

ALL YOU NEED FOR WRITING YOUR MASTER THESIS IN THE ANCIENT LANGUAGES & CULTURES

This book contains material normally found in at least three separate volumes: research methodology, philosophy of science, and (simple) statistics. They are all compressed into one relatively small volume. All the basics you should know about carrying out scientific research and writing your Master Thesis about it, is presented in this book. The book is especially written for students of the ancient Languages and Cultures. The object is to make them acquainted with basic theory and methods used in social sciences. This should enable them to get more grasp on research that they will be carrying out for their Master Thesis. The author is acquainted with both research practice and the study of ancient Languages and Cultures, as he studied Egyptology after having experience as a scientific researcher in sociology and criminology. He found out that fellow students were not able to get the most out of their Master Thesis research simply because they were not sufficiently thought on the subject. Some students were displaying a slight fear of using mathematics too. That is the reason that this subject is *very delicately handled* in this publication and examples from their own study are used. The student is guided through the material in a relatively easy way. Although one chapter is devoted to the use of statistical software, this is done only for those who want to master this relatively complicated matter as well. In practice, it is not really necessary because a standard program like Excel allows use of the most basic statistical functions as well.

ABOUT THE AUTHOR



J. C. Colder was born in the Netherlands and first studied Personal Management in Rotterdam where he acquired his BA. He then studied sociology at the Erasmus University in Rotterdam and specialised in deviant behaviour and sociology of organisations. At this university, he also followed an intensive training program in research methodology and statistics. After acquiring his MSc (Dutch: Drs.) title he started his career as a Scientific Researcher for the Scientific Bureau of the Ministry of Justice (WODC) in the Netherlands. There he acquired practical knowledge of scientific research that enabled him to write this book. After the publication of his research in preventing Crime in Shopping centers he worked for a large Software House as a Business Consultant. Later in life he studied Egyptology at the University of Leiden and graduated with a MA in Egyptology. Currently he is working on a PhD-thesis “Explaining Crime in Ancient Egypt”.

**Writing a scientifically sound
Master Thesis
for the Ancient Languages and Civilizations**

J.C. Colder MSc MA
January 2019

J.C. Colder MSc MA

Writing a scientifically sound Master Thesis for the Ancient Languages and Civilisations

ALL RIGHTS RESERVED. This book contains material protected under International and Federal Copyright Laws and Treaties. Any unauthorized reprint or use of this material is prohibited. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording, or by any information storage and retrieval system without express written permission from the author / publisher.

Acknowledgements

This book could not have been written without the help of several people that I would like to thank in the following lines. I was very pleased that my old tutor in Methodology and Statistics Dr. Jan Hakvoort agreed to review my concepts. At least 30 years of experience, he was quick to reply and very thorough. Thanks very much, Jan! Dr. Rene van Walsem, my old teacher of Material culture, Egyptian Art and -Archaeology, was my second commentator. Rene suggested many subjects and examples and corrected some of my errors. My third commentator was Dr. Lex Cachet, my old teacher General Theoretic Sociology and my current co-promoter. Lex supplied several useful tips and stimulated me to keep on going.

Three experts in the field worked with me and provided the information that I needed in order to write the addenda on Dr. Jan Gerrit Dercksen at Leiden University, who is a specialist of Akkadian, supplied me with the material for the first Addendum on “Writing correct translations”. Jan Gerrit was a great help and we got together very quickly. This was probably so because translation of the old Akkadian and Egyptian language had much in common. Prof. Dr. Maarten Raven, former curator of the Ancient Egyptian collection in the National Museum of Antiquities (Rijksmuseum van Oudheden) at Leiden, was my informant of the second addendum about “Writing Catalogs of Objects.” Maarten supplied me with lots of information and performed some critical reviews on the text himself, which was greatly appreciated. In the period that we worked together, Maarten became a professor of Museology at Leiden University. Under the supervision of Prof. Dr. Jacques van der Vliet, who is an authority of Egyptology, Coptic and early Christianity at Leiden and Nijmegen University in the Netherlands, I wrote the third Addendum “Looking into the life of historical Persons”. We discussed the material together repeatedly until we reached the final text that is published in this book. Jacques was very helpful. All contributors to the addenda were very critical and made me work very hard. In addition, to avail because the result is of a quality we can all adhere too.

Having remarked upon quality, this takes me at the English Language I used in this book. Every contributor motioned above, suggested that I should have my work tested by native English speakers. They were right! My first correctors Koos & Carolien Hoogenboom from New Zealand adapted my use of the English language to a much more readable text. Koos even caught a hitherto unnoticed mistake in my explanation of the interquartile range. Thanks very much Koos! Last but not least, my UK neighbor Phil Guy reviewed the use of my English language over and over again. As a mathematician he knew much about the topic's I have discussed in this book and he had a profound influence on them. I am very much indebted to him.

January 2019

J.C. Colder

j.colder@hccnet.nl

CONTENTS

Chapter	Page
1. DOING RESEARCH FOR A MASTER THESIS IN ANCIENT SOCIETIES.....	8
1.1. INTRODUCTION.....	8
1.2. WHAT IS (SCIENTIFIC) RESEARCH?.....	8
1.3. WHAT IS SPECIAL IN RESEARCH OF THE ANCIENTS?.....	10
1.3.1. <i>The cultural background of the researcher</i>	10
1.3.2. <i>About our sources</i>	11
1.4. ON BECOMING A SCIENTIFIC RESEARCHER.....	12
1.5. ABOUT THE BOOK.....	14
2. REASONING IN SCIENCE.....	15
2.1. INTRODUCTION.....	15
2.2. SCIENTIFIC REASONING.....	15
2.2.1. <i>Introduction</i>	15
2.2.2. <i>Logic</i>	16
2.2.3. <i>Causality</i>	18
2.2.4. <i>Induction and deduction</i>	19
2.2.5. <i>Confirmation or falsification</i>	20
2.3. PROGRESS IN SCIENCE.....	22
2.3.1. <i>What constitutes progress?</i>	22
2.3.2. <i>History of science in the 20th Century</i>	24
2.3.3. <i>Logical positivism</i>	25
2.3.4. <i>Logical empiricism</i>	26
2.3.5. <i>Falsification</i>	28
2.3.6. <i>The scientific revolution</i>	29
2.3.7. <i>Ockham's razor</i>	30
2.3.8. <i>Criticism of science</i>	30
2.4. TAKING UP A POSITION.....	31
3. DOING RESEARCH.....	32
3.1. TYPES OF RESEARCH FOR ANCIENT AND DEAD CIVILIZATIONS.....	32
3.1.1. <i>Introduction</i>	32
3.1.2. <i>General research types</i>	33
3.1.3. <i>Specific research types</i>	35
3.2. METHODS OF DOING RESEARCH.....	36
3.2.1. <i>General research methods</i>	36
3.2.2. <i>Methods of qualitative research</i>	37
3.2.3. <i>Grounded theory</i>	39
3.2.4. <i>Methods of quantitative research</i>	40
3.3. DESIGNING YOUR RESEARCH.....	40
3.3.1. <i>Design of an explorative research</i>	41
3.3.2. <i>Design for testing a theory or theoretical elements</i>	44
4. GATHERING DATA.....	48
4.1. THE POWER OF THE NUMBER.....	48
4.2. THE RESEARCH CODEBOOK.....	49
4.2.1. <i>Purpose</i>	49

4.2.2.	<i>Variables, values and labels.....</i>	<i>49</i>
4.2.3.	<i>Types of variables.....</i>	<i>52</i>
4.2.4.	<i>Rules and conventions.....</i>	<i>55</i>
4.2.5.	<i>Layout of the codebook.....</i>	<i>56</i>
4.2.6.	<i>Example of a simple codebook for ancient statues.....</i>	<i>56</i>
4.3.	GATHERING RESEARCH DATA.....	61
4.3.1.	<i>Introduction.....</i>	<i>61</i>
4.3.2.	<i>Learning the truth.....</i>	<i>61</i>
4.3.3.	<i>Recording and coding your findings.....</i>	<i>62</i>
4.3.4.	<i>Example data-gathering form.....</i>	<i>63</i>
4.4.	THE DATA-MATRIX.....	64
4.4.1.	<i>Building the data-matrix.....</i>	<i>64</i>
4.4.2.	<i>Exemplary data-matrices.....</i>	<i>65</i>
4.5.	PREPARATIONS FOR ANALYSIS.....	68
4.5.1.	<i>Introduction.....</i>	<i>68</i>
4.5.2.	<i>Choosing your means of analysis.....</i>	<i>68</i>
4.5.3.	<i>Preparing data.....</i>	<i>68</i>
5.	BASIC ANALYSIS.....	69
5.1.	SIMPLE STATISTICS.....	69
5.2.	DATA CLEANING.....	69
5.3.	FREQUENCIES AND CROSS TABLES.....	71
5.3.1.	<i>Frequencies.....</i>	<i>71</i>
5.3.2.	<i>Cross tables.....</i>	<i>74</i>
5.4.	PRESENTATION OF DATA.....	76
6.	ADVANCED ANALYSIS.....	77
6.1.	CHANCE AND PROBABILITY.....	77
6.1.1.	<i>Inductive statistics.....</i>	<i>77</i>
6.1.2.	<i>Population and samples.....</i>	<i>77</i>
6.1.3.	<i>Chance.....</i>	<i>78</i>
6.1.4.	<i>Probability.....</i>	<i>79</i>
6.1.5.	<i>Frequencies.....</i>	<i>80</i>
6.2.	MEAN.....	81
6.3.	MEASURING SIMPLE ASSOCIATION.....	87
6.4.	DISTRIBUTION OF VALUES.....	91
6.4.1.	<i>Distribution of research data.....</i>	<i>91</i>
6.4.2.	<i>Special distributions: the normal distribution.....</i>	<i>96</i>
6.5.	SAMPLING.....	99
6.5.1.	<i>Evaluating the population.....</i>	<i>99</i>
6.5.2.	<i>Sampling criteria.....</i>	<i>101</i>
6.5.3.	<i>The taking of a sample.....</i>	<i>102</i>
6.6.	USING STATISTICAL SOFTWARE: SPSS.....	105
6.6.1.	<i>Introduction.....</i>	<i>105</i>
6.6.2.	<i>Entering a codebook.....</i>	<i>105</i>
6.6.3.	<i>Entering research data.....</i>	<i>110</i>

ADDENDA.....	114
ADDENDUM 1. WRITING CORRECT TRANSLATIONS.....	114
ADDENDUM 2. WRITING CATALOGUES OF OBJECTS.....	118
ADDENDUM 3. LOOKING INTO THE LIFE OF HISTORICAL PERSONS.....	124
ADDENDUM 4. RESEARCH PRACTICE, REPORTING AND FOLLOW-UP.....	132
CONSULTED LITERATURE & DIGITAL SOURCES.....	136
CONSULTED LITERATURE.....	136
DIGITAL SOURCES.....	137

1. Doing research for a Master Thesis in ancient societies

1.1. *Introduction*

Up until today, research in the ancient cultures and languages has been of a basically descriptive nature. Because we still do not know everything there is to know about these societies scientists are putting a lot of effort into finding out how things really were back then. In other words, we describe (parts of) the society in words and perhaps illustrate that by a certain amount of figures. Usually we put in a lot of nice pictures too, in order to make the things we found out a little more understandable. We have been doing this for a long time, but I think now the time has come to improve upon our working methods.

Before I studied Egyptology myself, I studied sociology and that was a whole different world. I grew up with utterly heavy methodology books and quite an intensive training in statistics, both descriptive and inductive. Although I finished my exams and thought I knew all there was to know about research, I felt a little betrayed when I got my first job. I became a researcher for a ministry that had its own scientific department. There I found out the hard way about what doing research really meant. Putting into practice what I had learned was not easy. Writing a codebook for research work seemed simple enough, but constructing it in such a way that it could overcome the problems that awaited me later, was another matter. In the end I got it right and used it later on again when I studied Egyptology.

You are probably thinking: “Ok. But what has that got to do with me?” Well for one thing, if you are studying an ancient culture and continue in the above mentioned way, you will cut yourself short of possibilities. And that is what this book is all about.

When I graduated, I noticed that my fellow students would have been able to do a great deal more with the material they had collection if only they had been taught a little about methodology and simple statistics. In doing so, the quality of their research could have improved considerably. I started helping one student, and made some simple statistical calculations with a regular computer and statistical software. This student was now able to show *trends* occurring over time. This is something that would have remained hidden, had I not made those calculations. What I intend to do in this book is teach you the basics of how I did that and what you need to know in order to perform your own simple statistical analyses. Don't be alarmed by the word 'statistical'. We will be doing things the easy way and there will be no complicated formulas. I know how students of ancient cultures react to these matters. Because I myself learned how to do practical research, I decided to write a book especially for students of ancient cultures and languages. By following these simple techniques, you will be able to get more out of your master thesis and become a more skilled researcher.

1.2. *What is (scientific) research?*

According to the Oxford dictionary of English, research is “... the systematic investigation into and study of materials and sources in order to establish facts and reach new conclusions.” This can be summarised as: “Doing research, is finding something out.” For me researching is just that: finding out what I want to know. Your doctor can examine your body and question you about a possible complaint that you have put before him or her. A fireman or police officer can research a burned down building in order to determine if the fire was intentional or not. The local mechanic of a garage can examine your car in order to determine why the en-

gine is not starting. There are a myriad of examples of people performing research in everyday life situations. I would not go so far as to say that the above mentioned research is not *scientific*. In most cases the *methods* that are used in the examples above, have their *origins* in science and certainly *logic* is applied. Actually there is a distinction to be made between scientific and 'everyday' research, but the line that separates the two is very fine indeed.

The word science stems from a Latin term *scientificus* meaning knowledge. Today 'knowledge' by itself is not enough to indicate 'science'. There is more associated with this term. For instance, the way in which this knowledge is obtained is quite important. For knowledge to be qualified as 'scientific' means that it has to be *systematic* and *replicable* research along internationally *approved* methods. The results should be, as far as possible, undisputed and (relatively) modern. Perhaps I should also add that the research should be *unbiased*, free from opinions and cultural factors, *systematic* and (has the potential to be) *tested* (by others). Notice that the real difference with the above examples of research I gave to you is only the way in which the results were *obtained*. How a doctor examines you, is learned in medical sciences at the university and derives from research practices on the human body. We call this *applied* science and these methods have proven themselves throughout time. The university will provide additional courses to medical staff constantly in which they teach improvements in methods to doctors and teach them additional data, like exceptions or contra-indications.

As you can see, practical research is often (deeply) rooted in modern scientific research. Even the 'research' of your garage mechanic is quite logical and thorough if he is good. I once brought my car to the garage, because the brake-assistance didn't work and I had to push the brakes like my life depended on it. I had already checked out the car myself, but everything seemed to be ok and I didn't understand the problem. After checking and much thinking, my mechanic discovered the cause: break-assistance only works in between certain margins of the rod connected to the vacuum drum of the assistance cylinder. Because the brake pads had a little too much wear, the braking rod now moved just beyond the operational limit point and no longer worked. Two pairs of new pads solved the problem. Perhaps you don't care how your car is fixed as long as it is, but I was interested because I didn't understand the problem and found the reasoning of the garage mechanic to be very clever indeed. That the fitting of two new pads solved my brake assistance problem, proved him right. This was also a lesson in humility for me. I have learned that garage mechanics can think very cleverly too. You don't have to be a university student for that. By the way, car manufacturers tackled this problem quite some time ago. My car mechanic had a method, was systematic, logical and arrived at a theory which was testable.

I would like to add another thing or two about science. It should be something *active*, that scientists all over the world are engaged in. There is more than just publishing your papers or research results. Within your own field of science there is an active community of people involved in answering and posing questions, solving problems and publishing as well as discussing their research results with fellow scientists. Actually you could say that science is also 'discussion'. It is much more than simply learning facts and the appliance thereof. The exchange of ideas is what makes science come to life. Last but not least, a scientist should be creative; always thinking of new ways to understand the world and why it is the way it is. You don't have to be a student of the ancient languages and cultures to appreciate and use everything that is written in this book. For example, I spoke with a lecturer of a modern language and discovered that he too could use almost everything in this book. He was collecting graffiti writings at certain places and wanted to know more about that. There is really no difference between counting graves in an antique necropolis and counting writings on the wall.

The same techniques apply. If you are creative enough almost everything is possible in science ...

1.3. What is special in research of the ancients?

Researching ancient and dead languages and matching civilizations is not easy. For a start, we have much less information compared to research of our current society. To begin with, the ancients lived a very long time ago and we can't carry out anything *active* anymore like observational research or gather their opinions by sending them questionnaires. Because of that we are at a loss about (a part of) their motives, why they carried out the things they did. That is a serious problem in our attempt to understand them. In this paragraph, I will deal with the main differences between research of current cultures and that of the ancients.

1.3.1. The cultural background of the researcher

Even today in sociology, the science that studies our present day society, the background of the researcher could have an influence on how the research is designed and carried out. Sociologists know this and are more or less aware of the fact that their background might play a role. However, in practice a researcher from one cultural background may have difficulty in understanding the motives of people from another cultural background. Such an admission rarely finds its way into a research report and that is why a scientific approach is so important; it removes (at least to some degree) our own inherent cultural bias.

In the studies of ancient languages and cultures, the cultural background of the researcher is considerably more important and potentially more influential. Contrary to sociology, this has nothing to do with class in society, but your cultural background as a whole. Being a member of modern society can itself present significant problems in understanding the ancient world. Life now evolves significantly differently than it did back in previous ages, especially where really old and far away countries and societies are concerned. Most students learn this from their university professors who, upon entering the lecture theatre, talk about this fact immediately. Unfortunately the differences are not at all that clear and completely obvious from the start. Even though we might know it, we sometimes forget this in situations where the reason to make a difference is not altogether that obvious.

Most students immediately learn that a grave field or necropolis is certainly *not* a churchyard! The term would be 'correctly translated' in the language of today, that is if you are living in a (former) Christian country, because everybody would know what you mean, but it is far from correct if we are referring to the ancient society itself. In most cases Christianity did not exist and consequently there were also no churches. In ancient times the grave fields were *not* necessarily connected to a house of religion. Our present day culture is *embedded* in our language and thinking and we have to realize that. Currently I am researching facts about crime in an ancient society. But what is the word "crime". What is meant by that *back then*? Did the ancients consider the same (f) acts as we to be "criminal"? In our society "crime" is defined in law books and there is a difference between a perpetrator and a criminal. However in ancient times there were no law books, so you need to figure out what people thought about criminal or bad behaviour back then and how they would have reacted to that. But perhaps these examples are only about the definition of things. The object for us is to discover how people really thought and acted and the best way of doing that is to picture ourselves in their shoes. To try

and shake off our modern thoughts and active thinking is perhaps one of the most difficult tasks we have to perform as a good researcher.

1.3.2. About our sources

I've already indicated that we don't have all the data that we would like at our disposal anymore. Let's go a little deeper into this and have a look at the sources that we can use. There are no *active* sources anymore so we will have to rely on material remains of the culture we are studying. We call these remains *material culture*. Almost everything that the ancients once *made* falls into this category. This includes sculpture, paintings, buildings of all sorts, stele's, household objects, weapons and documents (such as papyri, clay tablets and ostraca). There is something peculiar with these objects. Some survived and some did not and there are *various reasons* for that.

One of these is the *climate*. The climate within a certain country can vary considerably. For example, Lower Egypt was moist in the Delta area and not many organic objects survived there. Upper Egypt was hot and arid and that is one of the main reasons that sensitive and organic material, like for instance papyri, had a better chance of survival than in Lower Egypt. So there is a sort of 'unevenness' in the distribution of material culture that has survived and that which has not. That can be a problem if we want to make statistically correct statements about the *whole* of Egypt because objects/documents from Lower Egypt seem to be under represented.

Another reason is *culture*. In ancient Egypt houses and even palaces were meant to last as long as their inhabitants and were usually made of mud brick: for contemporary use only. Not so for graves. These were meant to be places where the bodies of the deceased were kept, embalmed and preserved, to last for eternity. Consequently these were built in stone or cut in the rocks to ensure their existence throughout time. Much of the literature that survived is found in those graves.

There is another *cultural bias* that applies to many ancient cultures; and I am referring to writing and reading. In most ancient societies, most people were illiterate and consequently we know little about the common folk. Because only the elite could write and read, the writings that we have could be coloured by *their* views and perceptions. In addition, there could be an emphasis on material that was meant to survive or not. In ancient Egypt there were no city archives so we know very little of (life in) the city. Religion was more important in ancient Egypt and since temples controlled a large proportion of the economy, we know at least something about that.

Music as a cultural expression mostly did not survive from the ancient world. The instruments typically did not make it through time and the ancients had no ways of putting music into writing, like we do today. Even if they had, it would probably be impossible for us to decode it.

Of course, there are more reasons for the survival or destruction of material culture like war and many other reasons in various cultures, which each had their differences regarding the (non) survival of material culture. If we know them and are aware of these facts we can account for them in our research, but unfortunately, we are not sure or even aware of them in many cases. There is a discussion among scholars whether we can attribute the survival of material culture, and hence important writings of the culture we study, to *chance*. If so, then

we would not *have* to take the reasons for their survival into account in our research. Others will tell you that certain objects did come into our possession by *coincidence*. I am not so sure about that too but, in my opinion, it is better to keep an open mind about these explanations. There are perhaps reasons for the (none) survival of certain material culture that we have yet to discover. Hence, I prefer to treat these matters with caution because you never know what might turn up in the future.

I prefer to look at a possible lack of data as an *opportunity* to test out my skills and creativity, and not to treat it as a *problem*, because that would start you off on the wrong foot. There are many ways to overcome our lack of data. I name a few of the possibilities here. Maybe you could invent some new ones too.

I would think of using my creativity:

- to invent possible *replacements* or *alternatives* for data that cannot be retrieved;
- to design *new instruments* for the data that we are after;
- to look for *alternative* sources;
- to make use of *statistics* to get more out of existing data.

1.4. On becoming a scientific researcher

Most universities are training their students how to become a scientific researcher. That is their end goal. You are acquiring knowledge in your field of study and you are learning how to think and act scientifically and how to carry out scientific research and report on that. In most cases the master thesis that you are obliged to produce at the end of your study, is proof of the fact that you have learned your facts and are capable of single-handedly carrying out your own research, under some (mildly intense) guidance.

The training you get from your university professors, regarding ancient languages and cultures, is usually quite informal but pretty intense as they try to hammer into your head what's so special about the things we are occupying ourselves with. The intention is that a professional attitude towards your field of study is your only attitude towards it. Your attitude towards the material under study and your thinking and acting has to become your *second nature*. My experience is what I have explained above. There are other academic studies that are more formal about these matters and describe exactly what you need to know in order to become a fully qualified scientific researcher. I myself like clarity on this subject so I will try to throw some light on this matter i.e. try to describe what is expected of you, when you have qualified yourself as a scientific researcher. In doing so I will draw up an eight-point profile;

A scientific researcher should be:

1. Well educated.
2. Critical.
3. Objective.
4. Open to ideas of others.
5. Creative.
6. Supportive.
7. Communicative.
8. Possessing integrity.

