

THESIS FOR THE DEGREE OF DOCTOR OF ENGINEERING IN
APPLIED MATHEMATICS AND MATHEMATICAL STATISTICS

On the Optimisation and Regulation of Clinical Trials

SEBASTIAN JOBJÖRNSSON

CHALMERS



GÖTEBORGS UNIVERSITET

Division of Applied Mathematics and Statistics

Department of Mathematical Sciences

CHALMERS UNIVERSITY OF TECHNOLOGY AND UNIVERSITY OF GOTHENBURG
Göteborg, Sweden 2018

This project has received funding from the European Union's Seventh Framework Programme for research, technological development and demonstration under the IDEAL Grant Agreement no 602552.

On the Optimisation and Regulation of Clinical Trials

Sebastian Jobjörnsson

© Sebastian Jobjörnsson, 2018

ISBN: 978-91-7597-804-8

Doktorsavhandlingar vid Chalmers tekniska högskola

Ny serie (ISSN 0346-718X)

Löpnnummer 4485

Department of Mathematical Sciences

Chalmers University of Technology and University of Gothenburg

SE-412 96 Göteborg, Sweden

Phone: +46 (0)31 772 1000

Typeset with L^AT_EX.

Department of Mathematical Sciences

Printed in Göteborg, Sweden 2018

On the Optimisation and Regulation of Clinical Trials

Sebastian Jobbjörnsson

Department of Mathematical Sciences
Chalmers University of Technology and University of Gothenburg

Abstract

The basic premise of this thesis is that Bayesian Decision Theory (BDT) can and should be used to solve clinical trial design problems. While the flexibility of the framework allows for accommodating a great variety of situations, it also requires an explicit consideration of the gains and costs associated with the trial. This leads to an increased understanding of how the optimal design depends not only on statistical considerations, but also on the consequences of the decisions made during and after the trial.

The main contribution of the thesis consists of the four papers appended. In Paper I, optimisation is done by a drug company, taking the approval decision of a regulatory authority and the reimbursement decision of a health care insurer into account. A particular point of interest in this model is the effect that the uncertainty surrounding the insurer's willingness to pay has on the company's optimisation. Papers II and III are both concerned with comparing a number of different designs in the special case where the patient population can be partitioned using a binary biomarker. While Paper II restricts the analysis to single-stage designs, Paper III also considers adaptive, two-stage designs. The main method of analysis in all these papers is backward induction. Paper IV revisits Anscombe's classical model on fully sequential trials and also considers a number of different extensions. Approximate solutions are obtained using continuous-time optimal stopping theory. In addition to the papers, the thesis includes a discussion of the problem of optimal regulation of clinical trials, and defines and solves two simple example models.

Since several of the analyses presented in this thesis provide a detailed demonstration of how to formulate and solve clinical trial design problems, it should be of interest to statisticians seeking to apply BDT to real-world problems. Further, since the implications that the solutions have for regulation and reimbursement are discussed at several places, it should also be of value to government agencies tasked with creating an efficient environment for drug development.

Keywords: Bayesian statistics, decision theory, clinical trials, drug regulation, subgroup analysis, optimal stopping.

Acknowledgements

First and foremost I would like to thank my thesis advisor Carl-Fredrik Burman. Your constant support and guidance during the last five years has made this thesis possible. Paper IV would not have been possible to write without the mathematical expertise provided by Sören Christensen, who was also my co-advisor. It was great to have the opportunity to learn more about the fascinating world of optimal stopping and free-boundary value problems with you by my side. I'm also very grateful to my second co-advisor Serik Sagitov, and my two examiners Olle Nerman and Staffan Nilsson.

I thank Martin Forster and Paolo Pertile for teaching me about economics, for the great hospitality shown when I visited in Verona, and for their tireless effort when working on Paper I. Many thanks also goes to my collaborators on papers II and III, Thomas Ondra, Franz König, Martin Posch, Nigel Stallard and Robert A. Beckman. I wish all of you great success in the trials ahead.

Sebastian Jobjörnsson
Gothenburg, August 13, 2018

List of appended papers

- Paper I** **Jobbjörnsson, S.**, M. Forster, P. Pertile, and C.-F. Burman (2016). Late-stage pharmaceutical R&D and pricing policies under two-stage regulation. *Journal of Health Economics* 50, 298-311.
- Paper II** Ondra T., **S. Jobbjörnsson**, R. A. Beckman, C.-F. Burman, F. König, N. Stallard, and M. Posch (2016). Optimizing Trial Designs for Targeted Therapies. *PLoS ONE* 11(9): e0163726.
- Paper III** Ondra T., **S. Jobbjörnsson**, R. A. Beckman, C.-F. Burman, F. König, N. Stallard, M. Posch (2017). Optimized adaptive enrichment designs. *Statistical Methods in Medical Research*.
- Paper IV** **Jobbjörnsson, S.** and S. Christensen (2018). Anscombe’s model for sequential clinical trials revisited. *Sequential Analysis* 37(1), 115-144.

My contributions

- Paper I** I did most of the model development and the proofs. I also implemented the code for the application of the model and did most of the writing.
- Paper II** The model was jointly formulated by all authors at the beginning of the project. I and Thomas Ondra contributed equally to the derivation of some simplifying formulas for computing expected utilities and to the implementation of the code for trial design optimisation. I did contribute to the writing, but most of it was done by the other authors.
- Paper III** The basic model framework was jointly developed by all authors. I and Thomas Ondra both contributed to the development of the algorithms used in the optimisation code, although Thomas did most of actual implementation. I did, however, write a number of testing procedures in order to confirm the results and also reviewed the code written by Thomas. Most of the writing of the article was done by the other authors.
- Paper IV** The formulation of Anscombe's model as an optimal stopping problem was mainly done by Sören Christensen. He also led the work on how to solve the problem using the free-boundary approach. The two model extensions were arrived at through joint discussions, but Sören provided the analytical solutions to the example models for an unknown patient horizon. All of the algorithms used to compute the solutions were implemented by me. Most of the writing was done by me, and I did most of the work on the proofs in the appendix.

List of abbreviations

BDT	Bayesian Decision Theory
FDA	US Food and Drug Administration
FDR	False Discovery Rate
FWER	Family-Wise Error Rate
HCI	Health Care Insurer
ICER	Incremental Cost-Effectiveness Ratio
IDEAL	Integrated DEsign and AnaLysis (an EU project)
MTP	Multiple Testing Procedure
NICE	National Institute of Clinical Excellence
PCER	Per Comparison-Error Rate
QALY	Quality Adjusted Life Year
RA	Regulatory Authority
WTP	Willingness To Pay

Contents

1	Introduction	1
1.1	A brief history of clinical trials	2
1.2	Clinical trials today	4
1.2.1	The phases of a clinical trial	5
1.3	Brief summary of the contributed papers	7
1.4	Optimal regulation	9
1.5	Overview of thesis	9
2	Bayesian decision theory	11
2.1	Decision problems	11
2.2	Sequential decision problems	12
2.3	Backward induction	13
2.4	Example: Optimisation of sample size	14
3	Optimal stopping problems in continuous time	17
3.1	Markov processes and optimal stopping	17
3.2	Girsanov's transform	19
4	Summary of papers	21
4.1	Paper I: Late-stage pharmaceutical R&D and pricing policies under two-stage regulation	21
4.2	Paper II: Optimizing trial designs for targeted therapies	22
4.3	Paper III: Optimized adaptive enrichment designs	23
4.4	Paper IV: Anscombe's model for sequential clinical trials revisited	24
5	Optimal regulation of clinical trials	27
5.1	Basic skeleton model	28
5.2	Model with pre-clinical investments	29
5.3	Model with unknown trial cost	32
5.4	Further work on optimal regulation	34

6	Discussion	37
7	Conclusions	43
A	Multiple testing procedures	45
A.1	The Spiessens-Debois test	47
B	Some basic tools from analysis	49
B.1	The implicit function theorem	49
B.2	The envelope theorem	50
	Bibliography	51

Chapter 1

Introduction

Suppose that a new medical treatment has been discovered. Based on chemical analyses, simulation models and tests on animals, there is some indication that the treatment has a positive effect for a population of patients afflicted with a certain disease. However, no data has been collected on the actual medical response of the patient population when given the new treatment. Therefore, a substantial amount of uncertainty remains regarding whether or not the good effects outweigh any bad side effects. To reduce this uncertainty, a clinical trial is performed. In such a trial, a certain number patients are recruited and given the new treatment. After having observed the trial results, the decision maker responsible for approving the treatment for distribution will be better informed, and the probability of approving a bad treatment or rejecting a good treatment is smaller.

The situation just described can be formulated as a two-stage decision problem. In the first stage, a sample size for the trial is chosen. Then the trial is performed, yielding data used to estimate the effect of the treatment. In the second stage, the treatment is either accepted or rejected based on the estimate. But which sample size is optimal? On the one hand, a large sample size leads to less uncertainty when the approval decision is made. On the other hand, a large trial requires more resources to perform, resources that may very well be of better use in some alternative project. Moreover, a larger trial takes a longer time to execute. If the treatment has a large positive effect, this means that the patient population will have to wait longer for a potentially life-saving treatment. If the treatment has a large negative effect, a large number of patients will suffer in the trial. Hence, there is a basic trade-off between the information provided and the costs (monetary and pure health costs) incurred. Optimising the sample size corresponds to finding the most beneficial balance.

By viewing the trial design problem as a decision problem, it is possible to

apply the same methodology in the search of an optimal design regardless of the specifics, provided that a framework for handling general decision problems is available. One such framework is Bayesian Decision Theory (BDT), the main tool used to formulate and solve the trial design problems considered in this thesis. It has many advantages. First, it is based on a mathematical, axiomatic foundation. Second, the theory is very general, and may be applied to any decision problem once some basic structure has been introduced. Third, its use is well established in the literature, implying that a great number of specific applications are available to draw inspiration from, and also that many methods have been devised to deal with the computational challenges involved when searching for the optimal solution. The main alternative to BDT is commonly referred to as the *classical* or *frequentist* approach. As noted by Senn (2007), the classical approach is really a hybrid one, employing the Neyman-Pearson framework during the design stage while being Fisherian during analysis. The purpose of this thesis is not to enter into a detailed discussion comparing these approaches. Much has been published on this issue and the interested reader will have no trouble finding excellent overviews in the literature (see, for example, Bayarri and Berger (2004)).

The reader will find that the papers included in this thesis sometimes take a special interest in what happens when the target population for the medical treatment is small. The reason is that most of the work presented was done as part of an EU project called Integrated D^Esign and AnaL^Ysis of clinical trials in small population groups (IDEAL). Small population groups may arise because the treatment is for a rare disease or is only expected to work in a small subset of the total population. Trial design in such situations often requires a more careful analysis of how to make the most of the necessarily small sample sizes. From a regulatory perspective, small population groups also raise questions that are not purely statistical. If a pharmaceutical company expects the final market to be small, will it view basic research in this area as a viable investment? If not, what needs to be changed in the regulatory structure so as to incentivise more research targeting small population groups? The economic resources of any society are limited, so how should the cost of implementing the incentives be balanced against the potential health gains?

1.1 A brief history of clinical trials

A *clinical trial* is an experiment performed on human subjects for the purpose of generating data that may be used to estimate the efficacy and safety level (i.e., any harmful side-effects) of a new medical treatment. Note that there is a slight difference in meaning between the terms *efficacy* and *effectiveness* in the clinical trials literature. Efficacy is the extent to which a treatment does

more good than harm under ideal circumstances. Effectiveness, on the other hand, assesses whether a treatment does more good than harm when provided under usual circumstances of healthcare practice (Haynes, 1999). Following the literature, the term *drug* is often used in this thesis to refer to a chemical substance meant to treat some medical condition. However, in many cases the discussion also carries over to other kinds of treatments, for example a particular surgical procedure.

The idea that some kind of testing under controlled forms on humans should be required for new drugs might seem obvious. However, the establishment and subsequent development of the various governmental bodies responsible for ascertaining adequate testing has been a gradual process. Often, the regulatory rules have been extended as a direct response to drugs with severe side effects having been marketed. A very important part of modern clinical trial methodology is the inclusion of a control group in addition to the group of test subjects given the medical treatment. The subjects in both groups are typically recruited from the same population and then assigned to one of the two groups by means of some randomisation procedure. This serves to ensure that the any effects observed are really due to the treatment given and not circumstantial. While various forms of medical experimentation in its most general sense have been performed for thousands of years, the crucial idea of comparing with a control group is of a later origin. One of the first proper clinical trials in this respect was performed in 1747 by the Scottish physician James Lind (Baron, 2009), who divided a number of sailors afflicted with scurvy into different groups and noted that the group given oranges and lemons fared much better than the others. The first randomised curative clinical trial tested the drug streptomycin, aimed at curing pulmonary tuberculosis, and was carried out 1946-1947. The trial compared the active treatment with a placebo, and was also double-blinded (Hart, 1999).

As the production of pharmaceuticals became industrialised and the consumer market grew during the 20th century, a number of medical disasters prompted the establishment of various regulatory authorities. A well-known example is the drug Elixir Sulfanilamide, which had never been tested in human subjects before market introduction. Its use eventually led to more than 100 deaths and to the passing of The Federal Food, Drug and Cosmetic Act in the US in 1938, requiring pharmaceutical companies to submit reports on the safety of new drugs. The later Kefauver-Harris Amendment of 1962 strengthened the safety requirements and also introduced requirements on efficacy for the first time (Chow and Liu, 2014).

1.2 Clinical trials today

Spiegelhalter et al. (2004) use the following classification for the stakeholders involved in the clinical development process. The *sponsor* is the one who pays for the trial, for example a pharmaceutical company. The *investigators* are responsible for the actual conduct of the clinical trial. Based on confirmatory trial results, *reviewers* evaluate efficacy and safety and decide on market approval, while *policy-makers* estimate the cost-benefit impact of introducing the new drug.

