the Eurocodes, the reliability requirement is given as a design requirement β_{req} . Therefore, it is useful to reformulate the above expression, i.e. to facilitate the identification of a combination of R and S that comply with, say, $\beta_{req} = 3.8$ such that $\beta \geq \beta_{req}$. The distance between the mean values that is in compliance with the required

reliability is determined as $\mu_R - \mu_S \geq \beta_{req} \sigma_M$, whereas σ_M can be splitted into the contributions from R and S as

$$\mu_R - \mu_S \geq \beta_{req} \frac{\sigma_R}{\sqrt{\sigma_R^2 + \sigma_S^2}} \sigma_R + \beta_{req} \frac{\sigma_S}{\sqrt{\sigma_R^2 + \sigma_S^2}} \sigma_S$$

Introducing the weighting factors α_R and α_S as

$$\begin{aligned} \alpha_R &= \frac{\sigma_R}{\sqrt{\sigma_R^2 + \sigma_S^2}} \\ \alpha_S &= \frac{\sigma_S}{\sqrt{\sigma_R^2 + \sigma_S^2}} \end{aligned} \quad (5.8)$$

and separating the resistance from the load side, the following expression is obtained:

$$\begin{aligned} \mu_R - \alpha_R \beta_{req} \sigma_R &\geq \mu_S + \alpha_S \beta_{req} \sigma_S \\ \mu_R (1 - \alpha_R \beta_{req} V_R) &\geq \mu_S (1 + \alpha_S \beta_{req} V_S) \\ r_d &\geq s_d \end{aligned} \quad (5.9)$$

with the coefficients of variation V_R and V_S and so-called design values r_d and s_d .

Note: Equation (5.9) looks very convenient for structural design as it suggests the separation of load and resistance variables. But this is not really the case as the weighting factors α_R and α_S contain the standard deviations of both variables. The implications of this are demonstrated by the following example.

Example 1. Consider the simple example, where a tension rod with material capacity C is exposed to a load S . Both variables are considered as independent and normal distributed with parameters given in Table 5.1. The cross section area a that complies with the reliability requirement $\beta_{req} = 3.8$ is to be specified by the application of Eq. (5.9). The required design value r_d

Table 5.1: Mean value μ and coefficient of variation V of the capacity C and the load S .

	μ	V
Capacity C [kN/mm ²]	1	0.10
Load S [kN]	1	0.34

for the resistance $R = a \cdot C$ can be expressed based on Equation (5.9):

$$\begin{aligned} r_d = a \cdot \mu_C (1 - \alpha_R \beta_{req} V_C) &\geq \mu_S (1 + \alpha_S \beta_{req} V_S) = s_d \\ a \cdot \mu_C \left(1 - \frac{a \sigma_C}{\sqrt{a^2 \sigma_C^2 + \sigma_S^2}} \beta_{req} V_C \right) &\geq \mu_S \left(1 + \frac{\sigma_S}{\sqrt{a^2 \sigma_C^2 + \sigma_S^2}} \beta_{req} V_S \right) \end{aligned}$$

, i.e. the cross section area a is chosen such that the equation above is fulfilled.

Unfortunately, the identification of a is not as straight forward as it looks, since both, α_R and α_S are dependent on a . However, a solution can be found with $a = 2.615 \text{ mm}^2$, which corresponds to $r_d = s_d = 2 \text{ kN}$ and $\alpha_R = 0.6153$ and $\alpha_S = 0.7882$.

Obviously, the same result is obtained based on Equation (5.7), i.e. without the splitting into design values.

$$a \text{ such that } \beta_{req} = 3.8 \leq \beta = \frac{a \cdot \mu_C - \mu_S}{\sqrt{a^2 \sigma_C^2 + \sigma_S^2}}$$

△

From the example it becomes obvious that the computation of the design value of the resistance, r_d , requires the evaluation of α_R and this in turn requires the full evaluation of the reliability problem. In order to circumvent this problem the EUROCODES introduce so-called *standardized α values*. In order to develop the background for the choice of these *standardized α values* it is useful to represent the considered simple reliability problem as it was first introduced by Hasshofer and Lind in 1974.

2.4.1. Alternative representation of the simple reliability problem

Hasofer et al. 1974 suggested an alternative representation of the reliability problem and a generalised definition of the reliability index that is also applicable to non-linear reliability problems with non-normal random variables. Applied to the simple reliability problem from above, i.e. one normal distributed resistance variable R is compared with a normal distributed load variable S , a so-called limit state function is introduced, as

$$g(r, s) = r - s = 0 \quad (5.10)$$

The limit state function can be illustrated together with the joint probability density function of R and S , $f_{R,S}(r, s)$, compare Figure 5.3, where it can be seen that the limit state is dividing the possible domain of r and s into a failure domain, $g(r, s) \leq 0$ and a safe domain, $g(r, s) > 0$. The probability of failure can be computed by integrating the joint probability distribution function over the failure domain, i.e.

$$p_F = \int_{g(r,s) \leq 0} f_{R,S}(r, s) dr ds \quad (5.11)$$

By transforming the normal distributed random variables, R and S into standard normal variables $U_R = (R - \mu_R)/\sigma_R$ and $U_S = (S - \mu_S)/\sigma_S$, the limit state function can be written as a function of u_R and u_S , as

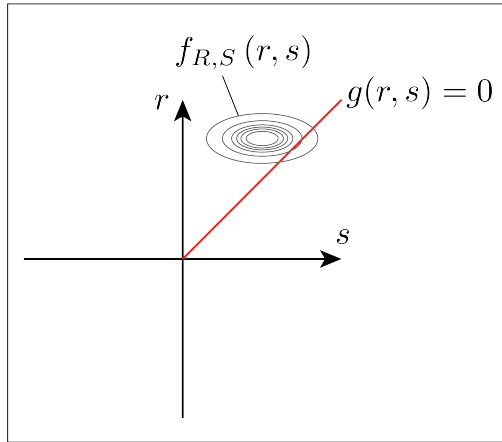
$$g(u_R, u_S) = (u_R \sigma_R + \mu_R) - (u_S \sigma_S + \mu_S) = 0 \quad (5.12)$$

Introducing the unit vector α , with the cosine of direction α_R and α_S , we find the reliability index β geometrically (compare Figure 5.3) as the shortest (orthogonal) distance between the limit state and the origin:

$$\begin{aligned} g(u_R, u_S) &= \beta - \alpha_R \cdot u_R - \alpha_S \cdot u_S = 0 \\ g(\mathbf{u}) &= \beta - \alpha^T \mathbf{u} = 0 \end{aligned} \quad (5.13)$$

For the simple reliability problem the α -values and β are identical to Equation (5.7) and (5.8), the different geometrical interpretation, however, is much more general, that is

real space



u-space

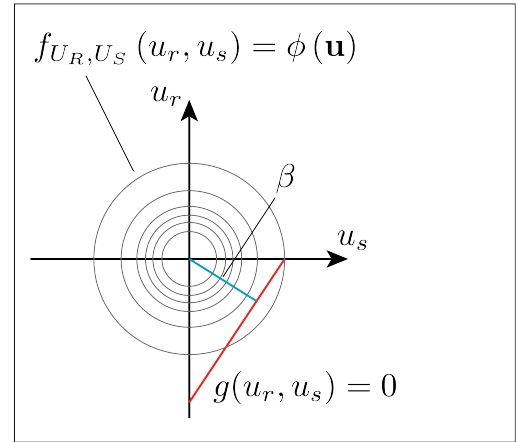


Figure 5.3: Left: The simple reliability problem with the resistance r and the load s . The topology of the joint probability density is indicated by the ellipses, the failure domain is below the limit state $g(r, s) = 0$, the safe domain is above.

Right: The same problem represented in the standard normal space and with the variables u_R and u_S . $f_{U_R, U_S}(u_R, u_S) = \phi(\mathbf{u})$ is the standard normal density.

- The design point $\mathbf{u}^* = \alpha\beta$ is the point on the limit state function $g(\mathbf{u}) = 0$ that is closest to the origin.
- α is the normal vector on the limit state in the design point.
- β is the distance between the origin and the design point.
- In correspondence to Equation (5.7) $r_d = \alpha_R\beta\sigma_R + \mu_R$ and $s_d = \alpha_S\beta\sigma_S + \mu_S$ are defining the coordinates of the design point and are referred to as the design values of the resistance and the load.

Example 2. The alternative representation is applied to the example above and obviously, similar results are obtained. The corresponding graphical representation can be inspected in Figure 5.4.

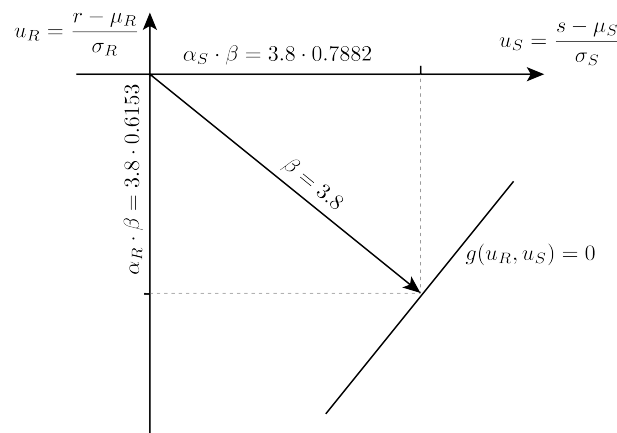


Figure 5.4: Illustration of Example 1 in the standard normal space.

△

2.5. Derivation of generalized α -values

In the following, the derivation of a set of generalised α -values that provide sufficiently accurate estimations of the design point for a relevant range of design situations is explored. Similar considerations have been done by König et al. 1981 and their results have been adapted in the current version of EN1990.

The idea is as follows: α is a unit vector of length 1. If a set of generalised α -values is selected such that the corresponding vector sum is larger than 1, these generalised α -values could be applied to a range of cases and still lead to design solutions that comply with the required reliability. The idea is illustrated in Figure 5.5.

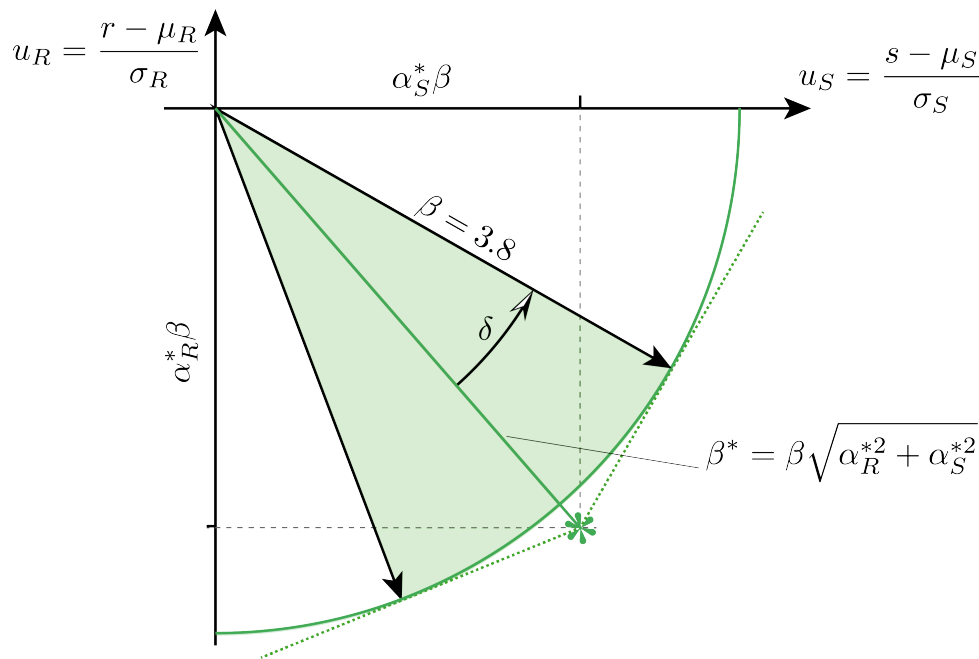


Figure 5.5: Illustration of the effect of the application of generalised α -values, α_R^* and α_S^* creating a range of design cases (indicated in green) for which compliance with the safety requirement can be achieved.

Accordingly and for linear limit state functions and normal distributed variables R and S , the range of cases for which the generalised α -values, if applied in combination both on the load and on the resistance side, result in design solutions that comply with the reliability requirement can be represented by the angle δ as

$$\cos \delta = \frac{\beta}{\beta^*} = \frac{1}{\sqrt{\alpha_R^{*2} + \alpha_S^{*2}}} \quad (5.14)$$

2.6. The generalised α -values in the Eurocodes

In the Eurocodes (prEN 1990:E2020 and EN1990:2002) the generalised α -values for the dominant resistance and load variable are

$$\begin{aligned}\alpha_{R,EC}^* &= -0.8 \\ \alpha_{S,EC}^* &= 0.7\end{aligned}\tag{5.15}$$

With $\beta_{req} = 3.8$ and following up the illustration in Figure 5.5, this corresponds to $\beta^* = 4.04$ and $\delta = 19.8^\circ$. When Equation (5.9) is applied for evaluating the design values r_d and s_d the design solutions identified based on the criterion $r_d = s_d$ have a reliability index larger than 3.8 for normal distributed variables and real α -values in the range of $\alpha_R > -0.956$ and $\alpha_S < 0.9$.

2.6.1. Generalisation of the Design Value Approach to other than Normal distributions

With proper transformation, the variables design value can be approximated.

Accordingly, the design value is defined as

$$y_d = F_Y^{-1}(\Phi(\alpha_Y \beta_t))\tag{5.16}$$

where y_d is the design value of a variable Y , F_Y is its cumulative distribution function, α_Y (with $|\alpha_Y| \leq 1$) is a sensitivity factor indicating the importance of Y in the reliability estimation, and β_t is the target value for the reliability index specifying the reliability requirement.

Design value y_d and characteristic value y_k for some common distributions are determined according:

Normal:	$y_d = \mu_Y (1 + \alpha_Y \beta_t V_Y)$ $y_k = \mu_Y (1 + \Phi^{-1}(p) V_Y)$
Log-Normal:	$y_d = \mu_Y \exp \left(-\frac{1}{2} \ln(1 + V_Y^2) + \alpha_Y \beta_t \sqrt{\ln(1 + V_Y^2)} \right)$ $y_k = \mu_Y \exp \left(-\frac{1}{2} \ln(1 + V_Y^2) + \Phi^{-1}(p) \sqrt{\ln(1 + V_Y^2)} \right)$
Gumbel:	$y_d = \mu_Y \left(1 - V_Y \frac{\sqrt{6}}{\pi} (0.5772 + \ln(-\ln(\Phi(\alpha_Y \beta_t)))) \right)$ $y_k = \mu_Y \left(1 - V_Y \frac{\sqrt{6}}{\pi} (0.5772 + \ln(-\ln(p))) \right)$

The FORM sensitivity factors α_Y are resulting from the reliability analysis of the the design situation at hand. If no reliability analysis is performed, the following Eurocode standardized values can be used for a 50 years reference period:

- If Y represents a strength related variable: $\alpha_Y = -0.8$
- If Y represents a load related variable: $\alpha_Y = 0.7$
- If Y is dominating the reliability problem: $\alpha_Y = (-)1$

- If Y represents a secondary strength or load related variable: $\alpha_Y = -0.8 \cdot 0.4$ or $\alpha_Y = 0.7 \cdot 0.4$ correspondingly.

Self-weight is usually represented by a Normal distribution; Resistance variables are often represented by a Log-Normal distribution; the extreme values per reference period of time-variable actions are represented by the Gumbel distribution.

3. Application examples

E5.1. Code calibration

Geometry

The simply supported beam in Figure 5.6 is to be designed with a given target reliability. The basic random variables that represent the situation are given in Table 5.3. The resulting beam span and cross-section modulus are dependent on the design philosophy that is followed. The following formats are to be considered: reliability based design, load and resistance factor design and global partial safety factor.

Table 5.3: Basic random variables (* $\equiv$ yearly maxima).

	G	Q^*	F_m
Distribution	Normal	Normal	Normal
Mean (μ_X)	2 N/mm	10000 N	200 MPa
Standard dev. (σ_X)	0.30 N/mm	4000 N	20 MPa
Coeff. of var. (COV_X)	15%	40%	10%
Characteristic fractile (k_X)	0.50	0.95	0.05

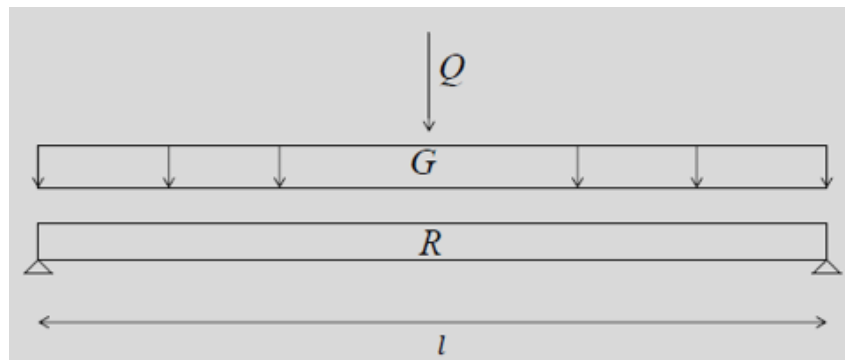


Figure 5.6: Beam geometry and loads.

The limit state function for bending moment at mid-span is:

$$g(f_m, g, q) = W \cdot f_m - g \frac{l^2}{8} - q \frac{l}{4} \leq 0 \quad (5.17)$$

The corresponding reliability index is:

$$\beta = \frac{W\mu_{F_m} - \mu_G \frac{l^2}{8} - \mu_Q \frac{l}{4}}{\sqrt{(W\sigma_{F_m})^2 + \left(\sigma_G \frac{l^2}{8}\right)^2 + \left(\sigma_Q \frac{l}{4}\right)^2}} \quad (5.18)$$

E5.2. Reliability based design

The beam cross-section W can be chosen such that the reliability index of the beam corresponds to the target reliability, i.e. $\beta = \beta_t$. This corresponds to the so-called **reliability based design**. E.g. for $l = 6$ m, the section plastic modulus giving $\beta = \beta_t$ is found to be $W = 3.01 \cdot 10^5$. The cross-section plastic modulus is plotted against the beam length in Figure 5.7.

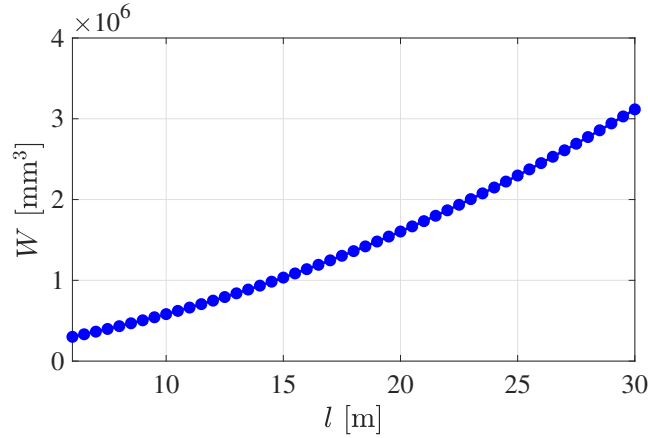


Figure 5.7: Reliability based design for different beam lengths for $\beta = \beta_t = 4.2$.

E5.2.1. Load and resistance factor design format

Although modern structural design codes (like Eurocodes) allow reliability based design, they propose simpler design procedures in the **semi-probabilistic design format**, where design equations are written in the **load and resistance factor design format** (LRFD). This format is referred to as **semi-probabilistic**, since random variables (represented by distribution functions) are reduced to design values, which are obtained from characteristic values as $r_d = r_k / \gamma_R$ (for resistance) and $s_d = s_k \gamma_S$ (for actions). The format is simpler, since the designer can check the design just by comparing design values, without the necessity of using reliability analyses.

The design equation of the beam above is:

$$W \cdot f_{m,d} - g_d \frac{l^2}{8} - q_d \frac{l}{4} \leq 0 \quad (5.19)$$

where:

- $f_{m,d} = \frac{f_{m,k}}{\gamma_R}$; $g_d = g_k \gamma_G$ and $q_d = q_k \gamma_Q$ are the **design values**;

- $f_{m,k} = F_{F_m}^{-1}(k_{F_m})$; $g_k = F_G^{-1}(k_G)$ and $q_k = F_Q^{-1}(k_Q)$ are the **characteristic values** corresponding to the characteristic fractiles k_{F_m} , k_G and k_Q , which are defined arbitrarily and
- γ_G ; γ_R and γ_Q are the **partial safety factors (p.s.f.)**.

The minimum cross-section plastic modulus satisfying Eq. (5.19) is:

$$W_{min} = \frac{g_K \gamma_G \frac{l^2}{8} + q_K \gamma_Q \frac{l}{4}}{\frac{f_{m,k}}{\gamma_R}} \quad (5.20)$$

3.0.1. Calibration of partial safety factors for a given span

The p.s.f. need to be **calibrated** so that the design given in Eq. (5.20) fulfills a target safety. This corresponds to answer the question *which values of the p.s.f. in Eq. (5.20) provides a design (i.e. W) associated with the level of risk $\beta = \beta_t$?*

The limit state function in Eq. (5.17) is rewritten with $W = W_{min}$:

$$g(f_m, g, q) = \frac{g_K \gamma_G \frac{l^2}{8} + q_K \gamma_Q \frac{l}{4}}{\frac{f_{m,k}}{\gamma_R}} f_m - g \frac{l^2}{8} - q \frac{l}{4} \leq 0 \quad (5.21)$$

Consequently Eq. (5.18) becomes:

$$\beta = \frac{\left(\frac{g_K \gamma_G \frac{l^2}{8} + q_K \gamma_Q \frac{l}{4}}{\frac{f_{m,k}}{\gamma_R}} \right) \mu_{F_m} - \mu_G \frac{l^2}{8} - \mu_Q \frac{l}{4}}{\sqrt{\left(\left(\frac{g_K \gamma_G \frac{l^2}{8} + q_K \gamma_Q \frac{l}{4}}{\frac{f_{m,k}}{\gamma_R}} \right) \sigma_{F_m} \right)^2 + \left(\sigma_G \frac{l^2}{8} \right)^2 + \left(\sigma_Q \frac{l}{4} \right)^2}} \quad (5.22)$$

where it is clear that the reliability of the design depends on the partial safety factors, i.e. $\beta = \beta(\gamma_R, \gamma_G, \gamma_Q)$.

The calibrate partial safety factors $(\gamma_R^*, \gamma_G^*, \gamma_Q^*)$ are found by setting $\beta(\gamma_R^*, \gamma_G^*, \gamma_Q^*) \equiv \beta_t$. Note that there are infinite sets of $(\gamma_R^*, \gamma_G^*, \gamma_Q^*)$ which satisfy the equality. By fixing one p.s.f. (e.g. $\gamma_R^* = 1.05$) a unique set is obtained.

By way of example, setting $l = 6$ m, the calibrated partial safety factors result in $(\gamma_R^* = 1.05, \gamma_G^* = 1.344, \gamma_Q^* = 1.437)$, yielding $W = 300610 \text{ mm}^3$. Numerical minimization in Matlab[®] has been used.

Design value method - FORM sensitivity factors method

An alternative method for p.s.f. calibration makes use of the FORM sensitivity factors, which are:

$$\alpha_{F_m} = \frac{-W\sigma_{F_m}}{k} ; \alpha_G = \frac{\left(\sigma_G \frac{l^2}{8}\right)}{k} ; \alpha_Q = \frac{\left(\sigma_Q \frac{l}{4}\right)}{k} \quad (5.23)$$

with:

$$k = \sqrt{(W\sigma_{F_m})^2 + \left(\sigma_G \frac{l^2}{8}\right)^2 + \left(\sigma_Q \frac{l}{4}\right)^2} \quad (5.24)$$

Introducing $W = 3.01 \cdot 10^5 \text{ mm}^3$ that was found for $l = 6 \text{ m}$, the sensitivity factors are ($\alpha_{F_m} = -0.699, \alpha_Q = 0.698, \alpha_G = 0.157$). The corresponding partial safety factors are:

$$\begin{aligned} \gamma_R^* &= \frac{f_{m,k}}{f_{m,d}} = \frac{1 + \Phi^{-1}(k_{F_m}) \cdot COV_{F_m}}{1 + \alpha_{F_m} \beta_t COV_{F_m}} = 1.18 \\ \gamma_Q^* &= \frac{q_d}{q_k} = \frac{1 + \alpha_Q \beta_t \cdot COV_Q}{1 + \Phi^{-1}(k_Q) COV_Q} = 1.31 \\ \gamma_G^* &= \frac{g_d}{g_k} = \frac{1 + \alpha_G \beta_t \cdot COV_G}{1 + \Phi^{-1}(k_G) COV_G} = 1.10 \end{aligned} \quad (5.25)$$

It is highlighted that the two sets of optimized p.s.f. ($\gamma_R^* = 1.05, \gamma_G^* = 1.344, \gamma_Q^* = 1.437$) and ($\gamma_R^* = 1.18, \gamma_G^* = 1.31, \gamma_Q^* = 1.10$) are equivalent in the sense that they lead to the same design (i.e. same W) and the same reliability index (i.e. $\beta = \beta_t = 4.2$) for $l = 6 \text{ m}$.

Calibration of partial safety factors for different scenarios

In general, the optimized p.s.f. obtained above using the FORM sensitivity factor method result in a reliability $\beta \neq \beta_t$, for a beam span $l \neq 6 \text{ m}$, cf. Figure 5.8. This is due to the fact that the fraction of the external moment induced by G (or Q) changes with the beam span. The reliability for $l = 6 \text{ m}$ is equal to the target, since that was the condition imposed for calibration. The reliability index is therefore dependent on the beam span (or equivalently the ratio between the moments induced by G and Q), i.e. $\beta = \beta(\gamma_R, \gamma_G, \gamma_Q, l)$.

In order to get a constant reliability, i.e. independent of the geometry, a new set of p.s.f should be calibrated for each beam length, see Figure 5.9. This is not practical and not implemented in design codes. The problem is solved by finding the set of p.s.f which ensures the reliability to be as homogeneous as possible over a range of possible scenarios. The set can be obtained by solving the minimization problem in Eq. (5.26).

$$\min_{\gamma_R^*, \gamma_G^*, \gamma_Q^*} \left\{ \sum_{i=1}^n (\beta_t - \beta_i(\gamma_R, \gamma_G, \gamma_Q, l_i))^2 \right\} \quad (5.26)$$

The set of p.s.f obtained is ($\gamma_R^* = 1.08, \gamma_G^* = 1.38, \gamma_Q^* = 1.35$). The corresponding reliability is plotted in Figure 5.10.