Ad 1. Well educated.

This is self-explanatory; you simply have to know your business as a researcher. Otherwise, you are no good. That does not mean that you have to know everything. That much is not required of you. However, you have to have the kind of attitude, where you grab a book immediately or go online, every time you feel that you do not know enough to make the things work. A good researcher knows his or her facts and brushes up on things that are stale. This also implies that you do not give away your opinion on academic matters, unless you are well informed and aware of what you do not know.

Ad 2. Critical.

In the fields where we are working, we have to be very critical of all the information that comes to us. In principle, *doubt everything*, unless you are very certain that the sources of the material are valid. There is nothing wrong in checking and double-checking. Better check one time too much than one time too little, has to be your motto.

Ad 3. Objective.

Objectivity has to be 'your middle name'. A good scientist is never biased. However, you can become an adherer of certain methodological schools if you wish. If you do so make that always, clear in advance. I would advise against it though. You keep your hands free if you do not and I do not see a good reason to do so anyway. Scientist will be *expected* to be bias free by the general public and you must be aware of this. You might damage the picture that people have of you, if you make mistakes in this field. One of the most important things is that you, the scientist, enter a study free of preconceptions about the outcome

Ad 4. Open to ideas of others.

Be aware that always there will be people or fellow scientists that differ in opinion from your own. They not only have a right to do so, they might even be correct! However, you do not know that so you do have to take up an academic discussion or an exchange of ideas. Remember that the object is always the search for the logical truth. A discussion on a scientific matter should never be a debate. It is not about winning! Being open to the ideas of others does not mean that you cannot challenge them. Do so and make your discussion public so that the entire scientific community can learn from it and maybe even give support. Ideally, a difference of scientific opinion should lead to something testable that would result in a conclusive proof one way or the other.

Ad 5. Creative.

You can only be a good researcher if you know how to employ your *creative* side. Do not be alarmed. We all have a creative side and you can learn how to use it. Creativity comes in many different forms. There is only one way to do this: think about your subject and write things down. Let it rest for a while and pick it up again. Look at it critically and try to improve. Carry a small notebook and pen everywhere you go or use your smart phone. Einstein was known to carry a notebook all the time. Every time he got an idea, he immediately jotted it down so that he could not forget. Not all the things that suddenly come into your head are useful, but even if you have only one good idea, than you are one-step further. I will not go much deeper into this but try to use your *intuition* as well. What does it tell you to do? Should you carry on like you intended or take up a new path? Sometimes starting anew is the best

thing to do although most people find that too rigorous or regard it as a loss of effort. Do what you *feel* is best. I had one basic ground rule in my old job: if my intuition and logic were not in harmony with each other, I did not do it. Both have to be positive.

Ad 6. Supportive.

Only in rare occasions do we scientists work alone. However, in such a situation we are still *dependent* on other people or fellow scientists. We are never alone. Archaeology only works in teams, so you better build up your skills to get along with your fellows workers or scientists. Be a good sport and encourage your colleagues! I am sure that you will get more in return, albeit not the next day ... My motto is one good turn deserves another.

Ad 7. Communicative.

A good researcher has to have very good communication skills. Not only to present your scientific findings in a clear and orderly way, but also to present your intentions to a potential funding partner. It is therefore necessary to acquire skills in presenting for a large audience and perhaps a smaller but very critical audience of a possible funding partner. Be prepared for critical questions and think about them beforehand. Do not ever make predictions for situations that you cannot foresee, especially where results are concerned. People funding your research will want results for their money. Sometimes you will not be able to provide them. Do not ever make statements that you cannot keep. In doing so you will only make matters worse for your team.

Ad 8. Possessing integrity

Integrity is what every researcher has to uphold under every circumstance. Only then can science really make progress. Never ever, come up with faked data! Every now and then, there is a 'bad apple' among us and stories appear in the newspapers about a corrupt scientist, data manipulations and so on. May you strive never to be on that path, no matter how difficult it is!

1.5. *About the book*

In this book, I have collected *the bare fundamentals* of all you need to know in order to write a scientifically sound MA Thesis. That entails at least, to be able to reason in a scientifically orderly way and to do some simple numerical calculations with the aid of a computer. To accomplish that, I have arranged all chapters logically to build up your knowledge cumulatively from Chapter one up until Chapter 6. I do consider it unwise to skip ahead and start with, for example, the "testing of hypotheses" if you do not already know what a hypothesis is or how to properly formulate one. You will *need* the information in each chapter to understand the subsequent chapter well. I therefore recommend that you read and use this book in the suggested order.

The first part of the book is dedicated to the basics for scientific reasoning, philosophy of science and doing research. The second part of the book is focused on how to prepare your research into a suitable format for (computer) calculations and the way in which to perform data analysis.

A *warning* is in place: this book only covers the very basics of scientific analysis and there is *much more* to be learned. I have added some hints in most chapters on how you can do that

and I encourage you to read more about the topics discussed in this book. It will broaden your knowledge of reasoning and eventually turn you into a better-equipped researcher. Someone better equipped is not only better prepared for the task, but is more confident in their ability to perform it. This in turn can lead to better research with improved results.

2. Reasoning in science

2.1. Introduction

In the sciences of ancient languages and cultures, scientific reasoning is currently taught by university professors as a byproduct. However as these sciences progress new study elements develop, like a philosophy of science, which are usually dedicated. This is an improvement, but not sufficient as students are required to be able to understand the logic and philosophies of *other* sciences as well. To fully comprehend a certain element within an ancient culture, more different and specialized sciences are needed.

Egyptology, for instance, could not be flourishing without the aid of: archeology, medical science, geologists, pharmaceutical sciences, physics and so on. A *general* understanding of the logic of these sciences is needed, to be able to -at least- *comprehend* how they work.

Luckily, scientific reasoning across these sciences is not altogether that different and they all follow the scientific process described in chapter 1. Basically the scientific standards are the same, and we will go into those below. I have limited the descriptions in this to what I consider to be really necessary in order to get a proper foundation of understanding on the situation in your own science and on other sciences as well. The field of discussion is very wide and details are many. Consult other literature or websites if you want to learn more¹.

2.2. Scientific reasoning

3.1. Introduction

In the sciences of ancient languages and cultures, scientific reasoning is currently taught by university professors as a byproduct. However as these sciences progress new study elements develop, like a philosophy of science, which are usually dedicated. This is an improvement, but not sufficient as students are required to be able to understand the logic and philosophies of *other* sciences as well. To fully comprehend a certain element within an ancient culture, more different and specialized sciences are needed. Egyptology, for instance, could not be flourishing without the aid of: archeology, medical science, geologists, pharmaceutical sciences, physics and so on. A *general* understanding of the logic of these sciences is needed, to be able to -at least- *comprehend* how they work.

Luckily, scientific reasoning across these sciences is not altogether that different and they all follow the scientific process described in chapter 1. Basically the scientific standards are the same, and we will go into those below. I have limited the descriptions in this to what I consider to be really necessary in order to get a proper foundation of understanding on the situa-

¹ Like for instance this university website: <http://plato.stanford.edu/>.

tion in your own science and on other sciences as well. The field of discussion is very wide and details are many. Consult other literature or websites if you want to learn more².

3.2. Logic

We have learned from the ancient Greek, that our *reasoning* must be logical and not based on *emotions*. In the modern western world, we are used to call our reasoning *logical* and are inclined to condemn reasoning that doesn't comply with our ideas of logic.

For instance, one of our main rules of logic is: two (exclusive) things cannot be true at the same time. So, if I state that a certain stick of wood is *long*, than it cannot be *short* at the same time. We take that for granted, but is it really true?

Based on what we have learned, and only considering, the rule: yes and no. The ancient Egyptians didn't have problems with two things being true at the same time, so why shouldn't we? If I have stated that a stick was long; that was only my *opinion*, based on the fact that I have known only short sticks in my lifetime. However another person could consider the same stick to be short, because this person was used to long sticks. Probably because he or she was acting in sports, like for instance pole vault.

Have I proven the rule to be faulty? Not at all. I have only shown that there can be multiple *opinions* that are true at the same time, for different people. I did not prove the rule to be faulty. But, have I *limited* the application of this rule, because it does not apply on opinions? No, not even that! Remember, the ancient Greek taught us to rely on *logic* and not on emotions. My opinion, that the stick is long, might well be classified as an *emotion*. The rule and its applicability therefore remain intact.

I wrote the paragraph above as an introduction into the world of logic. Many people take (their) logic for granted and this could lead to mistakes and an impaired vision on reality. We should not do that in science. We must be explicit in our reasoning and exclude emotion in (almost) all situations. Scientists are human beings too and there is nothing wrong in having or showing a bit of emotion. Emotion can be a powerful motivating force in us all, driving us to higher achievements. As such, it is invaluable.

In the past, scholars of the ancient languages and cultures have often been accused, by scholars of the "hard" sciences, of emotional behavior. Especially towards great finds of valuable ancient art. If this is true, then show me a scholar of physics who *isn't* exited and jumping up-and down because the formula that he has been working on for many years, now seems to be working! Showing emotion isn't bad. It's human. Just be careful *not* to let emotion have an influence in your *work*.³

The conditional statement

Logic is about statements. A conditional statement you can apply on real or fictional situations. A well-known example of such a statement is:

² The University of Hong Kong has a very nice and easy to comprehend website on these topics from which you can learn more. Visit: <http://philosophy.hku.hk/think/>.

³ The German philosopher / sociologist Norbert Elias wrote an entire article on this subject. Consult JSTOR. The British Journal of Sociology, Volume 7 Issue 3 (1956) *Problems of Involvement and Detachment*.

“When it rains, the streets will get wet.”

The first part of the statement is the *condition* and the second part of the statement is the *result*. If the condition is satisfied, the result will occur, i.e. it *has* rained and yes, the streets *are* wet.

Now if I have observed that the street in front of my house is wet, can I conclude that it has rained? This is equivalent to saying:

“When the streets are wet, it has rained.”

Is this necessarily true? No. My mother cleaned the street in front of our house! We call this the *inverse* statement. As you notice, the *original* statement might only work one way.

We may try to modify our first statement and claim, that:

“If it doesn’t rain, the streets will not get wet.”

It is basically the same statement but now the condition and the result are put in the negative. We call this a *converse* statement. It is the reverse of the first statement. As you noticed from my observation, this might *not* be true. This means that you have to be very careful with this kind of reasoning!

We can modify our first statement even further, and convert it. We first take the result and then put in the condition:

“If the streets aren’t wet, it hasn’t rained.”

We call this the *contrapositive* statement. If the first statement is always true, then so is the contra-positive one. As we have seen, this is not the case for both the inverse and the converse statements.

“The two parts of a conditional statement have specific terms with respect to logic. The first part is called a *premise*, and the second part is called a *conclusion*. Within a conditional statement, if a premise is true, the conclusion will be too, because it follows, or results from, the truth of the premise. ... In essence, the principles of conditional statements are the same for logical thinking”⁴.

Remember the term: “premise”. It is the *basis* of a statement but could also be the basis or *fundament* of an entire *theory*. As my old philosophy professor once told me during an oral examination, in which I tried to discover faults of reasoning within the writings of a very famous philosopher: “You can’t accuse these people to be faulty in reasoning. They are much too clever for that! Look at their *premises* instead! They could be at fault.”

Check out your skills on logic, now. Do the quiz on the Website of Hong Kong University⁵.

⁴ This is well put, so I quoted it from the website: <http://www.techrepublic.com/blog/10-things/10-tips-for-sharpening-your-logical-thinking/>. Consulted on 1-11-2018; my italic’s.

⁵ <http://philosophy.hku.hk/think/critical/> Note that these quizzes only tests *logic* and not *practice* in real life situations. For instance: if a sensor doesn’t report a volcano eruption, that doesn’t mean that there isn’t any. In practice the sensor could malfunction or may even be broken.

3.3. Causality

The ancient Greek already mentioned causality and the philosopher Aristotle even mentioned four different types of it. Causality revolves around two events that are *inevitably* linked together. The first event is the *cause* of the second event. The second event is therefore the *result* of this cause⁶.

For us scientists, it all revolves around the *proof* that the first event really triggered the second event. A good example, only slightly outside science, is the search that happens at all times, when a building catches fire in your own hometown. There are questions always about what started or *caused* the fire. The insurance company wants to know whether they can pay the insurance money. They want *proof* that the owner didn't start the fire himself in order to collect the insurance money. So the building and the background of the owner will be thoroughly researched. Firemen and experts will be looking for *clues* that can lead to the *cause* of the fire. Scientifically put: they are looking for *indicators* that could point to the cause. Indicators of a deliberate fire could be empty matchboxes and gasoline canisters found around or in the debris of the fire.

Indicators are objects, or perhaps even *events*, that could point to a cause. In this example, the objects or events that started the fire. There can be many indicators for one event and they needn't even be material. To continue our previous example: if you could prove that the owner of the building has serious money trouble, like debts, then this *could be* an indicator of foul play. It is possible but not necessary. If you also found half burnt canisters in the debris of the building *and* link that to the place where the fire probably started, then we have three different kinds of indicators regarding a probable fire starter and *probably* a strong case against the owner of the building.

I have used the word probably and that is because, although this looks like a strong case, it is not conclusive *proof*. It all depends on how well the investigation is performed.

In the research into the owner of the building, did you notice, for example, that he has a considerable number of enemies? There may be people who really hated his guts and maybe *set up* this fire in order to bring him definitely down. It is probable they are the same people that lie at the cause of his money problems. As we can learn from this example, the research of the fire has to be thoroughly and carried out well, just like your own scientific research. I hardly tell you a secret if I claim that a lot of people are in prison for crimes they did not commit because of badly performed research. New research methods, based on DNA-analyses, are improving the situation by providing more evidence and indicators.

The investigators of the fire in our example are actually using *scientific methods*. Scientists of almost all studies will be performing the same actions and perhaps even more than that. Causality is very important. If one event really can lead to another, then we are able to use that to construct *explanations* for events that we have witnessed in history. That is what we, as students of the ancient languages and cultures, are interested in.

As scientist we will be using *hypotheses* considering events that we want to research. A hypothesis is a *statement* about the situation we are researching. For instance, if we are researching a necropolis, our hypothesis may be: "This necropolis is used for the well to do people only." Actually, this hypothesis is not so good, because if I find one grave, belonging to a

⁶ Take note that there is *another* situation in which two events seem related but are not *causally* connected. These events only move together and are often confused with causality. We call this situation covariance.

commoner, I have refuted the hypothesis. Just drop the word “only” and the hypothesis becomes a bit less sharp and less easily refutable, but at the same time -unfortunately- less strong. The task that lies ahead for you, as a scientist, is to prove or refute this hypothesis.

This is where *logic* and *causality* come in. In the last two chapters of the book, you’ll learn how to deal with a hypothesis in a numerical / scientific way.

3.4. Induction and deduction

Reasoning in science has many forms. An important pair of opposite methods is whether to use *inductive* or *deductive* reasoning. Both forms have their own advantages. It is therefore not a matter of what you like or prefer, but simply using one of these methods in a situation where they can have an advantage.

Let us start with inductive reasoning. Induction is associated with *observation*. We depart from the detailed level by studying an item / thing / phenomenon or else and try to arrive at a higher level by turning several individual observations into a statement, rule or perhaps even a law.

Here is an example: If I am studying a swan; in many *individual* observations I have noticed that house swans are colored white. My observations lead me to proclaim the *general* statement that “Swans are colored white.” This is to show you how you can arrive at a higher class of statements, i.e. a general statement, that “Swans are colored white.” This seems simple, but what if I observe a swan with a *different* color than white; what happens then? There are roughly, two answers to this:

1. I *did not* observe a swan; it was a different kind of bird.
2. My statement that “Swans are colored white.” is *falsified* i.e. not valid.

For now you’ve learned what induction or inductive reasoning is. Let us continue with deductive reasoning. The *opposite* of induction is *deduction*. In this type of reasoning we depart from the higher level in the form of a general statement, rule or else and apply this to what we want on a lower level, usually the individual level.

Here is an example: I have learned, during my lessons in Material Culture, that ancient coffins of a certain time period of the people that I am studying are colored black (this conclusion has been arrived at by induction). During the excavations, that I am performing, I have just found inside an ancient tomb a black colored coffin. If other finds within the tomb confirm this, I can now date this coffin, and this formerly *undisturbed* tomb, to the exact period in time that I have learned!

What I am doing here is applying a *general* rule on a *specific* case. *Observations* may be used to create *generalizations*, such as “All swans that I have observed are white, therefore I shall generalize that all swans are white.” The generalization is a hypothesis, which is not proven, but which (for now) I believe to be true. The more white swans that I see, the more confident I am that my hypothesis is true. However, I need only to see a single black swan to disprove my hypothesis. I can apply my generalization to a specific instance to make a *deduction*. For example, if someone tells me that they have seen a swan then I can deduce, based on my hypothesis, that it was probably white.

That, in short, is deductive reasoning. By the way, if the grave I mentioned *was* formerly disturbed, than we could be talking here about a re-burial. It would in this case not be possible to date the tomb to the same period as the coffin inside. The tomb then probably belonged to an earlier time period.

3.5. Confirmation or falsification

Two important methods of verifying the statements that we have learned or used are confirmation or falsification. To continue the example of the swan: every time that I observe a swan colored white, I have increased confidence in my statement that “Swans are white.” Note that this is, as we have learned above, *inductive* reasoning. Confirmation is easy. You will only have to look for evidence that supports your statement, and the statement is confirmed.

There are scientists who consider this to be *too easy* and argue that signals about a possible non-validity of the statement, are overlooked. The confirming scientist is only occupied with the verification process and *ignores* other signals. This is known as positive bias. I suggest we take this critique by heart and try to make the confirmation process a little stronger. One method of doing this is, to actively look for *contra-indications* yourself during the validating of your statement. You have to perform this task with vigor, and you must *really* try to find indications that might or do stand in the way of the confirmation of your statement.

In the report or article, you are writing, you must report about this type of search. Particularly, *how* you performed this task. If you didn’t find anything that might challenge your confirmation, you actually made this a little stronger. But remember this is a *conscientious* task. You must be able to place yourself in the position of your opponent and reason like him / her. If you are writing an article, then you do have another option. Challenge your opponents! Describe in your article what you have done and how, and challenge anyone who thinks they know better. Allow them to send a letter to the editor and continue the discussion in the magazine. Scientific magazines do have this task. We can place these two suggested methods somewhere in-between confirmation and falsification, the next method.

In the next paragraph we will discuss a well-known theorist in science, Karl Popper. He considers *confirmation* not to be possible so as to achieve a scientific status for a statement. He suggests that, in order to be scientific, a statement has to be *falsifiable*. In other words, it must be possible to *test* the statement. If the statement is *not* true, then there must be a way to prove that. Actually this brings back *deductive* reasoning in the debate. Popper considers it *not* possible to derive a scientific statement based on several individual observations i.e. he does not believe in the *validity* of induction. However he does think it is possible to *falsify* a statement. Only *one* observation is needed to falsify a statement. My conclusion is that Popper needs induction for that too, in order to really falsify a statement.

Popper had quite an influence on science, because from the time of launching his ideas, scientist all over the world talked about falsification. In Poppers view, the statement doesn’t actually have to be falsified, in order to achieve the status of *falsifiable*. The possibility alone is enough.

What if your statement is actually falsified? Is there a way to avoid this or perhaps to adapt the statement? Yes there is, and this has to do with how you *formulate* your statement. In the

paragraph about causality, I used an example about a necropolis. I stated a hypothesis that I immediately rejected as not so good:

“This necropolis is used for the well to do people only.”

As you can see this hypothesis is formulated so sharply, that it takes only the detection of *one* grave of a less well to do person, and the hypothesis is *falsified*. The secret lies in the formulation of the hypothesis and I suggested, taking out the word “only”. The hypothesis is now less prone to falsification. However the force of the statement is also slightly diminished. The example of the swan that I have used is well known and is often used in this type of discussion. I made the statement that:

“All swans are colored white.”

This statement is falsified when I observe a *black* swan. Like above, you could argue that this was *not* a swan, but I am afraid that the odds are against me here.

How do we modify my statement so that it matches the new found reality? I could state:

“Not all swans are colored white.”

Or, alternatively:

“Most swans are colored white.”

Both reformulations of my original statement are correct and not easily falsifiable. However I have made a loss on the “power of expression” of my statement. So, yes, you can prevent your statement from being falsified, but it will cost you (perhaps) dearly.

2.3. *Progress in science*

It is not so easy to determine if science over a certain time-period has progressed. The reason for this is that there is a lot of division of opinion among scientist about how progress, looks like, how it is to be defined and measured, what progress really constitutes etc. There is also much discussion on whether science progresses gradually over time or in leaps (or ‘quantum steps’) at irregular intervals. There is an enormous amount of literature on the subject and we can only touch briefly upon the main subjects. Below you will find the major theories in a few words explained. For anyone who wants to read more, you don’t have to even leave your house anymore to visit a library. There are a number of nice sites on the web from which you can learn more. I recommend that you only have a look at websites belonging to a university⁷ because that offers more certainty, that what you’ll be reading has the required scientific level. There is some nice stuff on Wikipedia too, but I advise that you read that only *after* you have acquired enough knowledge to be able to judge the material for yourself. Before we have a look at the main theorists and their material, we will discuss the main topic of this paragraph, progress in science.

2.3.1. What constitutes progress?

Before we answer the question of what progress in science might be we first have to answer the question of what *science* is, or constitutes. This question too is not easy to answer. This is mainly because there are quite a few opposing views among scientists about the definition alone. Progress in science is even more debated. Let us start with an attempt to come a little closer to what “science” stands for.

Robert A. Nisbet’s (a sociologist) general idea is that science is a collective enterprise of researchers in successive generations and that it is characteristic of the Modern Age.”⁸ By taking up this position, Nisbet shares himself under the historicists and with that under one of two opposing basic schools of thought: *the scientific realists*. The scientific realists declare that science is searching for the truth and hold the view that scientific theories are true, roughly true, or very probably true.

This school of thought is diametrically opposed to *the scientific antirealist* or *instrumentalist*. They argue that science does *not* focus at least or succeed in, unraveling the truth. They regard it erroneous to view scientific theories as even potentially true.⁹ There are even antirealists that claim that scientific theories aim at being instrumentally useful and should only be regarded as useful, but not true, descriptions of the world.¹⁰

⁷ Like for instance this university website: <http://plato.stanford.edu/>. The University of Hong Kong maintains a very nice website on philosophy: <http://philosophy.hku.hk/>. The Internet Encyclopedia of Philosophy is not a university website but it maintains scientific standards. You can search for items within this website: <http://www.iep.utm.edu/>

⁸ the journal *Literature of Liberty: A Review of Contemporary Liberal Thought* , vol. II, no. 1, January/March 1979.

⁹ Levin, Michael (1984). "What Kind of Explanation is Truth?". In Jarrett Leplin.

¹⁰ van Fraassen, Bas (1980). *The Scientific Image*. Oxford: The Clarendon Press.

These seem to be radically opposed views. What are we to think of that? For the purpose clarity on the subject, that allows us to continue, maybe we can find some kind of definition that does justice to both opposing views.