A slightly different terminology is used in this thesis. It will be assumed that the investigator is the same as the sponsor. If the sponsor is a pharmaceutical company, it will be referred to as a *commercial sponsor*. A reviewer is referred to as a *Regulatory Authority* (RA). Using the clinical trial evidence provided by a sponsor, it decides on marketing approval for a proposed treatment. If approval is granted, a *Health Care Insurer* (HCI) then decides on the level of payment for the new treatment. This terminology can be applied to the regional EU and US markets. In the US, the Food and Drug Administration (FDA) takes on the role as the RA. Payment is typically provided by private insurance companies, which thus constitute the HCI's. In the EU, the RA is instead the European Medicines Agency, with the expressed purpose of harmonising the work of the national level regulatory bodies within the EU (EMA, 2016). The HCI's in the EU are the country-level health care authorities or insurance companies.

Even if the rules that the RA uses to decide on whether or not to approve a new drug vary considerably in practice, some conventions are in use. For example, in section 3.5 of the International Council for Harmonisation of Technical Requirements for Pharmaceuticals for Human Use's guidance on efficacy E9 (ICH, 2016), it is stated that

Conventionally the probability of type I error is set at 5% or less or as dictated by any adjustments made necessary for multiplicity considerations; the precise choice may be influenced by the prior plausibility of the hypothesis under test and the desired impact of the results. The probability of type II error is conventionally set at 10% to 20%; it is in the sponsor's interest to keep this figure as low as feasible especially in the case of trials that are difficult or impossible to repeat. Alternative values to the conventional levels of type I and type II error may be acceptable or even preferable in some cases.

The particular value of 5% was suggested by Fisher (1946) as a convenient cutoff level to reject a null hypothesis. However, Fisher did not intend that this level be fixed regardless of application. Rather, he recommended that a specific level be set according to circumstances (Fisher, 1956). In addition

to requirements on what constitutes a demonstration of efficacy, the RA also considers safety aspects, but the precise nature of these are often not as explicit as the requirements on efficacy.

The level of reimbursement decided by the HCI is typically based on the benefit-risk balance and/or the cost of the treatment. Clearly, the extent to which monetary and purely health related concerns should be combined is a complex ethical issue. It is therefore no surprise that the HCI conduct differs greatly between individual countries. One HCI of particular importance to this thesis is the National Institute of Clinical Excellence (NICE) in the UK. In Paper I, NICE serves as a specific example of a HCI that is willing to directly associate an incremental health benefit with a monetary cost. In order to do this, the positive and negative health effects of the drug are combined so as to form a quantity referred to as Quality Adjusted Life Years (QALYs).

1.2.1 The phases of a clinical trial

The entire drug development process can be divided into the following stages: (1) drug discovery, (2) laboratory development, (3) animal studies, (4) clinical trials, and (5) regulatory approval. The focus in this thesis is on the last two stages. It is a widespread convention in the industry to further divide the stage consisting of the clinical trials into five different phases:

Phase 0 First-in-human trials in which subtherapeutic doses are given to a small number of subjects. The aim is to confirm that the drug's pharmacodynamics (what the drug does to the body) and pharmacokinetics (what the body does to the drug) are as expected.

Phase 1 Small trial typically involving less than a hundred subjects. The main purpose of this phase is to determine which dosage levels are safe and to begin to register any short-term side effects associated with the drug.

Phase 2 The main objective of this phase is to find the proper dosing for the new drug, striking a balance between its beneficial effects and any potential side effects.

Phase 3 This is the final confirmatory phase before acceptance or rejection of the new drug, and consists of one or several relatively large trials. The purpose is to confirm the efficacy of the drug and to ensure that any short-term or long-term side effects constitute an acceptable risk.

Phase 4 This phase consists of studies made after market introduction of the new drug. If conducted, the data provided can help in adjusting the dosage for different subpopulations and also be used to discover any rare side effects which remained undiscovered in phase 3.

In principle, BDT can be used to optimise all of the phases above. However, this thesis is mainly concerned with the design of confirmatory phase 3 trials. There are two main reasons for this. Firstly, confirmatory trials are the largest and most expensive ones, which suggests that most of the potential gain that may be obtained through optimisation can be traced to improvements of this phase. Secondly, since the regulatory approval stage immediately follows phase 3, it is possible to analyse, in a relatively direct way, how changes to the RA's approval rule or the HCI's payment rule affects the optimal trial design.

The present day gold standard for a clinical trial is one which is placebo-controlled, randomised and double-blinded. It would be safe to say that the data produced by such a trial is viewed by the wide majority of clinical researchers as constituting a solid basis for estimating the medical value of the drug investigated. A trial is said to be *placebo-controlled* if the active treatment is compared with some alternative, called the placebo, with no therapeutic effect. *Randomisation* means that patients recruited to the trial are assigned either the active treatment or placebo via some procedure involving chance. The design of proper randomisation procedures is a challenge in itself and there is a large literature on the subject. However, such procedures are not a topic of this thesis. For our purposes, it is perfectly fine to employ the simple conceptual model that each recruited patient in a trial with two alternatives is assigned a group based on the flip of a (possibly weighted) coin. A *single-blinded* study is one in which the patients do not know which of the treatment alternatives they are given. Obviously, this will only be practical if the placebo can be made nearly indistinguishable from the active treatment. The trial is called *double-blinded* when also the investigators responsible for the trial are kept uninformed about the specific alternative given to each patient until the study has been finished and evaluation of the results is at hand. The reason for such blinding is to prevent the investigators from handling the two groups differently, even if such influence only occurs on an unconscious level.

The costs involved in bringing a new drug to the market today are quite substantial. Using data from 2004 through 2012, Sertkaya et al. (2016) estimate the average costs for phase 1, 2 and 3 to be in the ranges 1.4-6.6, 7.0-19.6 and 11.5-52.9 million US dollars, respectively, where the upper and lower limits of the ranges correspond to different therapeutic areas. While there are exceptions, it typically takes several years to go through all the phases. An estimate of the average times required for the total clinical programme and the subsequent approval is provided by Kaitin and DiMasi (2011). In the US, for the 5-year period 2005-2009, average times of 6.4 years and 1.2 years, respectively, are reported. Note that several years have typically been spent on pre-clinical research before a drug enters the clinical development stage.

There are a large number of excellent books and review articles available on the subject of clinical trials. Chow and Liu (2014) provide a broad overview of both the practical and theoretical issues of the design and analysis of clinical trials. They cover the regulatory structure for approval of new treatments as implemented by the FDA and the basic statistical methods available for different kinds of trials. Spiegelhalter et al. (2004) is another good reference. Although the focus is on the use of the Bayesian approach for clinical trial design, it also covers issues connected to health-care evaluation. An introduction to adaptive Bayesian designs is given by Berry et al. (2011). A recent publication by Ondra et al. (2016) provides a review of different trial designs involving biomarkers that have been proposed in the literature, which provides a context for the results presented in papers II and III.

1.3 Brief summary of the contributed papers

The common thread in all four papers is that the clinical trial design problem is formulated using BDT. The optimal solution is then found by maximising expected utility. Since the general approach is fixed, the most challenging part of the work has been to find a good balance between model complexity and realism.

Paper I is concerned with a two-stage problem in which a commercial sponsor first chooses a sample size for a confirmatory trial. Given that the results are good enough for market approval, the second choice for the sponsor is that of a price for the new treatment. The ratio between the price and the estimated effect from the trial is finally used by a HCI to determine if the new treatment should be reimbursed via a public health care system. The main point of interest is that the HCI's Willingness To Pay (WTP) for a unit increase in effectiveness may be unknown to the sponsor when it proposes a price for the new treatment. It is shown that the optimal design of the sponsor depends on the degree of this uncertainty, and that it may in some cases be beneficial for both the sponsor and the HCI to reduce it. In other words, for HCI's employing the kind of cost-effectiveness threshold assumed in the model, an increased transparency on its precise value may be warranted.

Papers II and III both consider trial optimisation in a situation where the total population can be divided into two subgroups defined by means of a *biomarker*. According to Biomarkers Definitions Working Group (2001), a biomarker may be defined as

A characteristic that is objectively measured and evaluated as an indicator of normal biological processes, pathogenic processes, or pharmacologic responses to a therapeutic intervention.

A biomarker can be either *prognostic* or *predictive* (or both). A prognostic biomarker is one which can be used to characterise the general outcome for a patient with a certain condition, independently of any specific treatment. In contrast, a biomarker is predictive only with respect to a specific treatment. In general, there are both discrete and continuous biomarkers. Papers II and III only consider the special case of a predictive binary biomarker.

Paper II compares three different design types. All of them results in single-stage decision problems, where the main question is which sample size to choose. The optimal type of design depends on the circumstances, such as the prior beliefs about the effect sizes of the treatment in the two subgroups and the specific cost structure of the design. Paper III extends Paper II by also optimising adaptive, two-stage designs. Such designs allow the trial designer to look at the data before the confirmatory trial has been completed, and adapt the design based on the information obtained so far. Moreover, the designs considered in Paper III allows for adjusting the patient recruitment so as to obtain a trial subgroup prevalence that differs from the prevalence in the target population. It is shown that this additional flexibility leads to superior designs in some situations.

Both Paper II and Paper III consider the trial design problem from two different angles, from the viewpoint of a commercial sponsor and from the viewpoint of a public health decision maker. Since these two decision makers have different goals, the former being driven by profit maximisation and the latter by public health concerns only, whether or not a given design is considered optimal will depend on the perspective adopted.

The last contribution of this thesis, Paper IV, stand out from the others in that the optimisation is performed over a large number of stages. Each stage corresponds to a pair of patients recruited to the trial, and the objective is to find the optimal rule for when to stop the trial. The basic model analysed is well known in the literature as Anscombe's model. Instead of directly solving this multi-stage decision problem using the standard method of backward induction, a connection is established between the discrete time problem and an optimal stopping problem in continuous time. The solution of the latter problem then constitutes an approximation to the solution of the former. In contrast to Paper I, where a specific form of HCI reimbursement was analysed, and papers II and III, in which trial designs were optimised in a quite specific setting of biomarker subgroups, the model analysed in Paper IV is not directly applicable to a particular, realistic trial setting. Instead, it is formulated using a minimum of assumptions, and the conclusions that may be drawn from the form of its solution can therefore be argued to be more generally applicable than the ones obtained for the models in the other papers.

1.4 Optimal regulation

The optimal designs obtained in the papers should be of interest to any trial designer facing a problem close to the models developed. However, perhaps even more valuable are the conclusions that may be drawn from these analyses regarding the behaviour of commercial sponsors facing different types of regulatory structures. As noted previously, the specific rules in place differ between countries, but they typically provide at least partial answers to the following questions:

- How much evidence must be provided by the clinical trial before deciding on market approval?
- After the trial, how large should the estimated effect be in order to approve the treatment?
- How should companies be reimbursed after approval?

Suppose that the trial designer is a pharmaceutical company. A given, fixed set of rules will then lead to a design which is optimal from the company's perspective, but may not be so from a public health perspective. By changing the regulatory rules, the effect on public health will change and one sees the possibility of establishing an optimal set of rules. None of the papers contains a full analysis of this optimal regulation problem. However, the analyses provided give a basis for performing a local analysis of how the behaviour of the company changes when the regulatory rules are perturbed. The optimal regulation problem is a natural next step after trial optimisation, and the thesis therefore includes a broader discussion of this issue in Chapter 5.

1.5 Overview of thesis

Chapter 2 contains an introduction to BDT. It covers the basic structure that needs to be specified before optimisation can begin, such as the prior and the utility function, and explains the distinction between single-stage and sequential decision problems. The main tool used to solve the design problems in the thesis, backward induction, is considered in some detail and applied to a simple two-stage example for illustration. The material is aimed at readers with some basic knowledge of probability theory, but with little prior exposure to Bayesian statistics or decision theory. However, the content of the chapter is completely standard, and contains no sophisticated theoretical results. Hence, it can be safely skipped by readers with some prior experience with BDT methods.

Since the theory of optimal stopping for continuous-time stochastic processes is used in Paper IV, a very brief introduction is provided in Chapter 3. It will

be useful for readers with some knowledge of stochastic processes that wants a quick review of the basic concepts before reading Paper IV. In particular, it states a result involving the Girsanov transform, which turns out to be a vital tool in our approach to solving Anscombe's problem.

Chapter 4 provides a more detailed summary of each of the papers than that provided in the introduction. It may be used to gain a quick overview of the papers, and hopefully serves to guide the reader to those of particular interest. The material in Chapter 5, on optimal regulation, is largely independent of the models analysed in the papers. Essentially, it consists of a short report on a few problems I've worked with after the publication of the papers on optimal trial design. The models analysed are very simple, and none of the results obtained have been published. However, I do think that the basic approach described shows some promise as a vehicle for solving more realistic problems. A general discussion of clinical trial optimisation that connects the contributions of this thesis to others in the literature is given in Chapter 6, and I summarise my conclusions in Chapter 7.

The two chapters of the appendix contains supporting material of a more technical nature that I hope will help to make the thesis more self-contained. Appendix A introduces the concept of the family-wise error rate, used in papers II and III, and describes how it may be controlled using a multiple testing procedure. In particular, it provides a rather detailed discussion of a procedure called the Spiessens-Debois test, a crucial part of the model in Paper II. Finally, Appendix B gives a brief review of two basic results from real analysis, the implicit function theorem and the envelope theorem, which are used when showing the more technical results in Paper I.

Chapter 2

Bayesian decision theory

A formal derivation of BDT can be given from a basic set of axioms of rationality. These are mathematical versions of rules that, presumably, any decision maker that wants to avoid certain inconsistent behaviour would want to adhere by. Proceeding from these, it is possible to show the basic tenets of BDT, namely, that (1) degrees of belief should follow the laws of probability theory, (2) preferences for different outcomes should be representable as a utility function, and (3) that the optimal action when faced with uncertainty is always to choose the alternative that maximises expected utility. Since the topic of this thesis is on applications of BDT rather than fundamentals, an exposition of these results will not be given here. Instead, the principle of expectation maximisation is accepted as a valid approach to making rational decisions, and we will proceed to outline the basic structure needed to formulate and solve applied problems. For a comprehensive reference on the foundations see, for example, Bernardo and Smith (1994).