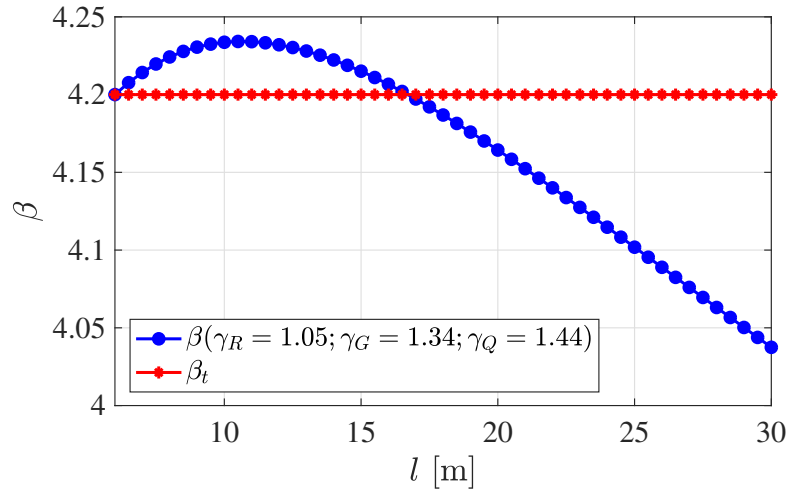


Figure 5.8: Reliability index of the beam designed with p.s.f optimized for $l = 6$ m.

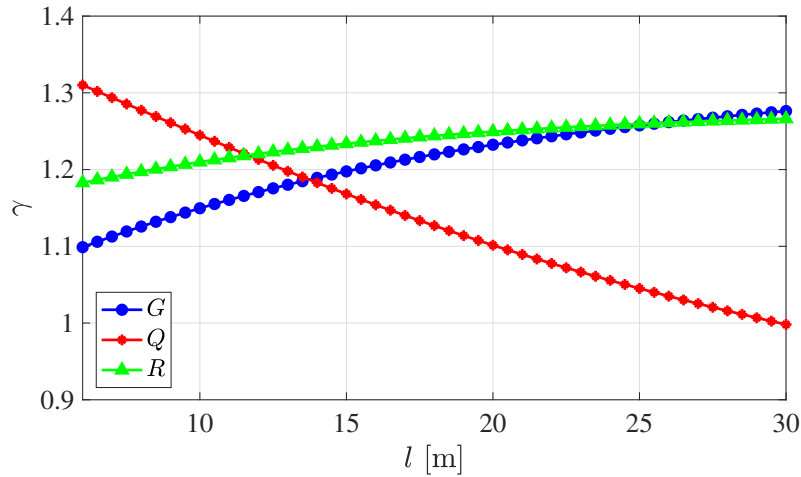


Figure 5.9: Partial safety factors giving constant reliability, calibrated with FORM α -method giving $\beta = \beta_t = 4.2$.

E5.2.2. Global partial safety factor

An alternative (and older) format of semi-probabilistic design equations makes use of the global safety factor (SF). The design equation with the global safety factor applied on the mean values reads:

$$\frac{W \mu_{F_m}}{l^2} \frac{l}{\mu_G \frac{8}{8} + \mu_Q \frac{4}{4}} > SF \quad (5.27)$$

where the numerator is the mean value of the moment resistance and the denominator is the mean value of the action induced moment.

The minimum cross-section plastic modulus satisfying Eq. (5.27) is:

$$W_{min} = \frac{SF}{\mu_{F_m}} \left(\mu_G \frac{l^2}{8} + \mu_Q \frac{l}{4} \right) \quad (5.28)$$

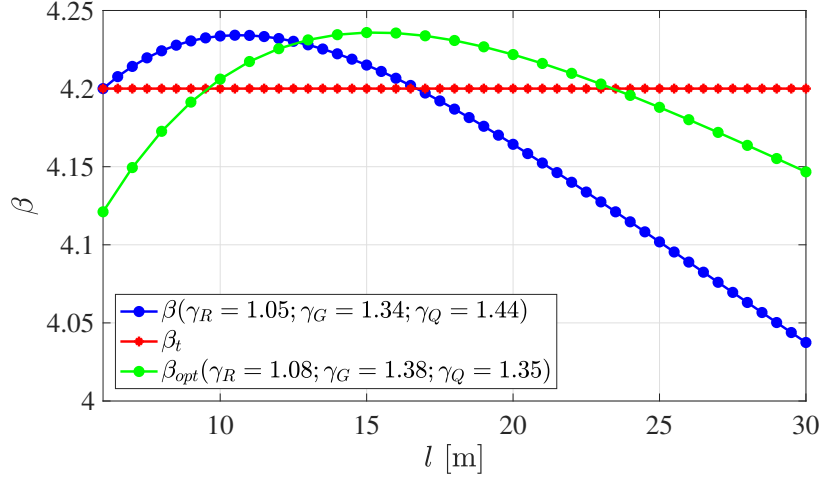


Figure 5.10: Comparison of the reliability index for p.s.f. calibrated for $l = 6$ m and for different scenarios simultaneously (β_{opt}).

Calibration of the global safety factor for a given span

As before, SF can be calibrated so that Eq. (5.28) provides a design with the target level of risk. The limit state function in Eq. (5.17) is rewritten with $W = W_{min}$:

$$g(f_m, g, q) = \frac{SF}{\mu_{F_m}} \left(\mu_G \frac{l^2}{8} + \mu_Q \frac{l}{4} \right) \cdot f_m - g \frac{l^2}{8} - q \frac{l}{4} \leq 0 \quad (5.29)$$

Consequently Eq. (5.18) becomes:

$$\beta = \frac{\frac{SF}{\mu_{F_m}} \left(\mu_G \frac{l^2}{8} + \mu_Q \frac{l}{4} \right) \mu_{F_m} - \mu_G \frac{l^2}{8} - \mu_Q \frac{l}{4}}{\sqrt{\left(\frac{SF}{\mu_{F_m}} \left(\mu_G \frac{l^2}{8} + \mu_Q \frac{l}{4} \right) \sigma_{F_m} \right)^2 + \left(\sigma_G \frac{l^2}{8} \right)^2 + \left(\sigma_Q \frac{l}{4} \right)^2}} \quad (5.30)$$

where it is clear that the reliability of the design is depending on the global safety factor, i.e. $\beta = \beta(SF)$.

The global safety factor is calibrated by setting $\beta(SF) \equiv \beta_t$. As an example, considering $l = 6$ m, the calibrated global safety factor is found with numerical minimization in Matlab[®] to be $SF^* = 2.51$, giving $W = 300610 \text{ mm}^3$.

Calibration of the global safety factor for different scenarios

As before, the design performed with the calibrated SF for $l = 6$ m results in a different reliability for different spans. The SF can be calibrated so to give a reliability as homogeneous as possible over the different design scenarios. Therefore, as previously, the SF can be calibrated by solving a minimization problem, cf. Eq. (5.31). The optimized SF obtained is $SF^* = 2.20$, see Figure 5.11.

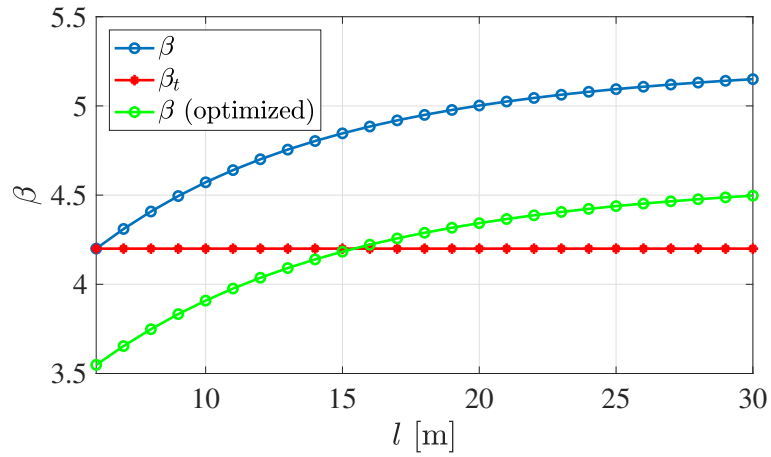


Figure 5.11: Reliability index for SF calibrated for $l = 6$ m (in blue) and for different scenarios simultaneously (in green).

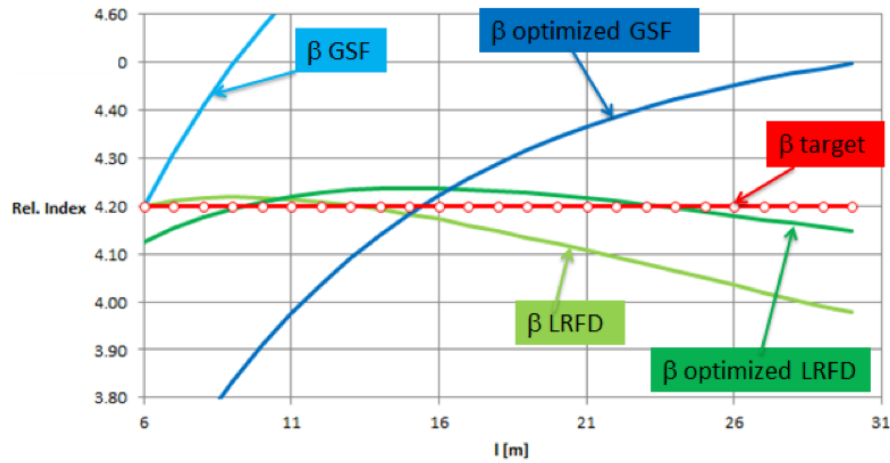


Figure 5.12: Comparison of β values for SF and $LRFD$ formats before and after optimization.

$$\min_{SF^*} \left\{ \sum_{i=1}^n (\beta_t - \beta_i(SF^*, l_i))^2 \right\} \quad (5.31)$$

Comments

It can be observed from Figure 5.10 and Figure 5.11 that, by changing the SF the curve is only shifted up and down while the three p.s.f. shift and distort the curve. This leads to higher levels of optimization (i.e. a flatter curve) with p.s.f.. The values of $\sum_{i=1}^n (\beta_t - \beta_i)^2$ show clearly this differences in the two formats. In fact, after optimization $\sum_{i=1}^n (\beta_t - \beta_i)^2 = 0.026$ for the p.s.f. format and $\sum_{i=1}^n (\beta_t - \beta_i)^2 = 1.88$ for the global safety factor format. The load and resistance factor design method offers therefore a more homogenized level of safety.

Question::

Structural Standards and Reliability

- What is the partial factor design concept?
- How does it correspond to reliability?
- Why are multiple partial factors better than a single global safety factor?
- What is the idea behind “generalized α -values”? What is the potential, what are the limitations?

Chapter 6

System Reliability

The fundamentals of system reliability analysis are introduced in this chapter.

1. Theoretical aspects

1.1. Introduction

So far, we considered failure as component failure, whereas the word "component" has no literal meaning; it is used in the sense of failure mode. Component failure is thus defined by the failure domain Ω_F that is specified by a single limit state function $g(\mathbf{x})$, which represents the corresponding failure mode as:

$$\Omega_F = \{g(\mathbf{x}) \leq 0\} \quad (6.1)$$

In the context of structural reliability theory, a system is defined as a combination of different failure modes. Correspondingly, the system state is described with a combination of several limit state functions $g_1(\mathbf{x})$, $g_2(\mathbf{x})$, ..., $g_n(\mathbf{x})$, each representing one possible failure mode.

It is common to distinguish between two basic system configurations that correspond to either the case where the realisation of a single failure mode equates to system failure, or the case where system failure requires all individual failure modes to realize. The former system configuration is in general referred to the *serial system*, the latter as *parallel system*. A combination of both is called mixed system and can be always represented as a serial arrangement of parallel sub-systems.

1.1.1. Serial Systems

The serial system fails if any of its components fails (if any of its possible failure modes realizes). A block diagram of a series system is illustrated in Figure 6.1. These systems are also called "weakest link" systems. Hence, the failure domain of a system composed of n elements is:

$$\Omega_F = \bigcup_{i=1}^n \{g_i(\mathbf{x}) \leq 0\} \quad (6.2)$$

Similarly, the probability of system failure is:

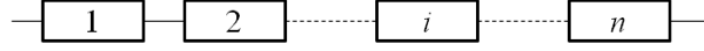


Figure 6.1: Block diagram of a series system.

$$\Pr(F_{sys}) = \Pr\left(\bigcup_{i=1}^n F_{comp,i}\right) = 1 - \Pr\left(\bigcap_{i=1}^n \bar{F}_{comp,i}\right) \quad (6.3)$$

where $\Pr(F_{sys})$ is the probability of occurrence of the system failure event, $\Pr(F_{comp,i})$ is the probability of failure of the component i and $\Pr(\bar{F}_{comp,i})$ the probability of survival or non-failure of the component i .

If the different failure modes are independent, Eq. (6.3) can be further elaborated:

$$\Pr(F_{sys}) = 1 - \prod_{i=1}^n \Pr(\bar{F}_{comp,i}) = 1 - \prod_{i=1}^n (1 - \Pr(F_{comp,i})) \quad (6.4)$$

It shall be noticed that $\Pr(F_{sys})$ increases with n . For small $\Pr(\bar{F}_{comp,i})$, Eq. (6.4) can be approximated by:

$$\Pr(F_{sys}) \approx \sum_{i=1}^n \Pr(F_{comp,i}) \quad (6.5)$$

1.1.2. Parallel Systems

A parallel system fails only if all of the failure modes of the structural system realize (all components fail). These systems are also called “redundant” systems. A block diagram of a parallel system is illustrated in Figure 6.2. The failure domain of a parallel system composed of n elements is thus:

$$\Omega_F = \bigcap_{i=1}^n \{g_i(\mathbf{x}) \leq 0\} \quad (6.6)$$

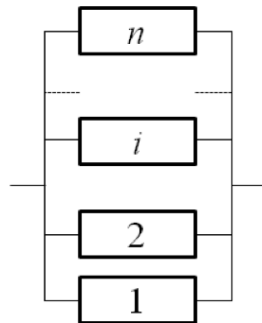


Figure 6.2: Block diagram of a parallel system.

Similarly, the probability of system failure F_{sys} can be related to component failure $F_{comp,i}$ as:

$$\Pr(F_{sys}) = \Pr\left(\bigcap_{i=1}^n F_{comp,i}\right) \quad (6.7)$$

For independent failure modes, it follows:

$$\Pr(F_{sys}) = \prod_{i=1}^n \Pr(F_{comp,i}) \quad (6.8)$$

Notice that $\Pr(F_{sys})$ decreases with n .

1.1.3. Mixed Systems

A combined system or mixed system is a system consisting of both series and parallel systems. In general, all kinds of combined systems can be reduced to a series of parallel sub-systems. For example, the combined system in Figure 6.3 may be used to represent the failure of a structural system, which fails if any of the sub-systems (parallel systems) representing failure of several failure modes fail.

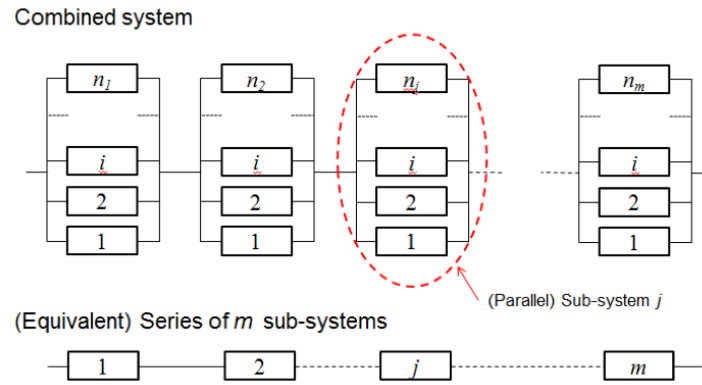


Figure 6.3: Block diagram of a mixed system.

The failure domain of a mixed system is given by:

$$\Omega_F = \bigcup_{j=1}^m \bigcap_{i=1}^{n_j} \{g_j(\mathbf{x}) \leq 0\} \quad (6.9)$$

where m is the number of parallel sub-systems (so-called cut-sets) connected in series and n_j is the number of elements connected in parallel in the j^{th} cut-set. If independence of the components failure events can be assumed, the failure of the subsystem (parallel system) j is calculated using Eq. (6.8) as:

$$\Pr(F_{cut-set,j}) = \Pr\left(\bigcap_{i=1}^{n_j} F_{comp,i,j}\right) = \prod_{i=1}^{n_j} \Pr(F_{comp,i,j}) \quad (6.10)$$

The failure of the whole system is calculated using Eq. (6.4) as:

$$\Pr(F_{sys}) = 1 - \prod_{j=1}^m (1 - \Pr(F_{cut-set,j})) \quad (6.11)$$

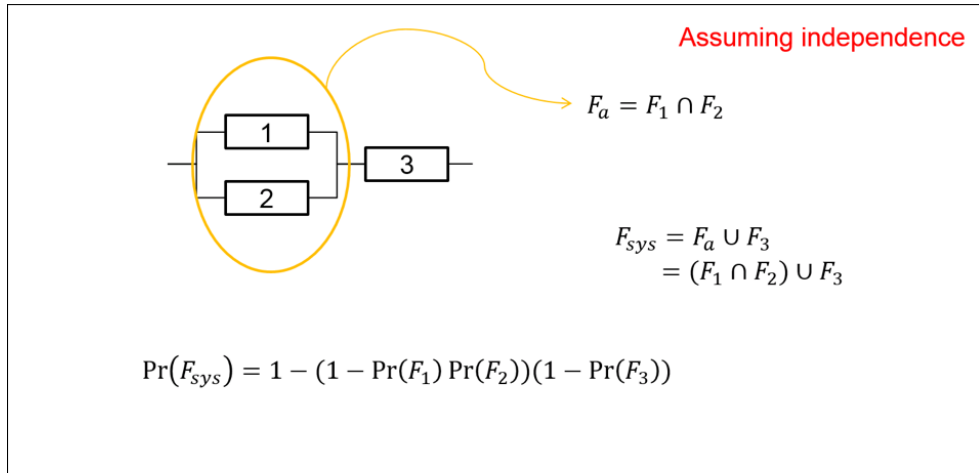


Figure 6.4: Example of mixed system.

1.2. Simple bounds for system reliability of failure

The above given results for serial-, parallel- and mixed systems are conditional to the assumption of *independent* failure modes. For many practical applications, however, this assumption is not appropriate. At the same time it is often difficult to quantify the dependence of different failure modes. For these cases the so-called simple bounds for the system reliability define an interval in which the system probability of failure is contained, i.e. the bounds are thereby defined by two extreme assumptions, i.e. considering the components as either fully independent or fully dependent.

Simple bounds for serial systems

The upper bound is given by the failure probability of a **series system** of **independent** components. The lower bound is defined by the failure probability of a **series system** of fully **dependent** components; the system failure probability is equal to the failure probability of the failure mode with the largest failure probability, i.e. the weakest link of the system. Therefore, when the correlation between the failure modes is somewhere between 0 and 1, the simple bounds on the system failure probability for a series system are given as:

$$\max_{i=1}^n \{\Pr(F_{comp,i})\} \leq \Pr(F_{sys}) \leq 1 - \prod_{i=1}^n (1 - \Pr(F_{comp,i})) \quad (6.12)$$

Simple bounds for parallel systems

The upper bound is defined by the failure probability of a **parallel system** of fully **dependent** components, i.e. the failure probability of the failure mode with the smallest failure probability. The lower bound is defined by the failure probability of a **parallel system** of **independent** components.

$$\prod_{i=1}^n (\Pr(F_{comp,i})) \leq \Pr(F_{sys}) \leq \min_{i=1}^n \{\Pr(F_{comp,i})\} \quad (6.13)$$

Simple bounds for mixed systems

For mixed systems Eq. (6.12) and (6.13) are applied in a sequential manner.

Second order bounds (not part of the pensum)

Second order bounds on the failure probability of systems with correlated failures of components can be established as well. See Chapter 5 in “Structural Reliability analysis and prediction”, Second edition –Robert E. Melchers. In general they are closer to the true system probability of failure.

1.3. System Reliability using FORM

System reliability problems can be represented in the standard normal space. The failure domain, as defined in Eqs. (6.2) and (6.6), can be graphically illustrated for serial and parallel systems as in Figure 6.5.

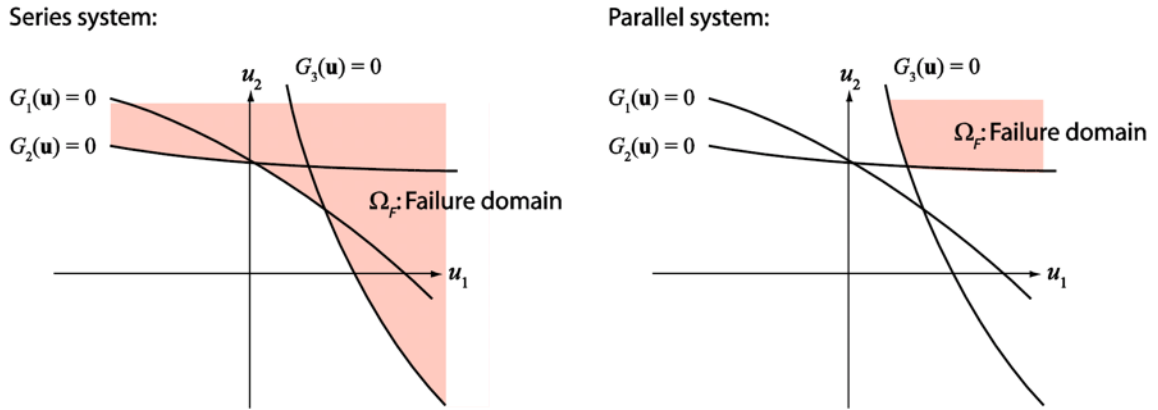


Figure 6.5: Graphical representation of the failure domain of a serial and a parallel system. (Picture adapted from Daniel Straub, TUM)

1.3.1. Serial Systems

For a serial system, the failure probability might be approximated as follows. A variable χ_i , $i = 1, 2, \dots, n$ is introduced and defined as in Eq. (6.14). χ_i represents a projection of the variables $\mathbf{U}$ on the straight line through the origin and the design point $\mathbf{u}_i^*$ associated to the failure mode/ limit state $g_i(\mathbf{u})$, and is standard normal distributed.

$$\chi_i = \boldsymbol{\alpha}_i \mathbf{U} \quad (6.14)$$

The component i failure event is thus computed defined as:

$$F_i = \{\beta_i \leq \chi_i\} \quad (6.15)$$

where β_i is the FORM reliability index of component i . The serial system probability of failure can thus be approximated as follows:

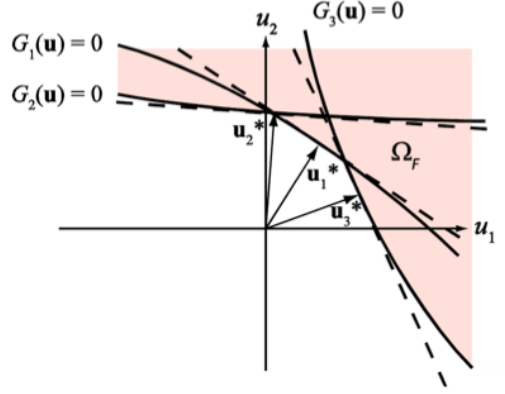


Figure 6.6: Graphical representation of the solution strategy for a serial system. (Picture adapted from Daniel Straub, TUM)

$$\begin{aligned}
 \Pr(F_{sys,series}) &= \Pr\left[\bigcup_{i=1}^n \{\beta_i \leq \chi_i\}\right] \\
 &= 1 - \Pr\left[\bigcap_{i=1}^n \{\chi_i < \beta_i\}\right] \\
 &= 1 - \Phi_n(\mathbf{B}, \mathbf{R})
 \end{aligned} \tag{6.16}$$

where Φ_n is the n -dimensional standard normal distribution with $\mathbf{B}$ the vector of component FORM reliability indices β_i and $\mathbf{R}$ the correlation matrix. $\mathbf{R}$ is equal to the identity matrix for the case of no correlation between the component failure events. However, correlation obviously exists for usual cases and the components of the correlation matrix $R_{ij} = \rho_{ij}$ are estimated by:

$$\rho_{ij} = \boldsymbol{\alpha}_i \boldsymbol{\alpha}_j^T \quad i = 1, 2, \dots, n; \quad j = 1, 2, \dots, n \tag{6.17}$$

The approximation in Eq. (6.16) was first introduced by Hohenbichler and Rackwitz in 1978.

Graphical interpretation of $\rho_{i,j}$

The correlation between two failure modes, has an illustrative graphical interpretation. Consider two linearized limit states $M_i = \beta_i - \boldsymbol{\alpha}_i^T \mathbf{U}$ and $M_j = \beta_j - \boldsymbol{\alpha}_j^T \mathbf{U}$ as shown in Figure (6.7).

It is seen that $\cos \theta_{i,j} = \boldsymbol{\alpha}_i^T \boldsymbol{\alpha}_j = \rho_{i,j}$ where $\theta_{i,j}$ is the angle between the $\boldsymbol{\alpha}$ -vectors or simply between the linearized safety margins. If now the two limit states are considered to have a shortest distance to the origin $\beta_i = \beta_j = 3.0$ the cases of $\rho_{i,j}$ equal to -1, 0, $\sqrt{0.5}$ and 1 are illustrated in Figure (6.8).

The generalized systems reliability index for the cases considered in Figure (6.8), $\beta^S = -\Phi^{-1}(1 - \Phi_2(3, 3; \rho))$ is $\beta^S = (2.782, 2.783, 2.812, 3)$. The general relationship between β^S and $\rho_{i,j}$ is illustrated in Figure (6.9).