According to the Oxford dictionaries¹¹, Science is “the intellectual and practical activity encompassing the systematic study of the structure and behaviour of the physical and natural world through observation and experiment.” This definition avoids the statement of whether or not there is *truth* in science. Maybe rightfully so, because the question of truth is extremely hard to answer; perhaps even, not at all.

Personally I considered the definition on Wikipedia of science¹² as a “... systematic enterprise that builds and organizes knowledge in the form of testable explanations and predictions about the universe.” Quoted from Wilson¹³ and The Oxford Companion to the History of Modern Science¹⁴, to be very practical and close to my own, because of the use of the words “systematic”, “organizes”, “testable explanations and predictions” in the definition. It nicely steers clear of some hotly debated scientific differences.

With this definition of science in mind, the next step is to define what actually constitutes *scientific progress*. To start with, the general view is that thinking in the western world, is influenced in considerable measure by the ancient Greek philosophers. Thales of Milete (c. 624 BCE – c. 546 BCE), is considered to be one of the “Founding Fathers” of Greek philosophy, at least as far as another great Greek philosopher, Aristotle, is concerned.

A modern philosopher, Bertrand Russell, states that “Western philosophy begins with Thales.”¹⁵ One of the reasons for western philosophers to be so enthusiastic about Thales is that he was one of the first to diverge from the path of his forefathers. Up to then it was customary to explain phenomena in the world, by referring at the will of the gods. Thales was the first to try to explain natural phenomena, by reasoning and thus constituting a *rational* explanation. Especially Aristotle (384 BC – 322 BC) was considered to lie at the heart of modern western thinking. Through his writing and thinking, an all-encompassing philosophical system became possible, including ethics, aesthetics, logic, science, politics, and metaphysics. His influence reached us via the medieval world, through the Renaissance up until the 19th century. Among scientist he is considered to be to first builder of a formal study of logic.

Not so long ago scientists attributed scientific *development* to a specific historical context (period, location and culture) which facilitated the change process. We call this view *historicism*. Hegel and Marx (the latter influenced by Hegel) were among the historicists. Other streams of thought oppose this view, like the *reductionists*, who attribute change to elementary principles or consider them product of random influences. Critique on historicists was also put forward by philosophers like P. Feyerabend and T. Kuhn. They considered mere historical circumstances not enough to explain the rather drastically changes that occurred in science in the last centuries.

¹¹ <http://oxforddictionaries.com/definition/english/science> (viewed on 1-11-2018)

¹² <http://en.wikipedia.org/wiki/Science> (viewed on 1-1-2018)

¹³ Wilson, Edward O. (1998). *Consilience: The Unity of Knowledge* (1st ed.). New York, NY: Vintage Books. pp. 49–71.

¹⁴ Refer to the remarks of the Editor in *The Oxford Companion to the History of Modern Science*. New York (2003) Oxford University Press, p.vii

¹⁵ Russell, Bertrand *The History of Western Philosophy*. New York (1945).: Simon and Schuster.

An article in the Encyclopedia Britannica about *change*, states one opinion that took hold some time ago¹⁶; "...a natural conception of scientific progress is that it consists in the accumulation of *truth*". In the same article¹⁷ this is immediately changed in: "... scientific progress consists in accumulating truths in the "observation language."¹⁸." Clearly, that addition is not enough and I am certainly *not* into a definition that has the word "truth" in it, because defining that is extremely difficult and hotly debated.¹⁹

I am afraid that this type of definition didn't bring us closer to a more modern and qualified solution for change. I am clearly in line with scientific anti-realists here that rejected this kind of reasoning too. Maybe a kind of definition on a lower level is better suited for our needs in this book right now. We could define "change" as a *succession of statements* that alter during time. And "growth" as an *improvement* in the way that these statements have over each other, in reaching for the best way to reflect reality, i.e. the subject we are researching or evaluating. And this brings us back in line with the scientific realists.

As you have noticed, this is a very difficult matter and up to now philosophers of science do not agree on the subject. I encourage you to form your own opinion and read more. However be careful, because this is a truly difficult matter and there is an abundance of literature, thinkers and opinions in existence. On many occasions it is hard to see the forest because there are so many trees.

Below I will discuss briefly a number of important scientists and thinkers and their thoughts on the subject of progress in science.

2.3.2. History of science in the 20th Century

Main author(s)

Richard Boyd

Background

In history, religion and technology did not occupy themselves with building up knowledge systematically. In this foregone period of time we cannot speak of "scientific progress". Most efforts were directed at passing on knowledge to the next generation, and, not in a -what we would call- scientific way. Some of the ancient Greeks did accumulate knowledge systematically but this came first to light regarding the accumulation of knowledge of the art of warfare. Boyd considers the history of science to be predominantly the study of the historical development of science and scientific knowledge.

Over the last ages, starting with the Age of Reason, the Enlightenment, the history of science is by and large the replacement of myths and superstition with more "knowledge of reason".

¹⁶ <https://www.britannica.com/topic/philosophy-of-science/Scientific-change> (viewed on 1-11-2018); my italics.

¹⁷ Ibid.

¹⁸ The philosophers Hanson, Popper, Kuhn, and Feyerabend agreed that all *observation is theory-laden*, therefore there cannot be a theory-neutral observational language.

¹⁹ When I read the word: "truth" I suddenly remembered part of a song, from my favorite musical: "Jesus Christ Super Star" by Andrew Lloyd Webber and Tim Rice (1970): When Jesus was arrested he was brought before Pontius Pilates and Pilates talked about truth. Whereupon Jesus sang: "But what is truth? Is truth unchanging law? We both have truths - are mine the same as yours?"

Main publication(s)

Co-editor of the book *The Philosophy of Science*, MIT Press, 1991

Theoretical concepts

Boyd is an adherer of scientific realism. In line of thought, the outcome of scientific research is knowledge of largely theory-independent phenomena. According to scientific realists, such knowledge is possible even when the phenomena are not actually observable. "Scientific realism is thus the common sense (or common science) conception that, subject to a recognition that scientific methods are fallible and that most scientific knowledge is approximate, we are justified in accepting the most secure findings of scientists at face value."²⁰

Boyd advocates a model of scientific progress: methods in science that came in existence are employed to produce scientific theories, which in turn are used to produce more methods, which lead to more theories and so on

Influence on science

The book "Philosophy of Science" is extensively used in undergraduate and graduate philosophy courses.

2.3.3. Logical positivism

Main author(s)

Ludwig Wittgenstein

Background

Roughly in-between the first and the second world war of the 20th century in Vienna Austria, there was a discussion group called the "First Vienna Circle". Members met at the "Café Central". In 1929 Otto Neurath, Hahn, and Rudolf Carnap wrote an article that briefly summed up the main thoughts of the Vienna Circle at that time. Their ideas were partly based on early work of Ludwig Wittgenstein. As the Second World War came more and more nearby, with corresponding political upheaval, members of the Vienna Circle dispersed. Some died and some migrated to England and some to the USA. Up until the 1950s, logical positivism was the leading school in the philosophy of science.

Main publication(s)

- The Tractatus Logico-Philosophicus (1922). (Latin for "Logical-Philosophical Treatise")
- Philosophical Investigations (posthumously published in 1953). Wittgenstein modified many of the ideas in the Tractatus Logico-Philosophicus.

Theoretical concepts

Logical positivists considered philosophy to be limited and only having a function on the organization of thought. In order to describe reality, they used *statements* about the reality.

²⁰ <http://stanford.library.usyd.edu.au/archives/spr2009/entries/scientific-realism/#6> (viewed on 1-11-2018)

These statements had to be *verified* in experiments to bring about their truthfulness. They rejected metaphysics and ontology because they considered those subjects to be meaningless.

In logical positivism, the idea that *observation* can generate knowledge (*empiricism*) is combined with mathematical or linguistic constructs or statements (*logic*). Logical positivists are using *deductions* of epistemology (i.e. the teaching of the nature, the methods and the boundaries of human knowledge).

Influence on science

The influence of logic positivism on science was, and still is, considerable to say the least. A school of thought that relies on formal logic and analysis of languages in the tradition of logic positivism is considered to be a (form of) analytic philosophy today. Logic positivism is still used in science these days despite a number of criticisms.

Main criticism

- Rejection of the inductive principle of verificationism as inadequate by Karl Popper and the anti-scientific realists.
- The philosopher C.G. Hempel considered the verification principle itself as not verifiable²¹.

2.3.4. Logical empiricism

Main author(s)

Carl Gustav Hempel

Background

Hempel studied mathematics, physics, and philosophy in Germany and Vienna. There he attended a few meetings of the Vienna Circle. Rudolf Carnap, also an important philosopher of the 20th century, was born in Germany but held an official US citizenship. He helped him migrate to the US before WWII started. There Hempel taught at a number of universities in the US: New York, Yale and Princeton.

Main publication(s)

- Hempel, C. G. (1934). *Beiträge zur logischen Analyse des Wahrscheinlichkeits-Begriffs*. Universitäts-Buchdruckerei G. Neuenhahn, Jena. (1934)
- Hempel, C. G., & Oppenheim, P. "Studies in the Logic of Explanation." *Philosophy of Science*, 15. (1948)
- Hempel, C. G. *Fundamentals of Concept Formation in Empirical Science*. Chicago: University of Chicago Press. (1952)
- Hempel, C. G. *Philosophy of Natural Science*. Englewood Cliffs, N.J.: Prentice-Hall. (1966)

²¹ This is an all time favorite question in philosophy: "Does a statement have to comply with its own condition?"

- Hempel, C. G. “Valuation and Objectivity in Science.” In R. S. Cohen & L. Laudan (Eds.), *Physics, Philosophy and Psychoanalysis*. Dordrecht, Holland: D. Reidel Pub. Co. (1983)

Theoretical concepts

Carl Gustav Hempel questioned regularly the theory of confirmation, which has importance in Logical Positivism. Together with fellow philosopher Paul Oppenheim²², “...he proposed a quantitative account of degrees of confirmation of hypotheses by evidence. His deductive-nomological model of scientific explanation put explanations on the same logical footing as predictions; they are both *deductive* arguments. The difference is a matter of pragmatics, namely that in an explanation the argument’s conclusion is intended to be assumed true whereas in a prediction the intention is to make a convincing case for the conclusion. ... He also emphasized the problems with logical positivism (logical empiricism), especially those concerning the verifiability criterion. Hempel eventually turned away from the logical positivists’ analysis of science to a more empirical analysis in terms of the sociology of science.”²³

Hempel liked paradoxes and he is famous for his formulation of the “Paradox of the ravens.” This is a quest into what exactly constitutes *evidence* in a statement. It is a problem of induction and deduction. Hempel uses statements, like a true logic positivist, to illustrate his paradox:

1. All ravens are black.

In strict logical terms, we can turn this statement around, like this:

2. Everything that is not black is not a raven.

Logically reasoning: if 1. is true than 2. must be true and if 1. is false than 2. must be false. For instance, if you could find a white raven, than these statements would be false. These two statements seem to be logically equal. We could make, *mutatis mutandis*, comparable statements about lemons. (Keep the first statements in mind.)

3. This yellow object (and not black) is a lemon (and thus not a raven).

Following the same kind of logic, it appears as if observing a yellow lemon supports the thesis that everything that is not black is not a raven. This appears paradoxical, because it seems that information about the raven was obtained by looking at a lemon.

Influence on science

Hempel was a significant advocate of logical positivism until later in life when he revised his opinions and is considered to be one of the prominent philosophers of science in the twentieth century.

²² Hempel and Oppenheim’s essay “Studies in the Logic of Explanation,” published in volume 15 of the journal *Philosophy of Science*,

²³ Adapted from the Internet Encyclopedia of Philosophy IEP <http://www.iep.utm.edu/hempel/> (viewed on 1-11-2018). My Italics.

2.3.5. Falsification

Main author(s)

Karl Popper

Background

Karl Popper was born in Vienna in 1902 and later in life immigrated to New Zealand. After WWII he moved to the United Kingdom where he taught Logic and Scientific method at the London School of Economics. In 1949 he was appointed professor at London University. K. Popper was knighted by Queen Elisabeth II in 1965. He died in 1994 in London.

Main publication(s)

- *The Logic of Scientific Discovery*, London and New York (1959). K. Popper.
- *Conjectures and Refutations*, London (1963).

Theoretical concepts

In his “Logic of Scientific Discovery”²⁴, Karl Popper clearly sketches what he identifies as: “The problem of induction”. He states that no matter how many observations you make of swans, you cannot conclude that a statement is *true* like “Swans are white”²⁵. This has serious implications on the validity or truth of universal statements. Popper introduces the term falsifiability. An almost infinite number of observations or verifies of one statement is not enough to confirm it, but it takes only one *opposite* outcome (the observation of a black swan) to *falsify* the statement (that all swans are white). Popper would therefore like to make demands on statements: A statement is only *temporarily* true until it has been falsified. In the eyes of Popper, statement(s) or theories are only scientific if they are able to be falsified.

Influence on science

The influence of Popper’s writings on science can hardly be overrated. He is generally regarded as one of the greatest philosophers of the 20th century. Scientists all over the world embraced his ideas. People now realized the shortcomings of Logical Positivism and -partly- how to deal with that. Falsifiability is still an import criterion in formulating hypotheses, statements or theories.

Main criticism

- The falsification is an *observation* too, and can be rejected as well.
- Confirmation and falsification rely on observation or instruments. Either can be *wrong*.
- Falsification does not *solve* the problem of induction.

²⁴ You can now read this book yourself by downloading it in PDF from the following link: <http://strangebeautiful.com/other-texts/popper-logic-scientific-discovery.pdf>.

²⁵ Perhaps only your confidence in the statement has increased.

2.3.6. The scientific revolution

Main author(s)

T.S. Kuhn

Background

Kuhn was born in Cincinnati Ohio USA. He earned a PhD in science from Harvard University in 1949. There he taught History of science until 1956. He then left for the University of Berkeley CA, where he became a professor of History in science in 1961. At this university he published his most influential work “The structure of Scientific Revolutions”. After that he joined two other universities until he died in 1996.

Main publication(s)

The Structure of Scientific Revolutions, Chicago (1962)

Theoretical concepts

Kuhn rejected the view that scientific progress moves along a *linear* path. Instead, science undergoes shifts in *paradigms*. In paradigm A, science is accumulating knowledge along a certain path. However, over time, *anomalies* are starting to build up, indicating that the trodden path might not be the correct one. The anomalies are growing stronger and stronger and eventually become so strong that paradigm A can no longer be held as just and must be overthrown. Then a new paradigm B is established, indicating the present state of science. This paradigm is also vulnerable to anomalies and exactly what happened to paradigm A eventually will happen to paradigm B. This goes on and on indefinitely. Kuhn hastens to add that this does not happen very often.

Kuhn recognized three phases: In the *pre-paradigmatic* phase, there is no agreement on a certain theory. There are also conflicting theories and opinions. Theories might not be completed, tested or both. The second phase is that of *normal science*. It is how we described paradigm A and B. Kuhn talks about *puzzles being solved* within the dominant paradigm. The third phase is where the normal science gets into trouble because of the anomalies accumulating. As a rule, crises can be resolved within normal science, but when that becomes impossible and the paradigm fails, a new paradigm is established. Kuhn refers to this phase as the revolutionary phase.

Influence on science

Kuhn's influence on science was huge. He contributed greatly to the thinking on the progress of science and the words “paradigm” and “paradigm-shift” has been established as buzz-words in many sciences. He is particularly popular in the social sciences.

2.3.7. Ockham's razor

A bit of an outlier, compared to the theorists above, is William of Ockham. He is a medieval philosopher and minimalist.²⁶ His contribution to the philosophy of science is well worth considering. According to his “philosophy” the simplest explanations, after testing, turned out to be better justified, than other more complex explanations.

His theory is focused on determining what *hypothesis* is the most likely to be true or false in the case of competing explanations or hypotheses. His razor is supposed to shave away unnecessary material on a hypothesis or else cut it in pieces. Because there are not many *competing* hypotheses or theories, his contribution is limited. To make a better use of Ockham's razor, it is possible to *quantify* an explanation by studying the length of the message. In this way we could construct a “measure of simplicity”.

A warning about Ockham's ideas is in place here. He is *not* indicating that the simplest explanation is always the best in every situation. What he does is demonstrating that out of two or more *alternative* explanations, the simplest one is *usually* the best suited.

2.3.8. Criticism of science

Criticism of science concentrate and hit on theories and methods applied within the scientific community in order to express their shortcomings with the intention of improvement. One of the main advocates of this stream of thought is Paul Feyerabend. In his book: *Against Method*, London (1993), he stated “... that there are no useful and exception-free methodological rules governing the progress of science or the growth of knowledge, and that the idea that science can or should operate according to universal and fixed rules is unrealistic, pernicious and detrimental to science itself.” This line of thought is an improvement on “Epistemological anarchism.” Feyerabend would like to limit the power of science in society. He accuses science of not having proof for its own philosophical guidelines.

In feminist philosophy the question has risen whether science can be considered as gender neutral. The answer was, probably not and some suggestions were made to come closer to the ideal of gender neutrality. One of these was improvement upon the way in which questions were asked in a survey, mainly because of the critique upon the gender bias in the language of science.

²⁶ Read more details on for instance: <http://math.ucr.edu/home/baez/physics/General/occam.html> (viewed on 1-11-2018).

2.4. Taking up a position

Knowing the above theorists and their ideas, might help you out taking up a position if you want, but you don't have to do that in order to be a good scientist. It is perfectly ok for you *not* to choose sides.

My personal view on this matter is rather practical. I have always chosen those theorist and theories which were *best suited* to the subject I was handling at a certain moment. I did choose another theorist or theory regarding another subject. Changing positions, methodologically, was rather easy for me. My audiences were not connected. Herein lays the problem if you strongly advocate yourself as an adherer of a certain line of thought. It is very likely that *your* audience will be the same throughout and switching positions probably can result in you losing your credibility. But that applies only if you are taking up position in an extremely outlined methodological camp. As long as you do not take up extreme positions, you can get away with a (perhaps) slight or moderate change in position. But be warned. But do not mix in one research! You will lose your credibility certainly if you do. You can however mix methodologically *only to demonstrate* how different streams of thought will work out on your research topic. Usually that will be a moment for you to actively choose sides within your paper. Applying theories and streams of thought, has to be done deliberately and with great forethought.

3. Doing research

3.1 *Types of research for ancient and dead civilizations*

3.1.1 Introduction

The object of this book is an introduction in *scientific research*. That indicates that we will be using academic methods for our research that are scientifically approved. You will find those further below in this book. As outlined before, not every type of research available to scientists can be used because of peculiarities of the situation. We are studying *ancient* and *dead* civilizations. Sciences that study present day societies, like sociology, political sciences, psychology and economics, are studying *dynamic* societies with *living* people they can question, observe etc. We, on the other hand, are dealing with societies in which its members are dead, and cannot be actively questioned, observed etc. This is a serious drawback, because we will have to rely on *secondary sources* to obtain our research material. These sources will be primarily:

- A. The *language* of these people in various forms of writings.
- B. The *material* culture, which the ancient society has left behind.

Ad. A. The Language

The language of dead societies comes to us in different forms.

- There are many *literary genres*, like letters, poems, biographies, hero's tales etc.
- State communications like announcements on steles and law books, recorded criminal process and documents but also administrative texts like accountancy reports.
- Religious and funerary texts in the form of inscriptions.

From each of these uses of language can we learn how society functioned or was supposed to function. The morals and thinking of these people during the period of investigation can be outlined using all texts available.

Ad. B. The Material culture

Every society in history has produced various kinds of objects that can still be studied today. Art is one category and so is clothing, weapons, domestic utensils etc. There is a *reason* why these people made the objects and we can learn from the various uses of language how and when these objects were used.

Material culture is also a great instrument for *dating* purposes. Studies that involve ceramic objects like pots indicate the use of certain types within known timeframes. This makes these pots potential ideal "time indicators". It is possible, for example, to date ancient gravesites by studying the grave goods. However, there are other material objects, like weapons or domestic tools that make good time indicators as well. Do not rely on *single* time indicators and gather all you can to get more certainty about the possible dating of a site.

3.1.2 General research types

We discern two basic types of doing research: qualitative and quantitative.

Qualitative research is focused on grasping elements of culture through methods like observation or questioning of participants. This type of research *describes* your subject in words. For instance, an ancient African tribe could be described as ‘peaceful’. Stoneware could be described as ‘intricately decorated’, a landscape or site could be hilly. It usually generates more detailed, and perhaps, subjective insights versus general insights in society. This type of research is habitually used on dynamic societies, but there are types that we can use on dead societies as well like text analysis.

Quantitative research is used on material we can put numbers on i.e. quantify. This research type is used when we want to make general and valid claims, usually numeric, resulting from our research. For instance, we can measure the height of statues. This type of research generally is based, and relies on, statistical principles and analysis.

Another central distinction in research is *applied* versus *basic* research.

Basic research is a research type that adds to the general knowledge of humankind or our scientific community. It increases the quantity/quality of observations or a theory. There is no money to be made from this kind of research. Basic research is largely fundamental to Applied Research.

Applied research is a research type that uses general and other applied research to create or invent something new. Theory can be used in experimentation or development of new products. For instance, like a new medicinal drug. The results are used to improve the basic (theoretical) research. Certainly, money is to be made with this kind of research.

The distinctions sketched above are largely academic and the borders between the types will be drawn less sharply in practice. Even basic research can have its advantages for the applied research like the discovery of new principles or methods.

Funding is not easy for basic research whereas funding for applied research is relatively easy. There is great return on investments for the last type. Generally speaking *our* research will fall in the basic research category. It is fantastic to add our knowledge of say a certain archaeological dig to the scientific community, but apart from our job as researcher or university professor, there is no money to be made. Maybe one can strike a small deal with the publishers of the find ...

There are still more (detailed) types of research to be mentioned:

Cross sectional versus *longitudinal* research.

Cross sectional research focuses on the study of one point in time; for example fourth dynasty China. By definition, they are not able to show *trends* over time i.e. changes in morphology or style in material culture.

Longitudinal research focuses on the study of change over time. For example, we can note differences in Aztec burial practices over time. An application of this type of study is the *trend study*. Samples are taken from different time-periods and compared to each other to dis-

cern a certain kind of trend. This type is very applicable to ancient civilizations and we can come to know all about changes in dress, pottery, burial customs, etc.

Cohort studies focus on populations of individuals with common experience. For instance, what happened to the students of class 2010? This type of study is extremely difficult if not impossible for studies of ancient and dead societies because this study usually involves questioning of the subjects. Participants in major historical events may thus be an exception here.

In *panel studies*, the same individuals are studied at different times or time intervals. This is not very easy for the students of ancient civilizations, however it is not impossible. We can study important persons and the course of their life, from descriptions of their lives by others or from autobiographies and other written documents like marriage records, tax records (inheritance) and the like.

A *case study* examines a certain person or organizational unit in depth. This type mostly belongs to the qualitative type of research and is not easy to perform by students of the ancient civilizations, because of the static nature of the society we study. Still a lot of these case studies are done by students of ancient civilizations who write descriptions of important persons like religious persons or military persons or persons of state.