2.1 Decision problems

An abstract decision problem may be described as an ordered list with four components, $(\mathcal{D}, \mathcal{Q}, \pi, u)$, where

1. $\mathcal{D}$ is a set of available decisions, from which a particular choice $d \in \mathcal{D}$ is to be made.
2. $\mathcal{Q}$ is a set of possible values for all the quantities that are unknown to the decision maker. It is assumed that a specific value $q \in \mathcal{Q}$ corresponds to the true state of the world, but that it is unknown at the time the decision

must be made. Note that q may be a single real number or a vector of numbers.

3. π is a probability distribution on the set $\mathcal{Q}$ which formalises the personal beliefs held by the decision maker about the true state of the world. In the language of probability theory, $\mathcal{Q}$ is a sample space and π is a probability measure on $\mathcal{Q}$.
4. $u = u(d, q)$ is a utility function mapping each possible combination of a decision d and state q into a real number. The interpretation is that $u(d, q)$ is the personal value that the decision maker places on making the decision d under the assumption that q is the true state of the world. The utility function formalises the decision maker's preferences.

Having fully specified the problem, the objective of a rational decision maker is to find the decision which maximises the expected utility, that is, the goal is to solve

$$\max_{d \in \mathcal{D}} \mathbb{E}[u(d, Q)], \quad (2.1)$$

where Q is a random variable with distribution π .¹

2.2 Sequential decision problems

Suppose that the decision maker faces a multi-stage problem in which n decisions are to be made in a sequence. At each stage, an element $d_i \in \mathcal{D}_i$ is to be selected, $i = 1, \dots, n$. It is assumed that this selection leads to an outcome $x_i \in \mathcal{X}_i$ which is observed by the decision maker before proceeding to stage $i + 1$. This has two consequences. First, the decision maker can condition its probability distribution for all future observations and other unknown quantities on the observations $x_1, \dots, x_i$ made so far. Second, the decision maker can also let the decision d_{i+1} depend on $x_1, \dots, x_i$, since these values are known at the time the decision is to be made. Hence, while d_1 is just selected as an element of $\mathcal{D}_1$, d_2 may be taken as the value of a function which maps each possible outcome x_1 into an element of $\mathcal{D}_2$. In other words, the search is not for an optimal sequence of pure decisions, but for a sequence of *decision rules* $\delta_1, \dots, \delta_n$ which maps previous observations into pure decisions at each stage. The entire sequence of decision rules is often called a *policy* or *plan*, the idea being that the decision maker can think about and commit to a course of action for the entire problem before making the first decision.

¹We follow the convention of denoting random variables by capital letters and specific realisations by the corresponding small letters.

This sequential setup can be formulated as an abstract decision problem of the form described in the previous section. Set²

$$\begin{aligned}\mathcal{D} &= \{(\delta_1, \dots, \delta_n) \mid \delta_1 \in \mathcal{D}_1, \dots, \delta_n : \mathcal{X}_1 \times \dots \times \mathcal{X}_{n-1} \rightarrow \mathcal{D}_n\}, \\ \mathcal{Q} &= \mathcal{X}_1 \times \dots \times \mathcal{X}_n \times \mathcal{T}.\end{aligned}$$

Here, $\mathcal{T}$ denotes a space containing all unknown quantities except the observations. An element of this space will be denoted by θ .

In this abstract formulation, π is a probability measure on $\mathcal{X}_1 \times \dots \times \mathcal{X}_n \times \mathcal{T}$ which will be updated using Bayes' rule after each observation is made. The abstract form of the utility function is $u = u(d, q)$, as before. However, if we make the very natural assumption that the utility should only depend on the pure decisions made at each stage, and not on the parts of δ_i which are not used given the previous observations, then we may write

$$u(d, q) = u(\delta_1, \dots, \delta_n, x_1, \dots, x_n, \theta) = u(d_1, \dots, d_n, x_1, \dots, x_n, \theta),$$

where

$$d_1 = \delta_1, d_2 = \delta_2(x_1), \dots, d_n = \delta_n(x_1, \dots, x_{n-1}).$$

So the utility function only needs to be specified for every sequence of pure decisions, observations and possible value for θ .

2.3 Backward induction

It should be intuitively clear that in order to find the optimal decision today, the decision maker must consider which situations that are possible tomorrow and how these should be optimised. Hence, to find the optimal plan in a sequential decision problem, it is necessary to look ahead and think about the future. This basic idea may be formalised and turned into a method for solving sequential decision problems which is called *backward induction*. Conceptually, the method works for all decision problems with a finite number of stages, although the computations required grows with the number of stages and may eventually make the problem intractable. For simplicity, the method will now be illustrated for a two-stage problem. It works exactly the same for any finite number of stages.

Backward induction finds the optimal plan $d^* = (\delta_1^*, \delta_2^*(x_1))$ by proceeding backwards in the decision problem while computing a sequence of induced utilities. The first step is to compute the expected utility of choosing d_2 given that d_1 led to x_1 in the first stage,

$$\bar{u}(d_1, d_2, x_1) \equiv \mathbb{E}[u(d_1, d_2, x_1, \Theta) \mid d_1, d_2, x_1].$$

²As usual, the symbol $\times$ denotes the Cartesian product operator.

This is done for all triples $(d_1, d_2, x_1) \in \mathcal{D}_1 \times \mathcal{D}_2 \times \mathcal{X}_1$. Since the decision maker is free to choose any $d_2 \in \mathcal{D}_2$, the rational choice is to select the one maximising $\bar{u}$. Therefore, the optimal decision rule is

$$\delta_2^*(d_1, x_1) = \arg \max_{d_2 \in \mathcal{D}_2} \bar{u}(d_1, d_2, x_1),$$

with a corresponding induced expected utility

$$\bar{u}^*(d_1, x_1) \equiv \bar{u}(d_1, \delta_2^*(d_1, x_1), x_1) = \max_{d_2 \in \mathcal{D}_2} \bar{u}(d_1, d_2, x_1).$$

The next step is to find the pure decision $\delta_1 = d_1$ that maximises the expected utility of observing the outcome x_1 , given that the optimal rule $\delta_2^*(d_1, x_1)$ is subsequently followed. The induced expected utility corresponding to a specific d_1 may be written as

$$\tilde{u}(d_1) \equiv \mathbb{E}[\bar{u}^*(d_1, X_1)|d_1].$$

The optimal choice of $\delta_1 = d_1$ is therefore

$$d_1^* \equiv \arg \max_{d_1 \in \mathcal{D}_1} \tilde{u}(d_1),$$

with a corresponding optimal expected utility

$$\tilde{u}^* \equiv \tilde{u}(d_1^*) = \max_{d_1 \in \mathcal{D}_1} \tilde{u}(d_1).$$

It may be shown that the optimal plan $d^* = (\delta_1^*, \delta_2^*(x_1))$ constructed in this way solves the problem in Eq. (2.1).

2.4 Example: Optimisation of sample size

In order to illustrate an application of the method of backward induction to a problem in the area of clinical trial design, this section presents an example of sample size optimisation in the context of a parallel-group trial comparing a new treatment A with a placebo alternative B . A balanced randomisation of n patients to the two groups is assumed, implying a per-group sample size of $n/2$. The unknown, incremental efficacy of A vs. B is denoted by θ . The individual responses in the two arms are assumed to be independent and identically distributed random variables which are combined into two sample means, $\bar{X}_A$ and $\bar{X}_B$. The variance of each individual response is for simplicity assumed to be known and equal to σ^2 . Based on the central limit theorem, the difference

$\bar{X} = \bar{X}_A - \bar{X}_B$ is used to estimate θ assuming a normal distribution according to

$$\bar{X} \mid n, \theta \sim \mathcal{N}(\theta, 4\sigma^2/n).$$

To simplify the calculations, a normal conjugate prior with zero mean is assumed for θ ,

$$\theta \sim \mathcal{N}(0, 4\sigma^2/n_0).$$

This normal conjugate model implies that the posterior distribution of θ given an observed value $\bar{x}$ is also normal (Raiffa and Schlaifer, 1961),

$$\theta \mid n, \bar{x} \sim \mathcal{N}\left(\frac{n\bar{x}}{n_0 + n}, \frac{4\sigma^2}{n_0 + n}\right).$$

Writing the prior variance as $4\sigma^2/n_0$ makes it easy to compare the amount of information about θ in the prior with the information about θ provided by the trial sample.

The decision procedure is as follows. First, a non-negative sample size $n \in [0, N]$ is selected, where $N > 0$ denotes the total size of the target population for the new treatment. After the trial, a decision $d \in \{0, 1\}$ is taken on which treatment to give to the patients in the target population, where $d = 1$ and $d = 0$ correspond to treatment A and B , respectively.

The utility of choosing a sample size n , observing a sample mean $\bar{x}$, and taking the post-trial decision d , given that θ is the true incremental efficacy, is taken to be the aggregated, incremental efficacy for both in-trial and post-trial patients:

$$u(n, d, \bar{x}, \theta) = \frac{n\bar{x}}{2} + d(N - n)\theta.$$

It follows that the induced utility of choosing n , observing $\bar{x}$ and subsequently selecting a treatment option d is

$$\bar{u}(n, d, \bar{x}) = \mathbb{E}\left[\frac{n\bar{x}}{2} + d(N - n)\theta \mid n, \bar{x}\right] = \frac{n\bar{x}}{2} + d(N - n)\left(\frac{n\bar{x}}{n_0 + n}\right).$$

Maximising over $d \in \{0, 1\}$ gives

$$\bar{u}^*(n, \bar{x}) = \frac{n\bar{x}}{2} + \frac{(N - n)n}{n_0 + n} \max(0, \bar{x}).$$

The prior predictive distribution for $\bar{X}$ is given by $\bar{X} \sim \mathcal{N}(0, 4\sigma^2/n_0 + 4\sigma^2/n)$. This implies that³

$$\mathbb{E}[\max(0, \bar{X}) \mid n] = \frac{1}{\sqrt{2\pi}} \sqrt{4\sigma^2 \left(\frac{1}{n_0} + \frac{1}{n}\right)}.$$

³It is straightforward to show that if Y is a random variable distributed as $Y \sim \mathcal{N}(0, \sigma_Y^2)$, then $\mathbb{E}[\max(0, Y)] = \sigma_Y/\sqrt{2\pi}$.

It follows that

$$\tilde{u}(n) = \frac{(N-n)n}{n_0+n} \mathbb{E}[\max(0, \bar{X})|n] = \frac{N-n}{\sqrt{2\pi}} \sqrt{\frac{4\sigma^2}{n_0}} \sqrt{\frac{n}{n_0+n}}.$$

Maximisation of this function gives

$$n^* = \frac{n_0}{4} \left(\sqrt{9 + \frac{8N}{n_0}} - 3 \right),$$

which in turn leads to the asymptotic result

$$n^* \sim \sqrt{\frac{n_0}{2}} \sqrt{N}, \quad N \rightarrow \infty.$$

A recent contribution by Stallard et al. (2017) contains a detailed analysis of a generalised version of the example considered in this section. They show that the optimal sample size is $O(\sqrt{N})$ as $N \rightarrow \infty$, under the assumptions that the distribution for the primary endpoint has a one-parameter exponential family form and that the utility for each patient is a continuous function of the parameter. In addition to the result for a fixed value of N , they also extend the asymptotic analysis to the case in which N is unknown and has a geometric distribution.

Chapter 3

Optimal stopping problems in continuous time

Paper IV stands out from the other three since the analysis of Anscombe's problem requires results from the theory of optimal stopping for continuous-time Markov processes. In order to make this thesis a bit more self-contained, this chapter therefore provides a very brief review of the basic concepts.

3.1 Markov processes and optimal stopping

The material in this section is mostly based on the more comprehensive introduction given in Christensen (2016). A *filtered probability space* $(\Omega, \mathcal{F}, \mathbb{P}, (\mathcal{F}_t)_{t \geq 0})$ consists of a sample space Ω (with outcomes $\omega \in \Omega$), a σ -algebra $\mathcal{F}$ of subsets of Ω that is interpreted as the set of events, a probability measure $\mathbb{P} : \mathcal{F} \rightarrow [0, 1]$ satisfying Kolmogorov's axioms and a filtration $(\mathcal{F}_t)_{t \geq 0}$ of σ -algebras satisfying $\mathcal{F}_{t_1} \subseteq \mathcal{F}_{t_2} \subseteq \mathcal{F}$ for all $t_1 \leq t_2$. A *continuous-time process* $(X_t)_{t \geq 0}$ is defined as a collection of random variables $X_t : \Omega \rightarrow E \subseteq \mathbb{R}^n$, indexed by the time parameter t . The *state space* of the process is the pair $(E, \mathcal{B}(E))$, where $\mathcal{B}(E)$ is the Borel σ -algebra of E .

The dynamics for many kinds of processes (e.g., diffusion or Markov processes) are determined by the starting state x , an element in E such that $\mathbb{P}(X_0 = x) = 1$. It is then of interest to consider the entire family of processes (X_t^x) corresponding to a set of such starting states. However, it is often more convenient to consider a single process (X_t) and handle the different starting states by introducing a family of probability measures $\mathbb{P}_x$ that conditions the process to start in x . Formally, we then have $\mathbb{P}_x((X_t) \in \cdot) = \mathbb{P}((X_t^x) \in \cdot)$. The

process (X_t) is said to be *adapted to the filtration* $(\mathcal{F}_t)_{t \geq 0}$ if X_t is a $(\mathcal{F}_t, \mathcal{B}(E))$ -measurable function for each $t \geq 0$.

We are now ready to define the central concepts of *stopping time* and *Markov process*.

Definition 3.1.1 (Stopping time): A random variable $\tau : \Omega \rightarrow [0, \infty]$ is a stopping time with respect to the filtration $(\mathcal{F}_t)_{t \geq 0}$ if the event $\{\tau \leq t\} \in \mathcal{F}_t$ for all $t \geq 0$.