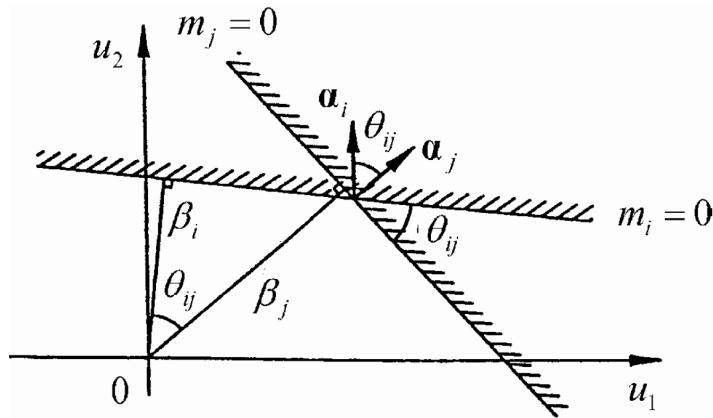


Figure 6.7: Graphical interpretation of the correlation coefficient as $\rho_{i,j} = \cos\theta_{i,j}$. (Picture adapted from J. D. Sørensen, AAU)

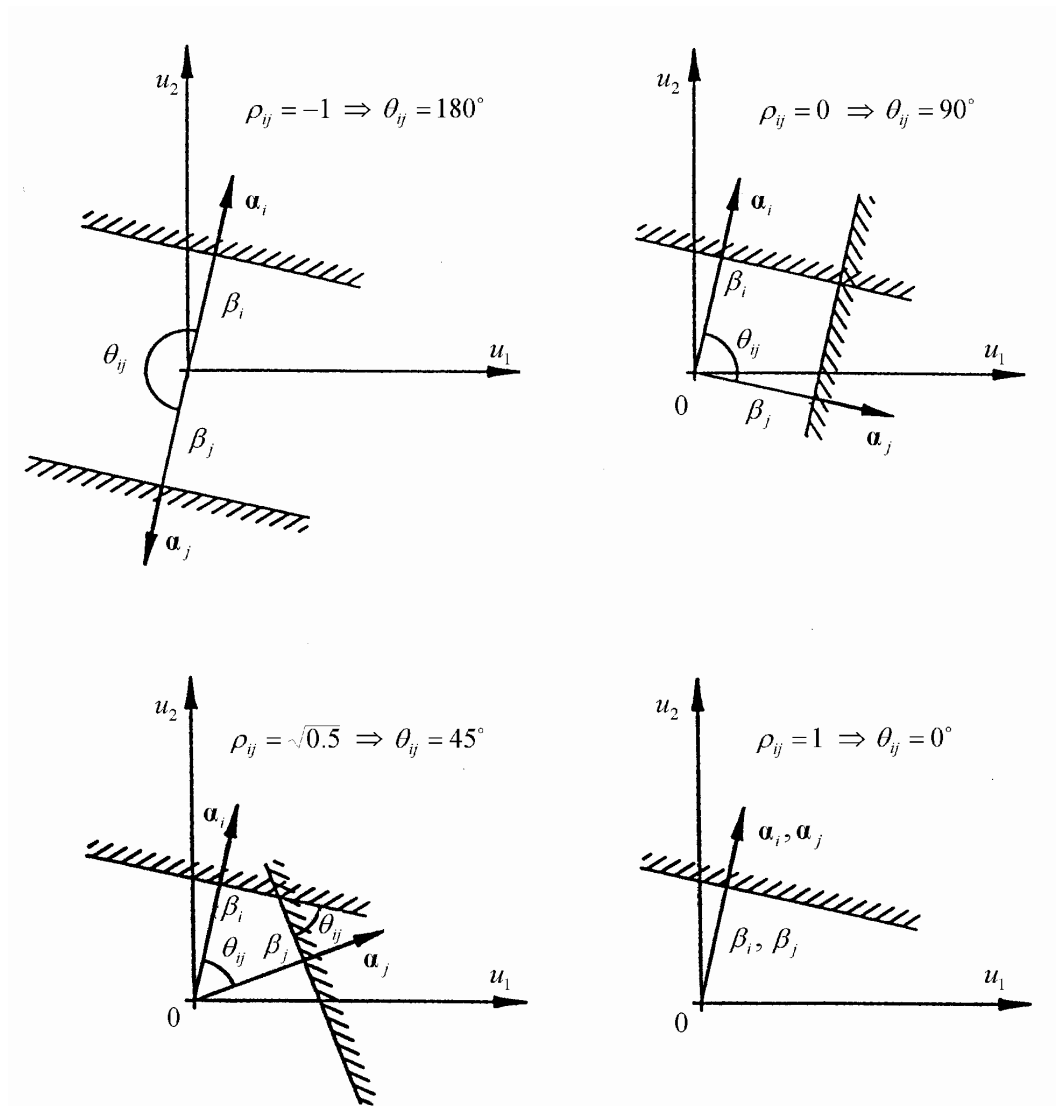


Figure 6.8: Illustration of different $\rho_{i,j} = \cos\theta_{i,j}$. (Picture adapted from J. D. Sørensen, AAU)

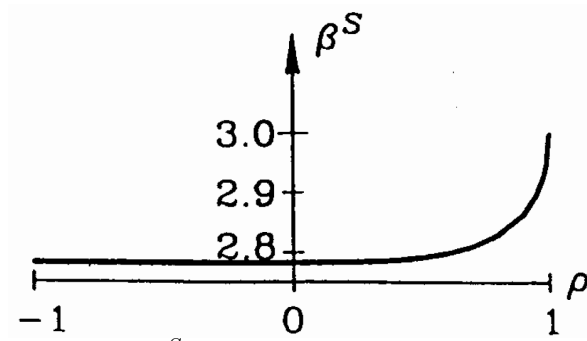


Figure 6.9: Illustration of different β^S dependent on $\rho_{i,j} = \cos\theta_{i,j}$. (Picture adapted from J. D. Sørensen, AAU)

1.3.2. Parallel Systems

For parallel systems, a similar approach might be followed. The solution in Eq. (6.16) is adapted as:

$$\begin{aligned}
 \Pr(F_{sys,parallel}) &= \Pr\left[\bigcap_{i=1}^n \{\beta_i \leq \chi_i\}\right] \\
 &= \Pr\left[\bigcap_{i=1}^n \{\chi_i < -\beta_i\}\right] \\
 &= \Phi_n(-\mathbf{B}, \mathbf{R})
 \end{aligned} \tag{6.18}$$

The accuracy of the estimate for the probability of failure of parallel systems is increased by including information about the position of the joint design point $\mathbf{u}^*$, cf. Figure 6.10. The joint design point is determined by the solution of the optimization problem in Eq. (6.19).

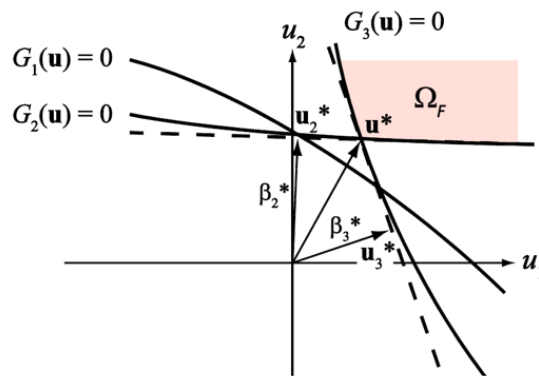


Figure 6.10: Illustration of the FORM estimation of the failure probability of a parallel system considering the joint design point $\mathbf{u}^*$. (Picture adapted from Daniel Straub, TUM)

$$\mathbf{u}^* = \arg \min \|\mathbf{u}\| \quad \text{subject to } g_i(\mathbf{u}) \leq 0, \quad i = 1, 2, \dots, n \tag{6.19}$$

The α -vectors for the different failure modes are then found by linearization in that joint

design point as:

$$\alpha_i^* = -\frac{\nabla g_i(\mathbf{u}^*)}{\|\nabla g_i(\mathbf{u}^*)\|} \quad (6.20)$$

The β values, however, are found as in usual FORM as $\beta_i^* = \alpha_i^* \mathbf{u}_i^*$, see Figure 6.10. The correlation matrix is estimated with $\rho_{ij}^* = \alpha_i^* \alpha_j^{*T}$. The improved FORM estimate for the parallel system reliability is then estimated with $\Phi_n(-\mathbf{B}^*, \mathbf{R}^*)$.

2. Application examples

E6.1.

Calculate the failure probability of the system in Figure 6.11 with $\Pr(F_1) = 0.10$, $\Pr(F_2) = 0.20$, $\Pr(F_3) = 0.30$ and $\Pr(F_4) = 0.40$. Failure events of different elements are independent.

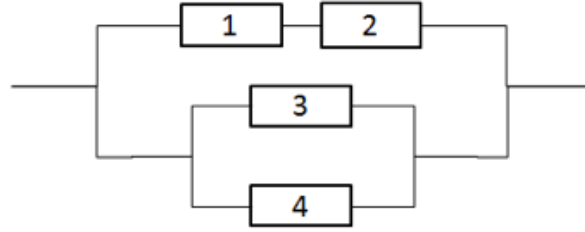


Figure 6.11: Combined system.

An equivalent system is represented in Figure 6.12 with:

$$\Pr(F_{\{1 \cup 2\}}) = 1 - (1 - \Pr(F_1))(1 - \Pr(F_2)) = 1 - (1 - 0.1)(1 - 0.2) = 0.28 \quad (6.21)$$

$$\Pr(F_{\{3 \cap 4\}}) = \Pr(F_3) \cdot \Pr(F_4) = 0.30 \cdot 0.40 = 0.12 \quad (6.22)$$

Considering the (parallel) system in Figure 6.12, the failure probability of the system is then:

$$\Pr(F_{sys}) = \Pr(F_{\{1 \cup 2\} \cap \{3 \cap 4\}}) = \Pr(F_{\{1 \cup 2\}}) \cdot \Pr(F_{\{3 \cap 4\}}) = 0.28 \cdot 0.12 = 0.0336 \quad (6.23)$$

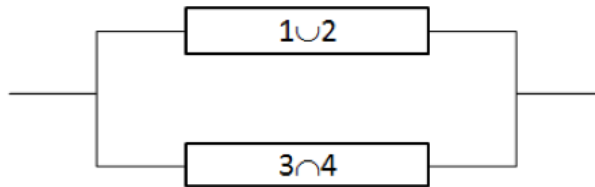


Figure 6.12: Equivalent system of parallel elements.

E6.2.

A fire protection system is installed in an industrial facility, containing two fire detectors (D1 and D2) and a sprinkler arrangement (S). The detectors are connected to the sprinkler arrange-

ment. If fire is indicated by one of the detectors, the sprinkler system is activated. Consequently, the system fails if the two detectors fail and/or the sprinkler system fails, see Figure 6.13.

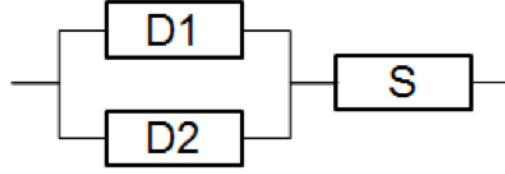


Figure 6.13: Block model of the system.

In case of fire, the probability that one fire detector fails is $\Pr(F_{Di}) = 0.05$ ($i = 1, 2$); and the probability that the sprinkler system fails is $\Pr(F_S) = 0.001$.

E6.2.1.

Assuming independent events, what is the probability that the fire protection system fails in case of fire?

If the failure events of D1 and D2 are independent, the probability that both detection systems fail simultaneously is given by:

$$\Pr(F_{\{D1 \cap D2\}}) = \Pr(F_{D1}) \Pr(F_{D2}) = 0.05 \cdot 0.05 = 0.0025 \quad (6.24)$$

The system failure probability is:

$$\begin{aligned} \Pr(F_{sys}) &= \Pr(F_{\{(D1 \cap D2) \cup S\}}) = 1 - \Pr(F_{\{(\overline{D1 \cap D2}) \cap \bar{S}\}}) \\ &= 1 - (1 - \Pr[F_{\{D1 \cap D2\}}]) (1 - \Pr(F_S)) = 1 - (1 - 0.0025)(1 - 0.001) = 0.0035 \end{aligned} \quad (6.25)$$

E6.2.2.

Assuming now that the detector failure events are dependent, what is the probability that the fire protection system fails in case of fire? The conditional probability that the detector 1 fails given detector 2 has failed (and vice versa) is given as $\Pr(F_{D1}|F_{D2}) = \Pr(F_{D2}|F_{D1}) = 0.50$.

The probability that D1 and D2 fail simultaneously is now given by:

$$\Pr(F_{(D1 \cap D2)}) = \Pr(F_{D1}|F_{D2}) \Pr(F_{D2}) = \Pr(F_{D2}|F_{D1}) \Pr(F_{D1}) = 0.50 \cdot 0.05 = 0.025 \quad (6.26)$$

The system failure results in:

$$\begin{aligned} \Pr(F_{sys}) &= \Pr(F_{\{(D1 \cap D2) \cup S\}}) = \\ &= 1 - (1 - \Pr(F_{\{D1 \cap D2\}})) (1 - \Pr(F_S)) = 1 - (1 - 0.025)(1 - 0.001) = 0.026 \end{aligned} \quad (6.27)$$

It should be noticed that neglecting the dependency of two events is not necessarily conservative. Figure 6.14 shows the system probability of failure for different conditional probabilities. In this case, neglecting the dependence between failures of D1 and D2 is not conservative, i.e. it leads to an underestimation of the probability of failure.

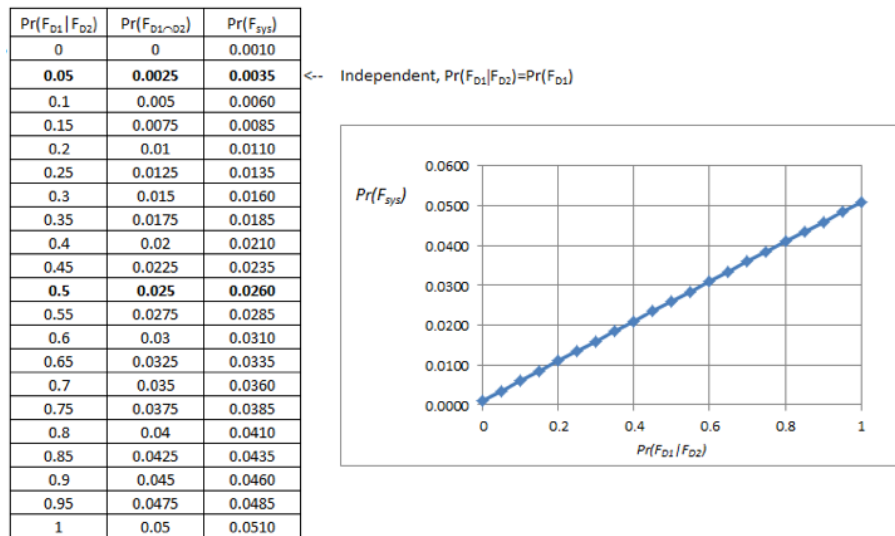


Figure 6.14: System failure probability for various values of the conditional probability of the failure of the detection systems.

The previous observation can be further elaborated. The system probability of failure for series systems is overestimated when dependence is not considered. In other words, **neglecting dependencies between failure modes is conservative in series systems**. On the opposite, the system probability of failure for parallel systems is underestimated when dependence is not considered. In other words, **neglecting dependencies between failure modes is not conservative in parallel systems**. This is illustrated in Figure 6.15.

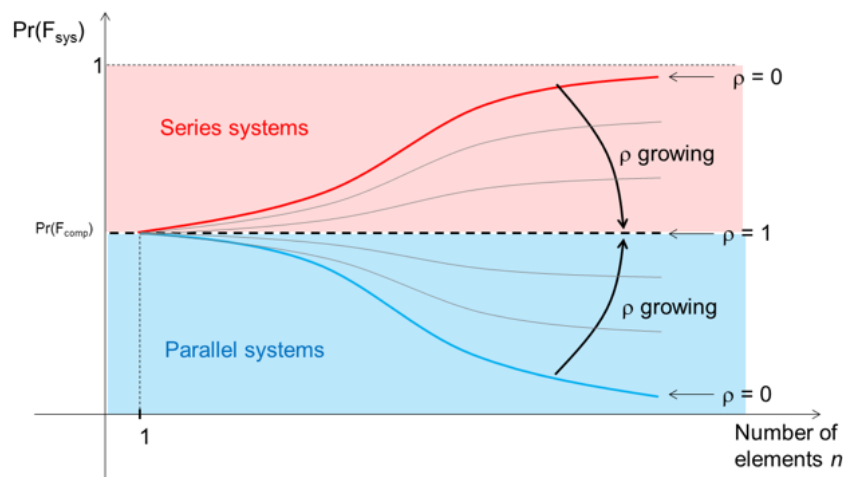


Figure 6.15: Qualitative relation between number of elements and correlation between failure events, and system probability of failure.

E6.2.3.

Calculate the simple bounds of the system probability of failure.

For the lower bound, events are assumed independent for components connected in parallel

and fully dependent for serial system.

$$\Pr(F_{\{D1 \cap D2\}}) = \Pr(F_{D1}) \cdot \Pr(F_{D2}) = 0.005 \cdot 0.05 = 0.0025 \quad (6.28)$$

$$\Pr(F_{sys}) = \max\{\Pr(F_{\{D1 \cap D2\}}); \Pr(F_S)\} = \max\{0.0025, 0.001\} = 0.0025 \quad (6.29)$$

For the upper bound, fully dependent events are assumed for parallel systems and independent events for serial systems.

$$\Pr(F_{\{D1 \cap D2\}}) = \min\{\Pr(F_{D1}); \Pr(F_{D2})\} = \min\{0.05; 0.05\} = 0.05 \quad (6.30)$$

$$\Pr(F_{sys}) = 1 - (1 - \Pr(F_{\{D1 \cap D2\}})) \cdot (1 - \Pr(F_S)) = 1 - (1 - 0.05) \cdot (1 - 0.001) = 0.05095 \quad (6.31)$$

A simplified approach can also be employed to estimate the simple bounds:

Considering independent events, i.e. $\rho = 0$, $\Pr(F_{sys})$ resulted in $\Pr(F_{sys, \rho=0}) = 0.0035$

Considering fully dependent events, i.e. $\rho = 1$, $\Pr(F_{sys})$ yields:

$$\Pr(F_{\{D1 \cap D2\}}) = \min\{0.05, 0.05\} = 0.05$$

$$\Pr(F_{sys, \rho=1}) = \Pr(F_{\{(D1 \cap D2) \cup S\}}) = \max\{0.05, 0.001\} = 0.05$$

Hence, the simple bounds are:

$$0.0035 \leq \Pr(F_{sys}) \leq 0.05$$

As expected, the value of 0.025 obtained for $\Pr(F_{D1}|F_{D2}) = \Pr(F_{D2}|F_{D1}) = 0.50$ is between the bounds.

E6.3.

Consider the system in Figure 6.16, where the probability of failure of the components is given in Eq. (6.32). Determine the simple bounds of the system probability of failure.

$$\begin{aligned} \Pr(F_{comp,1}) &= \Pr(F_{comp,2}) = \Pr(F_{comp,4}) = 1 \cdot 10^{-2} \\ \Pr(F_{comp,3}) &= \Pr(F_{comp,5}) = \Pr(F_{comp,6}) = 1 \cdot 10^{-5} \end{aligned} \quad (6.32)$$

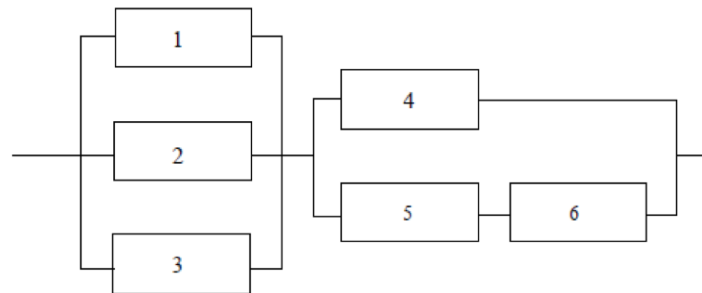


Figure 6.16: Mixed system with correlated failure modes.

The mixed system may be successively reduced to a series system as illustrated in Figure 6.17.

The failure probabilities $\{1 \cap 2 \cap 3\}$ and $\{4 \cap (5 \cup 6)\}$ are to be calculated to build the reduced system. To calculate the simple bounds, the system probability of failure first is to be computed under two extreme conditions. First, assuming uncorrelated failure events and second, assuming fully correlated failure events.

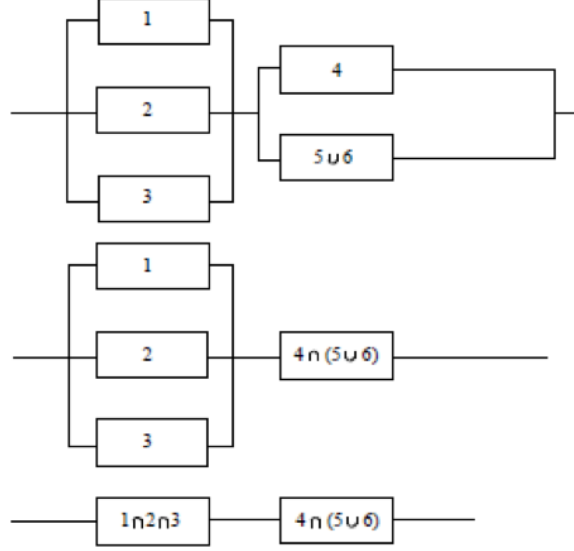


Figure 6.17: Illustration of scheme for successive reduction of a system based on the simple bounds.

1. Assuming uncorrelated failure events:

$$\Pr(F_{\{1 \cap 2 \cap 3\}}) = (1 \cdot 10^{-2})^2 (1 \cdot 10^{-5}) = 1 \cdot 10^{-9}$$

$$\Pr(F_{\{4 \cap (5 \cup 6)\}}) = \Pr(F_4) \cdot \Pr(F_{\{5 \cup 6\}}) = 1 \cdot 10^{-2} \cdot \left(1 - (1 - 1 \cdot 10^{-5})^2\right) = 2 \cdot 10^{-7}$$

$$\Pr(F_{sys, \rho=0}) = \Pr(F_{\{(1 \cap 2 \cap 3) \cup (4 \cap (5 \cup 6))\}}) = 1 - (1 - 2 \cdot 10^{-7})(1 - 1 \cdot 10^{-9}) = 2.01 \cdot 10^{-7}$$

2. Assuming fully correlated failure events:

$$\Pr(F_{\{1 \cap 2 \cap 3\}}) = \min\{1 \cdot 10^{-2}, 1 \cdot 10^{-2}, 1 \cdot 10^{-5}\} = 1 \cdot 10^{-5}$$

$$\Pr(F_{\{4 \cap (5 \cup 6)\}}) = \min\{\Pr(F_4), \Pr(F_{\{5 \cup 6\}})\} = \min\{1 \cdot 10^{-2}, \max\{1 \cdot 10^{-5}, 1 \cdot 10^{-5}\}\} = 1 \cdot 10^{-5}$$

$$\Pr(F_{sys, \rho=1}) = \Pr(F_{\{(1 \cap 2 \cap 3) \cup (4 \cap (5 \cup 6))\}}) = \max\{1 \cdot 10^{-5}, 1 \cdot 10^{-5}\} = 1 \cdot 10^{-5}$$

Therefore, the simple bounds are:

$$2.01 \cdot 10^{-7} \leq \Pr(F_{sys, \rho}) \leq 1 \cdot 10^{-5}$$

3. Practical exercises

P6.1. Reliability-based design of an elevator

Consider the elevator system in Figure 6.18 composed of:

- A steel support for the top sheave, with net resisting cross-section $A_{S, res}$ and material strength $F_{S, y}$, that is rotated by an electric motor;

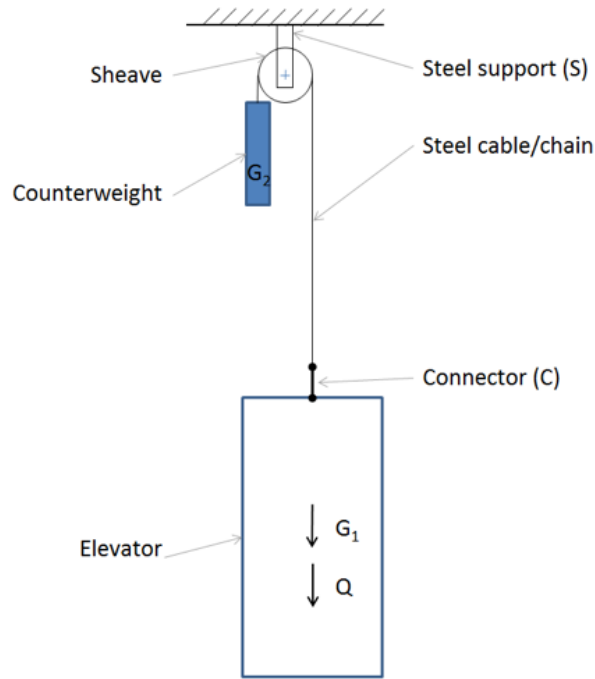


Figure 6.18: Elevator model.

- A steel cable/chain connecting the counterweight to the elevator, with net resisting cross-section $A_{T,res}$ and material strength $F_{T,y}$. The elevator is installed in a 30-story building so the cable/chain is approximately 100 m long;
- A connector between the cable/chain and the elevator, with net resisting cross-section $A_{C,res}$ and material strength $F_{C,y}$;
- The elevator with self-weight $G_1 = 300 \text{ kg} \cdot g = 2.94 \text{ kN}$;
- A counterweight of weight $G_2 = 10 \text{ kN}$.

The yearly maxima live load on the elevator is Q . The basic variables are given in Table 6.1.

Table 6.1: Basic variables.

Symbol	Distribution	Mean	C.o.V. [-]
Q	Normal	11 kN	0.4
G_1	Deterministic	2.94 kN	/
G_2	Deterministic	10 kN	/
$A_{S,res}$	Deterministic	150 mm ²	/
$F_{S,y}$	Normal	400 MPa	0.07
$A_{T,res}$	Deterministic	Decision parameter	/
$F_{T,y}$	Normal	630 MPa	0.05
$A_{C,res}$	Deterministic	100 mm ²	/
$F_{C,y}$	Normal	400 MPa	0.07

P6.1.1. Design of a chain

Design the chain, i.e. choose the resisting cross-section $A_{T,res}$, that gives the probability of failure of the supporting system, i.e. steel support + steel chain + connector equal to $P_{f,S-T-C} = 2.5 \cdot 10^{-6}$ ($\beta = 4.56$).