The *survey* makes use of questionnaires to quantify events in a civilization. Sociology and political sciences make active use of this kind of research. A survey can be used in basic as well as applied research. This type of research is based on statistics and can be broadly applied. It can generate general and fundamental knowledge. Surveys are generally held by political scientists during election time to predict which turn voters will take. Commerce also participates regularly in this kind of studies to determine which new products to fabricate and sell. This is called *marketing research*. We deal with the prerequisites of this type of research later in this book.

Experiments are often used for testing. They are both used in the “hard” sciences like physics and chemistry as in the “softer” sciences as psychology and sociology. We may discern between a *true experiment* and a *quasi-experiment*.

In a *true experiment*, the circumstances of the experiment can be 100% controlled and the variables can be fully manipulated or controlled. This means that the cause (or lack) of change can be directly attributed to the variables that were tested. An example is the testing of a new type of medicinal drug. A control group of say 100 ill persons, *the experimental group*, is given the actual drug and a different *control group* of a 100 ill persons is given a non-functional harmless substitute. The reason for the substitute or *placebo* is to remove the variable of people’s psychological reaction to receiving medication (functional or otherwise). The researcher tells the groups *after* the experiment.

In a *quasi-experiment*, there is less control over circumstances and or variables. In a functioning society, a researcher cannot control every variable. The society functions largely as it is. The object is to test a hypothesis and keep it as free from intervening variables as you can. The experimental design is measuring of the situation ‘as is’ → Appliance the experiment → Measuring of the new situation. The object of the last measuring is to detect the results of your experiments by comparing it with the first measuring. It is very difficult to *prove* that that an altered new situation is the *result* of your experiment. There *could* be other mechanisms at work you have not reckoned with. At best, you are able to show that the results of

your intervention *might* be due to your efforts. I myself used such an experiment²⁷. In trying to fight crime in shopping centres, I conducted an experiment involving, among other things, extra surveillance in the centre and the shops. I first measured the rate of crime in the shopping centres before the experiment. This we call a *zero-measure* and then again after the experiment, the *one- or experimental measure*. Certain aspects of the crime rate did drop but not all. You cannot simply control society in this period that lasted about nine months of the experiment. It is also not easy to attribute the measured changes to the experimental measures (the variables) I took. For example, there was publicity involved, and that might have had an influence. I could have used an experimental shopping centre and control shopping centre, but considered that option non-*ethical* in view of the willingness of the owners of the shopping centres and shop owners, that participated in my research.

There is also *content analysis* or research. This is used to discern and interpret events in all sorts of texts. This type is very useful in studying ancient civilizations as we have only a limited amount of sources from which we can learn about.

Then there is *comparative analysis*, which tries to find similarities and differences across -in our situation- multiple cultures, locations, situations etc.

3.1.3 Specific research types

One specific type of research is the *replication research*. This type of research can only be done after *other* research has been carried out first. Generally, that specific research has some characteristics that make it interesting for a fellow scientist. Several reasons for this apply. For instance, the research can have some unknown or untested methodology, or the research findings seem improbable. The conditions were not specified enough or were unlikely, or there are other reasons of scientific methodology that makes the research interesting. A researcher can decide to *replicate* that research to find out whether he or she arrives at the same conclusion. In other words, the research is done again using all (the same) data and conditions that the first researcher has used.

Prerequisite for this type of research is the availability of data, methodology used, research instruments etc. This type of research is often done in the *basic research* type to make sure that the findings are sound, because future research will be based upon it. However, it is rarely done in our type of research, because of the scarcity of funding. It is difficult to explain to funding parties why an already finished research should be done again. The expected outcome should be the same (?). The possibility of replication research is an academic prerequisite. It hinges on clear and comprehensive documentation. I consider replication research to be one of the most powerful means of testing within science. The scientist can be reviewed thoroughly on his honesty, integrity and the use of instruments (like dropping cases that contain contra information against a wanted research outcome). If you carry out academic research, you will have to ensure that other independent researchers can replicate your results. You will therefore save your data and decisions regarding your analysis for replication purposes.

²⁷ “Het winkelcentra project” WODC - SDU, Den Haag 1988

3.2 Methods of doing research

3.2.1 General research methods

In principle, every method known in present day research methodology applies to the research of ancient societies, except for those where the active involvement of participants of those societies is required. We will consider that later on, but first we shall consider the ordering of research. There are many ways to do this. A much-used method is to categorize research according to the intended results:

- A. Explorative research.
- B. Explanatory research.
- C. Testing.

Ad. A. Explorative research.

Researchers doing explorative research are *exploring* the civilization they are interested in. Explorative research is usually *descriptive*. This type is focused on “reconstruction” of parts of the civilization. In other words, we are trying to find out how things were and make a description of that. In Assyriology, Egyptology etc. much is still unknown and we are still completing the puzzle of how this civilization once looked like and functioned. Many authors make their contributions and together they are completing the puzzle. This research is *basic* and meant for other users to be used for their purposes. Because of that, this research has to be *very good* and tested, because other scientists will build their work on this previous research. Remaining errors will persist and compound.

Unfortunately, I have not seen much testing of this research and I consider that dangerous for the science as a whole. Of course, scholars *are* critical towards each other’s work and their scientific contributions *are* evaluated by others. However, there is *no real testing* here! There is no replication research. Remember that this research is *basic*, which is fundamental for fellow scientists.

One of the reasons for that is already mentioned: lack of funding. Another reason is that there are simply too many ‘white spots’ i.e. uncharted terrain, for us to discover first in the societies we are studying. Testing is considered to be a luxury. Also new areas are more interesting, associated with more prestige than following someone else’s footsteps. Most scientists seem to reason that the white spots must be covered first, and after that, there might be room for testing. Personally, I do not agree and prefer a small solid base above a not so solid broader base like there is now.

In practice, explorative research focuses on a certain topic of interest to the researcher. He/she gathers all the information there is on that topic from other publications, *explores* the field and makes a description of this new topic. This new description in turn can be of interest to another researcher. The question remains: how solid was the basis for this research? I make a plea here for as much *independent* research as possible. Fortunately, this *is* happening, as there are many countries that own archaeological institutes in the countries under surveillance and fund the researches tied to that institutes. However more is needed and we do not have to be in the spot ourselves in order to perform original research. Even *copies* of original material like papyri and ostraca, can be studied at home without losing their original value. This study can

make a valid contribution to the grand total of basic research that we have. Even the *testing* of such research can be done in a valid way at home and at low costs.

Explorative research makes use of ordering or classification techniques. It compares and tries to explain differences and relationships.

Ad. B. Explanatory research

Explanatory research is used to unravel working methods, political and organizational questions, etc. and intends to solve questions about *how* and *why* something is done. This type of research is often used for theory building. Although many researchers are truly interested in the how and why, they usually *describe* their findings as soon as they have uncovered the answer to their research question and leave it at that. Unfortunately there is little or no attempt to theorize about the subject on a *higher* level, apart from the actual unravelling itself which can only be considered to be a theory on a *micro* level. In a predominantly *descriptive* science like Assyriology and Egyptology, theory usually does not go *beyond* the micro level. A theory with a scope further than the micro level, like a *medium* level theory that works on the plane of organizations or parts of society or a *grand* theory that has an explanatory level of societies, is very rarely used or constructed in research on ancient civilizations.

Ad. C. Testing

Next to the categories explorative and explanatory research, testing can be indicated as a third category although strictly it can be applied to all forms of research. One of the most prominent representatives in this category is *replication research*. This type is solely developed for the testing of existing research. Within the mentioned categories, another type of testing is possible. I am referring to the testing of research statements or *hypotheses*, what could be, and often is, an integral part of a quantitative oriented type of research. Research results that may indicate, for instance, a difference between male and female variants of statues and can be tested against a so-called *null-hypothesis*. This type of hypothesis states that, there is *no* difference and that all the differences found can be attributed to *chance*. We describe this kind of hypothesis testing in the last part of this book. Considerably more complex is the testing of a theory.

Globally stated, the theory has to be “translated” or put into operational terms that can be measured. These terms are then measured by way of indicators. Their association with the main topic is established and the outcome then evaluated

3.2.2 Methods of qualitative research

In the previous paragraph, qualitative research was indicated as focusing on details of cultural elements, mostly through active methods like observation and questioning of the participants. But there are types that we can use on dead societies. Most of the qualitative research that we will be performing is related to an analysis of the *sources* that we have left from the ancient culture we are studying. As you will recall there were only two main sources: language and material culture. In present day science, by and large, the last category is not considered to be a part of qualitative research. Normally the focus is on *active* participating members of society and their thoughts and actions. As we cannot actively study living members of society anymore, we must fall back on language i.e. written text. Because this reduces our view on

society as a whole, we need to “back-up” the information to be gathered from language, with additional information. I suggest that we make use of material culture to elaborate on our findings from studying language and vice versa. Anything that we can use to enhance our limited sources is fine by me²⁸. For example, if during our research, we detect a change in posture of a certain type of statues over time; we must fall back on language to try to recover its origin. In my view, both sources must be confronted with each other, to arrive at a higher level of knowledge. As we have only limited sources, we must make the very best of those. Below we will briefly discuss a number of techniques used in this type of research.

Literature research

Literature research is very important for students of the ancient cultures, because much of our information about these cultures will have to come from analysis of literary sources. Literature research focuses on only *one* specific aspect of language, namely the *artistic* side of it. Examples are poems, song texts, (heroic) tales etc. There is much literature about this use of language around and some websites can help you out as well. Modern methods still apply. There are about seven important categories in this kind of analysis that you should pay attention to when analysing a text. These are:

- Symbolism.
- Character development.
- Plot.
- Author background.
- Location.
- Setting.
- Writing style.

However, this way of analysing only applies to the artistic side of language. In order to analyse a broader category of the use of language, we can revert to content analysis.

Content analysis

Content analysis is a kind of methodology that is used for analysing the content of communication. Often this is narrowed down to written or oral communication. However, because we are studying a dead society, we need all the information we can gather. Earlier on, I indicated the study of material culture in order to complete the study of language and vice versa. Therefore, I opt for a much broader definition. Fortunately, I am not alone in this. Even modern scientists consider other forms of “communication” as well. Earl Babbie defines it as “the study of recorded human communications, such as books, websites, paintings and laws.”²⁹ However, in Babylonian times, there were no websites and little or no *written* laws. We can replace them by filling in all the other “communication” there was. I suggest even a broader definition than that. As far as I am concerned, the design of a temple or a fortress or in fact any other piece of material culture, carries a message inside and can be analysed with content analysis. If your research can handle it, you may make use of this kind of analyses. A temple can carry a message. It says for example “look how large and impressive I am ... I demand a place in your mind ... Remember the gods of this temple ... Pay your respect etc.” However, these mes-

²⁸ In the social sciences Denzin promoted ‘triangulation’. This is using more than two methods to check up on the results and thus make up for a more solid base of your data. Well worth considering if our sources are deficient like in the research of ancient languages and civilizations. Denzin, N. (2006). *Sociological Methods: A Sourcebook*. Aldine Transaction.

²⁹ Earl Babbie, 'The Practice of Social Research', 10th edition, Wadsworth, Thomson Learning Inc.

sages are not enough. They have to be coupled to other elements of culture in order to get significance. Here O.R. Holsti³⁰ has an important message. He has designed a few simple and amazingly powerful tools. Holsti makes a difference between the *content* and *reference* aspect of communication. Remember that. It is vital in understanding communication.

A modern day example: a woman tells her husband “It is eight o’clock, Henry!” So, does she want to tell the time to her husband? No, he can tell the time as well. Therefore, there has to be something else. In fact, it turned out that Henry normally takes the dog out for a walk at eight o’clock in the evening. Henry’s wife simply wanted him to take the dog out for a walk. This is a nice illustration, that the *content* of the communication “It is eight o’clock.” might not be what it looks like. The *reference* is “the habit of taking the dog for a walk at eight o’clock”.

Social Network Analysis

Within a community or civilization, people interact with each other. When people join certain organizations, groups, or other organizational entities, they become members of a so called “social network”. Within such a network the interactions of the people within, will take up a certain kind of pattern. These patterns can have an influence over their behaviour and hence are important in understanding these people and their work or behaviour. Interaction between individuals and their interactional institutions can be analysed. A method for this is formalized and is called social network analysis or SNA. The relationships between people are viewed in terms of *network* theory. A network consists of *nodes*, the actors themselves, and *ties*, the relationships. Points and the relationships represent the actors by lines in between the nodes in network diagrams. This type of analysis can become complicated. Fortunately, there is software in the public domain to support this. There is a free plug-in for Microsoft Excel to support the analyses, which may be downloaded from this website³¹.

Participative observation

One last technique of our qualitative research repertoire, which we sadly cannot use, is participative observation. This is a technique³² in which the researcher observes one or more individuals. This might be done in an explorative way to gather information or to test a theory. Since the civilizations we are studying are dead, the chances of using this method are slim indeed but not impossible. I can imagine us using this technique regarding for instance a re-enactment of a ceremony. If this is done properly with attention to all the details, I can imagine that observing techniques like this can bring us new information.

3.2.3 Grounded theory

I already mentioned the fact that sciences of the ancient and dead civilizations usually do not theorize beyond the micro level and hence stay, primarily, descriptive. In the mid 1960’s, a couple of authors designed a way of generating theory bottom up and by relying on research techniques described above. These authors were Glaser and Strauss and they argued in favour

³⁰ Ole R. Holsti: Content Analysis for the Social Sciences and Humanities. Reading, MA: Addison-Wesley 1969.

³¹ <http://nodexl.codeplex.com/>.

³² Want to know more? Read the following article on the web: <http://www.qualitative-research.net/index.php/fqs/article/view/466> (viewed on 9-1-2019)

of *inductive* analysis, i.e. theory derived from the lowest possible level and aimed at a higher level. This was new, because up to that point, scientist argued that mid-range and top-level theory was only possible by way of statistical i.e. *quantitative* methods and not by *qualitative* research. Glaser and Strauss³³ defined grounded theory as “the discovery of theory from data systematically obtained from social research”³⁴ These authors were influenced by the line within sociology of “Symbolic-Interactionism”. Main authors of this line of thought are George Herbert Mead and Charles Horton Cooley and important interpreter H. Blumer³⁵ Grounded theory has two main features³⁶: it is *inductive* as I explained above (bottom up) and it is *iterative*. The analyses is a continuing process whereby data and analysis sort of ‘fall back’ on each other in such a way that the outcome is growing better with each cycle until no new information can be extracted from the data. That is the ultimate goal and is called *theoretical saturation*. A new “grounded” theoretical outline should than be available.

3.2.4 Methods of quantitative research

Quantitative research is a type of research in which *numbers* play an important part. This type of research is often displayed as the exact opposite of qualitative research, but actually, it is not. The only real difference with qualitative research is that there are different research methods being employed in this type. These methods predominantly have their basis in statistics. The object is to test a hypothesis or even complete theories and to make predictions about phenomena in the society. These methods are supposed to be objective in their working sphere, but caution must be taken towards this kind of research because it falls or stands given the exactness all the other research prerequisites. The exactness is the strength of this kind of research and the bridge between the empirical world of phenomena and the theoretical world can be bridged give a solid research design. In statistical terms, the quality of outcome of this kind of research can be given within certain, usually predefined, mathematical limits. Usually this type of research is being employed in sciences like sociology, psychology, political sciences and physics but also in marketing research. To these sciences, we can now add the sciences of ancient cultures and civilizations. The last part of this book is entirely devoted to this type of research.

3.3. Designing your research

In this paragraph we discuss the prerequisites of a solid scientific research design for the sciences of the *ancient cultures and civilizations*. In the last paragraph, I took up the position that these sciences could be added to the list of sciences studying a civilization. A civilization still is a civilization even if it is thousands of years old. The ancient Egyptians, Babylonians and lots of other ancient civilizations had for instance a division of labour, a hierarchy ruling the land etc. All the elements we still see today within our civilizations. Therefore, I see no reason to treat the science of ancient civilizations any different from sociology, psychology or even economics. The same rules and prerequisites for research apply and that is why I am going to rely on *modern* research methods in my exposition below.

³³ Glaser B.G. ,Strauss A.L. *The discovery of grounded theory; strategies for qualitative research*. New York (1967)

³⁴ Ibid (1967:2)

³⁵ Blumer, H. *Symbolic Interactionism; Perspective and Method*. (1969) Englewood Cliffs, NJ: Prentice-Hall

³⁶I have based my explanation on a wiki from the university of Amsterdam (UVA): http://wiki.uva.nl/kwam-cowiki/index.php/Gefundeerde_theoriebenadering. (Dutch! viewed on 1-11-2018)

I will discuss two types of research, going from general to specific:

1. Design for an explorative research.
2. Design for testing a theory or theoretical elements.

The processes of performing such types of research are roughly the same. They differ only in the *method* of research. With a different subject comes a different process, but regarding the design, only for that (relatively small) part of the research. Most of the research *steps* remain the same. I will therefore start with the description of a complete set of research steps for an explorative design. Then I will handle the design for testing a theory or theoretical elements, but only the steps in which it differs from the explorative design.

3.3.1 Design of an explorative research

Most students of the ancient civilizations will be performing a more *general* type of research. They usually are interested in a certain *subject* on which they wish to learn more. After stating their object of study, they scan all available literature; make a global design of their report and start reading and writing. That is how it is today, but hopefully *not* tomorrow. Although generally viewed, this kind of design is not entirely wrong; it is not *specific* enough to meet the scientific standards of today. Below, I will outline a set of *mandatory* standard research steps i.e. steps that every research performing student has to follow through, *before* the final research report can be written. The *steps* are the same for both explorative and testing research, but the *contents* of certain steps may differ. The process within the steps will be dynamic and cyclic as each step interacts with the previous step or the step(s) to come. During your design and carrying out of the research, you may find out that you will have to make some modifications to the decisions in the steps that you took earlier on or in the next. This is completely normal and part of the dynamics of your research.

The research steps to be taken in each research design are:

1. Introduction.
2. Research question.
3. Delineation.
4. Method & Operation.
5. Data collection.
6. Organizing and analysing data.
7. Drawing conclusions: answering the research question.

Ad 1. Introduction

In the *introduction* of your research, you give insight to the reader about the subjects that you are interested in and why. You may point out what is already done in this area and your scientific opinion about that; theoretically or empirically, perhaps even both. Do mention authors, theories and/or research fields. Also, lift up a tip of the blanket about you. Who are you and what is your background and what are your specialties? What do you hope to accomplish, or what are you striving at? Make sure that the reader knows everything there is to know so that he/she is able to follow your reasoning throughout your entire research process.

Ad 2. Research question

One of the main focal points in your research is the *research question*. Your entire research will evolve around this question, which is why it will have to be phenomenally good. In this part of your research design, you pose the question that the rest of your research activities will have to give an answer to. The question you formulate will have to be one-dimensional, clear, simple and, for every other researcher in the field, comprehensible. That is not to say that you may not pose more research questions. You may pose several if you wish, but they must all comply with what is stated above. However, do be careful! It is not easy to give a scientifically satisfactory answer to one question, let alone to two or more. When you do ask more research questions, I recommend that the second and following questions are in line with the first and *main* question of your research.

So if I state the following, example, research question: “What is the purpose of the Epic of Gilgamesh?” I might then ask, as follow-up questions: “By whom was this probably written?” And: “Why was this written?” Or, alternatively: “What is the best remaining copy of this epic and why?”

Another possibility is “What is currently considered to be the most valid translation of: ‘The Descent of the Goddess Ishtar into the lower world’?” “Is it possible to improve upon the translation; and if so, why and how? The possibilities are endless, but the example questions above may not be very easy to answer. Be careful with what you ask!

Ad 3. Delineation

The next part of your research design is to determine exactly what part of the field you are going to study and what not. The skill of designing the delineation is that you focus, as sharp as possible, on all of these things that are related to a possible answer on your research question. The intention is for you to narrow down the field of your research, in order to not get lost in details that drive you away from your path. Besides that, you usually will have a limited time to carry out your research and you will have to cut yourself short. For instance, if you are going to research a certain type of statue, select only one *time-period* (for example, in Egypt, in the Middle Kingdom) and one *category* of places (for example, temples). You may go even further than that and specify the *material* the statues must have been made of in order for you to study them. If that does not specify the research field enough, try more factors like *gender* or *location*.

This research step, like the most of the research steps, is *dynamic*. It may turn out, during the course of your research process that you will have to limit yourself even further or, conversely, remove some limitation. It may be even so that you will have to focus your limits on different aspects, because the initial path you projected did not bring you where you wanted to go. Do not get upset about this, because it is quite natural and common. However, be careful. Research is not about tuning your research parameters to get the results you *want*. It is about scientifically reporting *facts*, based on *traceable* and *replicable* research. The sheer reality of things is wayward and may not turn out to be what you expected.

I have called this step ‘delineation’ because that is what I am used to do. Alternatively, you could also refer to this step as determine the ‘scope’ or ‘scoping’. In many instances, the scope is often limited by funding. If you are writing a Master Thesis, then it is limited by the time and resources you have available. Usually a reduced scope or a narrowly defined delineation will have a profound influence on your research: Reduced scope means fewer research

variables which results in less available evidence, which gives a lower confidence level in the results. The scope should be driven by and appropriate to the question. Conclusions must be limited by the scope, but prompt follow-on questions for further research.

Ad 4. Method & Operation

Apart from your conclusion, this is the most central and technical part of your research. In this research step, you must shed light on exactly *how* you are going to answer the research question(s) that you have formulated in the second step.

To begin with, if you intend to make use of some *theoretical notions* or parts of a theory, explain that briefly, or project a chapter in your research report devoted to that. Explain *why* you are using those theoretical notions and *how* they affect your research. Do not forget to give definitions of the subjects that will be under research. To give an example: in my own master thesis, I discussed crime in ancient Egypt. I suggested some modern sociological theoretical notions about crime in society that may also be applied in the ancient Egyptian situation and explained some theoretical notions about the way in which people in a society viewed crime. This formed the starting point for my research to follow.

In the theoretical notions above and/or in your research question, you may be using some technical or theoretical *terms* that need to be explained. For example, in my own master thesis about “crime” in ancient Egypt, I had to clarify, what I meant by using the term “crime”.

Among other things, I wrote “In our present day society, what *we* consider to be a ‘crime’ is defined in *laws*. There were ancient societies that had their laws written down and enforced. In ancient Egypt, there was no codification of law until the late period. However, even if there was such a law book then their definition of crime, according to the law than, does not have to correspond to *our* definitions and thoughts about this subject. ...”

In the end, I came down to the term “deviant behaviour” as a working definition, because this can be evaluated against the existing values and norms in a society, which can be traced in ancient literature.

In this step, you should next answer the following questions: What are you going to do (your research method)? And: how are you going to do it? Back to my example: I am going to study ancient literature from which I can distil the values and norms people once had, *and* evaluate reports of criminal trials from the period. From these we can learn how people considered the *seriousness* of the events committed; this was also a variable in my theory.

Summarizing: in this step, you give some theoretical notions, provided you have these, you mention the way in which you want to work and you indicate the literature you will be studying. Finally yet importantly, you set out your research: *what*, *where* and *how*.