Definition 3.1.2 ((Strong) Markov process): $(X_t)_{t \geq 0}$ is called a time-homogeneous Markov process on $(\Omega, \mathcal{F}, (\mathbb{P})_{x \in E}, (\mathcal{F}_t)_{t \geq 0})$ if

1. The function $x \mapsto \mathbb{P}_x(X_t \in B)$ is measurable for all $t \geq 0$ and $B \in \mathcal{B}(E)$.
2. For all $\omega \in \Omega$ and $t > 0$, there exists $\omega' \in \Omega$ such that $X_{t+s}(\omega) = X_s(\omega')$ for all $s \geq 0$.
3. $\mathbb{P}_x(X_0 = x) = 1$ for all $x \in E$.
4. For all $s, t \geq 0$, $x \in E$ and $B \in \mathcal{B}(E)$, we have

$$\mathbb{P}_x(X_{t+s} \in B \mid \mathcal{F}_t) = \mathbb{P}_{X_t}(X_s \in B), \quad \mathbb{P}_x \text{ almost surely.}$$

If, in addition, we have

$$\mathbb{P}_x(X_{\tau+s} \in B \mid \mathcal{F}_\tau) = \mathbb{P}_{X_\tau}(X_s \in B), \quad \mathbb{P}_x \text{ almost surely on } \{\tau < \infty\},$$

for all $s \geq 0$, $B \in \mathcal{B}(E)$ and stopping times τ , then (X_t) is a strong Markov process. It is property 4 that is most important in this definition. It may be interpreted as stating that the future of a Markov process from time t and onwards only depends on the state x at time t , and not on the entire history up to time t .

Let $G : E \rightarrow \mathbb{R}$ be a measurable function such that

$$\mathbb{E}_x \left[\sup_{t \geq 0} |G(X_t)| \right] < \infty, \quad \text{for all } x \in E. \quad (3.1)$$

This function G is referred to as the *reward function* of an optimal stopping problem. Having introduced an appropriate set of stopping times $\mathcal{T}$, the *value function* is defined as

$$V(x) = \sup_{\tau \in \mathcal{T}} \mathbb{E}_x [G(X_\tau)]. \quad (3.2)$$

The right hand side above essentially defines the optimal stopping problem. To solve the problem, we have to find the value function $V(x)$ and an optimal stopping time $\tau^* \in \mathcal{T}$ satisfying

$$V(x) = \mathbb{E}_x [G(X_{\tau^*})], \quad \text{for all } x \in E.$$

Eq. (3.1) ensures that the value function $V(x)$ is well-defined for all $x \in E$. Moreover, since $\tau = 0$ is a finite stopping time, it must satisfy

$$V(x) \geq \mathbb{E}_x[G(X_0)] = G(x), \quad \text{for all } x \in E.$$

The assumption that we are working with Markov processes can be exploited when searching for an optimal stopping time. For suppose that we have observed a path of the process corresponding to a specific outcome ω until $t > 0$. If we have not stopped so far and if $X_t(\omega) = X_0(\omega)$, then since (X_t) is a Markov process we are at time t faced with the same optimal stopping problem that we were at time 0. Hence, if it was optimal not to stop at time 0, we should not stop at time t either. We can conclude that the decision of whether to stop (and get the reward $G(X_t)$) should only be based on the current state of the process at time t , not on the whole path up to this time point. This heuristic argument motivates the following definitions.

Definition 3.1.3 (Stopping region and continuation region): Let

$$\begin{aligned} \text{Stopping region} &= \{x \in E : V(x) = G(x)\}, \\ \text{continuation region} &= \{x \in E : V(x) > G(x)\}. \end{aligned}$$

The problem has been reduced to finding a way to characterise the stopping region. If this can be done, an optimal stopping time consists of observing the process until $X_t = x$ for some state x in the stopping region.

In Paper IV, we are mostly interested in applying the theory briefly reviewed here to problems with finite time horizons, in which the allowed stopping times are those which satisfy $0 \leq \tau \leq T$ for some constant T . Moreover, the reward function will depend on the current time t . This generalisation leads to no great difficulties, since we may consider the time-space process (t, X_t) , which is again Markov provided that (X_t) is. The value function is in this case defined as

$$V(t, x) = \sup_{0 \leq \tau \leq T-t} \mathbb{E}_{(t, x)}[G(t + \tau, X_{t+\tau})], \quad t \in [0, T], \quad x \in E.$$

3.2 Girsanov's transform

A simple version of this central result that is enough for our purposes is stated here, based on the formulation given by Øksendal (2003, Theorem 8.6.4).

Theorem 3.2.1: Let $Y_t \in \mathbb{R}$ be an Itô process of the form

$$dY_t = a(t)dt + dW_t, \quad 0 \leq t \leq T, \quad Y_0 = 0,$$

where $a(t)$ is a real-valued function and (W_t) is a standard Brownian motion. Define a process

$$M_t = \exp \left(- \int_0^t a(s) dW_s - \frac{1}{2} \int_0^t a^2(s) ds \right), \quad 0 \leq t \leq T.$$

Suppose (M_t) is a martingale with respect to $\mathcal{F}_t$ and define the measure $\mathbb{Q}$ on $\mathcal{F}_T$ by

$$d\mathbb{Q}(\omega) = M_T(\omega)d\mathbb{P}(\omega).$$

Then $\mathbb{Q}$ is a probability measure on $\mathcal{F}_T$ and (Y_t) becomes a standard Brownian motion with respect to $\mathbb{Q}$ for $0 \leq t \leq T$.

To illustrate the result, we here apply it to the situation of main interest in Paper IV.

Example 3.2.1: Suppose the process (Σ_t) is a diffusion of the form

$$d\Sigma_t = \mu dt + \sigma dW_t, \quad \Sigma_0 = 0,$$

where $\mu \in \mathbb{R}$ and $\sigma > 0$ are constants. Dividing through by σ yields

$$d\left(\frac{\Sigma_t}{\sigma}\right) = \left(\frac{\mu}{\sigma}\right) dt + dW_t.$$

An application of Girsanov's transform therefore implies that (Σ_t/σ) is standard Brownian motion with respect to $\mathbb{Q}$, where

$$d\mathbb{Q} = \exp \left(-\frac{\mu}{\sigma} W_t - \frac{1}{2} \frac{\mu^2}{\sigma^2} t \right) d\mathbb{P}.$$

Since

$$\Sigma_t = \mu t + \sigma W_t \iff W_t = \frac{\Sigma_t - \mu t}{\sigma},$$

$d\mathbb{Q}$ may also be written as

$$d\mathbb{Q} = \exp \left(-\frac{\mu}{\sigma^2} \Sigma_t + \frac{1}{2} \frac{\mu^2}{\sigma^2} t \right) d\mathbb{P}.$$

Chapter 4

Summary of papers

4.1 Paper I: Late-stage pharmaceutical R&D and pricing policies under two-stage regulation

Paper I investigates R&D incentives for the pharmaceutical industry in the presence of two exogenous regulatory stages. In Stage 0, a commercial sponsor deliberates on whether to run a phase 3 trial and, if it decides to go ahead, selects the sample size of the trial. The trial results in an estimate x of the incremental effectiveness. Upon trial completion, a RA in charge of granting access to a market considers the evidence produced by the trial. Approval for marketing is granted if the sample size is large enough and the new treatment shows superiority to the current standard alternative at a one-sided level of 2.5%. In Stage 1, a price is proposed by the sponsor for the new treatment. Combined with x , this price determines the Incremental Cost-Effectiveness Ratio (ICER) upon which a HCI bases its reimbursement decision. The optimal policy for the sponsor over both stages is found using backward induction.

From the perspective of the sponsor, the value of the HCI's maximum WTP for a unit increase in effectiveness is uncertain and is modelled using a continuous random variable W . It is assumed that W belongs to a location-scale family of random variables, implying that any member can be uniquely characterised in terms of a pair (m, s) , where m is the expected value (or location parameter) of W and the scale, s , can be considered a measure of how uncertain the sponsor is about the HCI's WTP. We identify three ranges for the uncertainty parameter s , in which increases in uncertainty have different effects. In the *low uncertainty* range, increases in s result in lower optimal prices, lower optimal expected profits and a smaller optimal trial sample size. In the *high uncertainty* range, the situation is reversed: greater uncertainty leads to higher prices, higher expected

profits and a larger trial sample size. For *intermediate uncertainty*, prices are increasing, expected profits decreasing and sample size decreasing in the degree of uncertainty. Hence, for the range of intermediate uncertainty, a smaller value for s benefits the sponsor, the HCI and the patients.

The framework is applied to a recent NICE appraisal of mannitol dry powder for treating cystic fibrosis. The status of cystic fibrosis as a rare disease means that the R&D decision could potentially be considered to be a marginal project, that is, one with a market size that is close to the minimum population size required for the investment to be deemed profitable. We investigate how the RA parameters defining the one-sided significance level and the minimum sample size required for marketing authorisation impact the minimum size of the target population that the sponsor requires in order to expect a positive profit when acting optimally.

4.2 Paper II: Optimizing trial designs for targeted therapies

Paper II is concerned with clinical trials in which the efficacy of a treatment is tested in an overall population and/or in a pre-specified subpopulation defined by a binary biomarker. Trial optimisation is done from two perspectives, that of a commercial sponsor and from the viewpoint of a public health decision maker.

For both perspectives, three different types of trial designs are considered. These are referred to as the *classical design*, the *stratified design* and the *enrichment design*. The classical design makes no use of the biomarker status and only tests for a treatment effect in the full population. This is done using a standard, parallel-group trial with equal group sizes. The stratified design also recruits patients from the full population, but the biomarker status of each patient is determined and the treatment effect is tested in the full population and in the subpopulation. This implies that the stratified design may lead to approval in either the full population or in the subpopulation only, which necessitates an appropriate control of the FWER. Such control is implemented using the closed Spiessens-Debois test (Spiessens and Debois, 2010). In the enrichment design, patients are screened for their biomarker status and only biomarker positive patients are included in the trial.

The use of the biomarker test in the stratified and enrichment designs implies that a fixed cost must be paid to develop the screening procedure. Moreover, a marginal cost must be paid for each patient screened. These biomarker costs are not present for the classical design. By comparing the optimal expected utilities for these three design types, the framework allows us to assess when it is favourable to determine the biomarker status of the patients in a clinical

trial and when it is actually more efficient to disregard the biomarker and to proceed with a classical trial design.

The sample size is optimised for each of the three design types. Moreover, for the stratified design, the two significance levels defining the Spiessens-Debois test are also optimised. This optimisation is done with respect to a prior that encodes the pre-trial knowledge about the efficacy of the treatment by means of a two dimensional distribution on the true effect sizes in the full population and the subpopulation. The utility functions considered account for the different costs of the design types as well as the expected benefit when demonstrating efficacy in the subpopulations.

Examples of trial designs obtained by numerical optimisation are presented for both perspectives. We find that the optimal type of design depends sensitively on the various parameters of the framework. A parameter of particular interest is the prevalence of the biomarker positive patients in the total target population, and we consider its impact in detail.

4.3 Paper III: Optimized adaptive enrichment designs

In this paper two types of clinical trial designs are optimised. Both may be referred to as *partial enrichment designs*, one single-stage and an adaptive, two-stage generalisation. The setting in which these designs are evaluated is one in which there is an a priori biological plausibility that the treatment effect is larger, or only present, in a subgroup defined by a binary biomarker. The term *partial enrichment* here refers to the method of adjusting the recruitment to the trial so that the trial subgroup prevalence differs from the prevalence in the total population. Hence, there may be over- or underrepresentation of the subpopulation in the trial. A full enrichment design is a special case of such a partial enrichment design where only patients from the subgroup are recruited.

Parallel group trials with normally distributed outcomes are assumed. As in Paper II, the optimisation is performed from the two perspectives of a commercial sponsor and a public health decision maker. The total population F is divided into a subgroup S and its complement S' . With a subgroup prevalence λ (assumed known), the treatment effects for F , S and S' are, respectively, $\delta_F = \lambda\delta_S + (1 - \lambda)\delta_{S'}$, δ_S and $\delta_{S'}$. The trial designs test the hypotheses $H_F : \delta_F \leq 0$ and $H_S : \delta_S \leq 0$, while controlling the FWER at a pre-specified one-sided level. Although more powerful procedures are available, for simplicity, and since it allows for the utilisation of numerical integration rather than simulation, the Bonferroni correction is used to adjust for multiplicity.

A complete specification of the problems in this paper depends on a number

of different parameters. Priors for effect sizes in the subgroups need to be fixed, the costs involved must be set to specific values, and the form of the utility functions must be determined before moving on to the optimisation. This leads to a large space of possible problem configurations. No attempt was made to perform a systematic numerical investigation of this space. Instead, we focused on a few specific scenarios while attempting to select realistic values for the parameters. For these scenarios, a rather detailed analysis is provided, showing the optimal expected utilities for both perspectives (commercial sponsor vs. public health) as well as the optimal trial sample sizes for the different designs. Our results suggest that partial enrichment designs can lead to substantial improvements of the expected utilities in some situations.

4.4 Paper IV: Anscombe’s model for sequential clinical trials revisited

The basic objective in this paper is to find the optimal sequential procedure for testing a new treatment A , and, subsequently, choosing either A or a standard alternative B . The goal is to maximise the expected utility for a finite population of size N , which includes both the patients recruited to the trial and the ones that are subsequently given the selected treatment once the trial is complete. In contrast to the more applied settings of Paper II and Paper III, we here wish to focus on the mathematical aspects of such a model. Hence, the model obtained is intentionally very simple. For example, it assumes that there are no trial costs and that the treatments are indistinguishable in terms of side effects. It has been studied in numerous other works in the literature since it was first formalised by Anscombe. Therefore, our aim in this article is not only to apply the modern optimal stopping theory to this well-known problem, but to also consider a number of generalisations.

The first generalisation allows for a prior on the incremental effect size that is not conjugate normal. A very important first step in handling the problem for such a general prior is to first apply Girsanov’s transform to the sum process that equals the total utility for all patients in the trial up to time point t . The second generalisation breaks the symmetry between the two treatments. The current standard B is assumed to be used in parallel with the ongoing trial, so that there are always more patients treated with this alternative as time goes. This changes the optimal stopping rule, and, in particular, we are able to generalise certain asymptotic result for the optimal stopping rule in the limit case of a vague prior and a large patient horizon. The third extension of the model considers the possibility that the patient horizon N is unknown, and hence must be modeled as a random variable instead of a known constant. It

turns out that certain specific distributional assumptions for N allows for an explicit solution for the optimal stopping rule.