Consider: 10,000 links composing the chain; for the worst scenario, i.e. the elevator in the lowest position, it can be assumed that 90% of the cable length is on the elevator side and 10% is on the counter-weight side.

P6.1.2. Design of a parallel wire cable

Design the parallel wire cable cross-section providing the same failure probability. Consider 20 similar parallel wires composing the cable.

Question::

System Reliability

- What are the different types of systems and what are the implications for structural reliability?
- How is dependence affect system reliability?
- What are simple bounds and how are they computed?
- How can system reliability assessed with FORM?

Chapter 7

Engineering Decision Making

Structural engineering and engineering in general involves decisions that have to be made in order to maximize the utility of structures in particular or the built environment in general. An obvious type of decisions are the decisions about physical configurations or changes related to the design and maintenance of structures. In chapter 1, we have already seen how risk based criteria can be utilized for structural design decision subject to uncertainty. A second less obvious but not less important type of decision is about what information and what level of modelling detail we should employ for representing the decision situations. As information comes at a cost, it has to be integrated in the utility maximisation as well.

Bayesian decision theory offers the theoretical framework for assessing both types of decisions and the basic concepts are introduced in this chapter.

1. Concept

The built environment, i.e. infrastructural elements, industrial buildings and facilities as well as residential buildings, constitutes the basis for our economy and the continuous development of our society. In this respect structures play an important role, since the primary purpose of structures is to provide the functionality of the built environment and the safety of its users. Consequently, a large proportion of the societal economic resources are invested into the continued development, maintenance and renewal of structures. Due to the high importance of structures and the large amount of economic resources invested for development, maintenance and renewal it is of utmost importance that decisions made in connection with structures are optimal. Optimal in the sense that the benefit of structures as well as the possible adverse consequences such as loss of lives, damage to the qualities of the environment and the direct and committed costs are considered and optimal engineering decisions are identified.

Normative decision theory, as introduced by von Neumann and Morgenstern already in 1943, (Neumann et al. 1943) provides the general framework for engineering decision making and provides the decision maker with a rational basis to act. Here, rationality is understood as “instrumental rationality”, i.e. the decision maker does evaluate and rank the utility of different decision outcomes in terms of expectations, (Weber 1921).

The evaluation of expected utilities for structural engineering decision problems is generally associated with large aleatory and epistemic uncertainties that have to be considered. A consistent framework for the account of information and uncertainty in decision making is Bayesian

Decision Analysis as presented in Raiffa et al. 1961. According to this, and using the same classical nomenclature, the principle layout of Bayesian Decision Analysis is illustrated in Figure 7.1.

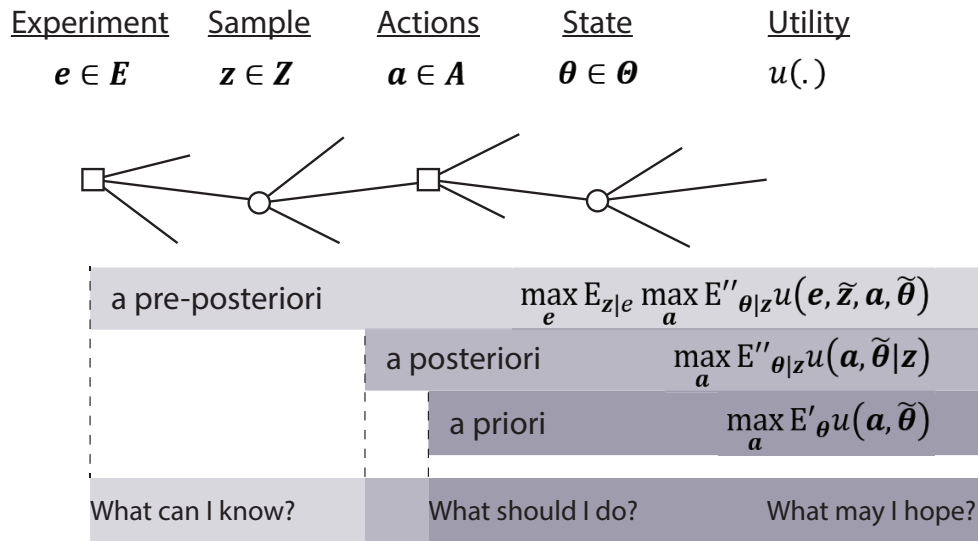


Figure 7.1: Basic decision problem facilitating for decisions about actions and experiments, according to Raiffa et al. 1961. In the graph, the squares represent decision nodes, and the circles represent chance nodes. It is indicated how the decision analysis might support the assessment of the three basic questions of Kant's Canon of Pure Reason (Kant).

The most simple form of decision analysis is the so called “prior decision analysis”. Here, the decision maker states his or her prior probabilities of state $P'(\theta)$ and his or her preferences among possible pairs of state θ and action a as expressed by the utility $u(a, \theta)$. The expectation over the different outcomes of the random state variable $\tilde{\theta}$ is then maximised subject to a . The posterior, or terminal analysis, does allow for the consideration of new information represented as the outcome of an experiment z . The new information is expressed in form of a likelihood $P(z|\theta)$ and the posterior probability is obtained by combining the likelihood with the prior probability, $P''(\theta|z) \propto P(z|\theta)P'(\theta)$. The utility to be maximised subject to a is then also conditional on z , i.e. $u(a, \theta|z)$. The pre-posterior analysis allows for the assessment of different possible experiments before they are chosen, i.e. this type of analysis includes a terminal analysis for each possible type of experiment e and the possible corresponding outcome z . The outcomes of the experiment are weighted according to the prior probabilities and the expectation of the utility $u(e, \tilde{z}, a, \tilde{\theta})$ can be maximised subject to e and a .

The concept of Bayesian Decision Analysis represents a logical and consistent framework for decision making under uncertainty, whereas decisions not only refer to actions that directly influence the state of nature, in Figure 7.1 this is referred to as a , but also our knowledge about the state of nature θ . In a broad sense also the choice of the modelling universe for the representation of the decision problem can be formally included in the framework, see Glavind et al. 2018 for an example in engineering decision making. Although the strict implementation of Bayesian Decision Analysis is challenging for some complex situations in practical engineering, the underlying logic may be used to sort and structure the engineering problem in a reasonable way. Before we start to illustrate the main principles of Bayesian Decision Analysis alongside an “engineering example”, we will recap the *Bayes' Rule* which is fundamental for the concept.

2. Bayes' Rule

This section is intended as a repetitorium from the statistics class.

2.1. Conditional Probability and Bayes' Rule

Conditional probabilities are of special interest in risk and reliability analysis as they form the basis of the updating of probability estimates based on new information, knowledge and evidence.

The conditional probability of the event E_1 given that the event E_2 has occurred is written as:

$$P(E_1 | E_2) = \frac{P(E_1 \cap E_2)}{P(E_2)} \quad (7.1)$$

It is seen that the conditional probability is not defined if the conditioning event is the empty set, i.e. when $P(E_2) = 0$.

The event E_1 is said to be probabilistically independent of the event E_2 if:

$$P(E_1 | E_2) = P(E_1) \quad (7.2)$$

implying that the occurrence of the event E_2 does not affect the probability of E_1 .

From Equation 7.1 the probability of the event $E_1 \cap E_2$ may be given as:

$$P(E_1 \cap E_2) = P(E_1 | E_2)P(E_2) \quad (7.3)$$

and it follows immediately that if the events E_1 and E_2 are independent, then:

$$P(E_1 \cap E_2) = P(E_1)P(E_2) \quad (7.4)$$

Based on the above findings, the important *Bayes' rule* can be derived. Consider the sample space Ω divided into n mutually exclusive events $E_1, E_2, \dots, E_n$ (see also Figure 7.2, where the case of $n = 8$ is considered).

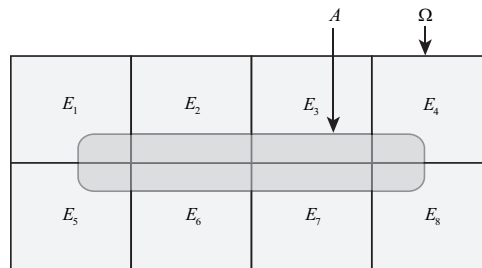


Figure 7.2: Illustration of the rule of Bayes.

Furthermore let the event A be an event in the sample space Ω . Then the probability of the

event A , i.e. $P(A)$, can be written as:

$$\begin{aligned}
 P(A) &= P(A \cap E_1) + P(A \cap E_2) + \dots + P(A \cap E_n) \\
 &= P(A | E_1)P(E_1) + P(A | E_2)P(E_2) + \dots + P(A | E_n)P(E_n) \\
 &= \sum_{i=1}^n P(A | E_i)P(E_i)
 \end{aligned} \tag{7.5}$$

this is also referred to as the *total probability theorem*.

From Equation 7.1 there is $P(A | E_i)P(E_i) = P(E_i | A)P(A)$ implying that:

$$P(E_i | A) = \frac{P(A | E_i)P(E_i)}{P(A)} \tag{7.6}$$

Now by inserting Equation 7.6 into Equation 7.5, the *Bayes' rule* results:

$$P(E_i | A) = \frac{P(A | E_i)P(E_i)}{\sum_{j=1}^n P(A | E_j)P(E_j)} \tag{7.7}$$

In Equation 7.6 $P(E_i | A)$ is called the *posterior probability* of E_i , the conditional term $P(A | E_i)$ is often referred to as the *likelihood* (i.e. the probability of observing a certain state given the true state). The term $P(E_i)$ is the *prior probability* of the event E_i (i.e. prior to the knowledge about the event A).

As mentioned previously, the Bayes' rule is extremely important, and in order to facilitate its appreciation a few illustrative applications will be given in the following.

2.1.1. Example 1 - Using Bayes' Rule for Concrete Assessment

A reinforced concrete beam is considered. From experience it is known that the probability that corrosion of the reinforcement has initiated (the event CI) is $P(CI) = 0.01$. However, in order to know the condition more precisely an inspection method (non-destructive) has been developed.

The quality of the inspection method may be characterised by the probability that the inspection method will indicate initiated corrosion given that corrosion has initiated $P(I | CI)$ (the probability of detection or equivalently the likelihood of an indication I given corrosion initiation CI) and the probability that the inspection method will indicate initiated corrosion given that no corrosion has initiated $P(I | \overline{CI})$ (the probability of erroneous findings or the likelihood of an indication given no corrosion initiation).

For the inspection method at hand the following characteristics have been established:

$$\begin{aligned}
 P(I | CI) &= 0.8 \\
 P(I | \overline{CI}) &= 0.1
 \end{aligned}$$

An inspection of the concrete beam is conducted with the result that the inspection method indicates that corrosion has initiated. Based on the findings from the inspection, what is the probability that corrosion of the reinforcement has initiated?

The answer is readily found by application of Bayes' rule:

$$P(CI|I) = \frac{P(I|CI)P(CI)}{P(I|CI)P(CI) + P(I|\overline{CI})P(\overline{CI})} = \frac{P(I \cap CI)}{P(I)} \quad (7.8)$$

With $P(I)$, the probability of obtaining an indication of corrosion at the inspection:

$$P(I) = P(I|CI)P(CI) + P(I|\overline{CI})P(\overline{CI}) = 0.8 \cdot 0.01 + 0.1 \cdot (1 - 0.01) = 0.107$$

and $P(I \cap CI)$ the probability of receiving an indication of initiated corrosion and at the same time to have initiated corrosion:

$$P(I \cap CI) = P(I|CI)P(CI) = 0.8 \cdot 0.01 = 0.008 \quad (7.9)$$

Thus, the probability that corrosion of the reinforcement has initiated given an indication of initiated corrosion by the inspection method is:

$$P(CI|I) = \frac{0.008}{0.107} = 0.075 \quad (7.10)$$

The probability of initiated corrosion, given an indication of initiated corrosion, is surprisingly low. This is due to the high probability of an erroneous indication of initiated corrosion by the inspection method relative to the small probability of initiated corrosion (i.e. the inspection method is not sufficiently accurate for the considered application).

2.2. Example 2 - Using Bayes' Rule for Bridge Upgrading

An old reinforced concrete bridge is reassessed in connection with an upgrading of the allowable traffic (see also Schneider¹¹). The concrete compressive strength class is unknown but concrete cylinder samples may be taken from the bridge and tested in the laboratory.

The following classification of the concrete is assumed:

$$\begin{aligned} B_1 &: 0 \leq \sigma_c < 30 \\ B_2 &: 30 \leq \sigma_c < 40 \\ B_3 &: 40 \leq \sigma_c \end{aligned}$$

Where σ_c is the compressive strength of concrete.

Even though the concrete class is unknown, experience with similar bridges suggests that the probability of the concrete of the bridge belonging to class B_1 , B_2 and B_3 is 0.65, 0.24 and 0.11, respectively. This information comprises the prior information - prior to any experiment result.

The test method is not perfect in the sense that even though the test indicates a value of the concrete compressive strength belonging to a certain class, there is a certain probability that the concrete belongs to another class. The likelihoods for the considered test method are given in Table 7.1.

It is assumed that one test is performed and it is found that the concrete compressive strength is equal to 36.2 MPa, i.e. in the interval of class B2.

Using Bayes' rule, the probability that the concrete belongs to one of the different classes may now be updated. The posterior probability that the concrete belongs to class B_2 is given by:

$$P(B_2 | I = B_2) = \frac{0.61 \cdot 0.24}{0.61 \cdot 0.24 + 0.28 \cdot 0.65 + 0.32 \cdot 0.11} = 0.40$$

The posterior probabilities for the other classes may be calculated in a similar manner, the results are given in Table 7.1.

Concrete Class	Prior Probability	Likelihood $P(I B_i)$			Posterior Probability
		$I = B_1$	$I = B_2$	$I = B_3$	
B_1	0.65	0.71	0.28	0.01	0.50
B_2	0.24	0.18	0.61	0.21	0.40
B_3	0.11	0.02	0.32	0.66	0.10

Table 7.1: Summary of prior probabilities, likelihoods of experiment outcomes and posterior probabilities given one test result in the interval of class B_2 .

2.3. Example 3 - Covid Test

Different test exist to indicate a COVID-19 infection. None of these tests are perfect and the uncertainty related to the tests is generally expressed as *sensitivity*, i.e. probability that a test correctly identifies an infection ($\Pr[I_p|P]$). *Specificity* is expressing that a test correctly identifies a non-infection ($\Pr[I_n|N]$).

It is hard to find some general guidelines which values could be taken, as both, sensitivity and specificity are highly dependent on the general constitution of patients and the time-stage of the possible infection.

The following example is thereby intended as indicative only:

Based on the assumption that swabs (molecular) tests correctly identify 70% of those with COVID-19 (sensitivity) and 95% of those without COVID-19 (specificity), the following likelihoods can be quantified:

$$\Pr[I_p|P] = 0.7$$

$$\Pr[I_n|N] = 0.95$$

Task:

1. Consider that you test a person in Trondheim at random. You get a positive test indication (I_p). What is the probability that the person is infected. (Hint: You might estimate the

a priori probability $\Pr'(P)$ from the current infection situation published in the media. Make some sound assumptions.

- Now you test not a random person but a person that shows specific symptoms, such that you would reevaluate the a priori probability $\Pr'(P) = 0.4$.

Discuss the results.

3. Engineering application

Bayesian decision theory, as indicated in Figure 7.1, is about the maximization of expected utility subject to the decision alternatives at hand. In this chapter the theoretical concept is further introduced and illustrated alongside a simplified example. The example follows actually the very classical one first introduced in Benjamin and Cornell in 1969.

Consider a construction site that involves the establishment of a pile foundation. The length of the piles has to be selected subject to uncertain knowledge about the soil condition and composition (namely, the depth of the load bearing rock bed that the pile have to reach is uncertain). The wrong choice of pile length is inducing some additional costs.

3.1. A priori analysis

Consider two possible uncertain conditions of the depth of the rock bed $\theta_0 = 40ft$ and $\theta_1 = 50ft$ and two choices of piles with length $a_0 = 40ft$ and $a_1 = 50ft$. The utility, in this simple example, is expressed in terms of negative costs and the costs assumed for the different combinations of θ and a .

Table 7.2: Costs of the different combinations of θ and a .

	a_0	a_1
θ_0	0	100
θ_1	400	0

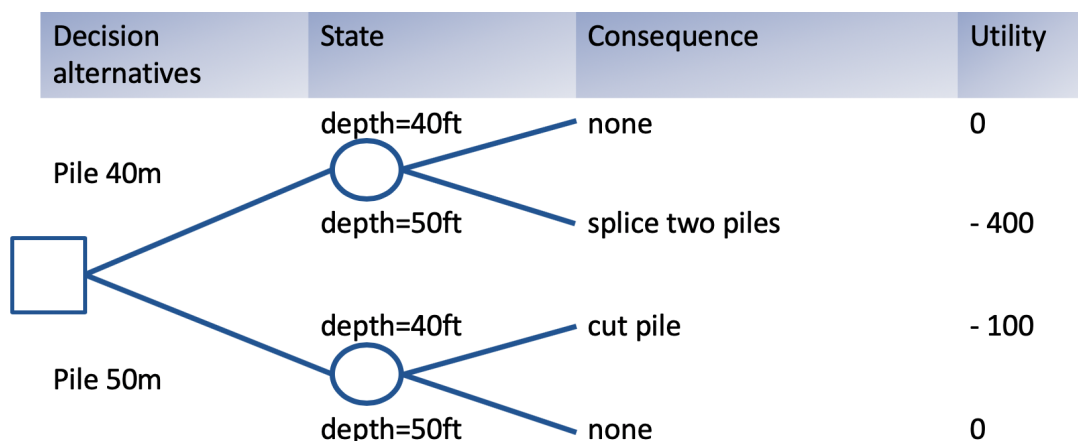


Figure 7.3: Decision tree for the a-priori decision analysis.

In order to employ an a-priori decision analysis the a-priori probabilities for the states θ_0 and θ_1 have to be specified and in this example they are estimated with: $\Pr'(\theta_0) = 0.7$ and $\Pr'(\theta_1) = 0.3$.

The expected utility is then assessed as follows:

$$\begin{aligned} \max_a E'_{\theta} [u(a, \tilde{\theta})] &= \max_a \left[\Pr'(\theta_0)u(\theta_0|a_0) + \Pr'(\theta_1)u(\theta_1|a_0); \right. \\ &\quad \left. \Pr'(\theta_0)u(\theta_0|a_1) + \Pr'(\theta_1)u(\theta_1|a_1) \right] \\ &= \max_a \left[\begin{array}{l} 0.7 \cdot 0 + 0.3 \cdot (-400) \\ 0.7 \cdot (-100) + 0.3 \cdot 0 \end{array} \right] \\ &= \max_a \left[\begin{array}{l} -120 \\ -70 \end{array} \right] = -70 \end{aligned} \quad (7.11)$$

3.2. A posteriori analysis

Suppose that additional test data on the soil condition becomes available. The data carries information about the true state of the soil, but, however, it is none-perfect information. The quality of the test method can be expressed in terms of the likelihoods, i.e. the probabilities to obtain a certain test result conditional on the true state of the soil. Let's consider a ultrasonic test device with three possible outcomes:

- Indication 40ft: z_0
- Indication 50ft: z_1
- Indication 45ft: z_2

The corresponding conditional probability table indicating $\Pr(z|\theta)$ is given as:

	θ_0	θ_1
z_0	0.6	0.1
z_1	0.1	0.7
z_2	0.3	0.2

If the test result indicates z_0 , i.e. 40ft the following updated posteriori probabilities for the states θ_0 and θ_1 can be obtained according to Bayes' Rule:

$$\begin{aligned} \Pr''(\theta_i|z_k) &= \frac{\Pr'(\theta_i) \Pr(z_k|\theta_i)}{\sum_j \Pr'(\theta_j) \Pr(z_k|\theta_j)} \\ \Pr''(\theta_0|z_0) &= \frac{\Pr'(\theta_0) \Pr(z_0|\theta_0)}{\sum_j \Pr'(\theta_j) \Pr(z_0|\theta_j)} \\ &= \frac{0.6 \cdot 0.7}{0.6 \cdot 0.7 + 0.1 \cdot 0.3} = 0.93 \\ \Pr''(\theta_1|z_0) &= \frac{0.1 \cdot 0.3}{0.6 \cdot 0.7 + 0.1 \cdot 0.3} = 0.067 \end{aligned} \quad (7.12)$$

The (updated) posterior probabilities can now be utilized for the maximization of the expected utility.

$$\begin{aligned}
 \max_a E''_{\theta|z} \left[u \left(\mathbf{a}, \tilde{\boldsymbol{\theta}} | \mathbf{z} \right) \right] &= \max_a \begin{bmatrix} \text{Pr}''(\theta_0|z_0)u(\theta_0|a_0) + \text{Pr}''(\theta_1|z_0)u(\theta_1|a_0); \\ \text{Pr}''(\theta_0|z_0)u(\theta_0|a_1) + \text{Pr}''(\theta_1|z_0)u(\theta_1|a_1) \end{bmatrix} \\
 &= \max_a \begin{bmatrix} 0.93 \cdot 0 + 0.067 \cdot (-400) \\ 0.93 \cdot (-100) + 0.067 \cdot 0 \end{bmatrix} \\
 &= \max_a \begin{bmatrix} -26 \\ -93 \end{bmatrix} = -26
 \end{aligned} \tag{7.13}$$

It can be seen that the additional information did increase the maximum expected benefit from -70 to -26. It is also seen that the corresponding decision alternative is now a_0 , the 40ft pile.

3.3. A pre-posteriori analysis

In the previous section it was demonstrated how new, relevant information can be utilized in a decision problem. That the decision benefits from this new information is thereby obvious. It is not the question whether the information should be used or not, of course, as the information was available already.

However, in practical engineering decision situations multiple options for gathering new information exist, e.g. resulting from different measurement techniques or advanced engineering modelling. The relevant information content of the different options is different and also is the price. The decision problem, where it is assessed which choice of possible information gathering technique is to be made based on the corresponding maximization of the expected utility is called the pre-posteriori analysis.

The present example is extended in order to demonstrate the concept as follows. Before any experimental results are available (in the a-priori state) we assess the effect of three different options for additional experiments:

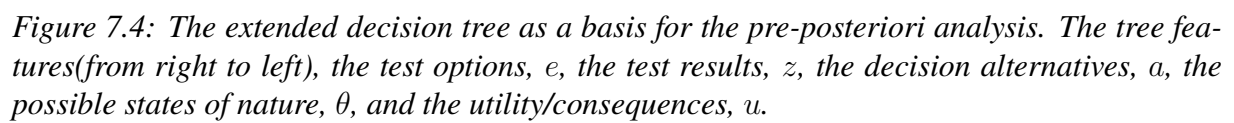
e_0 : No experiments - no additional costs.

e_1 : Ultrasonic tests - 20 MU.

e_2 : Probe - 80 MU.

The decision tree is extended with the corresponding options and the utility of the different end branches of the tree are indicated in Figure 7.4.

The pre-posteriori analysis evaluates (and maximizes) the expected utilities of the different test options e in terms of $\max_e E_{z|e} \left[\max_a E''_{\theta|z} \left[u \left(\mathbf{e}, \tilde{\mathbf{z}}, \mathbf{a}, \tilde{\boldsymbol{\theta}} \right) \right] \right]$. The term is solved by considering the decision tree in Figure 7.4 from the right to the left:



1. Assessment of the event probabilities $\Pr(\theta|e, z)$ (by using Bayes Rule):

$$\begin{aligned}
\Pr'(\theta_0|e_0) &= 0.7 & \Pr'(\theta_1|e_0) &= 0.3 \text{ (a-priori probabilities)} \\
\Pr''(\theta_0|e_1, z_0) &= 0.93 & \Pr''(\theta_1|e_1, z_0) &= 0.067 \\
\Pr''(\theta_0|e_1, z_1) &= 0.25 & \Pr''(\theta_1|e_1, z_1) &= 0.75 \\
\Pr''(\theta_0|e_1, z_2) &= 0.78 & \Pr''(\theta_1|e_1, z_2) &= 0.22 \\
\Pr''(\theta_0|e_2, z_0) &= 1 & \Pr''(\theta_1|e_2, z_0) &= 0 \\
\Pr''(\theta_0|e_2, z_1) &= 0 & \Pr''(\theta_1|e_2, z_1) &= 1
\end{aligned}$$

2. Computation of the expected utilities

$$E[u|e, z, a] = \sum_i \Pr(\theta_i|e, z) \cdot u[\theta_i|e, z, a]$$

, i.e. conditional on the chosen test, e , the observed test outcome, z and the decision on the length of the pile, a .