Ad 5. Data collection

The object of this step is to outline *how* you are about to collect the research data and *where* i.e. your sources. In the example I gave you about the research for my own master thesis, I did not study potsherds, steles, statues, grave goods or ornaments of temples and the like. However, I gather that perhaps you will and this is the step where you can outline your intentions. If for instance you are going to study the decorations of graves of nobles, what will you be looking at and how will you do it. Will you be recording the type of decoration, the use of colours, the size, the quantity, the place of the decorations etc. etc. Where are you going to find the data you need? In books, museums, on the spot itself? In addition, exactly how are you

going to gather your data? In the final report of your research, this has to be outlined in detail. Later on in this book, you will find out how to arrange this part of your research in practice.

Ad 6. Organizing and analysing data

In this step, you are really in the heart of your research. You will have a pile of information at hand and your job is to make meaningful information out of it, which at the same time is scientifically correct. The object of this step is to outline how you did this and inform your readers about it. How you can analyse your research results is outlined in the latter chapters of this book.

Ad 7. Drawing conclusions: answering the research question

As I outlined before, these research steps are all more or less *interactive* with the other steps. If you want to draw conclusions, you will have to make sure that the analyses you undertook are suitable to infer those from your research material. This is the time to look back at the research question(s) you posed and have a look if you are able to answer it or them. If not than possibly, you have to modify one of more steps so that you can. Perhaps you might even rephrase your research question, but I already warned you about the validity of your research. This could be at stake, if you are not able to show why you would take such a step. The explanations you *must* give in this instance, has to be phenomenally good. People who judge your research will question you about this too. However, there still is another possibility. The question cannot be answered because of serious limitations within the field. There is no shame in finding that out. The scientific world could still benefit from that, as long as you can *explain* exactly where all went wrong. Which paths cannot be trodden, and precisely what are the reasons for that?

3.3.2 Design for testing a theory or theoretical elements

The following topics are now discussed again in view of a different application of your research i.e. testing a theory or theoretical elements, in the empirical.

Ad 2. Research question

Sometimes a research question or problem definition can pop into your head because of something a fellow scientist states or from what comes to you by the media. In the example I am about to give you, this happened to me. In the first year of my study Egyptology, I read a statement from a scientist, to which I could not disagree more! This person claimed that the theory of Marx could *not* be applied on ancient Egypt. I was already a sociologist for a long time and had studied Marx quite well. I knew she was wrong and I wanted to prove it. Of course, you cannot apply *all* material of a theorist that wrote about the situation of turn of the 20th century England. That is unreasonable to expect, but Marx has some great economic theoretical notions about the exchange of goods in the economy and how this situation turns into a capitalist economy.

My research question was:

“Is (applicable) Marxist theory, suited to be used as explanation, next to other theory, for the economic relations in ancient Egypt during the Old, Middle and New Kingdom?”

That was quite a presumptuous research question, so I had to narrow that down a bit in my delineation. What I try to show you here is, that it is quite possible for you to test (part of) a theory or some theoretical notions that you think might be applicable. If it works, then that is great, if it does not than you have shown that it *might* not be done. I am cautious here. Remember the term *replication research* that I have told you about? Perhaps *you* could not but someone else could. It is also possible that you are proven wrong in your view that the theory was applicable. Other researchers might not agree and there can be a debate in the scientific magazines. I would *love* to see that, because that could be an indication of the growing maturity of our research field!

Ad 3. Delineation

Continuing my example, I had to limit myself, not only because of the rather presumptuous research question that I had stated, but also because of limitations on the size of my paper. I *selected a part* of Marxist theory that seemed the most applicable to me: chapter 1 thru 5 and chapter 12 of his famous book “Das Kapital”³⁷, which discusses barter and the division of labour and small workshops respectively. To rule out foreign influence in ancient Egypt, I carefully selected periods, where this was not the case i.e. the Old, Middle and New Kingdom. In this example of delineation, you can see some arguments and reasons and you can use in order to limit your research and rule out bad influences. The object is to have a good look at the theory and consider beforehand what might affect the application of it. You can switch off those factors, but the applicability of the theoretical notion might suffer a drawback at the same time. If you rule out certain periods, then the *applicability* of your theoretical notion will then automatically be *restricted* to the periods that you have researched.

4. Method & Operation

In this research step, you will find the most difference compared to the previous one. What you have to do is to *describe* the theoretical notions or theory briefly. In this description, you can *mark* the concepts or statements that you want to test against the empirical, by making them bold, cursive or underlined. *Explain* why these are the elements that are at the core of the theoretical notion or else why you want to test these.

Back to the example: In my own description of the theory of Marx, I made the words **bold** that I wanted to test against the empirical, i.e. a description of the economy of ancient Egypt, during the Old, Middle and New kingdom. I have also used the division in the theory that Marx made himself, but you can make your own division if that is better suited to your own purpose. Here is a fragment of the text I wrote:

³⁷ A word of caution, reading Marx. It is not easy! There are authors that have translated and explained his work. You might consider reading them instead. Look at the following website that may provide material for you in your own language: <http://www.marxists.org/>. I am no Marxist myself, but I do value his contributions to the understanding of economics in society, particularly of economies in ancient societies.

“... Chapter 1 The goods. Goods are objects outside of humankind... The usefulness of a thing gives it value of use ... If it is useful; you have to be able to trade it against something that is also useful. ... Apparently, two different things have something that is the same and that is the same size. Each of both must thus, as far as value of trade is concerned, be reduced to this third aspect ... The **value of trade** is marked without considering the usefulness ...”

What you can see, is that I consider “value of trade” to be a prime concept within this part of the theory. Now this theoretical notion has to be “translated” as it were in operational terms, i.e. in terms that I can *measure* or detect in the empirical. We call this *operationalizing* of the theory. The “things” that we can measure/detect in the empirical about our theoretical elements, we call *indicators*. A bunch of logically bundled indicators we can call a *testing instrument*.

5. Data collection empirical description of the civilization

In my example of the Egyptian economy, it was actually very simple. I first studied quite a number of texts that describe the economy of ancient Egypt. I then subdivided these texts into three levels. The *micro* level is about trade of *individuals* with each other. The *meso* level is about *organizations*, in this case: (temple) workshops. The *macro* level, is the level of the *country* itself or the state. An example is foreign trade, or state interventions in the economy. I then made a description of the ancient Egyptian economy according to this subdivision. This description in turn then served the purpose of the empirical. I had to reconstruct the empirical because there was no material at hand that was complete. I had to build one for my research purposes out of texts handling different details of the economy.

6. Organizing and analysing data

In this step of research we are going to actually test the operationalized theoretical notions or *indicators* against the description of the empirical. Continuing my example, I have studied the reconstruction of the economy that I have made on each level and tried to discover the evidence of trading. In ancient Egypt, there was definitely a form of exchange trade, so this indicator is confirmed. In my original text I have used ten indicators about exchange trade and three about the division of labour and the presence of workshops and I could confirm each of them. I described all levels and tried to discover the presence of the indicators that I have distilled from the theoretical notions of Marx’s theory. In the end I have made a table of all indicators and explained where I have found them. As you hopefully have learned from the previous chapter about reasoning in science, finding matching elements of a theory in the empiric is not so strong evidence for the support of it. If you look good enough, you could always find something that matches. One of my reasons for designing *many* indicators is exactly that. You might find one or two, may even three or more, but certainly not ten indicators! So this design of indicators makes my case stronger. What could make my finding of indicators in the empiric even better, is by trying to *reason away* the match you’ve found in each case of a confirmed indicator. This is a true test of your honesty and value freeness as a scientific researcher. If you can’t reason away the matches you’ve found then you have an even stronger case.

My conclusion in the end was, in research terms: “The hypothesis that the monetary theory of Marx was applicable on ancient Egypt, has not been disproved.” I.e. in my opinion: “Marx’s theory of exchange trade, the division of labour and presence of workshops, was sufficiently supported in the empiric.”

On top of that, it turned out that I was able to shed some light on another misconception about the ancient Egyptian economy. Namely that in ancient Egypt, there was no profit making. Of course there was! The problem was that you could not *detect it* easily and that you have to apply (Marx's / or some other) theory in order to locate it. Profit was usually made in the temple workshops at the time. In fact, this happens repeatedly in research. Results that you did not expect to find are thrown into your lap and these may turn out to be much more interesting than the thing you were looking for in the first place.

4. Gathering data

4.1. *The power of the number.*

Numbers are very important for a researcher, independent of what kind of a researcher you are. There will always be situations in which you are confronted with numbers. This is nothing to be afraid of. If you can manage your own home budget, make a few simple calculations and understand the basics of numbers, than you can become a researcher that is also skilled in numbers. It is not necessary to be an advanced mathematician (at least not in the fields that the readers of this book will be interested in).

We have all learned how to add, subtract, multiply and divide in primary school and that is quite enough to make a few nice tables for the presentation of our research findings. So do not worry, as you already know the basics. You just need to know how things are done and what the rules are. All you will need to know is explained in this book, particularly in this chapter.

In the sciences of ancient languages and territories (or “Area Studies” as they are now called) simple descriptive statistics are usually more than enough for the presentation of findings of your fieldwork to a selected public of fellow scientists and to a wider audience. You need to know the rules and have some understanding of chance (or probability), as well as some insight in sampling and samples, but no higher mathematical skills are needed for that. It is enough to simply know the rules and to follow them. However, you will also need to have some insight about how certain results are arrived at, but you do not need to make detailed calculations.

Learning the rules is primarily intended for judging your own work and that of others. Regarding the work of other scientists, there can be a lot of confusion. Much has to do with journalists eager to present “hot” research findings and omitting the necessary caution with which the quoted researchers worked. A heap of “hot” and rather blunt looking conclusions, presented in newspapers and on the internet are just actually that. These numbers are usually presented *without* their very necessary background, which would have put them in a completely different light. There is an important lesson in this. Trust no journalist whatsoever until you have got hold of the *original publication* from which you can decide for yourself if all the generated fuss was necessary and the research findings really are, what they appeared to be.

Remember that the number, i.e. research findings can have a profound influence. Not only on the public, but also on politicians who will grab hold on to everything that they can lay their hands on in order to prove their own right and another person’s wrong. So be very careful about the way in which you present your findings, especially to a wider audience. If you know how, you could really manipulate others with statistics!

Below you will find the way in which you can “convert” your research findings into numbers and how to prepare them for analysis. The best way to do that is to make a *research codebook*.

4.2. The research codebook.

4.2.1. Purpose.

Making a codebook is a good way to structure your thinking about the object(s) and the data to gather. A codebook translates everything you want to know about the subject you are re-searching in: *codes*. A code is just a simplified form of a fact you want to know and record, making it suitable for analysis (e.g. entering into a computer program). Usually a code is expressed in *numbers* because we can calculate with numbers but not with letters like a, b or c. The purpose of this chapter is to show how to make a codebook and teach you the *scientific terms* that researchers use when they prepare *data* for *analysis*. These terms will be in *italics*.

Ultimately, the codebook that you will be using should be definite. Which means that you have thought over all you want to know very well; played with the material, changed stuff repeatedly until you came up with the right and definite version.

4.2.2. Variables, values and labels.

Objects involved in our research will have a number of different qualities. They may be small or big, made from different kinds of material and so on. The qualities of the object(s) you are studying can be represented by what we call *variables*.

Example 1

Let us state that statues can be large or small, made from bronze, wood or clay. Therefore, your research objects, the statues in this example, will have *variation*. You can express this variation with one or more *variables*. A variable denotes one single *quality* of a (research) object. In this example we will differentiate between “size” (small or big) and “material” (i.e. bronze, wood or clay). In this way, we can make comparisons between the different variables in our analysis later on. When designing variables we have to make our *criteria* known beforehand. What do we mean by small and big in the first variable? People may differ on what they experience as “small” or “big”. To make thing perfectly clear you must define those qualities “small” and “big”. For instance, we will call all statues beneath the height of one meter “small”. Statues of one meter and more, we will call: “big”.

In this example, our analysis could arrive at the conclusion that that large statues were usually made from wood and that small statues were either bronze or clay. If we did *not* separate all the qualities of the research objects in different single variables beforehand, we would not be able to determine these results about our research objects.

Example 2

If you are studying statues, it is important to state whether these statues are representing males or females. Therefore, we will call this variable: “Gender” and for easy reference *label* it: V002 - “Gender”. The “V” represents the Variable and the number 002 the reference number of this variable. A *label* is short for the quality we are recording in a certain variable. We can use a label to quickly differentiate between the large number of qualities of an object that we will be recording and analysing. A label is used both for *variables* (to differentiate between them) and for (qualitative) *values*.

A variable has variation and the total number of variations that a variable has, we call: *values*. What we need to do now is to determine exactly what *values* a variable can have and assign *labels* to that variation. The variable V002 - "Gender" will have only two variations or values: male and female. Therefore we will *label* these two values of the variable V002 - "Gender": "Male" and "Female".

V002 (Gender)
Male
Female

Are we there yet? No. In order to make a computer understand our labels and calculate them for us, we need to add a "*code*" to it, i.e. assign a *number* to the values. To continue with our variable V002 - "Gender", most research institutes will code as follows: 1 = Male; 2 = Female (or the other way around, whatever you prefer). That is easy and simple. In this case, we do not need to write out "Male" or "Female" anymore. We simply assign a *code* that represents our values from which the computer can read: "Male" or "Female". This makes the data easier for a computer to read, but harder for a human.

V002 (Gender)	
1	Male
2	Female

That code must be put in a place where the computer can find it and use it. Usually this shall be a *data file* in which the whole of our research results shall be recorded. To make things easy, we will assume that our codes will be put in a *column* (a vertical line); namely, the column in which we "encode" or write our variable: V002 - "Gender". If we put all our codes regarding the gender of a statue, in the same column; than we only have to add up all the different codes of our research objects in order to state how many statues represent males and females. The computer can do this for us in the blink of an eye, so that's convenient and (hopefully) free from errors as well, provided we have made no mistakes. Because the variable V002 - "Gender", is rather limited in the number of values it will get, i.e. two, we will reserve one *position* per statue in our column for this variable.

Unfortunately doing research is not always that simple. There is a small but important catch: we have to provide for possible problems to be encountered later on. Let us say that there were a number of statues in the total number of statues we are studying, our so-called *research population*, that were badly damaged. It may be impossible to state whether these statues were once intended to be male or female. What are we to do about that? Well, there are more solutions to that problem. You could add another value to the variable V002 - "Gender" namely; "Undecided" and code it: 3. That could nicely solve our problem. However, it is not the way in which this universal problem: "lack of data", is generally solved by research institutes. The *convention* here is that we put the value "missing" in our variable. There is also a convention in the code itself, that we will assign to the value "missing" = 9. We have reserved only one column for recording the results of our variable V002 - "Gender" so we can put a single "9" in our column, to denote that we do not know the gender of this statue. Code "9" means no data available or "missing". In this manner, our problem is solved. Note that there could be more variation in statues, like androgyny, both male and female, etc.

V002 (Gender)	
1	Male
2	Female
9	Missing

Example 3

Another example of a variable may be the diversity of clay pots that were found near an ancient burial place. You can make a variable for the type of pot discovered, like: V033 – “Type of Pot” and the values are for example: for cooking, for storage, for display, for ceremonial purposes, for burial purposes, and so on. Many research institutes skip the values bit and mention the code first followed by the value label. They code as follows:

V033 (Type of Pot)	
1	For cooking
2	For storage
3	For display
4	For ceremonial purposes
5	For burial purposes
...	...
9	Missing

Do not forget to put in the option for missing data: 9 - “missing”. This variable also needs only one column for recording.

Perhaps you will need more variables in order to record supplementary data about the pots you have found, like: the condition of the pots and the dimensions.

V034 (Condition)		
1	Poor	More than 60% damage
2	Average	40% - 60% damage
3	Good	Less than 40% damage

Once again, define what you mean by: “Poor”, “Average” and “Good” or refer at international standards. For instance: I will call a pot “Poor”, if more than 60% of it is damaged. I will call a pot “Average” if between 40% and 60% is damaged. And I will call a pot “Good” if less than 40% is damaged. Draw your own lines here.

In this variable, it is not necessary to add a value “missing” because we will be able to judge all the pots. To be sure you could add extra values if you want to be more specific; for example: “above average” and “very poor”. So add more values in-between of what I suggested.

You might want to expand on qualities of the pots by adding extra variables. Next to the condition, you probably would like to say something about the dimensions of the pots. Now we are entering the realm of the quantitative variables. These are treated differently the qualitative ones above. In order to record the dimensions of the pots, make variables to account for: “length”, “width” and “height” in inches or centimetres. Probably add “content” in cubic inches or centimetres too³⁸.

In designing the variables make clear, what the boundaries of these variables are. What numbers are you expecting? How many numbers do you need for recording purposes? One number but ultimately four numbers should be enough. It seems a bit strange perhaps but do add the “9” or “99” for no data available or “missing”. The value “missing” might be recorded in case of a broken pot or knowledge about pots or other objects that you have heard about and, for the moment, cannot lay a hand on.

4.2.3. Types of variables.

Generally, variables are determined by the type of *data* you want to record. These data are translated into *numbers* so that a computer can make calculations for us, as explained above.

MEASUREMENT LEVEL

All types of variables correspond with a level of measurement.

Nominative Variables

There are nominative variables that only assign a *label* and no ordering whatsoever, like “gender” as discussed above, or “political party”, or “sort of pet” you may own and so forth. Thus you need only to assign the labels to get the variable right. Therefore, if your country has for instance 25 political parties; you will have to reserve 25 labels for the parties and one label (i.e. 99) for the missing value.

Ordinal Variables

In ordinal variables, there is some sort of ordering between the values, but the ordering is very loose. We only know that the second value is larger than the first and the third value is larger than the second is. However, we do not know by how much. There are no absolute distances among the valuables. A statue may be: 1 = small, 2 = medium and 3 = large as perceived by the general public, if you ask them that. In the same way a temple may be small, medium or large. Distances can be close, further away, far away or unknown. The differences between the smaller and the larger are not necessarily equal. It is therefore absolutely necessary to define your categories. What do you consider “small”, “medium” and “large”?

Scale Variables

These variables are also called quantitative variables. If you are using them, you have to be clear about the units of measurement. For instance if you are measuring and recording the heights of statues, state your unit of measurement. Are you using centimetres or inches? Are the weight measurements in kilogram’s, pounds or stones? These variables can consist of two

³⁸ If you are acquainted with this subject, you probably know that there are many, many more dimensions to be taken into account. If you need them, just add them. I try to keep it simple here.

types. In most statistical packages, they are treated equally because the differences between them are small. They are:

Interval variables

These variables have the same distance between the former and next value, but there is no definite beginning like a “0”. A much referred to example of such a variable is “Time”. We do not know when it started, but each second, minute, hour and so forth is the same. Perhaps ages ago there were other ideas about time than now. If you know these, you can use them in your research, but for now, we assume that a starting point is not within our knowledge.

Ratio variables

A ratio variable is only different from an interval variable, in the sense that here a true starting point or “0” is known. For example, a width, length, volume and mass of an object can be measured, calculated and compared with other objects.

EXAMPLES OF VARIABLES

For students of the ancient civilisations, it is not uncommon to date objects by the year in which they were produced, used, found, discarded or else. We are not palaeontologist who go back 60.000 year or more, so four figures, to determine the year, should be enough (i.e. variable V0018 – Year; values: 0000 – 9999). In the case of (quantitative) variable that expresses a number, like length, height, and year and so on, our values will equal the codes. So the value for the year 1088 on the variable: “Year of discovery” is coded: 1088. The value of “Missing” could be coded as “9999” but we have to be careful not to put in a “missing code” that could be a real figure or year. That is why I added one *extra* digit to the “missing” code, i.e. 99999. The values correspond with the year that we want to record. However, there is more. We should be able to fix the year as BC or AD, for the western world. An alternative for that is "BCE" (Before the Common Era). Here is a suggestion:

V019 (Time)	
1	BC
2	AD
9	Missing

V018 (Year)	
....	
1255	1255
1256	1256
....	
1813	1813
....	
99999	Missing

Please note that after a code: BC or BCE the objects get older as the numbers get bigger, but in AD or after the CE the object is dated as older as the numbers get smaller.

RECODING A VARIABLE

Above is a very simple example. Usually the dating will be more complicated. The estimated dating of the fabrication year of an artefact will lie within an *interval* or to put it another way in between two dates. You will have to double the variables above, because the dating could be between for example 0500 BC - 0300 AD. If you make a unique dating for every object and record that in this way, it preserves the most information about the object that is possible. This complies with one of the golden rules of data gathering: *“Collect as much information about a research object as you can with the minutest details.”* However, doing this will make it more difficult to make an analysis afterwards. It requires more skills to make an adequate presentation out of very detailed data. Therefore, I recommend something else. First, preserve as much detail as you can and record that for use afterwards. Second, to make your analysis somewhat easier, make relative broad dating intervals and record them in a simple variable called V0023 – “Dating interval”. In the case of very ancient artefacts, you will not be far off if you do the following:

V023 (Dating Interval)	
1	Older than 4000 BC
2	4000-3001 BC
3	3000-2001 BC
4	2000-1001 BC
5	1000-501 BC
6	500-251 BC
7	250-1 BC
8	1-250 AD
9	251-500 AD
99	Missing

You need only two positions in your database to make nearly a hundred intervals and that should be more than enough. It will make your analysis afterwards much simpler. In addition, in the analysis later on, you can reduce the number of intervals. If you learn the statistical package, with which you will be analysing, somewhat better, you will be able to use your first very detailed variable, better. A so-called “recode” into a new variable, suited to your purpose is a possible option. More about this will be discussed in the analysis section of this book.

In conclusion, do not make too many intervals because that will be difficult to present to a wider audience later on. Egyptologists can make use of the wide accepted and known era’s in the ancient Egyptian history like: Pre-dynastic, Old Kingdom, First Intermediate Period, Middle Kingdom, Second Intermediate Period, New Kingdom, Third Intermediate Period, Late Period, Greco-Roman and Coptic Period. Remember that these are just examples. Make the variable suitable to your wishes and you have just written a line in your codebook to be.

4.2.4. Rules and conventions

We indicate variables by giving them a number and the letter “V” put in front of them, in order to recognise a variable easily.

As a rule, the first variable in the codebook of our research should be the so-called “Case number” or *Casenum* for short. The purpose of this value is to uniquely identify each instance of data that we are observing. It is vital to ensure that the data from one observed object is not confused with another; it refers to all the observations that we have made and want to record. If we are researching statues, a casenum is simply the number of the statue we are studying and gathering data about. Therefore, if we are studying a hundred statues, the values of V001 “Case-number” will be 001 through 100.

Note that in this variable, the value 99 is a *valid value* and does not denote a missing value. Logically there cannot be missing cases. In addition, should there be any missing cases: just add them to the data! There is enough room in this variable to do that because three figures can hold up to 999 cases. (1000 actually, but do not use 000; see below.)