The most fundamental assumption in the model is that the response variables are normally distributed. This makes it possible to reformulate the problem as a Markovian optimal stopping problem in continuous time. The numerical solution of this latter problem, which proceeds by solving certain integral equations that follows from the free-boundary approach, then provides an approximation to the original discrete time problem (over patients $1, \dots, N$). This is a crucial step, because although the discrete time problem can in principle be solved using backwards induction, in practice this method becomes intractable as N grows.

The stopping rule derived is optimal from the perspective of a decision maker that is (1) able to perform a fully sequential trial, and (2) is only concerned with the public health aspects of the trial (i.e., that disregards any trial costs). This does not reflect the typical situation in drug development, where a commercial sponsor performs a trial and then submits the trial data to a regulator, which in turn decides on market approval. Nevertheless, we demonstrate that the solution obtained for the idealised model can be used to argue that the classical evaluation rule of comparing the p-value with a fixed significance level is suboptimal in certain situations.

Chapter 5

Optimal regulation of clinical trials

A pharmaceutical company optimises a trial relative to the environment shaped by the regulatory rules for acceptance of new drugs, together with the structure in place for monetary reimbursement. In this section, we'll refer to all these rules that in the end determine the profit for the company as the *market environment*. As discussed in Section 1.2, the market environment, and the different stakeholders involved in it, varies between different countries. However, in order to simplify the treatment of this topic, we will in this chapter assume that all governmental bodies (RA, HCI, etc.) are lumped together in a single entity. Following the game-theoretic literature, we will refer to this entity as the *principal*, and the pharmaceutical company will be called the *agent*. It is assumed that the principal's objective is to maximise the health gain of the patients in a certain population, say all patients afflicted with a certain disease or all patients within a certain country. For the principal, then, the basic question is: what is the optimal market environment? In other words, which incentives should be put in place in order to maximise the health gain of the target patient population over time, given that basic research and subsequent clinical trials are performed by profit seeking agents?

The central part of the models below is that the agent performs a phase 3, confirmatory trial that results in an estimate X of the true mean population effect Θ . For simplicity, safety issues are ignored. Letting N denote the size of the target population, not including the subjects enrolled in the trial, it follows that the total expected incremental health effect if the treatment is approved is $N\Theta$. We assume throughout that the sample size for the trial is fixed, so that the only decision that the agent makes is whether or not to perform the trial.

Before the trial, the mean effect Θ is unknown to the principal. This uncertainty is described by a prior density f_Θ . However, we will assume that the agent is much better informed and actually knows Θ when deciding on whether or not to go ahead with the trial.

5.1 Basic skeleton model

First, we'll define a basic decision procedure. The two specific models that will be solved in the following sections are both obtained as modifications of this procedure.

Step 1 The principal chooses an incentive structure $(a(x), t(x))$. $a : \mathbb{R} \rightarrow \{0, 1\}$ maps the observed effect estimate x into a binary decision, where $a(x) = 1$ means that the new drug is distributed to the patient population and $a(x) = 0$ means that it is not. $t : \mathbb{R} \rightarrow \mathbb{R}$ defines a monetary transfer from the principal to the agent based on the estimate x .

Step 2 Knowing Θ , the agent decides whether or not to perform a trial. The decision is a function $d : \mathbb{R} \rightarrow \{0, 1\}$, where $d(\theta)$ corresponds to a GO decision and $d(\theta) = 0$ to a NO GO decision. The trial cost is denoted by K .

If the agent makes a NO GO decision in Step 2, then neither the agent nor the principal will make any gain. If $d(\theta) = 1$, then the agent pays K to perform the trial and obtains a monetary reward $t(X)$. The utility for the principal equals the difference between $\gamma N\Theta$ and the payment $t(X)$ to the agent, where γ is a fixed multiplier which defines the monetary value that the regulator places on one additional unit of health. Hence, the utilities for the two actors in this model are

$$\begin{aligned} U_P &= d(\Theta) (a(X)\gamma N\Theta - t(X)), \quad \text{for the principal,} \\ U_A &= d(\Theta) (t(X) - K), \quad \text{for the agent.} \end{aligned}$$

Our general approach for solving models of this form proceeds as follows. First, we solve the decision problem for the agent for each fixed configuration of incentives, assuming it to be a rational expected utility maximiser. This will lead to a plan of actions and a corresponding optimal utility for the agent, both of which are functions of the incentives. Each such agent plan will determine an expected utility for the principal. In the second step, we search for the optimal incentives in some appropriate space, assuming that the agent always responds optimally according to the plan the was found in the first step.

Of course, the solution method just described is nothing else than backward induction, applied to a sequential game involving two players. It leads to a

Bayesian equilibrium, in which each strategy is an optimal response to the other. Let's illustrate this approach by applying it to the basic skeleton problem. In doing so, we make the following simplifying assumptions,

Perfect trial The trial is assumed to be so large that the remaining uncertainty regarding Θ after it is concluded may be neglected. Hence, $X = \Theta$ after the trial.

Known trial cost K is assumed known to both the principal and the agent.

The first step is to find the optimal d for the agent, given a fixed incentive structure (a, t) . Since a perfect trial is assumed,

$$\mathbb{E}[U_A \mid \Theta = \theta] = d(\theta)(t(\theta) - K) \implies d^*(\theta) = \mathbb{I}\{t(\theta) \geq K\},$$

where $\mathbb{I}$ denotes the indicator function. Assuming an optimal response by the agent, the perfect trial assumption implies that the principal's problem is

$$\max_{(a,t)} \mathbb{E}[d^*(\Theta)(a(X)\gamma N\Theta - t(X))],$$

where

$$\mathbb{E}[d^*(\Theta)(a(X)\gamma N\Theta - t(X))] = \int_{-\infty}^{\infty} \mathbb{I}\{t(\theta) \geq K\} (a(\theta)\gamma N\theta - t(\theta)) f_{\Theta}(\theta) d\theta.$$

It is immediately checked that the optimal incentives are given by

$$a^*(x) = \mathbb{I}\{x \geq 0\}, \quad t^*(x) = K\mathbb{I}\left\{x \geq \frac{K}{\gamma N}\right\},$$

implying a corresponding optimal expected utility of

$$\mathbb{E}[U_P] = \int_{-\infty}^{\infty} (\gamma N\theta - K)^+ f_{\Theta}(\theta) d\theta.$$

5.2 Model with pre-clinical investments

In this section, we assume that the size of the payment to the agent not only affects its willingness to go ahead with a trial, but also its willingness to invest in pre-clinical research. As before, we assume that the effect of the new treatment follows a density $f_{\Theta}(\theta)$. In the present model, this density is interpreted as arising from the frequency distribution of a stream of potential new drugs. When a new drug is discovered in the pre-clinical phase, the agent learns the true value of Θ and decides on whether or not to perform a perfect trial in order to obtain market approval.

The model differs in two important ways from the basic one. First, the agent can affect the probability that pre-clinical research leads to a drug discovery by adjusting an investment level $I \geq 0$. Second, instead of a general transfer function $t(x)$, we assume that payment must be proportional to the effect size demonstrated in the trial. We refer to the constant of proportionality as the WTP of the principal, and aim to find the optimal WTP level. Specifically, the decision process consists of

Step 1 The principal chooses an incentive structure $(a(x), \lambda)$. $\lambda > 0$ is interpreted as the payment per unit of incremental health per treated patient.

Step 2 Not yet knowing Θ , the agent decides on a level of investment $I \geq 0$. With probability $p(I)$, there is a discovery of a new drug candidate. Let B be a binary random variable, with $B = 1$ and $B = 0$ corresponding to discovery and no discovery, respectively.

Step 3 Having learned the true value of Θ , the agent decides whether or not to perform a trial.

Since it is a probability, the function p must always satisfy $0 \leq p(I) \leq 1$. In addition, we make the following assumptions:

$$p(0) = 0, \quad p'(I) > 0, \quad p''(I) < 0.$$

The utilities for the principal and the agent are defined as

$$\begin{aligned} U_P &= Bd(\Theta)(a(X)\gamma N\Theta - \lambda NX), \\ U_A &= Bd(\Theta)(\lambda NX - K) - I. \end{aligned}$$

As for the basic skeleton model, γ is a factor that converts health units into monetary units.

Before finding the WTP that maximises $\mathbb{E}[U_P]$, the optimal response of the agent for each specific value of λ must be found. Given that the investment made in Step 2 leads to a discovery of a new drug with effect θ , it is evident from the form of U_A that

$$d^*(\theta) = \mathbb{I} \left\{ \theta \geq \frac{K}{\lambda N} \right\}.$$

Hence, the expected agent utility of making the investment I and then continue

optimally is

$$\begin{aligned}\mathbb{E}[U_A \mid I, d^*] &= p(I) \int_{K/(\lambda N)}^{\infty} (\lambda N \theta - K) f_{\Theta}(\theta) d\theta - I \\ &= p(I) \left(\lambda N g_2 \left(\frac{K}{\lambda N} \right) - K g_1 \left(\frac{K}{\lambda N} \right) \right) - I, \\ g_1(z) &\equiv \int_z^{\infty} f_{\Theta}(\theta) d\theta, \quad g_2(z) \equiv \int_z^{\infty} \theta f_{\Theta}(\theta) d\theta.\end{aligned}$$

In Step 2, the agent's optimisation problem is $\max_{I \geq 0} \mathbb{E}[U_A \mid I, d^*]$. Since the derivative of $\mathbb{E}[U_A \mid I, d^*]$ with respect to I is given by

$$\frac{\partial}{\partial I} \mathbb{E}[U_A \mid I, d^*] = p'(I) \left(\lambda N g_2 \left(\frac{K}{\lambda N} \right) - K g_1 \left(\frac{K}{\lambda N} \right) \right) - 1,$$

and since we have assumed that $p''(I) < 0$, it follows immediately that there are only two cases to consider. If $p'(0)$ is sufficiently small, then it is optimal to choose $I = 0$. Otherwise, the unique optimal solution is given by the first order necessary condition for the maximisation problem. Specifically,

$$I^*(\lambda, N, K) = 0, \quad \text{if } p'(0) \leq \frac{1}{\lambda N g_2 \left(\frac{K}{\lambda N} \right) - K g_1 \left(\frac{K}{\lambda N} \right)},$$

and otherwise, $I^*(\lambda, N, K)$ is defined implicitly by the equation

$$p(I^*) = \frac{1}{\lambda N g_2 \left(\frac{K}{\lambda N} \right) - K g_1 \left(\frac{K}{\lambda N} \right)}.$$

Next, consider the selection of the optimal WTP λ and function $a(x)$ of the principal. Assuming an optimal response by the agent, the expected utility is given by

$$\mathbb{E}[U_P \mid \lambda, a] = p(I^*(\lambda, N, K)) \int_{-\infty}^{\infty} d^*(\theta) (a(\theta) \gamma N \theta - \lambda N \theta) f_{\Theta}(\theta) d\theta.$$

Clearly, $a^*(x) = \mathbb{I}\{x \geq 0\}$, which implies

$$\mathbb{E}[U_P \mid \lambda, a^*] = p(I^*(\lambda, N, K)) (\gamma - \lambda) N g_2 \left(\frac{K}{\lambda N} \right).$$

Since

$$\mathbb{E}[U_P \mid \lambda = 0, a^*] = \mathbb{E}[U_P \mid \lambda = \gamma, a^*] = 0, \quad \text{and } \mathbb{E}[U_P \mid \lambda, a^*] < 0 \text{ for } \lambda > \gamma,$$

the principal's optimisation problem is reduced to

$$\max_{0 \leq \lambda \leq \gamma} \left\{ p(I^*(\lambda, N, K))(\gamma - \lambda)Ng_2 \left(\frac{K}{\lambda N} \right) \right\}.$$

Existence of a solution is guaranteed, since a continuous function is maximised over a finite interval.

In order to illustrate a possible form of the solution to the model, we now consider a very simple numerical example. It is assumed that

$$\Theta \sim \mathcal{N}(0, 1), \quad p(I) = 1 - e^{-I}, \quad K = 1, \quad \gamma = 10.$$

That Θ is standard normal implies that $g_1(z) = 1 - \Phi(z)$ and $g_2(z) = \phi(z)$. The simple form assumed for $p(I)$ implies that an explicit expression is immediately obtained for the agent's optimal investment,

$$I^*(\lambda, N, K) = \begin{cases} 0, & g_2 \left(\frac{K}{\lambda N} \right) / g_1 \left(\frac{K}{\lambda N} \right) \leq \frac{K}{\lambda N}, \\ \ln \left(\lambda N g_2 \left(\frac{K}{\lambda N} \right) - K g_1 \left(\frac{K}{\lambda N} \right) \right), & g_2 \left(\frac{K}{\lambda N} \right) / g_1 \left(\frac{K}{\lambda N} \right) > \frac{K}{\lambda N}. \end{cases}$$

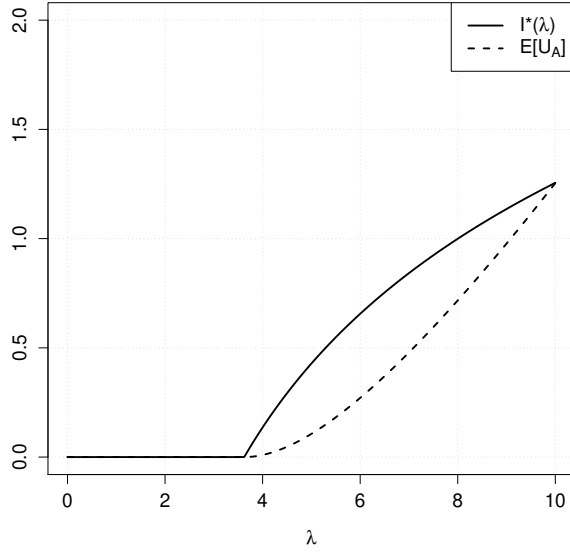
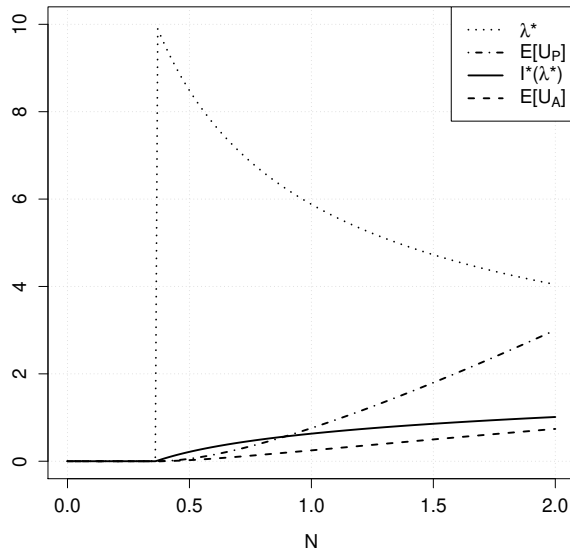
Figure 5.1 shows $I^*(\lambda, N = 1, K = 1)$ and the agent's corresponding expected utility as a function of λ . Figure 5.2 shows the principal's optimal WTP as a function of population size, together with the optimal investment level and the corresponding optimal utilities for the agent and the principal.