$\Pr'(\theta_0 e_0)$	$\cdot u[\theta_0 e_0, a_0]$	$+ \Pr'(\theta_1 e_0)$	$\cdot u[\theta_1 e_0, a_0]$	$= 0.7 \cdot 0 + 0.3 \cdot (-400)$	$= -120$
$\Pr'(\theta_0 e_0)$	$\cdot u[\theta_0 e_0, a_1]$	$+ \Pr'(\theta_1 e_0)$	$\cdot u[\theta_1 e_0, a_1]$	$= 0.7 \cdot (-100) + 0.3 \cdot 0$	$= -70$
$\Pr''(\theta_0 e_1, z_0)$	$\cdot u[\theta_0 e_1, z_0, a_0]$	$+ \Pr''(\theta_1 e_1, z_0)$	$\cdot u[\theta_1 e_1, z_0, a_0]$	$= 0.93 \cdot (-20) + 0.067 \cdot (-420)$	$= -48$
$\Pr''(\theta_0 e_1, z_0)$	$\cdot u[\theta_0 e_1, z_0, a_1]$	$+ \Pr''(\theta_1 e_1, z_0)$	$\cdot u[\theta_1 e_1, z_0, a_1]$	$= 0.93 \cdot (-120) + 0.067 \cdot (-20)$	$= -113$
$\Pr''(\theta_0 e_1, z_1)$	$\cdot u[\theta_0 e_1, z_1, a_0]$	$+ \Pr''(\theta_1 e_1, z_1)$	$\cdot u[\theta_1 e_1, z_1, a_0]$	$= 0.25 \cdot (-20) + 0.75 \cdot (-420)$	$= -320$
$\Pr''(\theta_0 e_1, z_1)$	$\cdot u[\theta_0 e_1, z_1, a_1]$	$+ \Pr''(\theta_1 e_1, z_1)$	$\cdot u[\theta_1 e_1, z_1, a_1]$	$= 0.25 \cdot (-120) + 0.75 \cdot (-20)$	$= -45$
$\Pr''(\theta_0 e_1, z_2)$	$\cdot u[\theta_0 e_1, z_2, a_0]$	$+ \Pr''(\theta_1 e_1, z_2)$	$\cdot u[\theta_1 e_1, z_2, a_0]$	$= 0.78 \cdot (-20) + 0.22 \cdot (-420)$	$= -112$
$\Pr''(\theta_0 e_1, z_2)$	$\cdot u[\theta_0 e_1, z_2, a_1]$	$+ \Pr''(\theta_1 e_1, z_2)$	$\cdot u[\theta_1 e_1, z_2, a_1]$	$= 0.78 \cdot (-120) + 0.22 \cdot (-40)$	$= -97$
$\Pr''(\theta_0 e_2, z_0)$	$\cdot u[\theta_0 e_2, z_0, a_0]$	$+ \Pr''(\theta_1 e_2, z_0)$	$\cdot u[\theta_1 e_2, z_0, a_0]$	$= 1 \cdot (-50) + 0 \cdot (-450)$	$= -50$
$\Pr''(\theta_0 e_2, z_0)$	$\cdot u[\theta_0 e_2, z_0, a_1]$	$+ \Pr''(\theta_1 e_2, z_0)$	$\cdot u[\theta_1 e_2, z_0, a_1]$	$= 1 \cdot (-150) + 0 \cdot (-50)$	$= -150$
$\Pr''(\theta_0 e_2, z_1)$	$\cdot u[\theta_0 e_2, z_1, a_0]$	$+ \Pr''(\theta_1 e_2, z_1)$	$\cdot u[\theta_1 e_2, z_1, a_0]$	$= 0 \cdot (-50) + 1 \cdot (-450)$	$= -450$
$\Pr''(\theta_0 e_2, z_1)$	$\cdot u[\theta_0 e_2, z_1, a_1]$	$+ \Pr''(\theta_1 e_2, z_1)$	$\cdot u[\theta_1 e_2, z_1, a_1]$	$= 0 \cdot (-150) + 1 \cdot (-50)$	$= -50$

3. Conditional on the type of experiment, e , and the corresponding outcome, z , which decision would be taken, i.e.:

$\arg \max_a E_{\theta z} [u(a e_0)]$	$= a_1$	with $E_{\theta z} [u(a_1 e_0)]$	$= -70$
$\arg \max_a E_{\theta z} [u(a e_1, z_0)]$	$= a_0$	with $E_{\theta z} [u(a_0 e_1, z_0)]$	$= -48$
$\arg \max_a E_{\theta z} [u(a e_1, z_1)]$	$= a_1$	with $E_{\theta z} [u(a_1 e_1, z_1)]$	$= -45$
$\arg \max_a E_{\theta z} [u(a e_1, z_2)]$	$= a_1$	with $E_{\theta z} [u(a_1 e_1, z_2)]$	$= -97$
$\arg \max_a E_{\theta z} [u(a e_2, z_0)]$	$= a_0$	with $E_{\theta z} [u(a_0 e_2, z_0)]$	$= -50$
$\arg \max_a E_{\theta z} [u(a e_2, z_1)]$	$= a_1$	with $E_{\theta z} [u(a_1 e_2, z_1)]$	$= -50$

4. Evaluate the probability of observing a test result z , $\Pr(z|e)$.

$$\begin{aligned}
\Pr(z_0|e_1) &= \Pr(z_0|e_1, \theta_0) \Pr'(\theta_0) + \Pr(z_0|e_1, \theta_1) \Pr'(\theta_1) = 0.6 \cdot 0.7 + 0.1 \cdot 0.3 = 0.45 \\
\Pr(z_1|e_1) &= \Pr(z_1|e_1, \theta_0) \Pr'(\theta_0) + \Pr(z_1|e_1, \theta_1) \Pr'(\theta_1) = 0.1 \cdot 0.7 + 0.7 \cdot 0.3 = 0.28 \\
\Pr(z_2|e_1) &= \Pr(z_2|e_1, \theta_0) \Pr'(\theta_0) + \Pr(z_2|e_1, \theta_1) \Pr'(\theta_1) = 0.3 \cdot 0.7 + 0.2 \cdot 0.3 = 0.27 \\
\Pr(z_0|e_2) &= \Pr(z_0|e_2, \theta_0) \Pr'(\theta_0) + \Pr(z_0|e_2, \theta_1) \Pr'(\theta_1) = 1 \cdot 0.7 + 0 \cdot 0.3 = 0.7 \\
\Pr(z_1|e_2) &= \Pr(z_1|e_2, \theta_0) \Pr'(\theta_0) + \Pr(z_1|e_2, \theta_1) \Pr'(\theta_1) = 0 \cdot 0.7 + 1 \cdot 0.3 = 0.3
\end{aligned}$$

5. Evaluate the expected utility per experiment option $E[u|e]$:

$$\begin{aligned}
E[u|e_1] &= & &= -70 \\
E[u|e_2] &= \Pr(z_0|e_1)E_{\theta|z}[u(a_0|e_1, z_0)] \\
&\quad + \Pr(z_1|e_1)E_{\theta|z}[u(a_0|e_1, z_1)] \\
&\quad + \Pr(z_2|e_1)E_{\theta|z}[u(a_0|e_1, z_2)] \\
&= 0.45 \cdot (-48) + 0.28 \cdot (-45) + 0.27 \cdot (-97) &= -60 \\
E[u|e_3] &= \Pr(z_0|e_3)E_{\theta|z}[u(a_0|e_3, z_0)] \\
&\quad + \Pr(z_1|e_3)E_{\theta|z}[u(a_0|e_3, z_1)] \\
&= 0.7 \cdot (-80) + 0.3 \cdot (-80) &= -80
\end{aligned}$$

The results are summarized in Figure 7.5.

3.4. Summary

Bayesian decision theory seems to be a valid framework for rational engineering decision making subject to uncertainty. Three types of decision analysis have been introduced and their practical application have been demonstrated. The example was kept simple in order to demonstrate the concept transparently, i.e. only discrete states of nature and discrete decision options and experimental outcomes have been considered. If some or all of these features are represented by continuous variables the summations turn into integrals and the solution of the problem becomes computationally demanding.

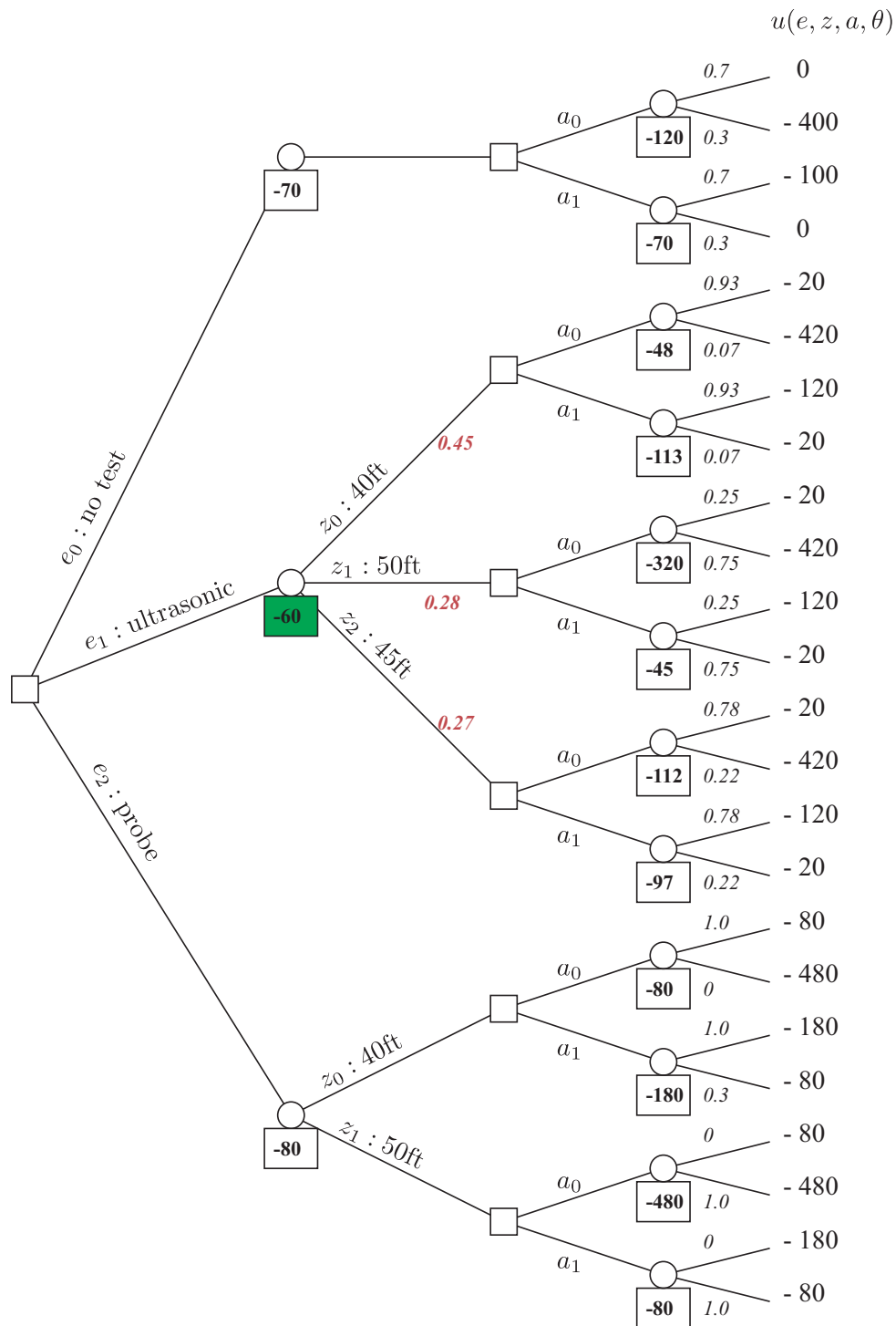


Figure 7.5: The extended decision tree as a basis for the pre-posterior analysis together with results. The optimal choice of experiment is e_1 ultrasonic test.

Question::

Engineering Decision Making

- Describe the three principle different types of decision analysis?
 - How is information used?
 - What are practical examples?
- Explain the rule of Bayes alongside a practical example.
- What are the different steps of a pre-posteriori decision analysis?

Chapter 8

Engineering uncertainty quantification

This chapter is not part of the pensum for the Exam 2021.

How do we translate information into probabilistic models? In the present chapter the very basics of engineering uncertainty quantification are addressed. It is demonstrated how the parameters of random variables can be estimated based on observations / data.

1. Engineering uncertainty quantification

1.1. Cumulative Distribution and Probability Density Functions

A random variable, which can take on any value, is called a *continuous random variable*. The probability that such a random variable takes on a specific value is zero. The probability that a continuous random variable, X , is less than or equal to a value, x , is given by the *cumulative distribution function*:

$$F_X(x) = P(X \leq x) \quad (8.1)$$

Generally speaking, capital letters denote a random variable and lowercase denote an outcome or a realization of a random variable. Figure 8.1 illustrates an example of a continuous cumulative distribution function. For continuous random variables, the *probability density function* is given by:

$$f_X(x) = \frac{dF(x)}{dx} \quad (8.2)$$

The probability of an outcome in the interval $[x; x + dx]$ where dx is small, is given by $P(X \in [x; x + dx]) = f_X(x) dx$.

Random variables with a finite or infinite countable sample space are called *discrete random variables*. For discrete random variables, the cumulative distribution function is given as:

$$P_X(x) = \sum_{x_i < x} p_X(x_i) \quad (8.3)$$

where $p_X(x_i)$ is the probability density function given as:

$$p_X(x_i) = P(X = x_i) \quad (8.4)$$

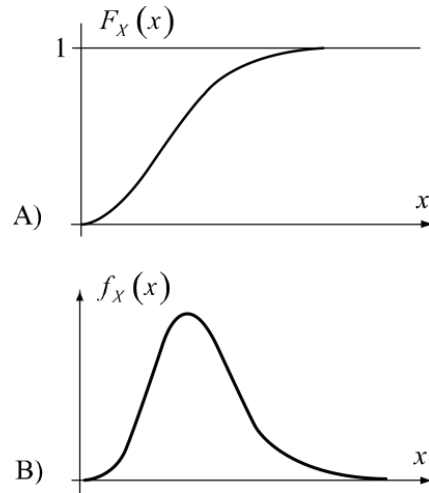


Figure 8.1: Illustration for a continuous random variable of A) a cumulative distribution function and B) a probability density function.

A discrete cumulative distribution function and probability density function is illustrated in Figure 8.2.

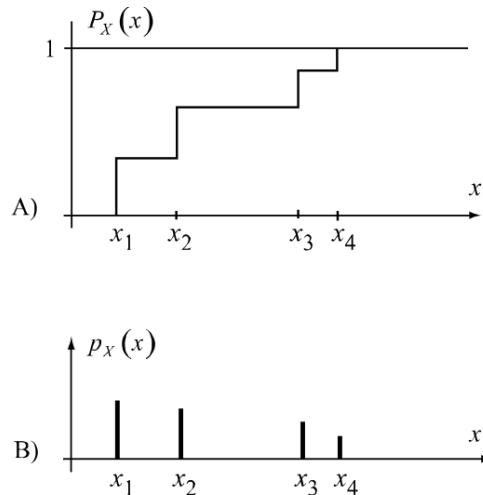


Figure 8.2: Illustration for a discrete random variable of A) a cumulative distribution function and B) a probability density function.

1.1.1. Moments of Random Variables and the Expectation Operator

Probability distributions may be defined in terms of their *parameters* or *moments*. Cumulative distribution functions and probability density functions are often written as $F_X(x, p)$ and $f_X(x, p)$, respectively, to indicate the parameters p (or moments) defining the functions.

The i^{th} moment m_i of a continuous random variable is defined by:

$$m_i = \int_{-\infty}^{\infty} x^i f_X(x) dx \quad (8.5)$$

and for a discrete random variable by:

$$m_i = \sum_{j=1}^n x_j^i p_X(x_j) \quad (8.6)$$

The mean (or *expected value*) of continuous and discrete random variables X is defined accordingly as the *first moment*, i.e.:

$$\mu_X = E[X] = \int_{-\infty}^{\infty} x f_X(x) dx \quad (8.7)$$

$$\mu_X = E[X] = \sum_{j=1}^n x_j p_X(x_j) \quad (8.8)$$

where $E[\cdot]$ denotes the expectation operator.

Similarly, the *variance* σ_X^2 is described by the *second central moment*, i.e.:

$$\sigma_X^2 = Var[X] = E[(X - \mu_X)^2] = \int_{-\infty}^{\infty} (x - \mu_X)^2 f_X(x) dx \quad (8.9)$$

$$\sigma_X^2 = Var[X] = \sum_{j=1}^n (x_j - \mu_X)^2 p_X(x_j) \quad (8.10)$$

where $Var[X]$ denotes the *variance* of X .

The ratio between the standard deviation σ_X and the expected value μ_X of a random variable X is denoted the *coefficient of variation* $CoV[X]$ and is given by:

$$CoV[X] = \frac{\sigma_X}{\mu_X} \quad (8.11)$$

The coefficient of variation provides a useful descriptor of the variability of a random variable around its expected value.

1.1.2. (Some relevant) Probability Density and Distribution Functions

In Table 8.1, a selection of probability density and cumulative distribution functions is given with the definition of their distribution parameters and moments. Note that the probability density and cumulative distribution functions are given as conditional distributions (conditional on the parameters), i.e. $F_X(x|\theta)$ and $f_X(x|\theta)$, where θ is the vector of parameters.

Table 8.1: Selection of probability distributions.

Distribution type	Param.	Mean and Standard deviation
<i>Uniform</i> $a \leq x \leq b$ $f_X(x a, b) = \frac{1}{b-a}$ $F_X(x a, b) = \frac{x-a}{b-a}$	u $a > 0$	$\mu = u + \frac{a+b}{2}$ $\sigma = \frac{b-a}{\sqrt{12}}$
<i>Normal</i> $f_X(x \mu, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{x-\mu}{\sigma}\right)^2\right)$ $F_X(x \mu, \sigma) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^x \exp\left(-\frac{1}{2}\left(\frac{t-\mu}{\sigma}\right)^2\right) dt$	μ $\sigma > 0$	μ σ
<i>Shifted Log-Normal</i> $x > \varepsilon$ $f_X(x \lambda, \zeta, \varepsilon) = \frac{1}{(x-\varepsilon)\zeta\sqrt{2\pi}} \exp\left(-\frac{1}{2}\left(\frac{\ln(x-\varepsilon)-\lambda}{\zeta}\right)^2\right)$ $F_X(x \lambda, \zeta, \varepsilon) = \Phi\left(\frac{\ln(x-\varepsilon)-\lambda}{\zeta}\right)$	λ $\zeta > 0$ ε	$\mu = \varepsilon + \exp\left(\lambda + \frac{\zeta^2}{2}\right)$ $\sigma = e^{\lambda + \frac{\zeta^2}{2}} \sqrt{e^{\zeta^2} - 1}$
<i>Shifted Exponential</i> $x \geq \varepsilon$ $f_X(x \lambda, \varepsilon) = \lambda \exp(-\lambda(x-\varepsilon))$ $F_X(x \lambda, \varepsilon) = 1 - \exp(-\lambda(x-\varepsilon))$	ε $\lambda > 0$	$\mu = \varepsilon + \frac{1}{\lambda}$ $\sigma = \frac{1}{\lambda}$
<i>Gamma</i> $x \geq 0$ $f_X(x b, p) = \frac{b^p}{\Gamma(p)} \exp(-bx) x^{p-1}$ $F_X(x b, p) = \frac{\Gamma(bx, p)}{\Gamma(p)}$ Where: $\Gamma(bx, p) = \int_0^b x t^{p-1} \exp(-t) dt$ $\Gamma(p) = \int_0^\infty t^{p-1} \exp(-t) dt$	$p > 0$ $b > 0$	$\mu = \frac{p}{b}$ $\sigma = \frac{\sqrt{p}}{b}$

<i>Beta</i>	a	$\mu = a + (b - a) \frac{r}{r + t}$
$a \leq x \leq b$	b	$\sigma = \frac{b - a}{r + t} \sqrt{\frac{rt}{r + t + 1}}$
$f_X(x a, b, t, r) = \frac{\Gamma(r + t)}{\Gamma(r)\Gamma(t)} \int_a^x \frac{(u - a)^{r-1}(b - u)^{t-1}}{(b - a)^{r+t-1}} du$	$r > 0$	
	$t > 0$	

The relevance of the different distribution functions given in Table 8.1, in connection with the probabilistic modelling of uncertainties in engineering risk and reliability analysis, is strongly case dependent. The reader is suggested to consult the specific literature for specific guidance on the application. Nevertheless, a few specific remarks will be provided in the following for the commonly applied Normal distribution and the Log-Normal distribution.

1.1.3. The Normal Distribution

The Normal probability distribution follows from the central limit theorem as a result of the sum of independent (or almost) random variables. Thus, it is applied very frequently in practical problems for the probabilistic modelling of uncertain phenomena which may be considered to originate from a cumulative effect of several independent uncertain contributions.

The Normal distribution has the property that the linear combination S of n Normal distributed random variables X_i , $i = 1, 2, \dots, n$, cf. Eq. (8.12), is also Normal distributed. The distribution is said to be *closed* regarding summation.

$$S = a_0 + \sum_{i=1}^n a_i X_i \quad (8.12)$$

One special version of the Normal distribution should be mentioned, namely the *Standard Normal distribution*. In general, a standardized (sometimes referred to as a reduced) random variable is a random variable which has been transformed such that it has an expected value equal to zero and a variance equal to one:

$$Y = \frac{X - \mu_X}{\sigma_X} \quad (8.13)$$

Y is a standardized random variable. If the random variable X follows a Normal distribution, then the random variable Y is standard Normal distributed. In Figure 8.3, the process of standardization is illustrated.

It is common practice to denote the cumulative distribution function for the standard Normal distribution by $\Phi(x)$ and the corresponding density function by $\varphi(x)$. These functions are broadly available in software packages such as Microsoft[®] Excel, Matlab[®] and Python.

1.1.4. The Log-Normal Distribution

A random variable Y is said to be *Log-Normal distributed* if the variable $Z = \ln(Y)$ is Normal distributed. It follows that if an uncertain phenomenon can be assumed to originate from a multiplicative effect of several uncertain contributions, then the probability distribution for the phenomenon can be assumed to be Log-Normal distributed.

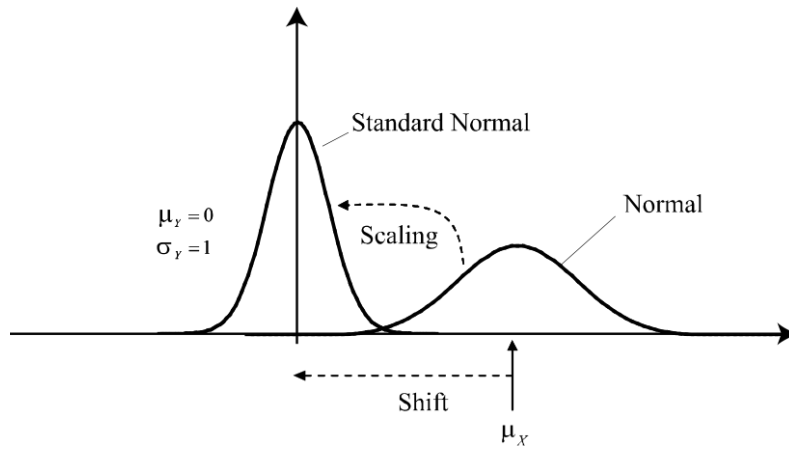


Figure 8.3: Illustration of the relationship between a Normal distributed random variable and a standard Normal distributed random variable.

The Log-Normal distribution has the property that if Eq. (8.14) is satisfied and all Y_i are independent Log-Normal random variables with parameters λ_i , ζ_i and $\varepsilon_i = 0$ (as given in the table above), then also P is Log-Normal with parameters:

$$P = \prod_{i=1}^n Y_i^{a_i} \quad (8.14)$$

$$\lambda_p = \sum_{i=1}^n a_i \lambda_i \quad (8.15)$$

$$\zeta_p^2 = \sum_{i=1}^n a_i^2 \zeta_i^2 \quad (8.16)$$

1.1.5. Properties of the Expectation Operator

The *expectation operation* has the following properties, where a , b and c are constants and X is a random variable:

$$\begin{aligned} E[c] &= c \\ E[cX] &= cE[X] \\ E[a + bX] &= a + bE[X] \\ E[g_1(X) + g_2(X)] &= E[g_1(X)] + E[g_2(X)] \end{aligned} \quad (8.17)$$

The implication of the last equation is that expectation, like differentiation or integration, is a linear operation. This property is useful since it can be used, for instance, to find the following formula for the variance of a random variable X in terms of more easily calculated quantities:

$$\begin{aligned} Var[X] &= E[(X - \mu_X)^2] = E[X^2 + \mu_X^2 - 2\mu_X X] \\ &= \mu_X^2 + E[X^2] - 2\mu_X E[X] = \mu_X^2 + E[X^2] - 2\mu_X^2 = E[X^2] - \mu_X^2 \end{aligned} \quad (8.18)$$

By application of Eq. (8.18), the following properties of the *variance operator* $Var[\cdot]$ can easily be derived:

$$\begin{aligned} \text{Var}[c] &= 0 \\ \text{Var}[cX] &= c^2 \text{Var}[X] \\ \text{Var}[a + bX] &= b^2 \text{Var}[X] \end{aligned} \quad (8.19)$$

where a , b and c are constants and X is a random variable.

From Eq. (8.18), it is furthermore seen that, in general, $E[g(X)] \neq g(E[X])$. In fact, for convex functions $g(x)$, it can be shown that the following inequality is valid (*Jensen's inequality*):

$$E[g(X)] \leq g(E[X]) \quad (8.20)$$

where the equality holds if $g(X)$ is linear.

Whether the cumulative distribution and density function are defined by their moments or by parameters is a matter of convenience and it is generally possible to establish the one from the other.

1.1.6. Random Vectors and Joint Moments (not part of pensum)

If an n -dimensional vector of continuous random variables $\mathbf{X} = (X_1, X_2, \dots, X_n)^T$ is considered, the *joint cumulative distribution function* is given by:

$$F_{\mathbf{X}}(\mathbf{x}) = P(X_1 \leq x_1 \cap X_2 \leq x_2 \cap \dots \cap X_n \leq x_n) \quad (8.21)$$

and the *joint probability density function* is:

$$f_{\mathbf{X}}(\mathbf{x}) = \frac{\partial^n}{\partial x_1 \partial x_2 \dots \partial x_n} F_{\mathbf{X}}(\mathbf{x}) \quad (8.22)$$

The covariance $C_{X_i X_j}$ between X_i and X_j is defined by:

$$\begin{aligned} C_{X_i X_j} &= E[(X_i - \mu_{X_i})(X_j - \mu_{X_j})] \\ &= \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} (x_i - \mu_{X_i})(x_j - \mu_{X_j}) f_{X_i X_j}(x_i, x_j) dx_i dx_j \end{aligned} \quad (8.23)$$

and is also called the *joint central moment* between the variables X_i and X_j .

The covariance expresses the linear dependence between two variables. On the basis of the covariance the *correlation coefficient* is defined by:

$$\rho_{X_i X_j} = \frac{C_{X_i X_j}}{\sigma_{X_i} \sigma_{X_j}} \quad (8.24)$$

It is evident that $C_{X_i X_i} = \text{Var}[X_i]$, i.e. $\rho_{X_i X_i} = 1$. The correlation coefficients can only take values in the interval $[-1; 1]$. A negative correlation coefficient between two random variables implies that, if the outcome of one variable is large compared to its mean value, then the outcome of the other variable is likely to be small compared to its mean value. A positive correlation coefficient between two variables implies that, if the outcome of one variable is large compared to its mean value, then the outcome of the other variable is also likely to be large compared to its mean value. If two variables are independent, their correlation coefficient is zero and the *joint density function* is the product of the 1-dimensional density functions. In many cases, it is possible to obtain a sufficiently accurate approximation to the *n-dimensional cumulative*

distribution function from the 1-dimensional distribution functions of the n variables and their parameters, and the correlation coefficients.