The function of this variable is that if there seems to be something wrong with data recorded within a certain *case*, a statue, pot or another research object, you can refer to the problematic case by its Casenum. Later on, you could make the decision to drop this case from your research because it could have a bad influence over the calculated results. However, you have to give a strong argument for doing so. It is not allowed to drop the case simply because you do not like the results that this case would give in your analysis. You have to have other strong content related arguments to drop a case. Unfortunately, in history too many scientists manipulated their data to get the results they wanted. However, those were not the *real* results. Sometimes the research results are very different to what we expected and we should seek explanation for that. Maybe you did something wrong. Do not ever manipulate your data to get nice results! It is *not* the truth a scientist is looking for.

Commonly, a “0” is not used in both qualitative and quantitative variables, because it denotes nothing and there could be obscurity if you are using this number. As already explained above, there is a standard way in which the common problem: “lack of data” is generally solved. The convention here is that we put a code 9 = “missing” in our variable for every position or column that our variable has. If a variable has only one column, a code “9” means no data available or “missing”. If a variable has two columns, code “99” means no data available or “missing” and so on. Remember that with *real data* like for instance “Age” 99 has the meaning of 99 years old. You could keep up with the standard coding of missing values and add one position or column extra to your variable. The missing value for V09 – “Age” becomes “999”. There cannot be confusion with the real age because living people will not become 999 years old.

Another convention already mentioned is the habit of placing the code before the label in the codebook like this:

V002 (Gender of Statue)	
Value	Label
1	Male
2	Female
9	Missing

4.2.5. Layout of the codebook.

You have to draw up a codebook consisting of all your variables, values and codes neatly in a row. You have to make specifically clear how you indicate a missing value within the different sorts of variables that you are presenting in your codebook.

The format is very simple like the example above.

V <number> Variable name [Let this stand out by giving it a bold type case for instance]

Values/Labels

<code> / <label>

The number of codes should make clear by itself how many positions or columns a variable needs in order to be recorded. If that is not clear because, you have entered “999” as missing for a variable “Real age”, you should indicate explicitly how many positions or columns the variable needs. Take note of the example codebook below.

4.2.6. Example of a simple codebook for ancient statues.

RESEARCH CODEBOOK STATUES

Variables / Values / Value labels:

V001 (Casenum.)
Values:
01
\
/
98

V002 (Gender of Statue)	
Value	Label
1	Male
2	Female
9	Missing

V003 (Position of the body)	
Value	Label
1	Standing
2	Seated
3	Squatting
4	Lying down
5	Other (specify)
9	Missing

V004 (Head position)	
Value	Label
1	Straight ahead
2	Right
3	Left
4	Looking down
5	Looking up
6	Lying down
7	Other (specify)
9	Missing

V005 (Status of statue)	
Value	Label
1	Royal
2	Sacred
3	Elite
4	Commoner
9	Missing

V004 (Construction material statue)	
Value	Label
1	Granite
2	Marble
3	Sand/Limestone
4	Basalt
5	Quartzite
6	Bronze
7	Wood
8	Other (specify)
9	Missing

V007 (Condition of the statue)	
Value	Label
1	Good
2	Average
3	Poor
9	Missing

V008 (Height of the statue in cm)	
Values:	
0001	
\	
/	
9998	
9999 (missing)	

V009 (Original location of the statue)	
Value	Label
1	Temple
2	Palace
3	Burial site
4	Other cult place
9	Missing

V010 (Dating of the statue)	
Value	Label
01	4000 BC or older
02	3000 through 3999 BC
03	2000 through 2999 BC
04	1000 through 1999 BC
05	0500 through 0999 BC
06	0250 through 0499 BC
07	0000 through 0249 BC
08	0000 through 0249 AD
09	0250 AD or younger
99	Missing

Alternative for Ancient Egypt

V010 (Dating of the statue)	
Value	Label
01	Pre-dynastic
02	Old Kingdom
03	First Intermediate Period
04	Middle Kingdom
05	Second Intermediate Period
06	New Kingdom
07	Third Intermediate Period
08	Late Period
09	Greco-Roman
10	Coptic Period
11	Post Coptic Period
99	Missing

Archaeological Alternative

V010 (Dating of the statue)	
Value	Label
01	Palaeolithic
02	Mesolithic
03	Neolithic
04	Copper Age
05	Bronze Age
06	Iron Age
99	Missing

Akkadian Alternative

V010 (Dating of the statue)	
Value	Label
01	Uruk
02	Proto-dynastic
03	Akkadian Empire
04	Guti
05	Sumerian Renaissance
06	Paleo-babylonian Empire
07	Cassite Dynasty
08	(Hittite Empire, Assyrian Empire)
09	Sea Peoples - New Deal
10	Neo-Assyrian Empire
11	Neo-Babylonian Empire
12	Persian Domination
99	Missing

Ancient Israel Alternative

V010 (Dating of the statue)	
Value	Label
1	Late Bronze Age (1550 - 1200 BCE)
2	Iron Age I (1200 - 1000 BCE)
3	Iron Age II (1000 - 586 BCE)
4	Babylonian and Persian Periods (586 - 333 BCE)
5	Hellenistic and Roman Periods (333 BCE - 70 CE)
9	Missing

4.3. Gathering research data.

4.3.1. Introduction

Now that our codebook is finished, we have to collect data on the subject that we will be researching. There are many ways to collect data, depending on the subject that you have chosen. The nice thing is that, since we have invested a lot of time in our codebook, we can build our data collection material on that same codebook. In fact, most computer programs, with which you will make a statistical analysis afterwards, can build a data gathering form for you based on your very own codebook. But then you will have to learn the computer program that codebook first. This step is not necessary and you can make paper forms based on the codebook or data collecting pages in your word-processing program first. You could also make a spread sheet to type in data directly, preparing for analysis at the same time. It is all up to you.

Your choice of gathering data will depend very much on your research subject and particularly how easy or difficult it is to gather your data. Are you actually entering ancient tombs or temples? Than perhaps, it is best to print same data-gathering forms and bring a pen or pencil to make notes. It may be even better to take an office dictation recorder with you and your codebook. Voice in your research data and feed them to your computer later on.

Pen and paper are very versatile and you could make extra notes on paper, to use afterwards or perhaps even modify your codebook. If you own a tablet computer or a smartphone, you could also use this.

4.3.2. Learning the truth

In the data-gathering process, you will learn the truth about your codebook. Is it really up to the task? Can my codebook accommodate for all the objects that I am encountering? Do I have to make extra values and labels? Perhaps even extra variables? It will all come to light when you start *using* your own codebook that you have constructed in advance.

Remember that doing research is an *interactive* practice. There is absolutely no shame in modifying your codebook when confronted with the “crude and brutal” reality that we are trying to research. In fact, this will usually be the case. In this process, you will become a much more experienced researcher and prepared for problems that may lie ahead. This means that you have to be able in the field to adapt your codebook and data-gathering forms. So bring a laptop or tablet computer.

After you have gathered data for a week or so, and made modifications, your codebook should be fine. Most problems are encountered in the first days that you are using your codebook.

Apart from data-gathering forms, you will also have to bring and use other equipment. Mostly tools like measuring tape and perhaps gloves for handling pots or statues. Moreover, everything you need to record your findings in the field.

4.3.3. Recording and coding your findings

During the process of finding and recording your research-data, many errors could be made. You have to be very alert in the recording process and make your data-gathering process as fool proof as can be. If you are using a computer data-entry program, like the ones provided by the manufacturers of statistical software, an error correction is built in. Based on your already entered codebook, the data you type in will be evaluated against the range that you have provided in the program, for your variables and values. However, that will only protect you from elementary mistakes. Typing in a '2' instead of a '1' could potentially change the gender of your research statue and this mistake could go unnoticed. If you are making pictures of your research objects, you could correct these mistakes later on in the process. Perhaps it is best to code for computer later on when you have returned home, based on your paper data recordings.

4.3.4. Example data-gathering form.

V001 – Casenum.

☐☐☐ Remarks _____

V002 Gender of statue

☐ Male ☐ Female ☐ Missing Remarks _____

V003 - Position of the body

☐ Standing ☐ Seated ☐ Squatting ☐ Other ☐ Not certain
Remarks _____

V004 – Head position

☐ Straight ahead ☐ Right ☐ Left shoulder ☐ Looking down
☐ Looking up ☐ Missing Remarks _____

V005 Status of statue

☐ Royal ☐ Sacred ☐ Elite ☐ Commoner ☐ Missing
Remarks _____

V 006 Construction Material of the statue

☐ Granite ☐ Marble ☐ Sandstone ☐ Limestone ☐ Basalt
☐ Quartzite ☐ Bronze ☐ Wood ☐ Missing Remarks _____

V007 Condition of the statue

☐ Good <40% Damage ☐ Average 40-60% Damage ☐ Poor > 60% damage
Remarks _____

V008 Height of the statue in cm. V 009 – Original location of the statue

☐☐☐☐ ☐ Temple ☐ Palace ☐ Burial site ☐ Other cult place
Remarks _____

V010 - Dating

☐ 4000 BC or older ☐ 3000 through 3999 BC ☐ 2000 through 2999 BC
☐ 1000 through 1999 BC ☐ 0500 through 0999 BC ☐ 0250 through 0499 BC
☐ 0000 through 0249 BC ☐ 0250 AD or younger ☐ Missing
Remarks _____

4.4. The data-matrix.

4.4.1. Building the data-matrix

Now that the data gathering process is over, we are going to prepare our data for computer analysis. As you have already noticed in the description of the variables and values, we have assigned columns or positions to our variables. These are necessary to construct our *data-matrix*. A data-matrix simply collates all the data we have gathered in the form of a table. We will place the variables we have worked with in the *columns* of the data matrix and the data we have gathered in the *rows*. Every object we have researched and recorded data from will have exactly one row in the data-matrix. These rows we call: *observations*. The place where a value of a variable will be placed in an observation, we call the *cells* of the matrix.

An empty data-matrix will look like this:

	Column1 V001 Casenum.	Column 2 V002 Gender of statue	Column 3 V003 Position ...
Row 1 Observation 1	Cell	Cell	Cell
Row 2 Observation 2	Cell	Cell	Cell
Row 3 Observation 3	Cell	Cell	Cell
Etc.			

The actual data we have gathered will be placed in the *cells* of the data-matrix. It is not necessary anymore to denote observation 1, observation 2 etc., because that information is already recorded in our first variable: **V001 Casenum.**

Our data-matrix will now look like this, with some fictional data added:

Column1 V001 Casenum.	Column 2 V002 Gender of statue	Column 3 V003 Position of the body
001	1	3
002	2	2
003	9	1

This is what our data-matrix will look like; when we have placed all the data we have recorded in the right spot. A computer program, usually an addition of the statistical package we will be using later on, can do that for us. Furthermore, modern computer programs can use data generated by another program. So it is quite possible to make your data-matrix in a program like Excel and import the data-matrix in a statistical program for analysis.

What we need to know in order to make our data-matrix, is remember the codes we have assigned to our values. If you know these well, and you should, because you have created them yourself, you must be able to read what is inside the data-matrix above. Look at paragraph 4.3 at the example codebook. In the first observation we can see that the gender of the statue is “male” and that the position of the body is “squatting”. In the second observation, the gender of the statue is “female” and the position of the body is “seated”. In the third line, there is something strange. You notice the figure “9” in the matrix. A code 9 in the variable V002 “Gender of statue”, means: “missing” or “no data” or “The gender could not be determined.”

This code is alarming. That is why we use a general code “9” for missing data. You can detect immediately that possibly something has gone wrong. I need not be the case, but you are now forewarned. Check if this code is correctly stated and continue your work. The position of the body is “standing”.

We can simplify the data-matrix by leaving out the labels of the variables. That leaves more space to work with and again, you should be familiar with what the V-numbers represent. Because you have labeled them yourself.

Then our data-matrix will look like this:

V001	V002	V003
01	1	3
02	2	2
03	9	1

4.4.2. Exemplary data-matrices

Example Full Data-Matrix with some fictional data added.

V001 Case- num.	V002 Gender	V003 Position body	V004 Head- position	V005 Status Statue	V006 Constr. Material	V007 Condit. Statue	V008 Height Statue	V009 Original Location	V010 Dating
01	1	3	1	2	4	3	124	1	4
02	2	2	3	3	3	2	93	3	6
03	9	1	2	2	5	3	155	4	8
04	2	4	9	9	8	3	31	4	99
05	1	1	1	1	3	3	1800	3	4

Try to read this data-matrix for yourself. Refer to the exemplary codebook in chapter 4.3. Do you want to keep observation (Casenum.) 04 or discard this? The condition is pretty bad. You should only keep it if it has additional value to the things you are after in your research. (If you are an Egyptologist ...) Have you guessed the identity of the statue in observation (Casenum.) 005? It is very tall: 18 meters in height. This is one of the colossi of Memnon Egypt.

If we leave out the labels of the variables, the matrix looks like this:

Example simplified Data-Matrix with some fictional data added.

V001	V002	V003	V004	V005	V006	V007	V008	V009	V010
01	1	3	1	2	4	3	124	1	4
02	2	2	3	3	3	2	93	3	6
03	9	1	2	2	5	3	155	4	8
04	2	4	9	9	8	3	31	4	99
05	1	1	1	1	3	3	1800	3	4

Example Data 1

Below you will find a data-matrix of 25 ancient Egyptian statues from places and museums all over the world, intended for practice later on in this book.

Full Data-Matrix Statue-Research with fictional data added.

V001 Case- num.	V002 Gender	V003 Position body	V004 Head- position	V005 Status Statue	V006 Constr. Material	V007 Condit. Statue	V008 Height Statue	V009 Original Location	V010 Dating
01	1	3	1	2	4	3	124	1	4
02	2	2	3	3	3	2	93	3	6
03	9	1	2	2	5	3	155	4	8
04	2	4	9	9	8	3	31	4	99
05	1	1	1	1	3	3	1800	3	6
06	1	3	1	3	1	1	31	4	3
07	1	2	1	3	3	1	57	9	3
08	1	1	1	1	8	1	56	3	4
09	2	2	1	1	4	1	195	3	4
10	2	1	1	3	4	1	158	3	4
11	1	1	3	2	3	3	134	4	3
12	1	1	1	2	1	1	75	1	4
13	2	2	1	1	1	1	99	1	4
14	2	1	1	1	4	1	48	1	7
15	1	3	1	3	6	1	74	1	5
16	1	1	1	9	5	2	9999	1	7
17	2	1	1	4	8	2	9999	3	4
18	1	2	1	3	4	1	120	3	3
19	2	2	1	3	4	1	120	3	3
20	1	3	1	3	4	1	22	3	3
21	1	1	1	1	3	2	252	3	3
22	1	3	4	3	1	1	131	9	4
23	2	1	1	1	4	3	28	9	4
24	1	1	1	1	8	1	170	3	4
25	1	2	1	1	4	1	142	3	3

Alternative Full Data-Matrix Statue-Research with fictional data added. *Egyptian Dating*

V001 Case- num.	V002 Gender	V003 Position body	V004 Head- position	V005 Status Statue	V006 Constr. Material	V007 Condit. Statue	V008 Height Statue	V009 Original Location	V010 Dating Egypt.
01	1	3	1	2	4	3	124	1	4
02	2	2	3	3	3	2	93	3	6
03	9	1	2	2	5	3	155	4	8
04	2	4	9	9	8	3	31	4	99
05	1	1	1	1	3	3	1800	3	6
06	1	3	1	3	1	1	31	4	2
07	1	2	1	3	3	1	57	9	2
08	1	1	1	1	8	1	56	3	4
09	2	2	1	1	4	1	195	3	6
10	2	1	1	3	4	1	158	3	4
11	1	1	3	2	3	3	134	4	3
12	1	1	1	2	1	1	75	1	6
13	2	2	1	1	1	1	99	1	6
14	2	1	1	1	4	1	48	1	7
15	1	3	1	3	6	1	74	1	5
16	1	1	1	9	5	2	9999	1	9
17	2	1	1	4	8	2	9999	3	5
18	1	2	1	3	4	1	120	3	2
19	2	2	1	3	4	1	120	3	2
20	1	3	1	3	4	1	22	3	2
21	1	1	1	1	3	2	252	3	2
22	1	3	4	3	1	1	131	9	4
23	2	1	1	1	4	3	28	9	6
24	1	1	1	1	8	1	170	3	4
25	1	2	1	1	4	1	142	3	2

4.5. Preparations for analysis.

4.5.1. Introduction

With our data-matrix established, we can now turn to the data-entry process. The data-matrix must be filled with our research data and coded in the way that our codebook describes. A computer data-entry program, like the ones provided by the manufacturers of statistical software, can fabricate our data-matrix for us, if we have entered our codebook already in the statistical program. Not everyone works like that and it is not necessary to do so. However it is a lot easier and the transference of our research data to the computer goes virtually unnoticed and, most important of all, free from errors. This requires advanced skills with the computer and the statistical program and is not necessary for the analyses. You do not need to have a statistical program in order to perform simple and straightforward analyses.

Most simple analyses, like tabulating frequencies or making cross tabulations, can be done by hand. We can turn our knowledge of a simple calculator or perhaps even spread sheets to our advantage and make some calculations that will enhance our research greatly.

4.5.2. Choosing your means of analysis

Usually your university will have bought licenses to statistical software and there will be courses in working with that software too albeit perhaps at another part of the university. The brand of statistical software is their choice, and you have to master the principles of it in order to work with it and perform some analysis of your research in the end.

However, the use of statistical software is not necessary. You do not even have to be a genius at mathematics. Even if you did not perform well on maths in high school, you need not worry. The maths you have learned in primary school are quite enough to make a few nice frequencies tables and cross tabulations. In addition, this is usually quite enough to fulfil the needs of your research. Actually, if you have learned to use spread sheets like Excel, you should be able to perform many statistical procedures, without the need of a statistical package. The nice thing is that you can make up a frequency count or cross table in your text-editor, copy that in your spread sheet, make your calculations and copy them from the spread sheet into your research rapport. A spread sheet can make the most statistical calculations you need, like percentages, totals, mean, variance and standard deviation from cross tables. You could also make very nice graphical representations of your research data. For the most research calculations, statistical software is not needed.

4.5.3. Preparing data

All you have to do is fill up our data-matrix with the research values that you have found during your fieldwork. Make sure that you work with the right codes and put them in the designated place i.e. the proper column. You have mentioned that column in your codebook.

If you are going to use a statistical package on a computer, be sure that your data-matrix is properly transferred to a file with which the statistical software can work. Most software can import spread sheets like Excel. In addition, if you already have your data-matrix in Excel, you can use that program to make the necessary calculations.

Make sure that the data-matrix you have looks like the “example simplified data-matrix” from paragraph 4.4.2.

5. Basic Analysis

5.1. Simple statistics.

Now that our research data have become available, the next step is to perform analysis upon it. Before we do that, it is necessary to explain a little further on what we are about to do. We are going to be occupied by what is called *descriptive* statistics. This means that we are going to extract useful information from our data that describe our research findings. In the previous chapter, we gathered data in a numerical format, which allows us to make *calculations* on it. As we will be presenting these data to fellow scientists and possibly to a wider audience, we need to do that in an orderly and scientifically accepted way. There are *rules* for that and I will explain these in this chapter below.

5.2. Data cleaning.

Data cleaning is an important process, by which you eliminate data that could potentially influence your research in a negative way. You have to be very careful in this process because there is a fine line between data cleaning and manipulating research results. Every data cleaning action you undertake will have to be motivated. There should be good reasons for you to legitimize your cleaning actions. You should note your actions and add them to your report.

In practice, the reality around us is much more obstinate than one can imagine. I am afraid that this is a fact with which all scientists will have to live with. Many scientists in the past have succumbed to the temptation to “adjust” their data or ignore data that did not fit their theory. There are many reasons why a scientist might do this; perhaps they have built their career based on a particular theory which the data refutes, or there may be financial motivations to arriving at a desired result.

Periodically, these facts come to the surface, and I am afraid that this will never end, but fortunately most researchers are honest and reliable. There is no shame in reporting irregularities in your research and present them to a scientific audience. Perhaps one day a colleague comes up with an explanation, or you will find the reasons later on yourself. As a whole the scientific community learns from these facts.

Before we go any further I have to explain two terms that might play a role in the data cleaning process. If for instance you are studying 25 statues of which 24 are male and you have found only one female statue, this might be an *anomaly*. Webster’s dictionary defines the term as “something that is unusual or unexpected”. You have to study all details of the anomaly very good to see whether it belongs in your research population. If you think it does, you have to prepare yourself to do some explaining. Studying an anomaly might turn out to be very enlightened. Conditions or events may surface that are beyond your imagination. If it does not seem to belong in your research, you should erase it from your research population on the bases of good motivation. Demonstrate to your audience what you did and why.

Another important term is: *outlier*. Suppose we produce a graphical plot of the users of flint stone within a certain population. One of the users is on this very edge of the plot while the rest clusters somewhere in the middle. We call the user at the edge of the plot an *outlier*. It literally lies way out of our normal plot. Handle this as you would do with an anomaly.

To illustrate the data cleaning process we continue with the exemplary database from the previous chapter.

One of the very first steps to undertake is to make a first “dry test-run” with the data, in order to check if there weren’t any *mistakes* made during the data-entry process. The best thing to do first is to go visually through all the columns of your data matrix to see whether of the codes you have put in them seem to be in order. Are there codes that should not be there? You know for instance that the variable V002 ‘Gender’ can only have the values 1 for male and 2 for female and possibly some values 9 for ‘missing’. If you can detect other values in the column of this variable like a 3 or another number that should not be there, check it! Try to find out what caused this error. Perhaps there are some numbers that ended up in the wrong column.

Make corrections for the data that were incorrectly fed into the computer and go over to the next step. Data cleaning is more than correcting data-entry mistakes and could be a rather long process before you get everything right.

The next step is to determine whether our codebook performed well. One of the first things to decide is if our variables *really have variation*. Variables that vary too little may have little or no contribution to our research. Perhaps they will have to be removed or modified, so that they do vary but that is not necessarily so.

You can check even without the use of a computer because we don’t have a large research population. There are only 25 cases. Take a good look at the values that you have entered per variable. Are there values in your codebook that you did not use? Is there one or are there many? If there are many, you might reconsider the buildup of your variable and design new values that are better suited. Take a good look at the “Example Full Data-Matrix Statue-Research with fictional data added” table in paragraph 4.4.2. What do you see when your eyes are browsing the columns? Are there many *different* numbers in the columns?

V002 has only two numbers and one missing case (i.e. a ‘9’). There is not much variation here, but this variable records gender and there are only two of them. If there would only be males or females, code 1 or code 2, than there could be a problem. If you count the code 2 = females, you can see that there are 9. This is not a bad score on 25 statues considering that we have one missing value to; so this variable is ok.

V003 ‘the position of the body’ has many different codes. Value 4 ‘other’ is only used once and 5 ‘not certain’ has 0 entries. These last two values are actually *meant* to be used infrequently, because they have little or no content to add to our research. They provided a spot to put our observations of non-regular statues in our research, and we did not use them very much. So we did well on our codebook here!

V004 ‘head position’ actually has very little variation. Most of the statues got a code 1. They look ‘straight ahead’. Only the code 4 ‘looking down’ is used once and code 5 ‘looking up’ not. We only have 1 missing value. Like V003, this variable seems to be ok. Regardless of variable variation, we do have something to report later on, in our analyses. It appears that most statues in our research population are looking straight ahead!