5.3 Model with unknown trial cost

In this section, another generalisation of the basic model is investigated. The assumption of a pre-clinical investment is dropped. Now, however, it is assumed that the trial cost $K > 0$ is only known to the agent. Because the principal is uncertain about its value, it uses a distribution F_K , with corresponding density f_K . As before, the value of the true effect Θ is also unknown to the principal. It is assumed that K and Θ are independent from the principal's viewpoint, so that the joint density may be written as $f_K(k)f_\Theta(\theta)$. Further, instead of assuming a payment proportional to the perfect trial estimate X of the effect Θ , we consider arbitrary payment functions $t(x)$. The main interest now lies in determining how the principal's uncertainty regarding the trial cost impacts the optimal form of the payment.

Since $\mathbb{E}[U_A \mid \Theta = \theta, K = k] = d(\theta, k)(t(\theta) - k)$, the agent's optimal GO / NO GO decision is $d^*(\theta, k) = \mathbb{I}\{t(\theta) \geq k\}$. Hence, the principal's expected utility is

$$\mathbb{E}[U_P] = \mathbb{E}[\mathbb{I}\{t(\Theta) \geq K\} (a(\Theta)\gamma N\Theta - t(\Theta))],$$

Figure 5.1: $I^*(\lambda, N = 1, K = 1)$ with corresponding $E[U_A]$.Figure 5.2: $\lambda^*(N)$ and $I^*(\lambda^*, N, K = 1)$, with corresponding $E[U_P]$ and $E[U_A]$.

where the expectation is taken with respect to the joint distribution of Θ and K . The independence of Θ and K implies that $\mathbb{E}[U_P]$ may be rewritten according to

$$\begin{aligned}\mathbb{E}[U_P] &= \mathbb{E}[(a(\Theta)\gamma N\Theta - t(\Theta))\mathbb{E}[\mathbb{I}\{t(\Theta) \geq K\} \mid \Theta]] \\ &= \mathbb{E}[(a(\Theta)\gamma N\Theta - t(\Theta))\mathbb{P}(t(\Theta) \geq K \mid \Theta)] \\ &= \mathbb{E}[(a(\Theta)\gamma N\Theta - t(\Theta))F_K(t(\Theta))].\end{aligned}$$

It is immediately clear that $a^*(x) = \mathbb{I}\{x \geq 0\}$, and the problem for the principal has thus been reduced to

$$\max_{t(\theta) \in \mathcal{C}} \mathbb{E}[(\gamma N\Theta^+ - t(\Theta))F_K(t(\Theta))],$$

where $\mathcal{C}$ is some suitable space of functions.

Since $\mathbb{E}[U_P]$ is a one-dimensional integral over θ with respect to the density f_Θ , the maximisation problem boils down to finding the value $t = t(\theta)$ which maximises

$$(\gamma N\theta^+ - t)F_K(t), \quad \theta \in \mathbb{R} \text{ such that } f_\Theta(\theta) > 0.$$

Clearly, if $\theta \leq 0$, $t(\theta) = 0$ is an optimal choice. For $\theta > 0$, there exists an optimal $t(\theta) \in (0, \gamma N\theta)$. Uniqueness depends on the precise form of the distribution assumed for K . However, with sufficient regularity, each such optimal payment must satisfy the first order condition

$$-F_K(t) + (\gamma N\theta - t)F'_K(t) = 0 \iff \gamma N\theta = t + \frac{F_K(t)}{F'_K(t)}. \quad (5.1)$$

As an example, consider the case when $K \sim \text{Log-Normal}(\mu = 0, \sigma^2 = 1)$. Then

$$F_K(t) = \Phi\left(\frac{\ln t - \mu}{\sigma}\right), \quad F'_K(t) = \frac{1}{\sigma t} \phi\left(\frac{\ln t - \mu}{\sigma}\right).$$

Figure 5.3 shows $t^*(x)$ when $\gamma N = 1$. Note that, by Eq. (5.1), the optimal payment for any combination of values for γ , N and θ can be read from this curve.

5.4 Further work on optimal regulation

There are a number of quite obvious extensions that could be made to the basic skeleton model. Clearly, we would like to drop the unrealistic assumption of a perfect trial and introduce a sample distribution for the trial estimate X that depends on Θ . Further, even if it seems reasonable to assume that the agent

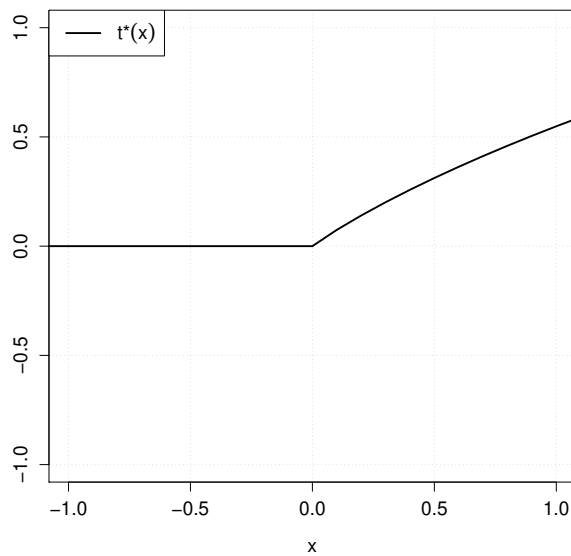


Figure 5.3: $t^*(x)$ for $K \sim \text{Log-Normal}(\mu = 0, \sigma^2 = 1)$ and $\gamma N = 1$.

might be better informed about Θ before the trial begins, this knowledge is hardly perfect. Hence, we would like to include an agent prior for Θ that might be different from the prior f_Θ used by the principal. While the issue of sample size selection permeates the rest of the thesis, it is missing in the simple models above, and the inclusion of such a choice for the agent seems both natural and interesting. Safety was also ignored, and might be included by letting the trial response be multivariate. Perhaps one of the most interesting extensions would be to generalise the trial design optimised by the agent to allow for an adaptive phase 3 trial, or to include earlier phases in a sequential optimisation.

The simple models covered here is just a first step towards an adequate framework for solving optimal regulation problems in the area of clinical trials. They are examples drawn from a number of different model variants that have been investigated to various extent during the last year of my thesis work. However, I decided to include them here since I believe that they are enough to demonstrate that the basic method of solving a game with two actors using backward induction can be a workable approach to the optimal regulation problem. Moreover, even these primitive models lead to some interesting conclusions. The model in Section 5.2 showed that, if a component describing pre-clinical investments is included, then it is suboptimal for the principal to use the same WTP level for all population sizes. That is, even if society values each individual QALY the same, the optimal regulatory scheme is to pay more for rare diseases. The model in Section 5.3 showed that, if the principal is uncertain about the trial cost, then it is not optimal to pay the agent in proportion to the trial estimate X .

Formally, the basic model framework analysed can be classified as a Bayesian game (see, e.g., Fudenberg and Tirole (1991) or Myerson (1997)). More precisely, it belongs to a specialised class of game-theoretic problems known as *principal-agent* problems. A comprehensive account of the principal-agent approach is provided by Laffont and Martimort (2002). This theory has been widely applied in different areas of economics, including the general problem of regulation. However, there seems to be relatively little published literature on specific applications to optimal regulation of clinical trials.

Chapter 6

Discussion

To what extent are Bayesian methods used in practice? Based on a literature review, Lee and Chu (2012) note that the frequentist design paradigm has dominated the field since the first modern trials started in the 1940's. They mention two main barriers associated with the Bayesian approach. The first is the often demanding computations required, and the second is the use of subjective information for the construction of priors. The computational challenge essentially restricted application to trials in which conjugate models could be used. However, the situation has been improved due to the invention of Markov Chain Monte Carlo (MCMC) methods and the present widespread availability of high-speed computers. Very few applications of Bayesian design methods were found before 1995. After this date, the number of trials in which Bayesian methods were used have been steadily increasing (with 74 applications in 2005-2011). Few of these, however, applied the full decision-theoretic approach. While there is a growing trend in the application of Bayesian methods, they are still only used in a small fraction of all trials performed.

Even if the BDT approach is not widely applied in practice, there are numerous earlier contributions in the literature that apply it to various theoretical clinical trial models. A classical, detailed account of the BDT approach is given by Raiffa and Schlaifer (1961). A shorter introduction to the methodology can be found in Lindley (1997), where a sample size selection problem is analysed. Gittins and Pezeshk (2000) discuss an approach that they refer to as "Behavioral Bayes". A pharmaceutical company performs the optimisation, while the actions of the other stakeholders involved, the regulator and the potential users of the treatment, are derived by making direct assumptions regarding their behaviour. In this sense, the model is very similar to the ones analysed in papers II and III. The major difference is that the behaviour of the regulator is not assumed to be based on traditional frequentist rules, but instead what the authors

believe are plausible assumptions about the likely decisions. It is stated that this fully Bayesian methodology was first introduced by Grundy et al. (1956).

Another BDT application is given by Stallard (1998). He considers the problem of phase 2 sample size optimisation, for both fixed size studies and sequential trials. The optimisation takes place from a commercial sponsor's perspective. The trial response is assumed to be binary, so that the outcome for each patient is either success or failure of the treatment tested. A result is given that is similar to the continuation region dynamics of Paper IV; letting s_{n_i} be the number of successes given that n_i responses have been observed up to stage i , there exist boundary functions $c(i)$ and $d(i)$ such that it is optimal to stop for futility if $s_{n_i} < c(i)$, to proceed to phase 3 if $s_{n_i} \geq d(i)$, and to continue the phase 2 trial otherwise. Backward induction is used to solve the problem.

In the papers included in this thesis, the statistical model was based on the normal distribution. This particular choice can be motivated by the central limit theorem and it often greatly simplifies the expected utility computation for a given design. If a normal sample distribution is not appropriate, then some other candidate must be found in the vast space of available statistical models. The choice made will determine how difficult the resulting optimisation problem will be. Sometimes, long-term experience may suggest that some particular distribution is a good choice. In other cases the process of finding a good statistical model may require substantial work. However, in this respect the BDT method is not different from any other approach involving statistics, such as, for example, the classical frequentist method. Therefore, we turn to the two features that makes the BDT framework stand out from other approaches to clinical trial design, the prior and the utility function.

The prime decision problem studied in this thesis is that of designing a phase 3 trial. For such trials, one would like to be able to form a prior by combining the data obtained from previous phases with any information derived from pre-clinical studies. This can be challenging since the form of the phase 3 study may be different from earlier trials. Moreover, even if pre-clinical evidence provides a substantial amount of information about a given drug, there may be no obvious way of converting it into a probability distribution for the parameters of the phase 3 model. Because of this, the choice of prior will often be quite subjective. This subjectivity is often criticised on the grounds that results from a scientific study should be of an objective nature, and should be influenced as little as possible by anything else than the actual trial data.

One solution that has been suggested in the literature is to use so called *objective* or *reference prior distributions*. See, for example, Bernardo and Smith (1994, Section 5.4) for an extensive discussion of the concept. Such distributions are defined so as to have a minimal effect on the final analysis, relative to the data. From the perspective of a RA or HCI in the process of reviewing the

results of a confirmatory trial, the use of such priors seems like a reasonable alternative. However, if we are concerned with clinical trial *design*, rather than analysis, then restricting the choice to the class of reference priors would be paramount to throwing away information that could have been used to optimise the design. My view is that the subjectivity issue is best resolved at the design stage by carefully motivating the form of the prior leading to the optimal design, and, in addition, to study the impact on the solution when perturbing the prior in various ways. Once the design has been fixed and the data collected, any prior can, in principle, be mapped uniquely to a corresponding posterior.

In general, the utility function of a BDT model has a big impact on whether or not a given design is optimal. This point has been made at several places in this thesis, and the issue was explored in some detail in papers II and III where separate utility functions were studied depending on whether or not the trial designer is a commercial sponsor or a public health decision maker. The problem of choosing an appropriate utility function is similar to the problem of subjectivity involved in the choice of prior. Once the choice has been, maximisation of expected value yields an optimal design. However, this design is only relevant to decision makers that accepts the utility specification. Hence, if BDT is used to derive an optimal design for a specific, real-world trial, then a comprehensive recommendation on trial design should include an analysis of the effect of different perturbations of the utility function used.

To solve the decision problems in papers I, II and III, backward induction combined with numerical integration was used. In Paper IV, optimal stopping theory led to an integral equation that could be solved numerically. Common to all these problems is that the form of the numerical computations depended heavily on the model specifics. An alternative route is to set up a simulation framework for estimating the expected utility for a given design, and then search for the optimal design directly. A major benefit of simulating the trial results is that the statistical model can easily be changed, that is, the same numerical method can be used to explore a much larger space of possible statistical models. Whether or not the simulation approach is feasible for solving multi-stage problems will of course depend on the number of stages involved. With too many stages, the variance (due to simulation) of the expected utility estimates might make the following optimisation result in large errors. Bayar et al. (2016) investigated the impact of the α -level (one-sided type I error) and the sample size in a simulation study. A sequence of trials, with the same, fixed design parameters, were simulated over a 15-year period. They were able to conclude that, on average, performing a series of smaller trials with relaxed α -levels (as compared to the traditional 2.5% level) in some scenarios lead to larger survival benefits over a long research horizon. Since traditionally sized trials typically require a large number of patients for moderate effect sizes, this insight is par-

ticularly important in the area of rare diseases. From the perspective of this thesis, it would be interesting to pursue a similar simulation approach while allowing for multiple hypothesis testing and adaptive, multi-stage trials.