If Y is a linear function of the random vector $\mathbf{X} = (X_1, X_2, \dots, X_n)^T$, i.e.:

$$Y = a_0 + \sum_{i=1}^n a_i X_i \quad (8.25)$$

using Eqs. (8.18), (8.19) and (8.23) it can be shown that the expected value $E[Y]$ and the variance $Var[Y]$ are given by:

$$\begin{aligned} E[Y] &= a_0 + \sum_{i=1}^n a_i E[X_i] \\ Var[Y] &= \sum_{i=1}^n a_i^2 Var[X_i] + 2 \left(\sum_{i,j=1; i \neq j}^n a_i a_j C_{X_i X_j} \right) \end{aligned} \quad (8.26)$$

1.1.7. Conditional Distributions and Conditional Moments (not part of pensum)

The *conditional probability density function* for the random variable X_1 , conditional on the outcome of the random variable X_2 is denoted $f_{X_1|X_2}(x_1|x_2)$ and defined, in accordance with the definition of conditional probability given previously, by:

$$f_{X_1|X_2}(x_1|x_2) = \frac{f_{X_1, X_2}(x_1, x_2)}{f_{X_2}(x_2)} \quad (8.27)$$

As for the case when probabilities of events were considered, two random variables X_1 and X_2 are said to be independent when:

$$f_{X_1|X_2}(x_1|x_2) = f_{X_1}(x_1) \quad (8.28)$$

By integration of Eq. (8.27) the conditional cumulative distribution $F_{X_1|X_2}(x_1|x_2)$ is obtained:

$$F_{X_1|X_2}(x_1|x_2) = \frac{\int_{-\infty}^{x_1} f_{X_1, X_2}(z, x_2) dz}{f_{X_2}(x_2)} \quad (8.29)$$

and finally, by integration of Eq. (8.29) weighed with the probability density function of X_2 $f_{X_2}(x_2)$, the unconditional cumulative distribution $F_{X_1}(x_1)$ is achieved by applying the *total probability theorem*:

$$F_{X_1}(x_1) = \int_{-\infty}^{\infty} F_{X_1|X_2}(x_1|x_2) f_{X_2}(x_2) dx_2 \quad (8.30)$$

The *conditional moments* of jointly distributed continuous random variables follow straightforwardly from Eq. (8.7), by use of Eq. (8.28). If X_1, X_2 are jointly distributed random variables, the conditional expected value $\mu_{X_1|X_2}$ of X_1 given X_2 is evaluated by:

$$\mu_{X_1|X_2} = E[X_1|X_2 = x_2] = \int_{-\infty}^{\infty} x f_{X_1|X_2}(x|x_2) dx \quad (8.31)$$

1.2. Sampling and representation of data

1.2.1. Sampling of data

Data is either observed directly from natural phenomena, or from controlled experiments. The information content in the data is relevant for the specification of the variable of interest if the character and origin of the magnitude and variability in the data is in correspondence with the magnitude and variability of the variable to be represented in the engineering problem at hand. In this regard, it is often stated that the “data has to be representative”.

Example/Question: The distribution function for the concrete compression strength of the strength grade C20/25 has to be represented for general structural design in Norway. Test data from one producer is available for the specification of the distribution parameters. Explain, why the data is not representative.

1.2.2. Describing data

In the following *numerical summaries* will first be introduced. These can be considered to be numerical characteristics of the observed data containing important information about the data and the nature of uncertainty associated with these. These are also referred to as *sample characteristics*. Thereafter *graphical representations* are introduced as means of visual characterisation and as a useful tool for data analysis. Descriptive statistics play an important role in engineering risk analysis as this forms a standardized basis for assessing and documenting data obtained for the purpose of understanding and representing uncertainties in risk assessment.

Numerical Summaries

One of the most useful numerical summaries is the *sample mean*. If the data set is collected in the vector $\hat{\mathbf{x}} = (\hat{x}_1, \hat{x}_2, \dots, \hat{x}_n)^T$ the sample mean $\bar{x}$ is simply given as:

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n \hat{x}_i \quad (8.32)$$

The sample mean may be interpreted as a central value of the data set. If, on the basis of the data set, one should give only one value characterising the data, one would normally use the sample mean. Another central measure is the mode of the data set i.e. the most frequently occurring value in the data set. When data samples are real values, the *mode* in general cannot be assessed numerically, but may be assessed from graphical representations of the data as will be illustrated below in connection to histograms.

It is often convenient to work with an ordered data set which is readily established by rearranging the original data set $\hat{\mathbf{x}} = (\hat{x}_1, \hat{x}_2, \dots, \hat{x}_n)^T$ such that the data are arranged in increasing order as $\hat{x}_1^o \leq \hat{x}_2^o \leq \dots \leq \hat{x}_i^o \dots \leq \hat{x}_{n-1}^o \leq \hat{x}_n^o$. The i^{th} value of an ordered data set is denoted by $\hat{x}_i^o$.

The *median* of the data set is defined as the middle value in the ordered list of data if n is

odd. If n is even the median is taken as the average value of the two middle values (see also the examples below).

Example - Concrete Compressive Strength Data

Consider the data set given in Figure 8.4 corresponding to concrete cube compressive strength measurements. In the table the data are listed both unordered, e.g. in the order they were observed and ordered according to increasing values.

The sample mean for the data set is readily evaluated using Eq. (8.32) and found to be equal to 32.67 MPa. All the observed values are different and therefore the mode cannot be determined without dividing the observations into intervals as will be shown below. However, the median is readily determined as being equal to 33.05 MPa.

i	unordered $\hat{x}_i$	ordered $\hat{x}_i^o$
1	35.8	24.4
2	39.2	27.6
3	34.6	27.8
4	27.6	27.9
5	37.1	28.5
6	33.3	30.1
7	32.8	30.3
8	34.1	31.7
9	27.9	32.2
10	24.4	32.8
11	27.8	33.3
12	33.5	33.5
13	35.9	34.1
14	39.7	34.6
15	28.5	35.8
16	30.3	35.9
17	31.7	36.8
18	32.2	37.1
19	36.8	39.2
20	30.1	39.7

Figure 8.4: Concrete Compressive Strength

The variability or the *dispersion* of the data set around the sample mean is also an important characteristic of the data set. This dispersion may be characterised by the *sample variance* s^2 given by:

$$s^2 = \frac{1}{n} \sum_{i=1}^n (\hat{x}_i - \bar{x})^2 \quad (8.33)$$

The *sample standard deviation* s is defined as the square root of the sample variance. From Eq. (8.33) it is seen that the sample standard deviation s is assessed in terms of the variability of the observations around the sample mean value $\bar{x}$.

Thus, the sample variance is the mean of the squared deviations from the sample mean and is in this way analogous to the moment of inertia as used in structural engineering.

As a means of comparison of the dispersion of different data sets, the dimensionless *sample*

coefficient of variation ν is convenient. The sample coefficient of variation ν is defined as the ratio of the sample standard deviation to the sample mean, i.e. given by:

$$\nu = \frac{s}{\bar{x}} \quad (8.34)$$

The *sample variance* for the concrete cube compressive strengths of Figure 8.4 may be evaluated using Eq. (8.33) and is found to be 16.36 MPa². The sample standard deviation is thus 4.04 MPa. For the considered concrete cube compressive strength data the sample coefficient of variation is equal to 0.12.

Whereas the sample mean, mode and median are central measures of a data set and the sample variance is a measure of the dispersion around the sample mean it is sometimes also useful to have some characteristic indicating the degree of symmetry of the data set. To this end, we introduce the sample coefficient of *skewness*, which is a simple logical extension of the sample variance. The sample coefficient of skewness η is defined as:

$$\eta = \frac{1}{n} \frac{\sum_{i=1}^n (\hat{x}_i - \bar{x})^3}{s^3} \quad (8.35)$$

This coefficient is positive if the mode of the data set is less than its mean value (skewed to the right) and negative if the mode is larger than the mean value (skewed to the left). For the concrete cube compressive strengths (Figure 8.4) the sample coefficient of skewness is - 0.12.

In a similar way the sample coefficient of *kurtosis* κ is defined as:

$$\kappa = \frac{1}{n} \frac{\sum_{i=1}^n (\hat{x}_i - \bar{x})^4}{s^4} \quad (8.36)$$

which is a measure of how closely the data are distributed around the mode (peakedness). Typically one would compare the sample coefficient of kurtosis to that of a normal distribution (introduced in Module D), which is equal to 3.0. The kurtosis for the concrete cube compressive strength (Figure 8.4) is evaluated as equal to 2.23, i.e. the considered data set is less peaked than the normal distribution.

Sample Moments and Sample Central Moments

In Eqs. (8.33), (8.35) and (8.36) the sample central moments are used. These are denoted as “central” as they are always taken relative to the mean. The equations show that the sample mean $\bar{x}$ is subtracted from the particular observations $\hat{x}_i$. In general the central sample moments can be defined as:

$$z_j = \frac{1}{n} \sum_{i=1}^n (\hat{x}_i - \bar{x})^j \quad (8.37)$$

In Eqs. (8.35) and (8.36) the 3rd and 4th central sample moment are divided by the quadratic standard deviation to assess the skewness and the kurtosis. The “simple” (non-central) sample

moments do not use the sample mean in their definition. Hence, the general definition is:

$$m_j = \frac{1}{n} \sum_{i=1}^n \hat{x}_i^j \quad (8.38)$$

Graphical Representations

Graphical representations provide a convenient and strong basis for assessing data and to communicate these to other persons. There exist a relatively large number of different possible graphical representations of data, of which some are better suited than others depending on the purpose of the representations. Some are better for representing the characteristics of data sets containing observations of one characteristic, like e.g. the concrete compressing strength and others are better for representing the characteristics of two or more data sets (e.g. the simultaneously observed material properties). In the following, the most frequently applied graphical representations are introduced and discussed with the help of examples.

One-Dimensional Scatter Diagrams

The simplest graphical representation is the *scatter diagram* which provides a means to represent observations contained in one or more data sets. The scatter diagram may be constructed by plotting the observed values of the data set along an axis labeled according to the scale of the observations. In a *one-dimensional scatter diagram* the minimum and maximum values of the data set can be readily observed. Furthermore, as long as the number of data points is not very large, the central value of the observed data may be observed directly from the plot. In the case where a data set contains a large number of data, some of these may be overlapping and this makes it difficult to distinguish the individual observations. In such cases it may be beneficial to apply another graphical representation such as histograms, as described subsequently.

Consider again the concrete strength data from Figure 8.4. A one-dimensional scatter diagram can be produced by plotting the data along one axis. In Figure 8.5 the resulting scatter diagram is shown. It can be seen from Figure 8.5 that the lowest value of the data lies close to 24 MPa while the highest lies close to 40 MPa. Moreover, a high concentration of observations is observed in the range 32 to 34, indicating that the mode of this data is in that range.

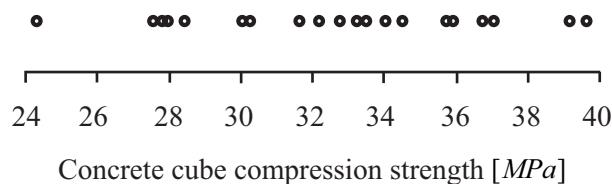


Figure 8.5: Scatter plot of the Concrete Compressive Strength

Histograms

A frequently applied graphical representation of data sets is the *histogram*. Consider again, as an example, the concrete strength data from Figure 8.4. The data are further processed and the observed strength values are subdivided into intervals. For each interval the number of observations within each interval is counted. Thereafter, the *frequencies* of the measurements

within each interval are evaluated as the number of observations within one interval divided by the total number of observations. Figure 8.6 show the graphical representation of the processed data of Figure 8.4. Both, the absolute and the relative amount of observations per interval is indicated.

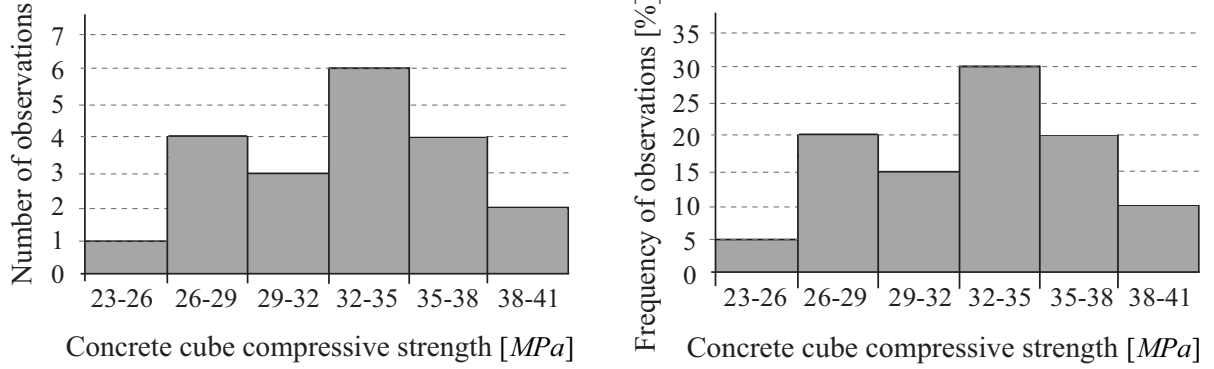


Figure 8.6: Histogram of the Concrete Compressive Strength, absolute (left) and relative(right) associations to the intervals.

It has to be noted that a histogram may reduce the information provided by the data examined. The interval width plays an important role for the resolution of the representation of the observations. It is suggested to subdivide the interval between the maximum and minimum value into k intervals where k is given by:

$$k = 1 + 3.3 \log_{10}(n) \quad (8.39)$$

where n is the number of data points in the data set.

Quantile Plots

Quantile plots are graphical representations containing information about the cumulative frequency of the data with fractiles. A fractile is related to a given percentage, e.g. the 0.65 quantile of a given data set of observations is the observation for which 65% of all observations in the data set have smaller values. The 0.75 fractile is also called the *upper quantile* while the 0.25 fractile is called the *lower quantile*. The median is thus equal to 0.5 fractile.

In order to construct a quantile plot the observations in the data set are arranged in ascending order $\hat{x}_i^o$ and associated to q_i which is given by:

$$q_i = \frac{i}{n + 1} \quad (8.40)$$

Consider again the concrete strength data from Figure 8.4. The corresponding quantile plot is shown in Figure 8.7.

1.3. Model building and parameter estimation

The information contained in the data has to be represented with random variables in order to be utilized in risk or reliability analysis. Typically the initial choice of the model i.e. underlying assumptions regarding distributions and parameters may be based mainly on subjective information whereas the assessment of the parameters of the distribution function and the verification

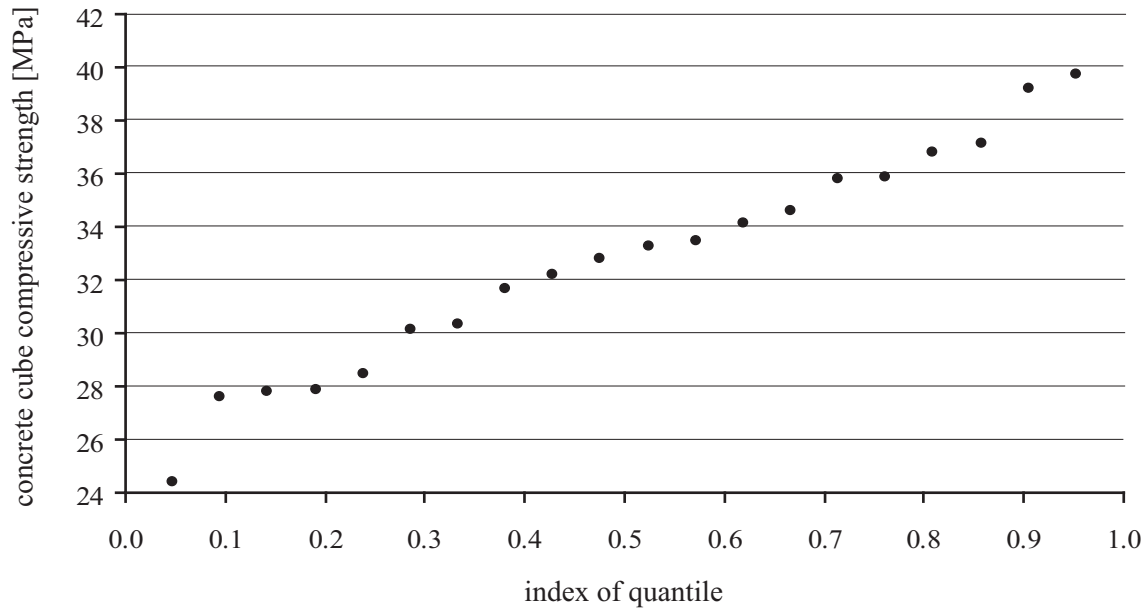


Figure 8.7: Quantile plot of the Concrete Compressive Strength.

of the models is performed on the basis of the available data. As the probabilistic models are based on both frequentistic information and subjective information these are Bayesian in nature.

In the following, only the probabilistic modelling of random variables will be considered, but the described approach applies with some extensions also to the probabilistic modelling of random processes and random fields.

Note that a quantile plot with swapped axes is called *cumulative frequency plot* and much more commonly used in engineering. The good thing about the cumulative frequency plot is that the shape that is created by the data entries can directly compared with the shape of a cumulative probability distribution function that is intended to be used for the representation of the data, see also Figure 8.1 as an example for a cumulative frequency plot.

1.3.1. Probability Paper

Probability paper is a useful tool for the selection of the distribution family that represents the data and the underlying phenomenon with sufficient accuracy.

Probability paper for a given distribution family is constructed such that the cumulative probability distribution function (or the complement) for that distribution family will have the shape of a straight line when plotted on the paper. A probability paper is thus constructed by a non-linear transformation of the cumulative distribution function (y -axis in the CDF diagram).

For a Normal distributed random variable the cumulative distribution function is given as:

$$F_X(x) = \Phi\left(\frac{x - \mu_X}{\sigma_X}\right) \quad (8.41)$$

where μ_X and σ_X are the mean value and the standard deviation of the Normal distributed random variable and where $\Phi(\cdot)$ is the standard Normal probability distribution function. Eq. (8.41) can be re-written as:

$$x = \Phi^{-1}(F_X(x))\sigma_X + \mu_X \quad (8.42)$$

Now by plotting x against $\Phi^{-1}(F_X(x))$, see also Figure 8.8, it is seen that a straight line is obtained with the slope depending on the standard deviation of the random variable X and the intercept on with the y -axis depending on the mean value of the random variable. Such a plot is sometimes called a quantile plot.

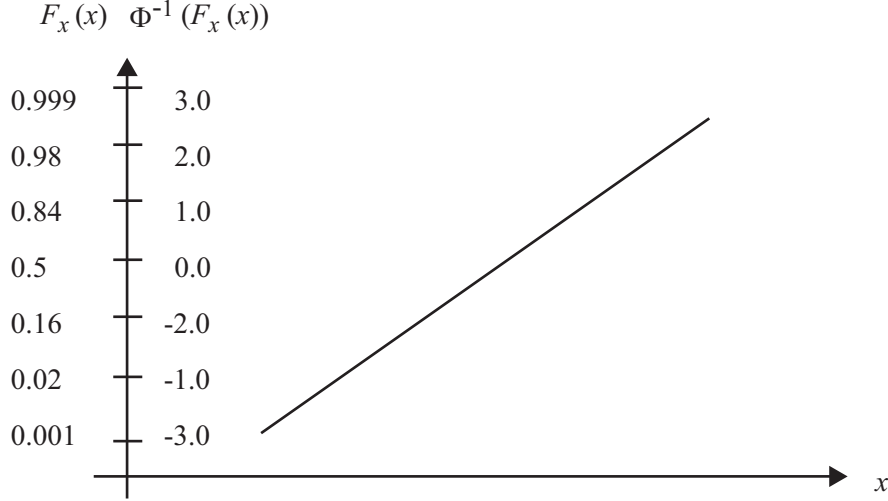


Figure 8.8: Illustration of the non-linear scaling of the y -axis for a Normal distributed random variable.

Also in Figure 8.8 the scale of the non-linear y -axis is given corresponding to the linear mapping of the observed cumulative probability densities. In probability papers typically only this non-linear scale is given.

Probability papers may also be constructed graphically. In Figure 8.9 the graphical construction of a Normal probability paper is illustrated.

Various types of probability paper are readily available for use. Given an ordered set of observed values $\hat{\mathbf{x}}_i^o = (\hat{x}_i^o, \hat{x}_i^o, \dots, \hat{x}_N^o)^T$ of a random variable the cumulative distribution function may be evaluated as:

$$F_X(\hat{x}_i^o) = \frac{i}{N+1} \quad (8.43)$$

In Figure 8.10 an example is given for the set of observed concrete cube compressive strength values together with the cumulative distribution function values as calculated using Eq. (8.43). In Figure 8.11 the cumulative distribution values are plotted on a Normal distribution probability paper.

A first estimate of the distribution parameters may be readily determined from the slope and the position of the best straight line through the plotted cumulative distribution values. In the next section the problem of parameter estimation is considered in more detail.

From Figure 8.11 it is seen that the observed cumulative distribution function fits relatively well with a straight line. This might also be expected considering that the observed values of the concrete compressive strength are not really representative for the lower tail of the distribution, where due to the non-negativity of the compressive strength it might be assumed that a Log-normal distribution would be more suitable.

The estimation of the probabilities of values which are outside the measured range can be done by extrapolation (see Figure 8.11). The probability paper may also be used for extreme

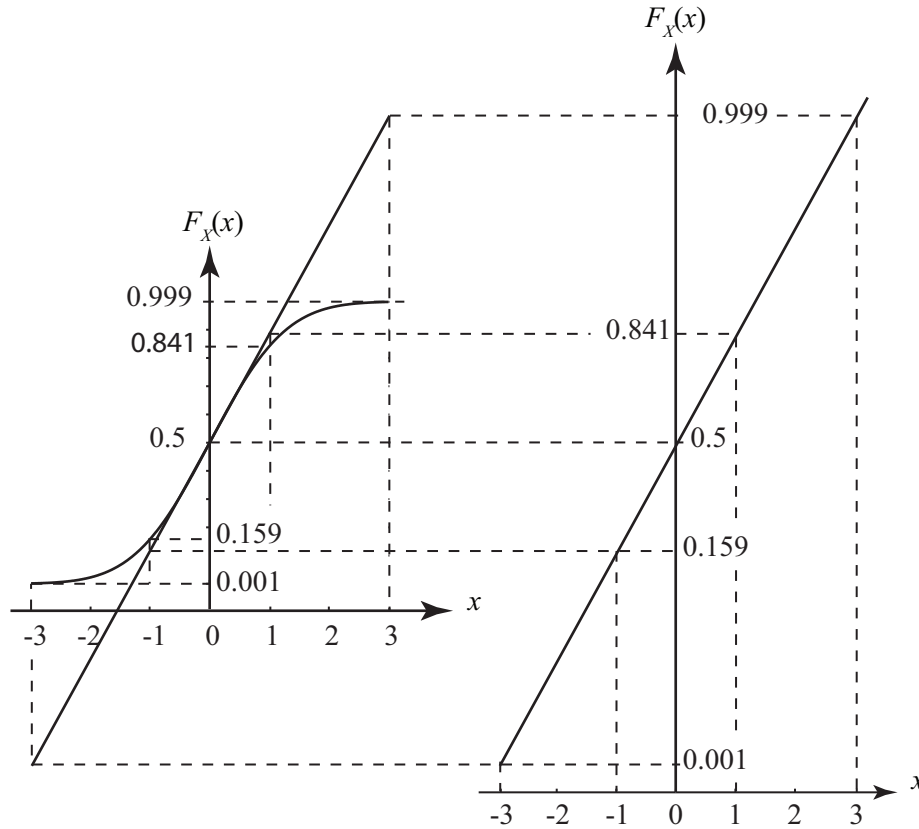


Figure 8.9: Illustration of the graphical construction of a Normal distribution probability paper.

phenomena such as the maximum water level, to estimate the values of the water level with a certain return period (see e.g. J.Schneider 2006). However, as always when extrapolating, extreme care must be exercised.

1.3.2. Parameter estimation with the Method of Moments

Using the method of moments, the parameters of the distribution can be estimated by equating the moments obtained from the sample and the moments obtained from the distribution.

Assuming that the considered random variable X may be modelled by the probability density function $f_X(x; \Theta)$, where $\Theta = (\theta_1, \theta_2, \dots, \theta_k)^T$ are the distribution parameters, the first k moments $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_k)^T$ of the random variable X may be written as:

$$\begin{aligned} \lambda_j(\theta) &= \int_{-\infty}^{\infty} x^j f_X(x | \theta) dx \\ &= \lambda_j(\theta_1, \theta_2, \dots, \theta_k) \end{aligned} \quad (8.44)$$

If the random sample, from which the distribution parameters $\Theta = (\theta_1, \theta_2, \dots, \theta_k)^T$ are to be estimated, is collected in the vector $\hat{\mathbf{x}} = (\hat{x}_1, \hat{x}_2, \dots, \hat{x}_n)^T$ the corresponding k sample moments may be calculated as:

$$m_j = \frac{1}{n} \sum_{i=1}^n \hat{x}_i^j \quad (8.45)$$

i	$\hat{x}_i^o$	$F_X(\hat{x}_i^o)$	$\Phi^{-1}(F_X(\hat{x}_i^o))$
1	24.4	0.048	- 1.668
2	27.6	0.095	- 1.309
3	27.8	0.143	- 1.068
4	27.9	0.190	- 0.876
5	28.5	0.238	- 0.712
6	30.1	0.286	- 0.566
7	30.3	0.333	- 0.431
8	31.7	0.381	- 0.303
9	32.2	0.429	- 0.180
10	32.8	0.476	- 0.060
11	33.3	0.524	0.060
12	33.5	0.571	0.180
13	34.1	0.619	0.303
14	34.6	0.667	0.431
15	35.8	0.714	0.566
16	35.9	0.762	0.712
17	36.8	0.810	0.876
18	37.1	0.857	1.068
19	39.2	0.905	1.309
20	39.7	0.952	1.668

Figure 8.10: Ordered set of observed concrete cube compressive strengths and the calculated cumulative distribution values.