V005 ‘status of the statue’ has much variation, so that is ok.

V006 ‘construction material’ also has much variation too and seems to be ok.

V007 ‘condition of the statue’ in addition, has enough variation, although 15 appearances of code 1 i.e. ‘good’ on 25 statues, indicates another fact to report later on. About $15/25 = 60\%$ of the statues in our research population, are in good condition. In the analyses later on, you might want to report this and try to explain why this is so. You didn’t go for only the nice ones in your research now, did you?

V008 ‘height of the statue’ has tremendous variation, from 28 cm till 18 meters, so that is ok, but you will have to try to explain these large differences later on.

V009 ‘original location of the statue’ varies quite enough. All the values are there and in enough quantities.

V010 ‘dating’ might pose a problem here, but it is not so apparent. In order to fill the data-matrix, I browsed the internet and added the data of ancient Egyptian statues, to be found in museums all over the world and in the open air in Egypt. However I have used non-Egypt related time intervals. A massive 10 statues out of 25 or 40% now fall into value 4 (‘1000-2000 BC’). This is a very big problem for the Egyptologists among us. They cannot decide which of these statues are from the Middle Kingdom, the Second Intermediate Period or the New Kingdom. I am sorry folks, but there is only one solution. You have to recode this variable with real ancient Egyptian related time intervals. And you should be able to do that. Remember that I have told you in Chapter 4 to preserve as many of the details in your research as you can? Use the “Alternative Full Data-Matrix Statue-Research with fictional data added *Egyptian Dating*” from paragraph 4.4.2 in chapter 4.

5.3. Frequencies and Cross tables.

Now that we have cleaned up our research data, it is time to go over them in order to see what’s inside. The best thing to do, in order to have an overview of all the data we have gathered, is to count all the values. The results of the counts we call *frequencies*. That is because these data tell us how ‘frequent’ a certain value appears in our data. Let us have a look at the “Example Full Data-Matrix Statue-Research with fictional data added” table that I have provided in paragraph 4.4.2 in chapter 4.

We are going to do a little handy work and calculate some values for ourselves. It is by no means necessary for our data analysis to use the computer although it is handy and could save us some time. If we have fed everything into the computer, we can have the statistical package make the calculations for us. But for now it is best to do this a couple of times manually, because you will have to be able to understand how everything works and what the computer and the statistical program does.

5.3.1. Frequencies

Let us start with something simple, like variable V002 ‘gender’. We have coded 1 – male, 2 – female and 9 – missing. Look at the second column of the table and count all the different numbers that represent the values of the variable.

The results are:

V002 Gender	
Code	Frequency
Male	15
Female	9
Missing	1
Total (N)	25

I have counted all the values of this variable and presented them in a table. This is the way that the presentation of your research results should be done. Notice the letter “N” in the table below. This letter always displays the total number of cases you are reporting about. But this table is still incomplete. There is more to be told about the frequencies. We not only want to know how many males and females there were, we want to know how many *males regarding the total research population* there were and how many females. It is best to express this in a *percentage* so that we can compare the numbers as follows. We have a total number of cases N=25 and that equals 100%.

If we want to know the percentage of males, we simply divide the number of males (i.e. 15) by the total number of cases (i.e. 25) and multiply by a hundred percent.

$$\begin{aligned}
 \text{Percentage Male Statues} &= (\text{Number of Male Statues} / \text{Total Number of Statues}) \times 100 \\
 &= (15 / 25) \times 100 \\
 &= 60\%
 \end{aligned}$$

Similarly, the percentage of females is: $9/25 \times 100 = 36\%$.
And finally the percentage missing is: $1/25 \times 100 = 4\%$.

Our new table now looks like this:

V002 Gender		
Code	Frequency	Percentage
Male	15	60%
Female	9	36%
Missing	1	4%
Total (N)	25	100%

This table is better, because now we can compare the percentage of women against the percentage man. Or can we? There is still one missing case and this should not interfere with our calculations. You will have to leave the missing case out and recalculate the percentage in order to arrive at the *valid percentage*.

The total number of valid cases is now $25-1 = 24$
The valid percentage males is $15/24 \times 100 = 62.5\%$
The valid percentage females is $9/24 \times 100 = 37.5\%$

Our new table with valid frequencies looks like this:

V002 Gender			
Code	Frequency	Percentage	Valid Percent.
Male	15	60%	62.5%
Female	9	36%	37.5%
Missing	1	4%	
Total (N)	25	100%	N=24 (100%)

Statistical packages, like SPSS or STATA, always display frequencies like this (though sometimes the N after subtracting the number of missing cases is left out). Now you can build your own frequency tables from the rest of the data. You don't have to make a frequency table for V001 'Casenum'.

However you will probably have a problem with variable V008 'Height' because only two of all the measurements are the same. We have to come up with a solution here. Might I suggest the same as we did in V010 'Dating'. Instead of stating a lot of individual numbers, we made *intervals* in order to reduce the large number of all different and individual data. When you make (new) intervals you are actually *recoding* the variable. This can also be done with the statistical package from your university. But we are going to do this by hand as the thinking that is required in order to do so, *cannot* be done with that statistical package. You have to think it over and make the new intervals yourself.

Let us have a look at the "Example Full Data-Matrix Statue-Research with fictional data added" table in that I have provided in paragraph 4.4.2 in chapter 4 again. Look at the numbers under V008 'Height'. What could be nice intervals? I count 11 numbers under 1 meter, 11 numbers in between 1 and 2 meters and 1 number (way) above 2 meters. Apparently we are having:

- model statues
- lifelike statues
- above lifelike or grotesque statues (as done by a lot pharaoh's of ancient Egypt).

Two height measurements are missing. I suggested that we convert, or actually recode our observations into a new (recoded) variable R008, with the following labels/values:

R008 Height intervals of the statue	
Value	Label
1	Model height
2	Lifelike height
3	Oversized
9	Missing

The conversion will be:

Numbers below 1 meter →	1 - Model height
Numbers in between 1 and 2 meters →	2 - Lifelike height
Numbers above 2 meters →	3 - Oversized
9999 Missing →	9 - Missing

Case num: 2,4,6,7,8,12,13,14,15,20,23 1 - Model height
Case num: 1,3,9,10,11,18,19,22,24,25 2 - Lifelike height
Case num: 5,21 3 - Oversized
Case num: 16,17 9 - Missing

Our new table with valid frequencies now looks like this:

R008 Height intervals of the statue			
Code	Frequency	Percentage	Valid Percent.
Model height	11	44%	47.8%
Lifelike height	10	40%	43.5%
Oversized	2	8%	8.7%
Missing	2	8%	
Total (N)	25	N=25 (100%)	N=23 (100%)

5.3.2. Cross tables

Designing and making cross tables is one of the most powerful tools you have for analyses and presentation of your research findings. A cross table is simply two variables put together. You must let your imagination work and choose those variables for the cross tables you want to make, that might have a relation to each other or perhaps can reveal an *association*. If the association is established statistically, we call this a *correlation*. Relating types of statues, their head position or perhaps even their height to time or ‘dating’ is usually insightful. Maybe you can spot a tendency. Are the statues getting larger over time? Has the place of origin something to do with the head position of the statue, or perhaps with the height of the statue? The possibilities are endless, but *you* should be the one detecting a certain kind of association. But you can’t really do that without some *explanation*.

It is often said that “everything has an association with something else”, but the association has to be meaningful and we must not rely on coincidence. Association can be tested statistically (like we do in the next paragraph), but even that could prove to be wrong. There once was a sociologist who could prove that the number of births in a certain spot statistically coincided with occurrences of stalks. This coincidence was later reasoned away by the introduction of a third variable which controlled the areas the ‘measured’ people lived in (urbane-non urbane). This is called a *spurious correlation*. Be careful what you are doing. You always have to give a reason for your way of analyses.

I was interested to see whether male statues were more frequently found in temples than female statues, and I formulated the following hypothesis: “There is a higher number of male statues in temples, than in other sites.” I have made a cross tabulation between the variables V002 ‘Gender’ and V009 ‘Original location of the statue’. The general rule for the analyses of cross tables is: “make *vertical percentages and compare horizontally*”, or “column percentages are to be compared row by row”.

Without the need of calculating a statistical association, you can see that the percentages of male and female statues are roughly the same, whether they are found at a temple, burial site

or other cult place. The conclusion is clear: no differences between male and female statues in this sample. However, if you have collected a sample on different grounds it might lead to a different outcome.

Although my hypothesis about the male-female statue differences is not wrong, researchers prefer to work with a so called *null hypothesis*. They argue that the differences you might observe in your data, is a result of *chance*. If this is not the case you might be able to show a correlation between variables in the end. I will go into that in the next chapter.

Cross table: Location \ Gender (N=21)

Location\Gender	Male	Female	Total
Temple	4 (30.8%)	2 (25.0%)	6 (28.6%)
Palace	0 (00.0%)	0 (00.0%)	0 (00.0%)
Burial site	7 (53.8%)	5 (62.5%)	12 (57.1%)
Other cult place	2 (15.4 %)	1 (12.5%)	3 (14.3%)
TOTAL	13 (100%)	8 (100%)	21 (100%)

I was then anxious to see whether statues showed up at strange unexpected places, so I have made a cross tabulation of V009 'Original location of the statue' and V005 'Status of statue'.

Cross table: Location \ Statues of Statue (N=20)

Location\Status	Royal	Sacred	Elite	Commoner	Total
Temple	2 (25.0%)	2 (50.0%)	1 (14.3%)	0 (00.0%)	2 (25.0%)
Palace	0 (00.0%)	0 (00.0%)	0 (00.0%)	0 (00.0%)	0 (00.0%)
Burial site	6 (75.0%)	0 (00.0%)	5 (71.4%)	1 (100%)	12 (60.0%)
Other cult place	0 (00.0%)	2 (50.0%)	1 (14.3%)	0 (00.0%)	3 (15.0%)
TOTAL	8 (100%)	4 (100%)	7 (100%)	1 (100%)	20 (100%)

As we can see from this table (but this is only a *created* example), these statues have been found in the locations we would have expected in ancient Egypt. There are not many palaces left and statues have rarely been found there. Burial sites and temples are often the places of origin of statues in ancient Egypt and this is reflected in our cross table. As royalty has a profound influence over temple life in ancient Egypt, their statues (2) were found in temples too, next to statues of gods (2) and a little less statues of high officials or elite (1). As we can see from the row with the value 'burial site' in this cross table, the most statues were discovered here.

Another trend in statues of ancient Egypt is that males are generally depicted larger than females and that 'royals' are the only ones with grotesque statues. To find out whether my sample reflected this trend I made a cross table of the variables V002 'Gender' and the recoded variable R008 'Height intervals'.

Cross table: Gender \ Height intervals of Statue (N=22)

Gender\Height intervals	Model height	Lifelike height	Grotesque	Total
Male	6 (54.5%)	6 (66.7%)	2 (100%)	14 (63.4%)
Female	5 (45.5%)	3 (33.3%)	0 (00.0%)	8 (36.4%)
Total	11 (100%)	9 (100%)	2 (100%)	22 (100%)

Looking at the frequencies in this table, both trends tend to be confirmed.

5.4. *Presentation of data.*

When you are about to present your findings in a master thesis or to your colleague researchers, you must tell your audience exactly and completely honest what it is that you did. Only presenting your research data in the form of frequencies or cross tables is not good enough. That would mean that other significant persons, perhaps those judging your work, are devoid of context, and therefore are not able to consider your research truly on its merits. In your research report you should tell the whole story.

Pay attention to the following points in the list:

- What was the goal of your (data) research?
 - In what way did it fit in your general research?
 - Exactly what did you study or collect?
 - What data did you sample and exactly how did you do that?
 - What were the circumstances of your data gathering process?
 - Did you have problems collecting your data and if yes, what were they?
 - Could the data gathering process be influenced in any way and how?
 - How much research objects you studied are there in existence?
 - Which part of the research population did you actually study?
 - What type of analysis did you undertake?
 - Did you discard data during the data cleaning process? If so, how much and why?
 - Did you mention all details of the data gathering-, cleaning- and analyses process?
-
- Advanced (these topics are covered in the next chapter)
 - Did you make a selection or a sample?
 - If so how big and what margin of error and confidence interval did you choose?
 - What type of statistical analysis did you perform?
 - Were your results statistically significant or not?
 - Did you confirm or refute your hypotheses and why?
 - How statistically strong were your results?

6. Advanced Analysis

In this chapter we will go a little deeper into the matter, and in order for you to follow what is inside, you will need a *basic knowledge of mathematics*. In particular you will need to know the *general rules* of calculation. Most of this is learned in the first years of high school.

6.1. Chance and probability

6.1.1. Inductive statistics

Descriptive statistics paint a picture of the data that you have been gathering. They describe what you have been doing and observed in your research. The next step in statistics is *inductive* statistics. Here a scientist determines what *sample* to take from a *research population* in order to make a meaningful estimate of a fact he or she is researching. Political science, sociology, psychology, econometrics and statistical bureaus are examples of such sciences and users of that kind of statistics. Students of area studies can make use of these kinds of statistics too. For example, if they wanted to do a prediction about a necropolis containing a large number of graves. You usually don't study *all* the graves, although that would be the best thing to do, but you take a representative *sample* instead.

6.1.2. Population and samples

One of the main and important terms to know is the term *population* and *research population*. A population contains *all* the items there are. Remember the example of the 1100 graves in the necropolis I presented earlier on. If we study this necropolis, those 1100 graves are our *population*. One grave in this necropolis is one *member* of that population. Usually we don't study an entire population because there is too much work involved and generally research budgets are limited. So what we need to do is to take a *sample* of the population. This sample contains our *research population*. Take note that the necropolis, which we are studying as a whole, could be in turn the *research population* of all the necropolises in the designated area.

Unfortunately, you can't just take the first 10% or 110 graves in the vicinity of your camp as a sample. If you do that you would be taking a so called: *non-random* sample. You can't do that if you want to make statements about the population as *a whole*. If you want to do the last, then you will have to comply with the *rules of statistics*. One of the main rules is that you take a *randomly* based sample of the entire research population. Only then are you allowed to make statements about the *entire* population however with a certain (chosen) margin of error.

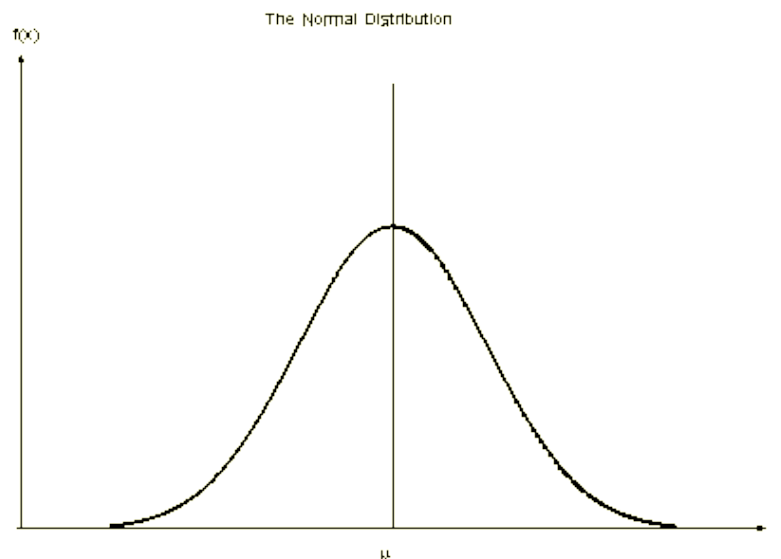
The problem with a *nonrandom* chosen sample is that this sample may not contain all the *variation* there is within the population. Certainly within a population of graves, like a necropolis, this can be a major problem. What we know about a necropolis, that it grows in the course of time because new graves are added continuously. It started off as a spot to dig graves and grows according to a certain pattern. From one or perhaps even more directions, graves are added over time. Perhaps even in different era's graves were added from different directions. So if we take a sample from, let us say, the south side, we could end up having for instance the almost complete newest part of the necropolis or else maybe even the oldest part. We simply don't know and that is why we cannot take just *a* sample.

A wise thing to do in this case, is to *keep the spread of graves constant* by virtually laying a grid over the whole terrain of let's say squares 50 x 50 meters and select randomly one or more graves per square of the grid. We should take a sample in this way because we have already knowledge of how a necropolis comes into being. If we act in this way, we may be able to make estimates of how the necropolis is built up (i.e. what parts contain graves from what period).

If you make yourself acquainted with the rules of sampling and the consequences of altering the 'parameters', then sampling can be a convenient and powerful way of learning about our object of study. Use your creativity and try to gain as much information as possible under the given circumstances.

6.1.3. Chance

The rules of statistics are primarily based on the calculation of chance and most commonly on the so called *normal distribution* or *Gauss curve* (also popularly referred to as a "Bell Curve"). This is a *continuous* probability distribution with a relation to chance and probability and an example of it is shown below. Continuous means that there are an infinite number of numbers, values etc. The line in the graph below shows all *real* interconnected values. This stands (more or less) in contrast to *discrete* values. When we roll a dice we can have only 6 outcomes and there are no values in-between. We can draw a graph of these 6 values but they may not be connected by a line. Although we see this often in reality, it is *not correct* because if there really would be values in-between, we are not certain that they lie on the straight drawn line that is usually drawn in-between.



Let us start with discussing *chance*. What *is* chance? When I flip a coin, what is the chance for heads or tails? Remember the term population. How many sides of a coin are there? Two (not counting the side of the coin). So there are *two sides* in our population. If we flip our coin there is a maximum of two *probable* outcomes: heads or tails. We know that if we flip the coin, only one side will be in sight.

Ok, now what is the chance on only heads or tails? This is very simple and basic mathematics. There are two possible outcomes and only one of the two can show up, upon tossing the coin.

One outcome divided by two possibilities = $\frac{1}{2}$. In the previous chapter, the so-called *percentage* was discussed. A per - cent - age. The word says it all, 1 part of hundred is one percent. Nowadays chances are expressed in percentages. So instead of our calculated chance of $\frac{1}{2}$ that either heads or tails of a flipped coin turn up we say: the chance is $\frac{1}{2} \times 100\% = 50\%$.

Now let us make it a little more difficult. Suppose we have a mixed collection of potsherds. 25 potsherds of a large colored and oven fired pot and 25 potsherd of a plain and large sun baked pot. They are all together in one large wooden box. What is the chance of me grabbing a colored potsherd, if I close my eyes and take one out of the box? [We call this an *a-select draw*.] Let us see. Since the potsherds are all together in one box, the population of the *box* is: $25 + 25 = 50$ potsherds. But actually, we have only two populations here: the large colored oven fired potsherds of 25 and the large plain sun baked potsherds also of 25. So there are really only two possibilities.

I grab either a plain potsherd or a colored one. The maths: 25 colored potsherds divided by the total number of potsherds $25 / 50 = \frac{1}{2}$. $\frac{1}{2} \times 100\% = 50\%$.

In *theory* we will still have the same chance on a colored potsherd as we have on heads of a coin. However in *practice*, things can turn out pretty different. It is possible to get three tails before you get one head if you continue to toss the coin. Only when you flip the coin a large number of times, say a thousand times, than we will be more confident that there appear to be really 50% heads and 50% tails. The same goes for the potsherds, provided you put the one drawn potsherd back in the box, every time that you took one out. Otherwise the populations could become uneven. However this is only the case if your number of potsherds is relatively small (say below 50).

6.1.4. Probability

In the paragraph about chance, we described *known* populations and the chance we have if we made an a-select draw within them. We described *chance* as a *percentage*. Suppose we turn everything around. Now, we do *not* know the build-up of our population, but we perform quite a number of a-select draws. We put the drawn object back every time so as not to alter our unknown population and make an inventory of the outcomes. At another dig we discovered a big box with potsherds. We made a 100 a-select draws from it. Every time we put the drawn item back in the box and we didn't look at the box when we made our draws.

The drawing results are:

Number	Type
19	Black potsherds
31	Yellow clay potsherds
38	Colored Potsherds
12	Red clay Potsherds

What can we learn from that? First of all, there appears to be only *four types* of potsherds. We have made a hundred draws, so if there were still other types in the box we would at least have drawn one of them. That is not the case, so our confidence is heightened that our hitherto unknown population has become a little less unknown and contain four types of potsherds.

What else do we know? In terms of *chance* we know what chance we have to draw a certain type of potsherd, namely:

Number	Type
19/100 = 19%	Black potsherds
31/100 = 31%	Yellow clay potsherds
38/100 = 38%	Colored Potsherds
12/ 100 = 12%	Red clay Potsherds

But that is *not* entirely true because we don't know this population well enough to establish the findings of our draws and trust them for real. We could be unlucky and still need 500 or more draws to have certainty about the numbers we have found. So that is not what we will be reporting about our finds. We will report in terms of *probability*, because we don't know for sure. This term is expressed differently. We express probability, indicated by the letter P, of finding a certain kind of potsherd in our box, as a part of the number 1.

Based on our sample population, we may estimate the proportions of coloured potsherds as:

Type	P.
Black potsherds	.19
Yellow clay potsherds	.31
Colored Potsherds	.38
Red clay Potsherds	.12
Total	1.00

6.1.5. Frequencies

The 100 draws that we made above can also be assumed to be a *sample* from a *survey*. Each item has a probability to be drawn equal to the figure in the table. If we had more boxes with more potsherds from exactly the same archaeological dig on the very same spot, we could multiply our calculated probability with the total number of potsherds and predict how many of each type of potsherd there will be. Let us say that we have 2000 potsherds. Or *prediction* is:

Type	P.	Number
Black potsherds	.19 x 2000	380
Yellow clay potsherds	.31 x 2000 = 620	620
Colored Potsherds	.38 x 2000 = 760	760
Red clay Potsherds	.12 x 2000 = 240	240
Total	1.00	2000

So because we know the probability, we can *assume* a table of frequencies that will be equal to the probability of being found in the sample. Reasoning like that, probability, as described above, and (a distribution of) *frequencies*, are about the same. In both cases, the object is to report on the findings in your research. Frequencies are simply the *occurrences of items* you have measured in your research. In this case the “measurement” consisted of a 100 a-select draws of potsherds in a box. The total number of observations will be indicated by the letter N. We express frequencies, indicated by the letter F, in plain numbers, like in the table below:

Type	F.
Black potsherds	380
Yellow clay potsherds	620
Colored Potsherds	760
Red clay Potsherds	240
Total	N=2000

One thing could arise in the example that I gave you about the box with potsherds, but that is practical and not mathematical. Remember I told you that we found no *other* type of potsherd in the box we have used to make draws from. Well, if you have twenty or more boxes, there could be a real chance that you do find another type in it. But in that case we were not allowed to make draws from just one box and make prediction from that box only. We should have made equal draws from all boxes, according to the rules of statistics. Always remember the rules; there is actually no real problem here.