None of the papers included in this thesis presents a full analysis of the optimal regulation problem. However, it should be possible to generalise the models in papers I, II, and III by assuming that the RA and/or the HCI can choose some of the parameters of the sponsor's decision problem. We would then arrive at a framework similar to the one presented in Chapter 5. An example along these lines is analysed by Miller and Burman (2018). Their framework includes both phase 2 and phase 3 of the drug approval process. In phase 2, the sponsor selects a sample size n_2 and a threshold, corresponding to a significance level α_2 , where the latter determines a threshold for progressing to phase 3. The regulator selects a significance level α_3 in phase 3 which determines if the new drug receives market authorisation. They solve this model in two steps. First, the optimal pair (n_2^*, α_2^*) is found by maximising the sponsor's expected utility for each admissible regulator choice of α_3 . Then, society's utility is optimised by the regulator, under the assumption that the sponsor responds optimally to each candidate value α_3 . Results are given for a number of different numerical examples. These indicate that sponsors are more restrictive in starting phase 3 if trial costs increases, the regulator requires stronger evidence for approval, or the market potential decreases (for example, due to a smaller target patient population). Importantly, they note that it will not be optimal for society to use a fixed type I error α_3 for all trial situations. In particular, the burden of evidence should typically be relaxed for rare diseases. On a qualitative level, this is the same result as that derived from Anscombe's fully sequential model in Paper IV, even though the decision process analysed by Miller and Burman (2018) is quite different.

An important aspect of the regulation of drug research is the pricing mechanism. If it is too strict, then there will be insufficient investments in basic research by the industry. On the other hand, high prices becomes a burden on the public health system, and may restrict access to consumers in markets dominated by private HCI's. This trade-off is sometimes referred to as that between *static efficiency*, which means ensuring access to the innovation, and *dynamic efficiency*, which means incentivising innovation in the future. Most of the literature has focused on analysing static efficiency rather than its dynamic counterpart. Babar (2015, Chapter 21) provides an introduction to the ideas behind some different pricing policies and reviews some of the recent contributions to the literature on innovation incentives for the pharmaceutical industry. In connection with a discussion on value based pricing, which here means payment to a commercial sponsor proportional to the estimated effect, it is noted that using the same constant of proportionality for all reimbursement decisions

can be problematic. Such a scheme leads to different drugs being granted the same price, regardless of development costs and market size (i.e., size of the target population).

Jena and Philipson (2008) provide an analysis of the impact of cost-effectiveness (CE) rules on static and dynamic efficiency. They argue that CE thresholds are concerned with maximising consumer surplus at the cost of reduced producer surplus. As an illustrative example they mention vaccines, which although they are often very cost-effective lack any appreciable research investment by producers. Their framework is reminiscent of the regulation model solved in Section 5.2. In particular, they include a factor in their social welfare function representing the probability of discovery for a given level of R&D undertaken. Their aim is not to derive the optimal incentivising mechanism from a public health perspective, but instead to underline that reimbursement based on cost-effectiveness is essentially a price control procedure, and hence the well-known fact that price controls affect both dynamic and static efficiency must be taken into account if one wants to find, say, an optimal CE threshold.

Levaggi et al. (2017) compares two schemes for regulating drug reimbursement: CE thresholds and risk-sharing agreements. The optimal behaviour of the firm under these two schemes is derived in a three stage model consisting of (1) discovery of new drug, (2) development and (3) commercialisation. At stage 1 and 3, the effectiveness of the new drug is modeled using a geometric Brownian motion. While there is no explicit attempt to define a societal welfare function and derive the optimal incentivising structure, an extensive discussion is provided regarding to what extent four different policy goals are achieved. There has been some concern in the literature that risk-sharing agreements, which as defined by Levaggi et al. (2017) shifts some of the risk from the HCI to the company, might disincentivise investments into R&D, as compared to the more standard CE threshold scheme. However, their results show that this need not necessarily be the case; it depends on the specifics of the implementation of the risk-sharing agreement, and the nature of the trade-off between paying more for new treatments and reducing the time patients have to wait before approval.

There are many interesting potential extensions for the models analysed in this thesis. For example, one possibility is to include earlier phases in the models of papers I, II, and III. Another natural extension is to consider multivariate outcomes in the trials. Although none of the papers considers safety issues in any detail, safety endpoints are very important in clinical trials and could be included as additional variables. On a conceptual level, both extensions are straightforward, but there will be a limit to the number of stages that can be handled numerically using backward induction.

Chapter 7

Conclusions

If one accepts the BDT approach as a general method for clinical trial design, then a natural direction for further work is to investigate more complex models. The utility function or the statistical model can be made more detailed, and more stages can be added to achieve a greater flexibility. The main difficulty associated with such extensions is the complexity of the numerical computations required when finding the optimal design. Indeed, numerical issues are starting to become troublesome even in the relatively simple two-stage designs of Paper III. One possible way forward may be to abandon the goal of finding the optimal solution via backward induction. For example, it may very well be preferable to use a suboptimal design allowing for five stages instead of an optimal design allowing for two. Although this line of research has not been the topic of this thesis, it is a highly relevant issue. Presumably, the extent to which BDT methods are used by practicing statisticians is determined by the availability of efficient and user friendly software that can be used to obtain near-optimal designs for given priors, statistical models and utility functions.

The main conclusion of Paper I is that the extent of the uncertainty about the HCI's willingness to pay for a new treatment can have a major influence on the decisions made by a commercial sponsor. Although we stop short of optimising the regulatory rules, the basic model is richer than the principal-agent setting discussed in Chapter 5, since it includes three stakeholders: a commercial sponsor, a RA and a HCI. An interesting expansion of this model would be to introduce utility functions for both the RA and the HCI and then try to find the optimal regulatory and reimbursement rules by optimising a social welfare function. This social welfare function would be a linear combination of the utility function for the three stakeholders in the model.

The comparisons of optimised trial designs in papers II and III reveal that it is not always optimal to incorporate biomarker testing, even if there is a

possibility to do so. The form of the optimal design depends on the economics of the situation, such as trial costs and biomarker screening costs, and the prior information available. Moreover, these papers show how the utility functions associated with the two different perspectives, that of a commercial sponsor and a public health decision maker, will affect the optimality of the designs. In general, the investigations in these papers show that, since the form of the optimal trial design depend heavily on both the prior and the utility functions adopted, it is very important to think carefully about appropriate models for these before optimisation begins.

Anscombe's original model was intentionally made as simple as possible. For example, there are no trial or production costs, the randomisation was assumed to be perfectly balanced and the patient horizon N was assumed known. However, we showed that it is possible to make reasonable extensions to this model, such as including patients given the standard treatment in parallel with the trial in the utility function, or assuming that N is unknown, while still being able to obtain solutions. Hence, it may be conjectured that other interesting results can be obtained by applying modern optimal stopping theory in continuous time to sequential decision problems similar to Anscombe's. A particularly interesting result for this thesis is that Anscombe's model led to an argument for using a regulatory rule that is different from the classical one (i.e., frequentist hypothesis testing at a fixed significance level). An example was given that considered the case of a rare disease (small N) and a small trial, and it was shown that the solution to Anscombe's model led to a less conservative level for approving a new treatment.

The model in Section 5.2 shows that, if pre-clinical investments are included as a model component, then it follows that it is not optimal to pay the same for all QALYs, even if they have the same value to the HCl. To some extent, this model captures the trade-off between static and dynamic efficiency. It is a simple model, to be sure, but the structure that is present seems reasonable enough. This indicates that the result may be quite general, and that extended, more realistic models would lead to a similar conclusion.

Appendix A

Multiple testing procedures

Consider a statistical model in which the probability distribution of a random variable X is determined by the value of a parameter θ that belongs to some parameter space $\mathcal{T}$. In this setting, a subset H of $\mathcal{T}$ is referred to as a *null hypothesis* about the value of θ , and the complement H^c is called the *alternative hypothesis*. A *test* associated with a null hypothesis is then defined using a binary function $t(X)$ mapping X into a value that determines whether or not the null hypothesis is rejected.

Given a hypothesis $H \subseteq \mathcal{T}$ and an associated test t , the *type I error* is defined as the (maximum) probability of falsely rejecting H . Formally,

$$\text{type I error} = \sup_{\theta \in H} \mathbb{P}_{\theta}(t(X) \text{ rejects } H).$$

A statistical test is said to control the type I error at the level α if the error is less than or equal to α . We now generalise this notion of type I error control to the setting of multiple hypothesis testing.

Let $\{H_1, \dots, H_m\}$ be a family of $m \geq 1$ hypotheses. For any value of the parameter θ , let $I(\theta) = \{i \in \{1, \dots, m\} \mid \theta \in H_i\}$ be the index set corresponding to all hypotheses containing θ . Any fixed value of θ defines a subset of true hypotheses, $M_{I(\theta)} = \{H_i \mid i \in I(\theta)\}$, containing $m_{I(\theta)} = |M_{I(\theta)}|$ elements. In this setting of multiple hypotheses, a test function t maps an outcome X of the experiment into a subset of $\{H_1, \dots, H_m\}$, the elements of which correspond to the individual hypotheses that are rejected by the test.

Having a fixed test procedure t in mind, let R (an observable random variable) be the total number of hypotheses which are rejected by the test and let V (an unobservable random variable) be the number of true hypotheses which are rejected. There are several alternative definitions that may be used when generalising the type I error rate to the testing of several hypotheses (Bretz

et al., 2011, Chapter 2). The *Per Comparison-Error Rate* (PCER) is defined as the expected number of true hypotheses rejected per comparison,

$$\text{PCER} = \mathbb{E}_\theta [V] / m.$$

The *False Discovery Rate* (FDR) is defined as the expected proportion of falsely rejected hypotheses among the rejected hypotheses (V/R is defined as 0 for $R = 0$),

$$\text{FDR} = \mathbb{E}_\theta [V/R \mid R > 0] \mathbb{P}_\theta (R > 0).$$

The *Family-Wise Error Rate* (FWER) is defined as the probability that at least one hypothesis is falsely rejected,

$$\text{FWER} = \mathbb{P}_\theta (V > 0).$$

In the context of confirmatory clinical trials, a common contemporary regulatory requirement is that the FWER is controlled at a given significance level α (Bretz et al., 2011, p. 13). This is also the generalisation of type I error to multiple hypotheses used in this thesis.

A *Multiple Testing Procedure* (MTP) is said to control the FWER in the *weak* sense if the rate is controlled only under the *global null hypothesis*, which is defined as the intersection of all null hypotheses in the family. Hence, there is weak FWER control if

$$\sup_{\theta \in \cap_{i=1}^m H_i} \mathbb{P}_\theta (V > 0) \leq \alpha.$$

Control of the FWER is said to be *strong* if the type I error is controlled for any configuration of true and false hypotheses, which may be expressed as

$$\sup_{\theta \in \cup_{i=1}^m H_i} \mathbb{P}_\theta (V > 0) \leq \alpha.$$

In many situations of multiple testing, there are natural tests and corresponding *unadjusted p-values* associated with the individual null hypotheses H_i , $i = 1, \dots, m$, or simple combinations of them. These simple tests are often used as a basis for the construction of a MTP. An *adjusted p-value* can then be associated with each hypothesis, defined in such a way that a direct comparison with a prescribed FWER is possible. For example, the closed testing principle (Marcus et al., 1976) can be used to construct a MTP given that tests for rejection have been defined for all possible intersection hypotheses $H_I = \cap_{i \in I} H_i$, $I \subseteq \{1, \dots, m\}$. An elementary hypothesis H_i is then rejected by the overall test (i.e., by the MTP) if H_I can be rejected for all I containing i . It can be shown that, if the tests for the intersections control the error rate at level α , then so will the overall MTP.

One of the simplest and most well-known MTP is the Bonferroni procedure, which may serve to illustrate the general philosophy behind such methods. Suppose that individual tests $t_1, \dots, t_m$ have been defined for the hypotheses, with corresponding unadjusted p-values $p_1, \dots, p_m$. The Bonferroni procedure is then implemented by rejecting H_i if $p_i \leq \alpha/m$, $1 \leq i \leq m$. That the FWER is controlled in the strong sense at the level α follows directly by the Bonferroni inequality, since $\mathbb{P}_\theta(V > 0)$ may be written as

$$\mathbb{P}_\theta(\cup_{i \in I(\theta)} \{p_i \leq \alpha/m\}) \leq \sum_{i \in I(\theta)} \mathbb{P}_\theta(p_i \leq \alpha/m) \leq m_{I(\theta)} \left(\frac{\alpha}{m}\right) \leq \alpha. \quad (\text{A.1})$$

A.1 The Spiessens-Debois test

Since no distributional assumptions are used in the proof of Eq. (A.1), the Bonferroni MTP is completely general and may be applied in any situation. However, this generality comes at a price. More specific procedures that are only applicable if certain distributional assumptions are placed on the test statistics often lead to higher power. One such procedure, which is employed in Paper II, have been proposed by Spiessens and Debois (2010). They refer to it as the *general bivariate normal method*.

We now consider the construction of a MTP controlling the FWER in the strong sense by combining the closed testing principle with the method of Spiessens and Debois. The context is that of a clinical trial with two analyses, namely, an overall analysis of the efficacy in the total population and a subgroup analysis which only considers a subset of the population. The two test statistics corresponding to the overall and subgroup analyses are assumed to follow a bivariate normal distribution. Let Z_1 and Z_2 denote the standardised test statistics used to test the null hypotheses H_1 and H_2 of a zero treatment effect in the total population and the subgroup, respectively. It may be shown (Jennison and Turnbull, 2000) that under $H_{12} = H_1 \cap H_2$,

$$\begin{bmatrix} Z_1 \\ Z_2 \end{bmatrix} \sim \mathcal{N} \left(\begin{bmatrix} 0 \\ 0 \end{bmatrix}, \begin{bmatrix} 1 & \sqrt{\lambda} \\ \sqrt{\lambda} & 1 \end{bmatrix} \right), \quad (\text{A.2})$$

where $\lambda = I_S/I_T$ is the fraction of information¹ in the subgroup relative to the total population.