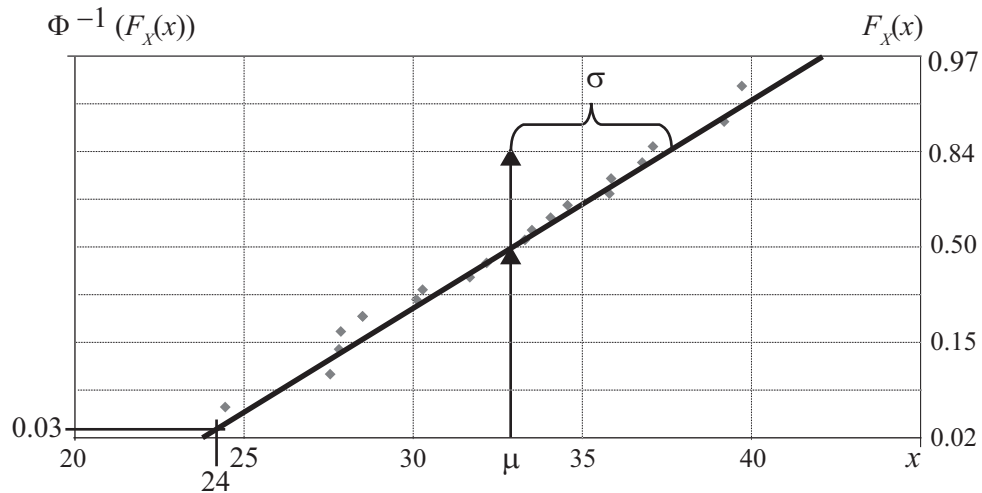


Figure 8.11: Concrete cube compressive strength data plotted in a probability paper for Normal distribution.

By equating the k sample moments to the k moments of the random variable X a set of k equations with the k unknown distribution parameters is obtained, the solution of which gives the *point estimates* of the distribution parameters.

1.3.3. The Maximum Likelihood Method

The principle of the Maximum Likelihood Method (MLM) is to find a set of parameters of an assumed probability distribution function, which most likely reflects the statistical behaviour of the underlying set of data (sample). The format of the method can be derived by means of the following considerations. Supposing that the parameters $\boldsymbol{\theta} = (\theta_1, \dots, \theta_n)^T$ of the distribution of X are known, the joint probability distribution of a (random) sample $X_1, X_2, X_3, \dots, X_n$ can be written as:

$$\begin{aligned} f_X(\mathbf{x}|\boldsymbol{\theta}) &= f_{X_1, X_2, \dots, X_n}(x_1, x_2, \dots, x_n|\boldsymbol{\theta}) = \\ &= f_{X_1}(x_1) f_{X_2}(x_2) \cdots f_{X_n}(x_n) = \prod_{i=1}^n f_X(x_i|\boldsymbol{\theta}) \end{aligned} \quad (8.46)$$

In other words, Eq. (8.46) can be seen as a relative measure for the belief that the sample $X_1, X_2, \dots, X_n$ belongs to the distribution function with given parameters $\boldsymbol{\theta} = (\theta_1, \dots, \theta_n)^T$. In general, the situation is obviously contrary. A sample $\hat{\mathbf{x}} = \hat{x}_1, \hat{x}_2, \hat{x}_3, \dots, \hat{x}_n$ is observed and the set of distribution parameters is not known. In that context, Eq. (8.31) can be similarly seen as a relative measure for the likelihood that the distribution determined by $\boldsymbol{\theta}$ is reflecting the statistical behaviour of the sample $\hat{\mathbf{x}}$. Over the entire domain of all possible parameters $\boldsymbol{\theta}$, the likelihood $L(\cdot)$ that the parameters belong to the sample is:

$$L(\boldsymbol{\theta}|\hat{x}_1, \hat{x}_2, \dots, \hat{x}_n) = \prod_{i=1}^n f_X(\hat{x}_i|\boldsymbol{\theta}) \quad (8.47)$$

The maximum likelihood estimators can now be defined as the set of parameters $\hat{\boldsymbol{\theta}}$ which are most likely representing the set of sample values, i.e. the set of parameters which maximize the likelihood function $L(\cdot)$ over the entire domain of $\boldsymbol{\theta}$.

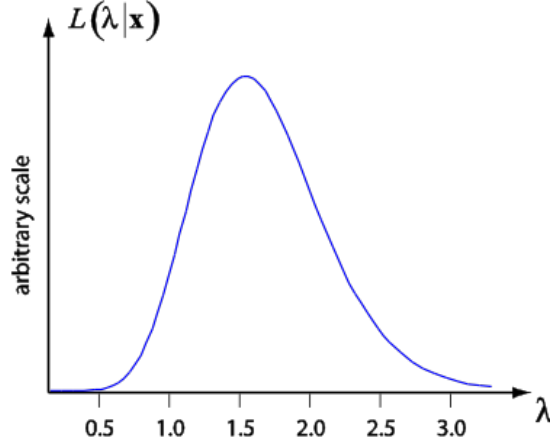
$$\hat{\boldsymbol{\theta}} = \arg \max_{\boldsymbol{\theta}} \{L(\boldsymbol{\theta}|\hat{x}_1, \hat{x}_2, \dots, \hat{x}_n)\} \quad (8.48)$$

To illustrate the nature of the likelihood function, a set of sample values $\hat{\mathbf{x}}$ is considered. The exponential distribution is chosen as a proper distribution function for the sample. The probability density function for the exponential distribution is:

$$f(x|\lambda) = \lambda \exp(-\lambda x) \quad (8.49)$$

By given realizations of X , i.e. $\hat{\mathbf{x}} = \hat{x}_1, \hat{x}_2, \dots, \hat{x}_n$, the joint probability density is defined as in Eq. (8.50), which is equivalent to the likelihood function given λ . For illustration, the likelihood function of the exponential distribution for a given sample is sketched in Figure 8.12.

$$\begin{aligned} f_X(\hat{\mathbf{x}}|\lambda) &= \prod_{i=1}^n f_X(\hat{x}_i|\lambda) = \lambda \exp(-\lambda \hat{x}_1) \lambda \exp(-\lambda \hat{x}_2) \cdots \lambda \exp(-\lambda \hat{x}_n) \\ &= \lambda^n \exp\left(-\lambda \sum_{i=1}^n \hat{x}_i\right) = L(\lambda|\hat{\mathbf{x}}) \end{aligned} \quad (8.50)$$


 Figure 8.12: Likelihood function on λ .

According to Figure 8.12, the likelihood function is attaining its maximum at $\lambda = 1.58$, i.e. the parameter which is most likely representing the considered sample. As an intuitive interpretation, the likelihood function might also be helpful to assess the likelihood of every possible parameter. In Figure 8.12 for example, 1.5 might be a more likely estimate for λ than 2.5. In general, it might be advantageous to consider the logarithm of the likelihood function:

$$\ln [L(\boldsymbol{\theta} | \hat{\mathbf{x}})] = l(\boldsymbol{\theta} | \hat{\mathbf{x}}) = \ln \left[\prod_{i=1}^n f_X(\hat{x}_i | \boldsymbol{\theta}) \right] = \sum_{i=1}^n \ln [f_X(\hat{x}_i | \boldsymbol{\theta})] \quad (8.51)$$

Then, the maximum may be obtained by solving a set of m equations:

$$\sum_{i=1}^n \frac{\partial}{\partial \theta_j} \ln [f_X(\hat{x}_i | \boldsymbol{\theta})] = 0 \quad j = 1, 2, \dots, m \quad (8.52)$$

However, the solution of the system might not always be possible in a closed form. For complex problems, computer-automated search functions like e.g. the simplex method (Nelder and Mead (1965)), can be employed to find the values $\hat{\boldsymbol{\theta}}$ which maximize the likelihood function.

As for the method of moments, the estimators are, before the experiment, random variables and can be studied as such. Their properties are well known and intensively studied in the literature (Lindley (1965)). For large samples ($n > 15$) the estimators $\hat{\boldsymbol{\theta}}$ are approximately normal distributed (due the summation form in Eq. (8.51)). Their means are asymptotically equal to the true parameter values $\boldsymbol{\theta}$; which means that the estimators are asymptotically unbiased. The covariance matrix $\mathbf{C}$ for the parameters $\boldsymbol{\theta} = (\theta_1, \theta_2, \dots, \theta_n)^T$ may be obtained through the inverse of the Fischer matrix $\mathbf{H}$:

$$\mathbf{C}_{\boldsymbol{\theta}\boldsymbol{\theta}} = \mathbf{H}^{-1} \quad (8.53)$$

with components given by:

$$H_{ij} = - \frac{\partial^2 l}{\partial \theta_i \partial \theta_j} \bigg|_{\boldsymbol{\theta} = \hat{\boldsymbol{\theta}}} \quad (8.54)$$

E.g. the variance of θ_i is asymptotically equal to:

$$\text{Var} [\hat{\boldsymbol{\theta}}] = -\frac{1}{nE \left[\frac{\partial^2 l}{\partial \theta_i^2} \right]} \quad (8.55)$$

Contrary to the method of moments, the estimations of the distribution variables are studied in regard to their uncertainty and not the sample moments. The main properties of the maximum likelihood estimates are:

1. For large data samples n , the distribution parameters $\boldsymbol{\theta}$ are Normal distributed.
2. The maximum likelihood estimates are consistent: For large data samples n , the estimates converge to the true value of the parameters of the distribution.
3. The maximum likelihood estimates are unbiased: For all sample sizes, the parameter of interest is calculated correctly.
4. The maximum likelihood estimates are efficient: The estimate has the smallest variance.
5. The maximum likelihood estimate is sufficient: All the information of the observations is utilized.
6. The solution of the MLM is unique.
7. The probability distribution for the problem at hand has to be known.

1.3.4. Predictive distribution

The parameters of distribution functions can be estimated as normal distributed random variables with the maximum likelihood method. Utilizing the total probability theorem, the unconditional or predictive distribution of the variable X can be formulated as:

$$f_X(x) = \int f_X(x|\boldsymbol{\theta}) f_{\boldsymbol{\Theta}}(\boldsymbol{\theta}) d\boldsymbol{\theta} \quad (8.56)$$

2. Application examples

E8.1.

E8.1.1. Data set

Records from the Norwegian Meteorological Institute (MET Norway) of the yearly highest 10 minutes mean wind speed U are utilized. Data from Torsvåg Fyr station are analyzed (WMO station number 01033). Analyze the data and estimate the parameters of the distribution that represents the data the best by following three different methods: probability plots, the method of moments (MoM) and the maximum likelihood method (MLM).

Table 8.2: Wind speed data (units: meter per second).

Year	$\hat{u}$	Year	$\hat{u}$	Year	$\hat{u}$	Year	$\hat{u}$	Year	$\hat{u}$	Year	$\hat{u}$
1957	26.06	1967	22.35	1977	25.59	1987	26.93	1997	34.74	2007	24.87
1958	29.43	1968	29.68	1978	28.61	1988	25.55	1998	27.81	2008	25.18
1959	33.01	1969	22.12	1979	23.94	1989	34.16	1999	21.98	2009	29.48
1960	24.01	1970	25.46	1980	24.80	1990	34.80	2000	29.19	2010	26.55
1961	23.42	1971	27.12	1981	29.87	1991	39.18	2001	26.49	2011	26.64
1962	24.20	1972	28.92	1982	27.74	1992	24.96	2002	24.06	2012	25.50
1963	27.73	1973	27.33	1983	25.92	1993	35.84	2003	27.00	2013	25.37
1964	23.80	1974	30.90	1984	27.78	1994	22.86	2004	23.15	2014	26.21
1965	26.42	1975	25.12	1985	26.80	1995	26.81	2005	27.06		
1966	24.00	1976	28.38	1986	24.01	1996	35.18	2006	23.25		



Figure 8.13: Weather station location.

Data reduction – Graphical display

Histograms With n values, the number of intervals k between the minimum and maximum value observed should be about $k = 1 + 3.3 \log_{10}(n)$ (Benjamin et al. 1970). For $n = 58 \rightarrow k = 1 + 3.3 \log_{10}(58) = 6.8$. Histograms with 6 and 7 bins are produced.

Cumulative frequency distribution

The wind speed data are sorted from the smallest to the largest $\hat{u}_1 \leq \hat{u}_2 \leq \dots \leq \hat{u}_i \leq \hat{u}_{58}$ ($i = 1, 2, \dots, 58$) and plotted against $i/(n + 1)$.

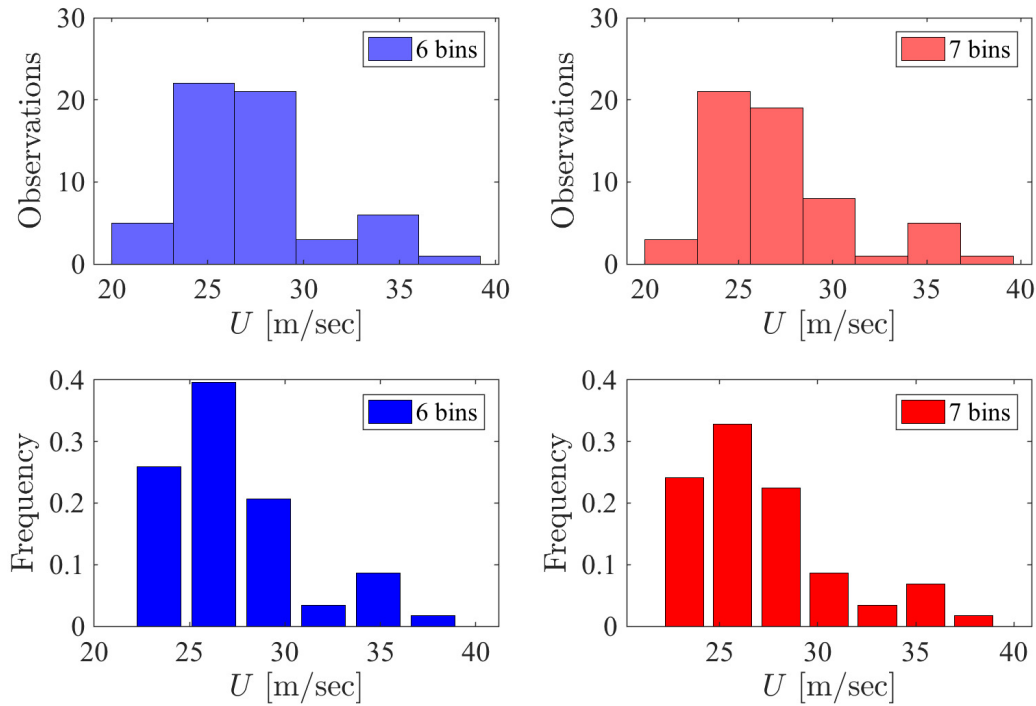


Figure 8.14: Histograms of the wind data.

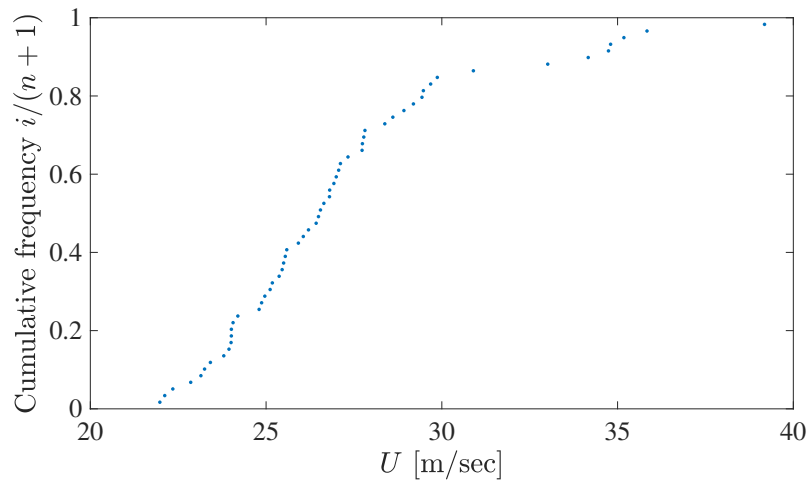


Figure 8.15: Cumulative frequency plot.

Numerical summaries

Sample mean (first moment)

The single most helpful number associated with a set of data is its average values, or arithmetic mean. The sample mean is:

$$m = \frac{1}{n} \sum_{i=1}^n \hat{u}_i = 27.16 \text{ m/s} \quad (8.57)$$

Sample variance (second central moment)

A measure of dispersion is the sample variance; it is analogous to the moment of inertia in the

sense that it deals with squares of distances from a centre of gravity, which in this case is the sample mean. The sample variance is defined here to be:

$$s^2 = \frac{1}{n} \sum_{i=1}^n (\hat{u}_i - m)^2 = 13.56 \text{ m}^2/\text{s}^2 \quad (8.58)$$

The sample standard deviation is then:

$$s = \sqrt{s^2} = 3.68 \text{ m/s} \quad (8.59)$$

And the sample coefficient of variation:

$$V = \frac{s}{m} = \frac{3.68}{27.16} = 0.136 \quad (8.60)$$

Sample coefficient of skewness (third central moment)

Another numerical summary of observed data is simply a logical extension of the reasoning leading to the formula for the sample variance. The sample coefficient of skewness is related to the third moment about the mean and measures the asymmetry of the data.

$$k = \frac{\frac{1}{n} \sum_{i=1}^n (\hat{u}_i - m)^3}{s^3} = 1.22 \quad (8.61)$$

The coefficient is positive for histograms skewed to the right (i.e. longer right tails) and negative for those skewed to the left.

Parameter estimates

Three different methods are used for estimating the parameters of the distribution representing the data: Probability plots, the method of moments (MoM) and the maximum likelihood method (MLM).

E8.1.2. Probability plots

The idea behind the probability plots is to transform the expression of the probability distribution so that it holds a linear relation with the data. First, this transformation is going to be elaborated for four distributions: Normal, Log-Normal, Gumbel and Weibull. After this, the linear relations are going to be used to fit the wind data and conclude which distribution provides the best fit.

Normal Distribution:

The normal cumulative density function (CDF) is:

$$F_X(x|\mu_X, \sigma_X) = \Phi\left(\frac{x - \mu_X}{\sigma_X}\right) \quad (8.62)$$

The equation can be rewritten as a linear relation between $\Phi^{-1}(F_X(x))$ and x :

$$\Phi^{-1}(F_X(x|\mu_X, \sigma_X)) = \frac{1}{\sigma_X}x - \frac{\mu_X}{\sigma_X} \quad (8.63)$$

Therefore, for normal variates the plot $[\hat{x}_i, \Phi^{-1}(i/(n+1))]$ is linear with intercept $-\mu_X/\sigma_X$ and slope $1/\sigma_X$.

Log-Normal Distribution:

The variate X is lognormal distributed if $Y = \ln(X)$ is normal distributed. The lognormal CDF is:

$$F_x(x|\mu_X, \sigma_X) = \Phi\left(\frac{\ln(x) - \mu_{\ln(X)}}{\sigma_{\ln(X)}}\right) \quad (8.64)$$

where the mean and standard deviation of the random variable Y are:

$$\Phi^{-1}(F_X(x|\mu_X, \sigma_X)) = \frac{1}{\sigma_{\ln(X)}} \ln(x) - \frac{\mu_{\ln(X)}}{\sigma_{\ln(X)}} \quad (8.65)$$

Therefore, for lognormal variates the plot $[\ln(\hat{x}_i), \Phi^{-1}(i/(n+1))]$ is linear with slope $1/\sigma_{\ln(X)}$ and intercept $-\mu_{\ln(X)}/\sigma_{\ln(X)}$.

Gumbel distribution:

The Gumbel CDF is:

$$F_x(x|a, b) = \exp\left(-\exp\left(-\frac{x-a}{b}\right)\right) \quad (8.66)$$

The distribution mean and standard deviation are:

$$\sigma_X = \frac{\pi}{\sqrt{6}}b; \quad \mu_X = a + b\gamma \quad (8.67)$$

where $\gamma \approx 0.577$ is the Euler's constant.

The equation can be rewritten as a linear relation between $-\ln(-\ln(F_X(x|a, b)))$ and x :

$$-\ln(-\ln(F_X(x|a, b))) = \frac{1}{b}x - \frac{a}{b} \quad (8.68)$$

Therefore, for Gumbel distributed variates the plot $[\hat{x}_i, -\ln(-\ln(i/(n+1)))]$ is linear with intercept $-a/b$ and slope $1/b$.

Weibull distribution:

The Weibull CDF is:

$$F_x(x|a, b) = 1 - \exp\left(-\left(\frac{x}{b}\right)^a\right) \quad (8.69)$$

The equation can be rewritten as a linear relation between $\ln(-\ln(1 - F_X(x|a, b)))$ and $\ln(x)$:

$$\ln(-\ln(1 - F_X(x|a, b))) = a \ln(x) - a \ln(b) \quad (8.70)$$

Therefore, for Weibull distributed variates the plot $[\ln(\hat{x}_i), \ln(-\ln(1 - i/(n+1)))]$ is linear with intercept $-a \ln(b)$ and slope a .

Probability plots with data

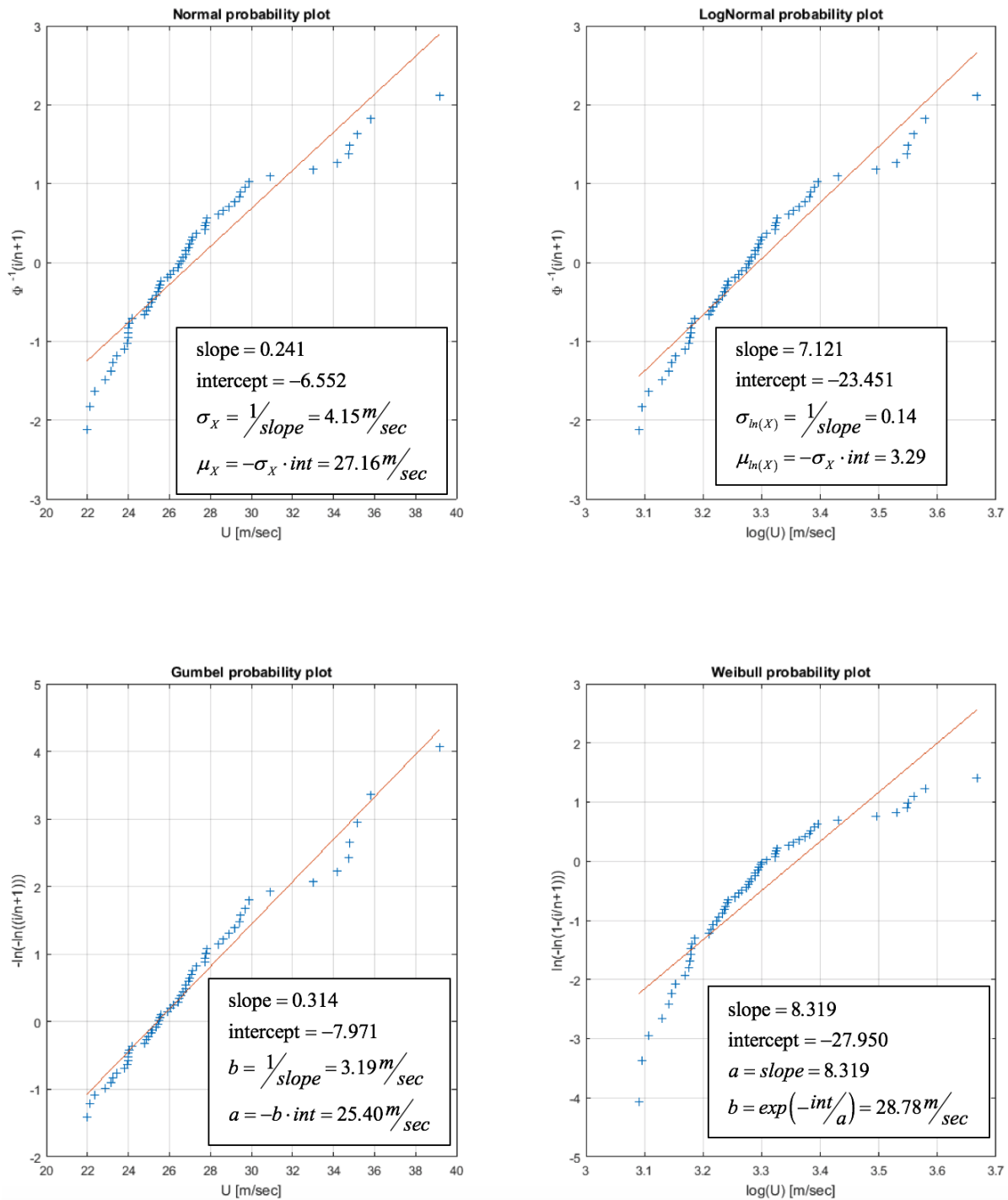


Figure 8.16: Data in probability plot.

The data is plotted in the form of probability plots in Figure 8.16. It can be observed that the “most linear” plot is given by the Gumbel distribution. This means that data are best represented by a Gumbel distribution. This was already expected since it can be proved analytically that the maxima of n independent identically distributed realizations tend to a Gumbel distribution for n growing.

The slope and intercept of the straight line, fitted with the least square method, are respec-

tively 0.3138 and -7.9709. It follows that the Gumbel parameters are $b = 3.19$ m/s; $a = 25.40$ m/s. The distribution mean and standard deviation result in:

$$\sigma_X = \frac{\pi}{\sqrt{6}} 3.187 = 4.09 \text{ m/s}; \mu_X = 25.40 + 3.187\gamma = 27.24 \text{ m/s}$$

The distributions are plotted in Figure 8.17.

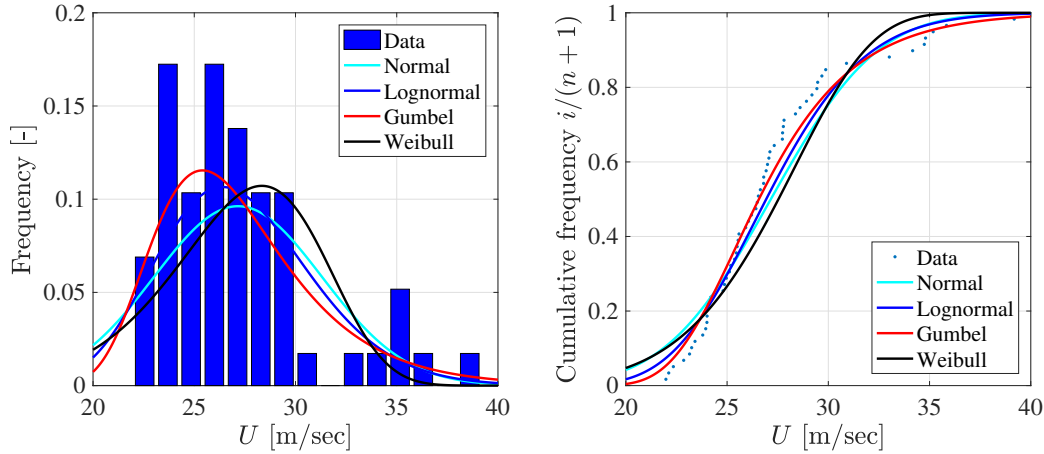


Figure 8.17: Distributions pdf and cdf fitted on probability plot.