If H_{12} is true, then it is unlikely that extreme values will be observed for Z_1 or Z_2 . Hence, a reasonable test is one which rejects H_{12} if either $Z_1 > z_{\alpha_1}$ or $Z_2 > z_{\alpha_2}$, where z_{α_1} and z_{α_2} are the critical values corresponding to some significance levels α_1 and α_2 . Suppose that the level α_1 is fixed at some specific

¹Information is here defined as the reciprocal of the variance of the respective parameter.

value. Then in order to control the FWER at an overall significance level α , the other significance level α_2 must be chosen in such a way that

$$\mathbb{P}_\theta (Z_1 > z_{\alpha_1} \text{ or } Z_2 > z_{\alpha_2}) = \alpha, \quad \text{for } \theta \in H_{12}.$$

Since

$$\begin{aligned} \mathbb{P}_\theta (Z_1 > z_{\alpha_1} \text{ or } Z_2 > z_{\alpha_2}) &= \mathbb{P}_\theta (Z_1 > z_{\alpha_1}) + \mathbb{P}_\theta (Z_1 \leq z_{\alpha_1} \text{ and } Z_2 > z_{\alpha_2}) \\ &= \alpha_1 + \mathbb{P}_\theta (Z_1 \leq z_{\alpha_1} \text{ and } Z_2 > z_{\alpha_2}), \end{aligned}$$

this is equivalent to finding a value of α_2 satisfying

$$\mathbb{P}_\theta (Z_1 \leq z_{\alpha_1} \text{ and } Z_2 > z_{\alpha_2}) = \alpha - \alpha_1. \quad (\text{A.3})$$

Using the density in Eq. (A.2), it can be seen that solving Eq. (A.3) is equivalent to solving

$$\int_{-\infty}^{z_{\alpha_1}} \Phi \left(\frac{z_{\alpha_2} - \sqrt{\lambda} z_1}{\sqrt{1 - \lambda}} \right) \phi(z_1) dx = 1 - \alpha. \quad (\text{A.4})$$

Given a selected value of α_1 , the solution of Eq. (A.4) provides a value of α_2 such that the test which rejects the intersection hypothesis H_{12} if $Z_1 > z_{\alpha_1}$ or $Z_2 > z_{\alpha_2}$ controls the type I error at level α . Clearly, the canonical univariate tests which reject H_1 if $Z_1 > z_\alpha$ and H_2 if $Z_2 > z_\alpha$ also control the type I error rate for the rejection of H_1 and H_2 , respectively. Hence, by the closed testing principle, the following MTP used in Paper II controls the FWER in the strong sense:

1. Specify values of α and α_1 .
2. Reject H_1 if $Z_1 > z_\alpha$ (unadjusted rejection of H_1) and $Z_1 > z_{\alpha_1}$ or $Z_2 > z_{\alpha_2}$ (unadjusted rejection of H_{12}).
3. Reject H_2 if $Z_2 > z_\alpha$ (unadjusted rejection of H_2) and $Z_1 > z_{\alpha_1}$ or $Z_2 > z_{\alpha_2}$ (unadjusted rejection of H_{12}).

Appendix B

Some basic tools from analysis

B.1 The implicit function theorem

In Paper I we are faced with the problem of maximising a certain objective function (an expected utility) with respect to a price variable. In addition to the price variable, the objective function also depends on a set of parameters and it is of interest to analyse how the optimal price depends on these parameters. The implicit function theorem (see, for example, Rudin (1976, Chapter 9)) is used in Paper I to establish that the optimal price function is continuously differentiable, given that this holds for the objective function that is maximised. Moreover, the theorem provides formulas for computing the partial derivatives of the optimal price function with respect to the parameters. For reference, a version of the theorem will now be stated.

Let f be a continuously differentiable function mapping points $(\mathbf{x}, y) \in \mathbb{R}^{n+1}$ into $\mathbb{R}$ and suppose that the point $(\mathbf{a}, b) = (a_1, \dots, a_n, b) \in \mathbb{R}^{n+1}$ satisfies $f(\mathbf{a}, b) = 0$. Then, if $(\partial f / \partial y)(\mathbf{a}, b) \neq 0$, there exists an open set U containing $\mathbf{a}$, an open set V containing b , and a unique continuously differentiable function $g : U \rightarrow V$ such that the local graph of g , $\{(\mathbf{x}, g(\mathbf{x})) \mid \mathbf{x} \in U\}$, coincides with the local level set $\{(\mathbf{x}, y) \in U \times V \mid f(\mathbf{x}, y) = 0\}$. Moreover, the partial derivative of g with respect to the component x_i in the point $\mathbf{a}$ is given by

$$\frac{\partial g}{\partial x_i}(\mathbf{a}) = -\frac{\partial f}{\partial x_i}(\mathbf{a}, b) \bigg/ \frac{\partial f}{\partial y}(\mathbf{a}, b).$$

B.2 The envelope theorem

A general result that is often useful when analysing how changes to parameters influence the optimised value of an objective function is the envelope theorem (Varian, 1992, Chapter 27). There are many versions of this theorem, but the result will only be stated here under rather strong smoothness assumptions. The theorem is used in Paper I to analyse how the optimal expected utility responds to changes in various parameters.

Suppose that f is a differentiable function mapping points $(\mathbf{x}, y) \in \mathbb{R}^{n+1}$ into $\mathbb{R}$, and consider the maximisation problem

$$f^*(\mathbf{x}) \equiv \max_y f(\mathbf{x}, y).$$

Assume further that the maximising argument $y^*(\mathbf{x})$ is differentiable for $\mathbf{x}$ in some region U of interest and that $(\mathbf{x}, y^*(\mathbf{x}))$ corresponds to an interior global optimum for $\mathbf{x} \in U$. The envelope theorem then states that

$$\frac{\partial}{\partial x_i} (f^*(\mathbf{x})) = \frac{\partial f}{\partial x_i}(\mathbf{x}, y) \Big|_{y=y^*(\mathbf{x})}. \quad (\text{B.1})$$

Hence, changes in the optimal value of the objective function may be analysed in terms of partial derivatives of the original objective function. With the stated assumptions, this result is easily shown. Since $f^*(\mathbf{x}) = f(\mathbf{x}, y^*(\mathbf{x}))$, the left hand side of Eq. (B.1) is

$$\frac{\partial}{\partial x_i} (f(\mathbf{x}, y^*(\mathbf{x}))) = \frac{\partial f}{\partial x_i}(\mathbf{x}, y^*(\mathbf{x})) + \frac{\partial f}{\partial y}(\mathbf{x}, y^*(\mathbf{x})) \frac{\partial y^*}{\partial x_i}(\mathbf{x}).$$

But from the assumptions of differentiability and an interior optimum, $(\partial y^* / \partial x_i)(\mathbf{x}) = 0$ and the result follows.

Bibliography

- Babar, Z.-U.-D. (Ed.) (2015). *Pharmaceutical Prices in the 21st Century* (First ed.). Switzerland: Springer International Publishing.
- Baron, J. H. (2009, June). Sailors' scurvy before and after James Lind — a reassessment. *Nutrition Reviews* 67(6), 315–332.
- Bayar, M. A., G. L. Teuff, S. Michiels, D. J. Sargent, and M.-C. L. Deley (2016, March). New insights into the evaluation of randomized controlled trials for rare diseases over a long-term research horizon: a simulation study. *Statistics in Medicine* 35(19), 3245–3258.
- Bayarri, M. J. and J. O. Berger (2004, Feb). The Interplay of Bayesian and Frequentist Analysis. *Statistical Science* 19(1), 58–80.
- Bernardo, J. M. and A. F. Smith (1994). *Bayesian Theory* (First ed.). 111 River Street, Hoboken, NJ 07030, USA: John Wiley & Sons.
- Berry, S. M., B. P. Carlin, J. J. Lee, and P. Müller (2011). *Bayesian Adaptive Methods for Clinical Trials* (First ed.). 6000 Broken Sound Parkway NW, Suite 300, Boca Raton, FL 33487-2742: CRC Press, Taylor & Francis Group.
- Biomarkers Definitions Working Group (2001, March). Biomarkers and surrogate endpoints: Preferred definitions and conceptual framework. *Clinical Pharmacology & Therapeutics* 69(3), 89–95.
- Bretz, F., T. Hothorn, and P. Westfall (2011). *Multiple Comparisons Using R* (First ed.). 6000 Broken Sound Parkway NW, Suite 300, Boca Raton, FL 33487-2742: Taylor & Francis Group.
- Chow, S.-C. and J.-P. Liu (2014). *Design and Analysis of Clinical Trials* (Third ed.). Hoboken, New Jersey: John Wiley & Sons.
- Christensen, S. (2016). *Stochastic Control with Financial Applications*. Unpublished lecture notes. Version: April 26, 2016.

- EMA (2016, April). History of EMA. Available at http://www.ema.europa.eu/ema/index.jsp?curl=pages/about_us/general/general_content_000628.jsp&mid=WC0b01ac058087add.
- Fisher, R. A. (1946). *Statistical Methods for Research Workers* (Tenth ed.). Edinburgh: Oliver and Boyd.
- Fisher, R. A. (1956). *Statistical Methods and Scientific Inference*. Edinburgh: Oliver and Boyd.
- Fudenberg, D. and J. Tirole (1991). *Game Theory* (First ed.). Cambridge, Massachusetts: MIT Press.
- Gittins, J. and H. Pezeshk (2000). A Behavioral Bayes Method for Determining the Size of a Clinical Trial. *Drug Information Journal* 34(2), 355–363.
- Grundy, P. M., M. J. R. Healy, and D. H. Rees (1956). Economic Choice of the Amount of Experimentation. *Journal of the Royal Statistical Society. Series B (Methodological)*. 18(1), 32–55.
- Hart, P. D. (1999). A change in scientific approach: from alternation to randomised allocation in clinical trials in the 1940s. *British Medical Journal* 319(7209), 572–573.
- Haynes, B. (1999, Sep). Can it work? Does it work? Is it worth it? The testing of healthcare interventions is evolving. *BMJ* 319(7211), 652–653.
- ICH (2016, April). Statistical principles for clinical trials (E9). Available at http://www.ich.org/fileadmin/Public_Web_Site/ICH_Products/Guidelines/Efficacy/E9/Step4/E9_Guideline.pdf.
- Jena, A. B. and T. J. Philipson (2008, September). Cost-effectiveness analysis and innovation. *Journal of Health Economics* 27(5), 1224–1236.
- Jennison, C. and B. W. Turnbull (2000). *Group Sequential Methods with Applications to Clinical Trials* (First ed.). 2000 N.W. Corporate Blvd., Boca Raton, Florida 33431: Chapman & Hall/CRC.
- Kaitin, K. I. and J. A. DiMasi (2011). Pharmaceutical Innovation in the 21st Century: New Drug Approvals in the First Decade, 2000-2009. *Clinical Pharmacology & Therapeutics* 89(2), 183–188.
- Laffont, J.-J. and D. Martimort (2002). *The Theory of Incentives* (First ed.). 41 William Street, Princeton, New Jersey: Princeton University Press.
- Lee, J. J. and C. T. Chu (2012). Bayesian clinical trials in action. *Statistics in Medicine* 31(25), 2955–2972.

- Levaggi, R., M. Moretto, and P. Pertile (2017). The Dynamics of Pharmaceutical Regulation and R&D Investments. *Journal of Public Economic Theory* 19(1), 121–141.
- Lindley, D. V. (1997, July). The choice of sample size. *The Statistician* 46(2), 129–138.
- Marcus, R., E. Peritz, and K. R. Gabriel (1976). On closed testing procedures with special reference to ordered analysis of variance. *Biometrika* 63(3), 655–660.
- Miller, F. and C.-F. Burman (2018). A decision theoretical modeling for Phase III investments and drug licensing. *Journal of Biopharmaceutical Statistics* 28(4), 698–721.
- Myerson, R. B. (1997). *Game Theory: Analysis of Conflict* (First ed.). United States of America: Harvard University Press.
- Øksendal, B. (2003). *Stochastic Differential Equations* (Sixth ed.). Berlin: Springer-Verlag.
- Ondra, T., A. Dmitrienko, T. Friede, A. Graf, F. Miller, N. Stallard, and M. Posch (2016). Methods for identification and confirmation of targeted subgroups in clinical trials: A systematic review. *Journal of Biopharmaceutical Statistics* 26(1), 99–119.
- Raiffa, H. and R. Schlaifer (1961). *Applied Statistical Decision Theory* (First ed.). Boston: President and Fellows of Harvard College.
- Rudin, W. (1976). *Principles of Mathematical Analysis* (Third ed.). New York: McGraw-Hill.
- Senn, S. (2007). *Statistical Issues in Drug Development* (Second ed.). The Atrium, Southern Gate, Chichester, West Sussex PO19 8SQ, England: John Wiley & Sons.
- Sertkaya, A., H.-H. Wong, A. Jessup, and T. Beleche (2016). Key cost drivers of pharmaceutical clinical trials in the United States. *Clinical Trials* 13(2), 117–126.
- Spiegelhalter, D. J., K. R. Abrams, and J. P. Myles (2004). *Bayesian Approaches to Clinical Trials and Health-Care Evaluation* (First ed.). The Atrium, Southern Gate, Chichester, West Sussex PO19 8SQ, England: John Wiley & Sons.
- Spiessens, B. and M. Debois (2010, November). Adjusted significance levels for subgroup analyses in clinical trials. *Contemporary Clinical Trials* 31(6), 647–656.

- Stallard, N. (1998, March). Sample Size Determination for Phase II Clinical Trials Based on Bayesian Decision Theory. *Biometrics* 54(1), 279–294.
- Stallard, N., F. Miller, S. Day, S. W. Hee, J. Madan, S. Zohar, and M. Posch (2017). Determination of the optimal sample size for a clinical trial accounting for the population size. *Biometrical Journal* 59(4), 609–625.
- Varian, H. R. (1992). *Microeconomic Analysis* (Third ed.). 500 Fifth Avenue, New York, N.Y. 10110: W. W. Norton & Company.