E8.1.3. Method of Moments

The parameters of the distribution can be estimated by equating the sample moments ($m, s, k, \dots$) to the distribution moments ($\mu_X, \sigma_X, k_X, \dots$). The number of conditions shall be equal to the number of parameters.

For a distribution with one parameter (e.g. exponential) only one condition is necessary:

$$m = \mu_X \quad (8.71)$$

For a two parameter distribution two conditions are necessary:

$$\begin{cases} m = \mu_X \\ s = \sigma_X \end{cases} \quad (8.72)$$

For a three parameter distribution (e.g. 3 parameter Log-Normal) three conditions are necessary, e.g. Eq. (8.73).

$$\begin{cases} m = \mu_X \\ s = \sigma_X \\ k = k_X \end{cases} \quad (8.73)$$

Solutions for some distributions

The distribution parameters are calculated below. The PDF and CDF of the distributions evaluated with the estimated parameters are plotted in Figure 8.18.

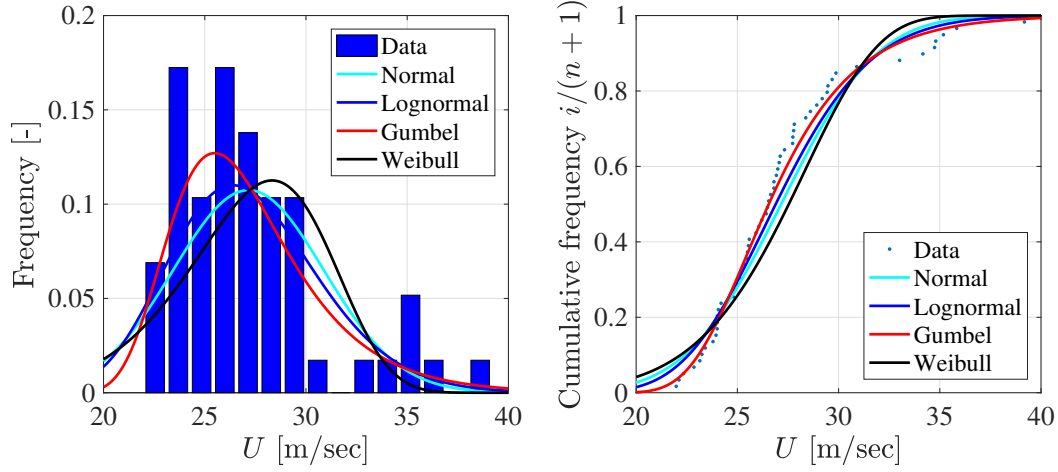


Figure 8.18: Fitted distributions with the method of moments.

Normal Distribution:

$$\begin{cases} m = \frac{1}{n} \sum_{i=1}^n \hat{u}_i = 27.16 \text{ m/s} = \mu_X \\ s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (\hat{u}_i - m)^2} = 3.71 \text{ m/s} = \sigma_X \end{cases} \quad (8.74)$$

Log-Normal Distribution:

$$\begin{cases} m = \frac{1}{n} \sum_{i=1}^n \hat{u}_i = 27.16 \text{ m/s} = \mu_X = \exp \left(\mu_{\ln(X)} + \frac{\sigma_{\ln(X)}^2}{2} \right) \\ s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (\hat{u}_i - m)^2} = 3.71 \text{ m/s} = \sigma_X = \mu_X \sqrt{\exp(\sigma_{\ln(X)}^2) - 1} \end{cases} \quad (8.75)$$

yielding:

$$\sigma_{\ln(X)} = \sqrt{\ln \left(1 + \left(\frac{s}{m} \right)^2 \right)} = 0.14; \quad \mu_{\ln(X)} = \ln(m) - \frac{\sigma_{\ln(X)}^2}{2} = -3.60 \quad (8.76)$$

Gumbel Distribution:

$$\begin{cases} m = \frac{1}{n} \sum_{i=1}^n \hat{u}_i = 27.16 \text{ m/s} = \mu_X = a + b\gamma \\ s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (\hat{u}_i - m)^2} = 3.71 \text{ m/s} = \sigma_X = \frac{\pi}{\sqrt{6}} b \end{cases} \quad (8.77)$$

yielding:

$$b = s \frac{\sqrt{6}}{\pi} = 2.90; \quad a = m - b\gamma = 25.49 \quad (8.78)$$

Weibull Distribution:

$$\begin{cases} m = \frac{1}{n} \sum_{i=1}^n \hat{u}_i = 27.16 \text{ m/s} = \mu_X = b\Gamma\left(1 + \frac{1}{a}\right) \\ s = \sqrt{\frac{1}{n-1} \sum_{i=1}^n (\hat{u}_i - m)^2} = 3.71 \text{ m/s} = \sigma_X = b\sqrt{\Gamma\left(1 + \frac{2}{a}\right) - \left(\Gamma\left(1 + \frac{1}{a}\right)\right)^2} \end{cases} \quad (8.79)$$

The parameters a and b can be obtained with numerical solver.

$$b = 28.7 \text{ m/s}; \quad a = 8.73$$

E8.1.4. Maximum Likelihood Method

The random variable of interest has a PDF $f_X(x|\boldsymbol{\theta})$, where $\boldsymbol{\theta} = (\theta_1, \theta_2, \dots, \theta_k)^T$ is the arrow of distribution parameters to be estimated. The random sample used to estimate the parameters is collected in the vector $\hat{\mathbf{x}} = (\hat{x}_1, \hat{x}_2, \dots, \hat{x}_n)$. The likelihood of the observed random sample is defined as:

$$L(\boldsymbol{\theta}|\hat{\mathbf{x}}) = \prod_{i=1}^n f_X(\hat{x}_i|\boldsymbol{\theta}) \quad (8.80)$$

The maximum likelihood point estimates of the parameters is the vector $\boldsymbol{\theta}^* = (\theta_1^*, \theta_2^*, \dots, \theta_k^*)^T$, which maximizes the likelihood function, i.e.:

$$\min_{\boldsymbol{\theta}} \{-L(\boldsymbol{\theta}|\hat{\mathbf{x}})\} \quad (8.81)$$

It is often advantageous to consider the log-likelihood, which is defined as:

$$l(\boldsymbol{\theta}|\hat{\mathbf{x}}) = \ln \left[\prod_{i=1}^n f_X(\hat{x}_i|\boldsymbol{\theta}) \right] = \sum_{i=1}^n \ln[f_X(\hat{x}_i|\boldsymbol{\theta})] \quad (8.82)$$

The maximum likelihood point estimate of the parameters is the vector $\boldsymbol{\theta}^* = (\theta_1^*, \theta_2^*, \dots, \theta_k^*)^T$ which maximizes the log-likelihood function:

$$\min_{\boldsymbol{\theta}} \{-l(\boldsymbol{\theta}|\hat{\mathbf{x}})\} \quad (8.83)$$

For sufficiently large n , the parameter distribution, i.e. the distribution representing their uncertainties, converges toward a Normal distribution with mean $\mu_{\boldsymbol{\theta}} = \boldsymbol{\theta}^* = (\theta_1^*, \theta_2^*, \dots, \theta_k^*)^T$ and covariance matrix $\mathbf{C}_{\boldsymbol{\theta}\boldsymbol{\theta}} = \mathbf{H}^{-1}$. $\mathbf{H}$ is the Fisher information matrix, whose elements are the second order partial derivatives of the log-likelihood function evaluated at $\boldsymbol{\theta}^*$:

$$H_{ij} = - \frac{\partial^2 l(\boldsymbol{\theta}|\hat{\mathbf{x}})}{\partial \theta_i \partial \theta_j} \bigg|_{\boldsymbol{\theta}=\boldsymbol{\theta}^*} \quad (8.84)$$

ML estimates for a Gumbel distribution:

The Gumbel CDF is shown in Eq. (8.66). The PDF is reported below:

$$F_x(x|a, b) = \exp \left(-\exp \left(-\frac{x-a}{b} \right) \right) \quad (8.85)$$

$$f_x(x|a, b) = \frac{1}{b} \exp \left(-\frac{x-a}{b} - \exp \left(-\frac{x-a}{b} \right) \right) \quad (8.86)$$

The vector of parameters θ^* has two elements: the location parameter $a \in (-\infty; +\infty)$, and the scale parameter $b \in (0; +\infty)$. The log-likelihood function is:

$$\begin{aligned} l(\theta|\hat{\mathbf{x}}) &= \ln [\prod_{i=1}^n f_X(\hat{x}_i|\theta)] = \sum_{i=1}^n \ln[f_X(\hat{x}_i|\theta)] = \\ &= \sum_{i=1}^n \ln \left[\frac{1}{b} \exp \left(-\frac{\hat{x}_i - a}{b} - \exp \left(-\frac{\hat{x}_i - a}{b} \right) \right) \right] \end{aligned} \quad (8.87)$$

The parameter vector θ^* maximizing $l(\theta|\hat{\mathbf{x}})$ is found numerically. In some cases, as for normal distribution, analytical solutions are obtained by setting the first derivatives of $l(\theta|\hat{\mathbf{x}})$ equal to zero.

The fisher information matrix is then:

$$\mathbf{H} = \begin{bmatrix} -\frac{\partial^2 l(\theta^*|\hat{\mathbf{x}})}{\partial^2 \theta_1} & -\frac{\partial^2 l(\theta^*|\hat{\mathbf{x}})}{\partial \theta_1 \partial \theta_2} \\ -\frac{\partial^2 l(\theta^*|\hat{\mathbf{x}})}{\partial \theta_1 \partial \theta_2} & -\frac{\partial^2 l(\theta^*|\hat{\mathbf{x}})}{\partial^2 \theta_2} \end{bmatrix}_{\theta=\theta^*} \quad (8.88)$$

Application of the method to the wind data

The application of the MLM to the data results in $\theta^* = \mu_\theta = (25.55, 2.65)^T$. The value of the log-likelihood function at maximum is equal to $l(\theta^*|\hat{\mathbf{x}}) = -149.78$, cf. Figure 8.19.

The `fminsearch` Matlab[®] function has been used with starting values $a_0 = 20$ and $b_0 = 3$. The distributions CDF and PDF are illustrated in Figure 8.20.

Integration of the parameter uncertainty

The distribution of the parameter estimates converges towards a bivariate Normal distribution with mean $\theta^* = \mu_\theta = (25.55, 2.65)^T$ and covariance matrix obtained by inverting the Fisher information matrix:

$$\mathbf{C}_{\theta\theta} = \begin{bmatrix} 0.1333 & 0.0309 \\ 0.0309 & 0.0786 \end{bmatrix} \quad (8.89)$$

The parameter joint density function $f_{A,B}(a, b)$ is illustrated in Figure 8.21 together with the marginal distributions.

The parameter uncertainties can be integrated in the conditional distribution $F_{U|A,B}(u|a, b)$ for obtaining an unconditional distribution $F_U(u)$ which includes the uncertainties on A and B .

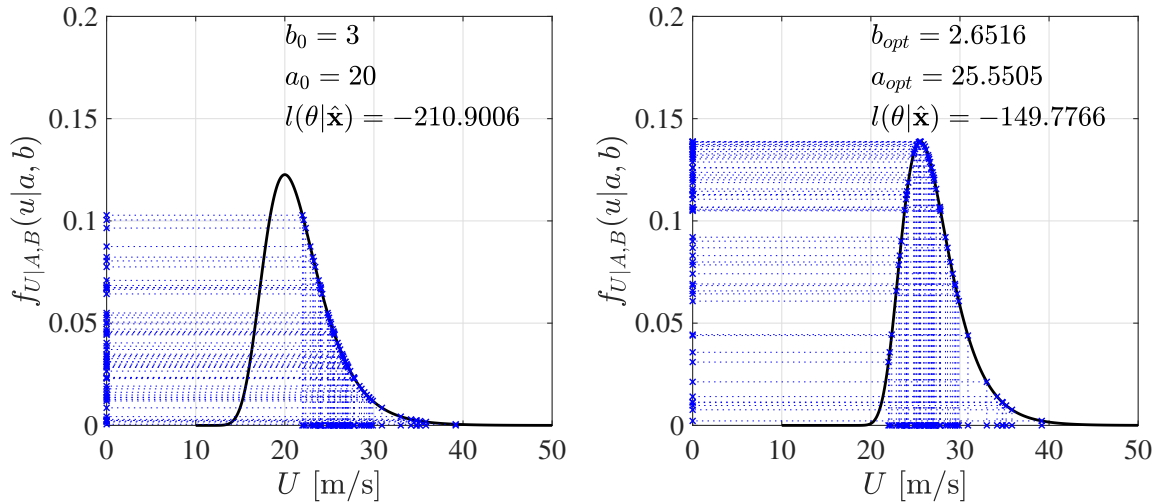


Figure 8.19: PDF and log-likelihood function for first guessed values of a and b (left) and for parameters maximizing the likelihood function (right).

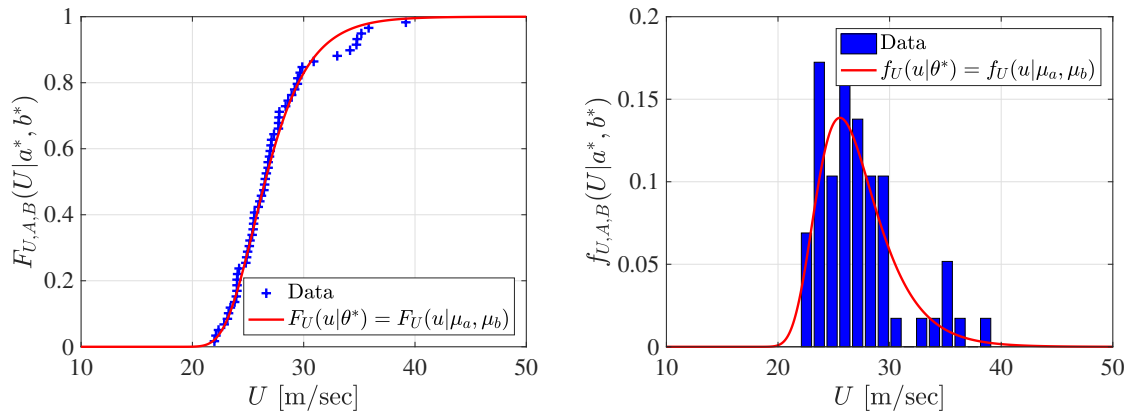


Figure 8.20: CDF and PDF of the Gumbel distribution with estimated parameters.

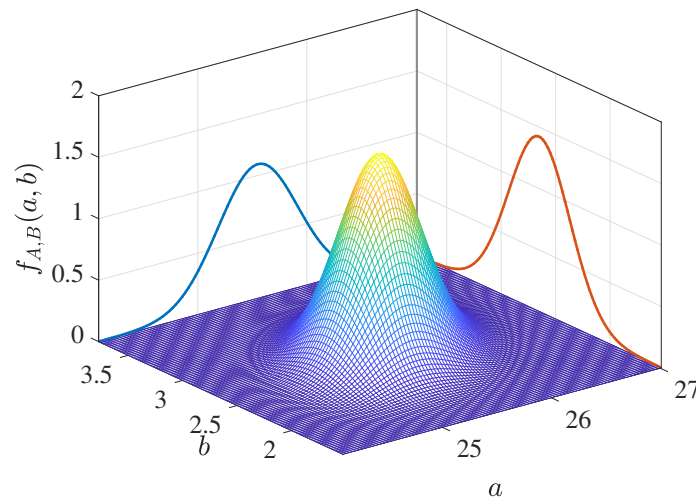


Figure 8.21: Joint PDF for the parameters A and B .

The integration reads:

$$f_U(u) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} f_{U|A,B}(u|a, b) f_{A,B}(a, b) da db \quad (8.90)$$

Or for the CDF:

$$F_U(u) = \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} F_{U|A,B}(u|a, b) f_{A,B}(a, b) da db \quad (8.91)$$

The integration is performed numerically; faster calculations are obtained by changing the integration limits to $\pm t \cdot \sigma$, with t equal to, for example, 4 or 5.

$$F_U(u) = \int_{\mu_A - t\sigma_A}^{\mu_A + t\sigma_A} \int_{\mu_B - t\sigma_B}^{\mu_B + t\sigma_B} F_{U|A,B}(u|a, b) f_{A,B}(a, b) da db \quad (8.92)$$

The conditional and unconditional (or predictive) distributions are illustrated in Figure 8.23.

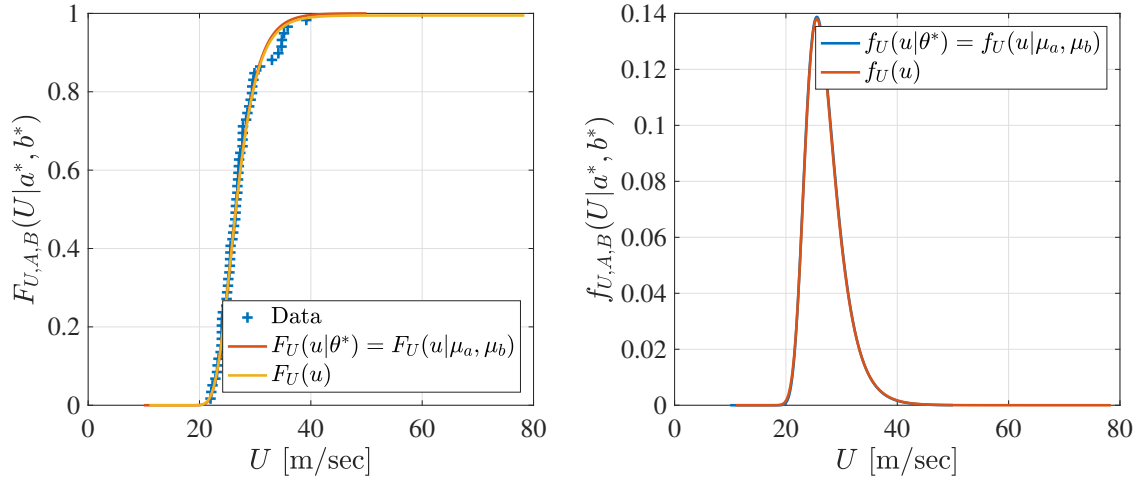


Figure 8.22: Predictive CDF and PDF. The numerical integration is performed over $\pm 3 \cdot \sigma$.

The parameter uncertainty is quite low due to the large amount of data in this example, as a consequence the predictive distribution looks similar to the conditional one. Nevertheless, the differences in the tail are very important for structural reliability problems, since the probability of failure is very sensible to the tails of the random variables.

If we imagine to have only the first 7 years of records, i.e. a sample of 7 values, the conditional and predictive CDFs will differ more since there is higher uncertainties on the parameters, see Figure 8.23.

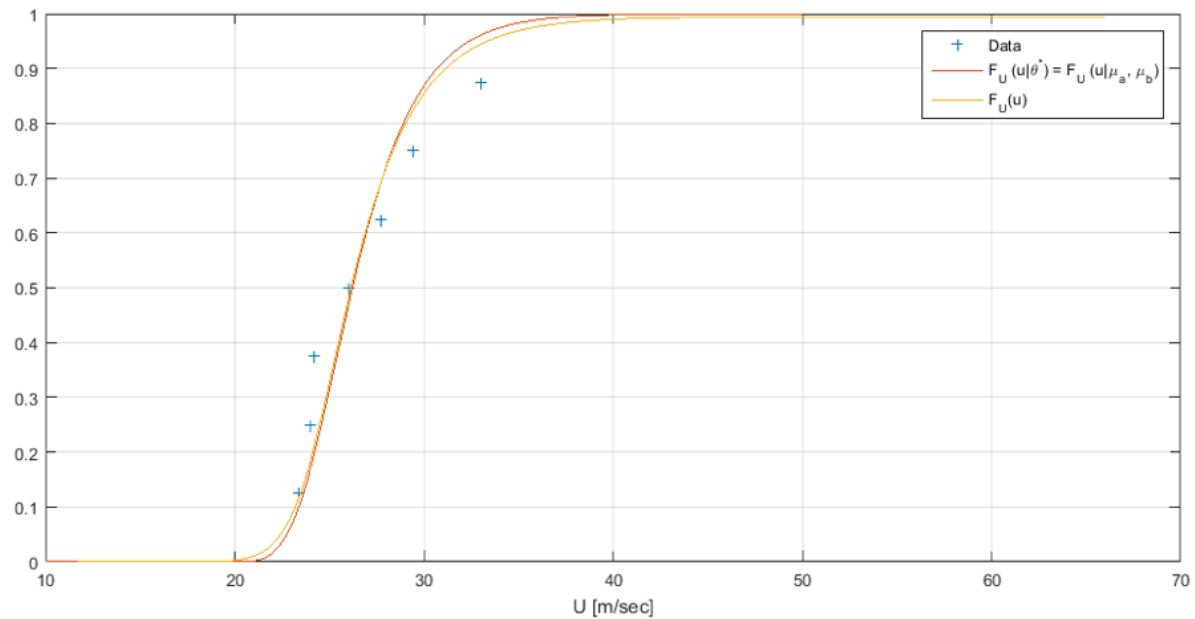


Figure 8.23: Conditional and predictive CDFs for a sample of 7 data.

Question::

Engineering uncertainty quantification

- Given data, how would you choose the distribution type and the parameters?
- Explain the principle of Maximum Likelihood Estimation.
- What is a predictive distribution, and how it is computed?

Bibliography

- Baravalle, M. and J. Köhler (2017). “A framework for estimating the implicit safety level of existing design codes”. In: *Proceedings of the 12th International Conference on Structural Safety and Reliability (ICOSSAR2017)*, Vienna, Austria.
- Benjamin, J. and C. Cornell (1970). *Probability, statistics and decision for civil engineers*. New York: McGraw-Hill.
- Blomquist, G. (1979). “Value of Life Saving: Implications of Consumption Activity”. In: *Journal of Political Economy* 87.3, pp. 540–558. ISSN: 00223808, 1537534X. URL: <http://www.jstor.org/stable/1832022>.
- CEN (2002). “Basis of Structural Design and other parts”. In: *Structural Standard EN1990 - EN1999*.
- EN 1990:2002 (2002). *Eurocode 0: Basis of Structural Design*. Standard. Brussels, Belgium: European Committee for Standardization.
- Faber, M. H., S. Miraglia, J. Qin, and M. G. Stewart (2018). “Bridging resilience and sustainability - decision analysis for design and management of infrastructure systems”. In: *Sustainable and Resilient Infrastructure* 0.0, pp. 1–23. DOI: [10.1080/23789689.2017.1417348](https://doi.org/10.1080/23789689.2017.1417348). eprint: <https://doi.org/10.1080/23789689.2017.1417348>. URL: <https://doi.org/10.1080/23789689.2017.1417348>.
- Glavind, S. and M. Faber (2018). “A framework for offshore environment modelling”. In: *Proceedings of the ASME 2018 37th International Conference on Ocean, Offshore and Arctic Engineering OMAE2018 June 17-22, 2018, Madrid, Spain*.
- Hasofer, A. M. and N. Lind (1974). “Exact and Invariant Second Moment Code Format”. In: *Journal Engineering Mechanics Division* 100.EM1, pp. 111–121.
- ISO (2015). *ISO 2394:2015: General principles on reliability for structures*. Tech. rep. Geneva, Switzerland: International Organization for Standardization.
- J.Schneider (2006). *Introduction to safety and reliability of structures*. Zürich: International Association for Bridge and Structural Engineering. [ebook at NTNU Library].
- JCSS (2001). *Joint Committee on Structural Safety - Probabilistic Model Code*. Standard. URL: <http://www.jcss.byg.dtu.dk>.
- JCSS (2008). *Risk Assessment in Engineering – Principles, System Representation & Risk Criteria*. ISBN 978-3-909386-78-9.
- Jordaan, I. (2005). *Decisions under uncertainty: probabilistic analysis for engineering decisions*. Cambridge University Press.

- Köhler, J. and M. Baravalle (2019). “Risk-based decision making and the calibration of structural design codes – prospects and challenges”. In: *Civil Engineering and Environmental Systems* 36.1, pp. 55–72. DOI: [10.1080/10286608.2019.1615477](https://doi.org/10.1080/10286608.2019.1615477). eprint: <https://doi.org/10.1080/10286608.2019.1615477>. URL: <https://doi.org/10.1080/10286608.2019.1615477>.
- Köhler, J., A. Frangi, and R. Steiger (2008). “On the role of stiffness properties for ultimate limit state design of slender columns”. In: *Proceedings of CIB-W18 Meeting*. Vol. 41.
- König, G. and D. Hosser (1981). *The Simplified Level II Method and its application on the derivation of safety elements for Level I*. Tech. rep. Comité Euro-International du Béton: Technische Hochschule Darmstadt.
- Madsen, H., S. Krenk, and N. Lind (1985). *Methods of structural safety*. Prentice-Hall.
- Nathwani, J. S., M. D. Pandey, and N. C. Lind (1997). “Affordable safety by choice: the life quality method”. In: *Waterloo, Canada: Institute for Risk Research, University of Waterloo*.
- Neumann, J. von and O. Morgenstern (1943). *Theory of games and economical behaviour*. Princeton, USA: Princeton University Press.
- Rackwitz, R. (2000). “Optimization — the basis of code-making and reliability verification”. In: *Structural Safety* 22.1, pp. 27–60. ISSN: 0167-4730. DOI: [https://doi.org/10.1016/S0167-4730\(99\)00037-5](https://doi.org/10.1016/S0167-4730(99)00037-5). URL: <http://www.sciencedirect.com/science/article/pii/S0167473099000375>.
- Raiffa, H. and R. Schlaifer (1961). *Applied statistical decision theory*. ISBN: 978-0-471-38349-9: Boston: Clinton Press, Inc.
- Ravindra, M. and T. Galambos (1978). “Load and resistance factor design for steel”. In: *Journal of the Structural Division* 104(9), pp. 1337–53.
- Weber, M. (1921). *Wirtschaft und Gesellschaft*. <http://www.textlog.de/7295.html>: 5. Aufl. Mohr Siebeck, Tübingen 1976.