6.2. Mean

In this paragraph we discuss some more statistical measures that we can use in order to display our research results. A very common and well known measure is the arithmetic *mean*. Suppose we have researched say ten tombs and have counted the number of grave goods that we have found inside. The results are:

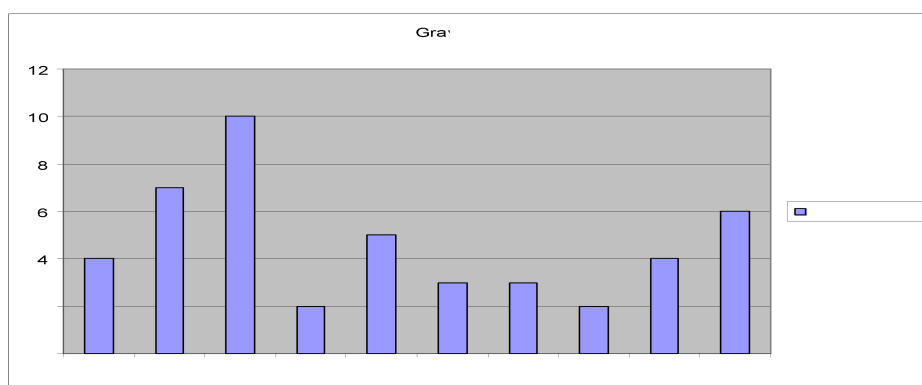
Table of tomb findings	
Tomb	Nr. of Grave goods
01	4
02	7
03	10
04	2
05	5
06	3
07	3
08	2
09	4
10	6
Total X	46

When we want to display the results of our tomb research, we might want to calculate how many items we have found on *average* in the tombs. The way of doing that is by adding up all the items found in every tomb and dividing that by the number of tombs. As a rule, we indicate the total number of objects we have found, with the letter X.

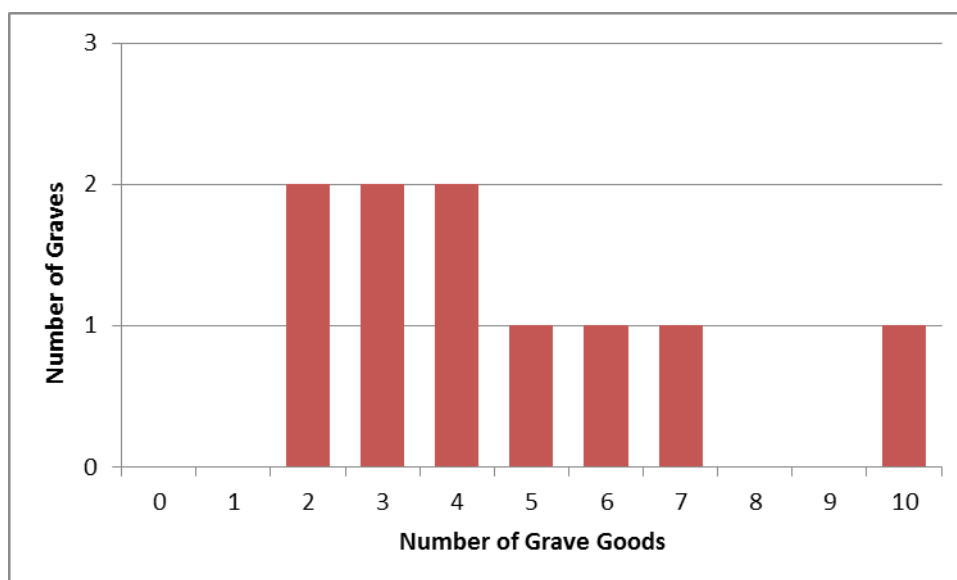
First, add up all the items: $4+7+10+2+5+3+3+2+4+6=46$; then divide this number by the total number of tombs ($=10$. $46/10=4.6$). The outcome of this sum, and therefore the total number of objects found in the 10 tombs, is 4.6.

To indicate that we are mentioning the mean value, the convention is that we put a line above the X that denotes the total number of values we have, like this: $\bar{x} = 4.6$.

But what does it really tell us, this *mean* of 4.6? For one thing, there is no real object of 0.6 because we are talking about *calculated* averages or mean. But if I am the one you are communicating your research results to, and you have told me that the average number of objects found in the tombs is 4.6, what does that signify? Am I to believe that every tomb has nearly 5 objects in them? Because that is what most people are inclined to believe. The best way of judging that number is by looking at how the objects are *spread* among the tombs. If you look at the graphics below, which displays the finds per individual tomb, you will notice that *the spread* of grave goods is *very uneven*! We have found 10 objects in tomb nr. 3 and only 2 at tomb nr. 4 and 8. Only in tomb 5 we find a number of 5 grave goods; more or less our calculated average of 4.6. This warns us to be *very careful* with the calculated means or $\bar{x}$.



Rather than looking at the individual members of sample population, you may look at the data in a more aggregated form:



Looking at the amounts of grave goods found per tomb, we can see that in almost every case our calculated $\bar{x}$ of 4.6 deviates significantly from the *real* values found in every grave. The *variation* is enormous! So is there a way for us to *judge* the spread of grave goods, or whatever we are measuring, among our research objects? Yes there is; the *standard deviation*!

Definition: Standard Deviation

“Standard deviation is a statistical measure of spread or variability. The standard deviation is the root mean square (RMS) deviation of the values from their arithmetic mean.”³⁹

Let us make that easier. What we want to know is *how much* our calculated mean or $\bar{x}$ differs from all the values we have measured. We return now to our example, the table of grave goods found in 10 tombs. What we are going to do first is to calculate the difference between every observation we have made i.e. the number of grave goods per tomb and $\bar{x}$. In order to do that we subtract the $\bar{x}$ from every observation like this:

Table 1 of differences observations and $\bar{x}$

Tomb	Grave goods	$\bar{x}$	Difference
1	4	4.6	- 0.6
2	7	4.6	2.4
3	10	4.6	5.4
4	2	4.6	- 2.6
5	5	4.6	0.4
6	3	4.6	- 1.6
7	3	4.6	- 1.6
8	2	4.6	- 2.6
9	4	4.6	- 0.6
10	6	4.6	1.4

As you can see from this table, many numbers are *negative* because of mathematical reasons. Since we are only interested in the *amount* of difference and not if the number is negative or not, we apply a second rule of arithmetic. We *multiply* the difference *by itself* (i.e. square²) the calculated values, thus making every value *positive*. We then add up all the squared differences.

Table 2 of differences observations and $\bar{x}$ squared

Tomb	Grave goods	$\bar{x}$	Difference	Difference ²
1	4	4.6	- 0.6	0.36
2	7	4.6	2.4	5.76
3	10	4.6	5.4	29.16
4	2	4.6	- 2.6	6.76
5	5	4.6	0.4	0.16
6	3	4.6	- 1.6	2.56
7	3	4.6	- 1.6	2.56
8	2	4.6	- 2.6	6.76
9	4	4.6	- 0.6	0.36
10	6	4.6	1.4	1.96
Total	46			56.40

Because we have squared all the differences we have to *make up* for the ‘enlargement’ of numbers that goes with that. We do that by taking *the root* of the number, this is exactly the *opposite* of what we did earlier on, namely taking the square of the number or multiplying by itself. However, since we are interested in the differences with the calculated means $\bar{x}$, we

³⁹ This definition comes from the website: <http://easycalculation.com/statistics/standard-deviation.php>.

first divide the total number of squared differences by the total number of observations. In this way we can see how much *variation* we have. Our Standard Deviation or SD thus becomes:

$$\text{Square Root of } \frac{\text{Sum of squared differences}}{N \text{ or number of observations } \textit{minus} 1}$$

“Some books show division by "N". However, when calculating the Standard Deviation of small sample, a better estimate (of the parent group) is obtained by dividing by (N-1) instead of dividing by "N". For large "n", the difference between using "N" or "N-1" is small.”

In either case, we can see that a larger values of N (i.e. larger sample sizes) result in a smaller standard deviation value. In other words, the more observations we make, the more confidence we have in our estimated value.

To continue our example:

$$\text{Square Root of } \frac{\text{Sum of squared differences}}{N \text{ or number of observations } \textit{minus} 1} = \text{Square Root of } \frac{56.40}{10-1} =$$

$$\text{Square Root of } 6.27 = \mathbf{2.5}$$

Ok, so now we know that the standard deviation calculated from our example is 2.5. What does that value mean? Unfortunately, the standard deviation is not a standardized measure. This means that there are no permanent values against which we can evaluate our calculated standard deviance. The number that we have found has a direct relation to the total number of cases or tombs in our example that we have measured. It is therefore a *relative* number. “A low standard deviation indicates that the data points tend to be very close to the mean; high standard deviation indicates that the data points are spread out over a large range of values.”⁴⁰ To give you an idea of the spread of our example, I have calculated three extra SD's on the same example of grave goods found in tombs but each with a very different *distribution*.

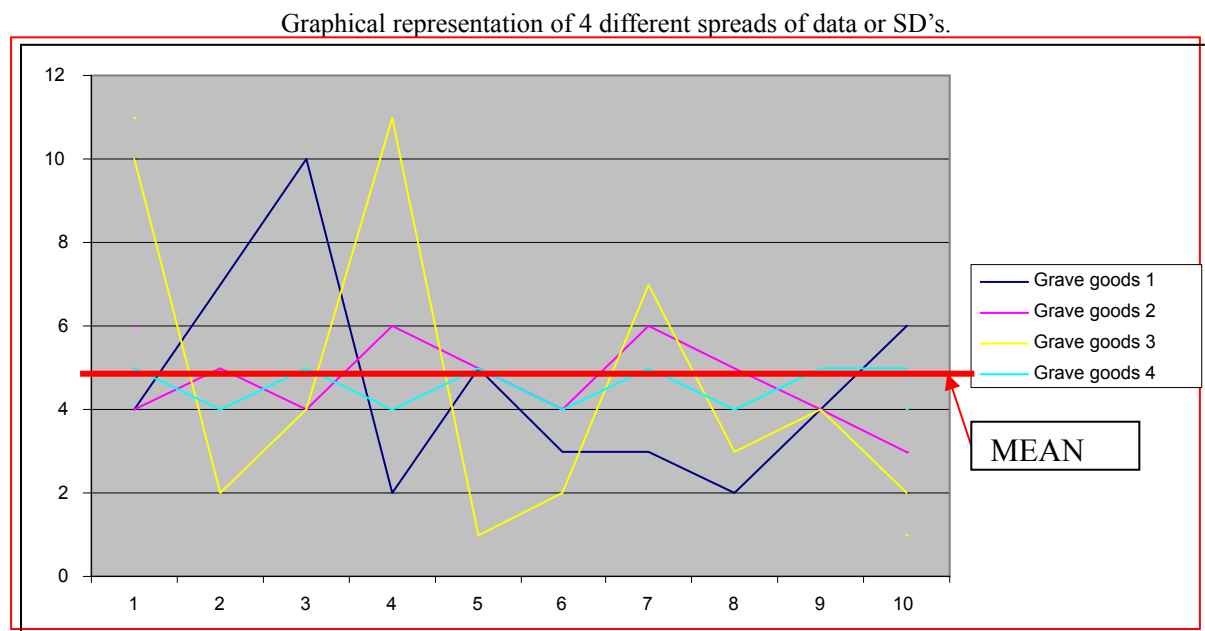
Table 3 Different distributions of grave goods, all with a $\bar{x}$ of 4.6

Tomb	Grave goods 1	Grave goods 2	Grave goods 3	Grave goods 4
1	4	4	10	5
2	7	5	2	4
3	10	4	4	5
4	2	6	11	4
5	5	5	1	5
6	3	4	2	4
7	3	6	7	5
8	2	5	3	4
9	4	4	4	5
10	6	3	2	5
SD	2.5	0.8	3.6	0.5

⁴⁰ This rule comes from Wikipedia: http://en.wikipedia.org/wiki/Standard_deviation

In the table above, our original example data is displayed in the first column under “Grave goods 1” We calculated the SD to be 2.5 and as you can see from the graphical presentation below, the spread is considerable. Next column, Grave goods 2, has a calculated SD of 0.8. According to the rule, given above: “A low standard deviation indicates that the data points tend to be very close to the mean; high standard deviation indicates that the data points are spread out over a large range of values.”

We can see from the graph below that the spread is small. The third column, with a SD of 3.6 has a wild fluctuation of values, as you can see in the graphical representation. The last fourth column has the least SD and has minimal variation, as displayed below. If there was no difference at all between the tombs and all would contain a theoretical 4.6 grave goods, which equals $\bar{x}$, the SD would be 0.0 or no spread at all.



Using Excel

I have calculated the mean and SD with the aid of my standard spreadsheet. If you are using Excel, first click on a cell that you want to contain your SD. Then click on the fx , just below the toolbar at the top of your screen. Select AVERAGE or STDEV from the submenu that appears and click OK. A new sub screen appears. Above in that screen the program suggests for you what cell to calculate. Mostly the suggested cells to use are *not* ok, so correct those by typing in the cells you want to use for the calculation, for example B2:B11. This means use cell B2 until B11. Press Ok and see some screens are flashing. The calculated value now appears in the cell you selected.

Table: Average

	A	B	C
1			
2		Tomb	Grave goods 1
3		1	4
4		2	7
5		3	10
6		4	2
7		5	5
8		6	3
9		7	3
10		8	2
11		9	4
12		10	6
13		Mean	=AVERAGE(C3:C12)

Variance

A closely related measure is *variance*. The definition of variance is:

“The square of the standard deviation. A measure of the degree of spread among a set of values; a measure of the tendency of individual values to vary from the mean value.”⁴¹

The calculation is exactly the same as a standard deviation. The only difference is that we do *not* take the square root of the fraction:

$$\text{Variance thus becomes: } \frac{\text{Sum of squared differences}}{\text{n or number of observations } \textit{minus} 1}$$

We could also say that variance is defined as: The average of the squared differences from the Mean. “Variance tells us the dispersion (i.e. Is it a fat or skinny bell shape?) of our Normal Distribution (bell curve)”⁴². To understand distribution, continue with the next paragraph of this book.

To continue our example above:

$$\text{Variance} = \frac{\text{Sum of squared differences}}{\text{N or number of observations } \textit{minus} 1} = \frac{56.40}{10-1} = \frac{56.40}{9} = \mathbf{6.27}$$

⁴¹ This definition comes from the website: <http://easycalculation.com/statistics/standard-deviation.php>.

⁴² Ibid.

6.3. Measuring simple association

There are a number of statistical test to apply to your research findings, to see whether a connection or *correlation* between variables you detected, is also *statistically significant*. A small variation in your numbers does not mean that there is a *real* statistical association. When working with cross tables one of the most common test, to see if there is a certain kind of association, is the Pearson's Chi-square test or χ^2 test. This test assumes that all counts in a cross table are equally distributed and can be used in any kind of measurement level. So if you are measuring Gender, for example, you are *expected* to find as many females as there are males. So in a research population of $N=22$ cases, 50% of the cases or 11 should be female and 50% or 11 should be male. Any *deviation* from this distribution of cell counts could denote a statistical association and this might be shown in the cell counts that we have *observed* in our research. But only by calculation we can find out if this is really the case.

The χ^2 test is particularly suited to indicate whether a *hypothesis* that we have is true or false.

χ^2 has the following form:

$$\chi^2 = \text{the sum of } \frac{(\text{Observed frequency} - \text{Expected frequency})^2}{\text{Expected frequency}}$$

This means that for *every observation* we have made, we should add the difference in (observed and expected frequencies) ² and divide that by the expected frequency.

We now need to calculate the so called *degree of freedom*.

$$\text{Degrees of Freedom (Df)} = \text{Number of columns being tested} - 1$$

Therefore in a table with 3 columns, the degree of freedom or $Df = 3 - 1 = 2$. You need a table with values with which to determine whether your hypothesis is confirmed or not.

Referring to the table at the end of this paragraph, the χ^2 in our example has to have a value of 6.27, which is over 5.99, with an error margin (P) of 5%. This means that we are 95% confident that our hypothesis is correct. However, our χ^2 value does not exceed the 1% error margin value of 6.64, meaning that we cannot state that we are 99% confident that our hypothesis is correct.

This value of 5% is the standard margin of error of most statistical research. We then should call our findings statistically *significant at the 95% confidence level*.

Example 1

Let us have a look at a table of statues from the previous chapter.

Location\Gender	Male	Female	Total
Temple	4 (30.8%)	2 (25.0%)	6 (28.6%)
Palace	0 (00.0%)	0 (00.0%)	0 (00.0%)
Burial site	7 (53.8%)	5 (62.5%)	12 (57.1%)
Other cult place	2 (15.4 %)	1 (12.5%)	3 (14.3%)
TOTAL	13 (100%)	8 (100%)	21 (100%)

We want to know whether there is a difference between male and female statues regarding the places where they have been found. Our hypothesis is, that there are *no differences* and that the differences that we have found, can be attributed to *chance*. This kind of hypothesis we call a *null-hypothesis*. Let us start first in making the cross table a little more sensible by removing the value 'Palace' of V009 'Original location of the statue'. This value was not used and can safely be removed. Our table now looks like this:

Location\Gender	Male	Female	Total
Temple	4 (30.8%)	2 (25.0%)	6 (28.6%)
Burial site	7 (53.8%)	5 (62.5%)	12 (57.1%)
Other cult place	2 (15.4 %)	1 (12.5%)	3 (14.3%)
TOTAL	13 (100%)	8 (100%)	21 (100%)

The next thing that we are going to do is to leave out the percentages and replace them by the values that we expected to find, in order to make a Chi² calculation. Remember that we are expecting an *equal* number of males and females. A total number of 6 statues in the temple will be evenly distributed between male and female statues i.e. 3 and 3 respectively. The 12 statues found in burial sites will be distributed accordingly in 6 male and 6 female statues. There are merely 3 statues found in other cult places and we can't physically split them in two. But since we are only making *calculations*, we can. We attribute 1 ½ statues to the male and 1 ½ statues to the female side. Our table now looks like this:

Location\Gender	Male	Female	Total
Temple	4 (exp. 3)	2 (exp. 3)	6
Burial site	7 (exp. 6)	5 (exp. 6)	12
Other cult place	2 (exp. 1 ½)	1 (exp. 1 ½)	3
TOTAL	12 (exp 10½)	8 (exp 10½)	21

Now we are going to calculate the Chi². As we have seen before, the way to do that is:

$$\text{Chi}^2 = \text{the sum of } \frac{(\text{Observed frequency} - \text{Expected frequency})^2}{\text{Expected frequency}}$$

This means that for *every observation* we have made, we should add the difference in (observed and expected frequencies) ² and divide that by the expected frequency. Let us start with the male column and calculate the first cell. We have observed 4 and expect to find 3. As we can see above we have to subtract the expected frequency from the observed frequency (4-3) than we have to take the square of that number (4-3)² and divide by 3. Our first number is:

$$\frac{(4-3)^2}{3} = \frac{1}{3} = 0.33$$

The next cell from the male side is 'the Burial site' value. We have observed 7 and expect to find 6. We have to subtract the expected frequency from the observed frequency (7-6) than we have to take the square of that number (7-6)² and divide by 6. Our second number is:

$$\frac{(7-6)^2}{6} = \frac{1}{6} = 0.17$$

The next cell from the male side is 'the other cult place' value. We have observed 2 and expect to find 1½. We have to subtract the expected frequency from the observed frequency (1½-2) than we have to take the square of that number (1½-2)² and divide by 1½. Our third number is:

$$\frac{(1\frac{1}{2}-2)^2}{1\frac{1}{2}} = \frac{(\frac{1}{2})^2}{1\frac{1}{2}} = \frac{\frac{1}{4}}{1\frac{1}{2}} = \frac{1}{6} = 0.17$$

Let us now start with the female column and calculate the first cell. We have observed 2 and expect to find 3. We have to subtract the expected frequency from the observed frequency (2-3) than we have to take the square of that number (2-3)² and divide by 3. Our first number is:

$$\frac{(2-3)^2}{3} = \frac{1}{3} = 0.33$$

The next cell from the female side is 'the Burial site' value. We have observed 5 and expect to find 6. We have to subtract the expected frequency from the observed frequency (5-6) than we have to take the square of that number (5-6)² and divide by 6. Our second number is:

$$\frac{(5-6)^2}{6} = \frac{1}{6} = 0.17$$

The next cell from the female side is 'the other cult place' value. We have observed 1 and expect to find 1½. We have to subtract the expected frequency from the observed frequency (1-1½) than we have to take the square of that number (1-1½)² and divide by 1½. Our third number is:

$$\frac{(1-1\frac{1}{2})^2 (-\frac{1}{2})^2}{1\frac{1}{2} \cdot 1\frac{1}{2}} = \frac{\frac{1}{4}}{1\frac{1}{2} \cdot 1\frac{1}{2}} = \frac{1}{6} = 0.17$$

Ch² is the sum of all the values we have calculated.

$$\text{Ch}^2 = 0.33 + 0.17 + 0.17 + 0.33 + 0.17 + 0.17 = \mathbf{1.34}$$

We know that our degree of freedom or Df.=1. If you look at the table below, the number in the table to meet, that would give us within a 5 % margin of error = 3.84. Our calculated Ch² of 1.34 is well below that figure. That means that there is *no statistical significance*. In other words: the differences that we have found between male and female statues *can* be attributed to chance. Our null hypothesis, that there are *no differences* and that the differences that we have found, can be attributed to *chance*, is hereby confirmed! We can say that with a confidence level of 95%.

Table of Chi2 statistics

Df.	P=0.05	P=0.01	P = 0.001
1	3.84	6.64	10.83
2	5.99	9.21	13.82
3	7.82	11.35	16.27
4	9.49	13.28	18.47
5	11.07	15.09	20.52
6	12.59	16.81	22.46
7	14.07	18.48	24.32
8	15.51	20.09	26.13
9	16.92	21.67	27.88
10	18.31	23.21	29.59

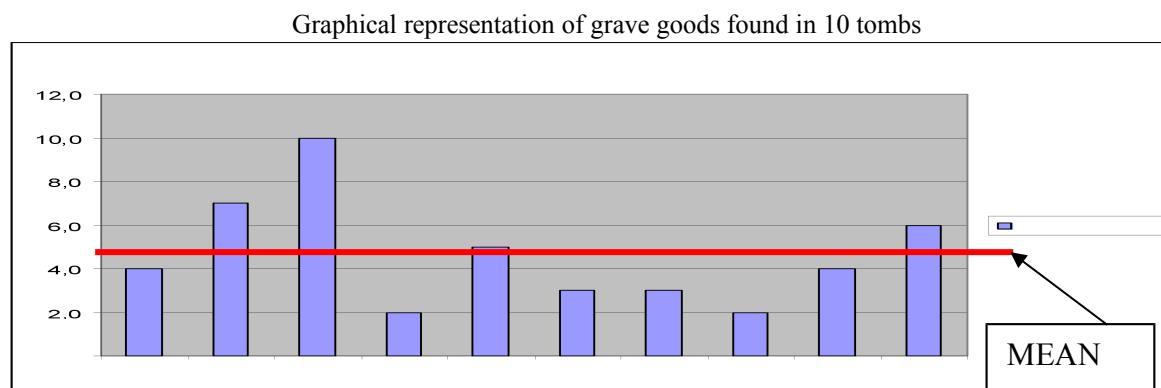
6.4. Distribution of values

One of the most important things that we are doing as a researcher is to collect *data*. Usually values of variables in our research. Above I have described some methods of gaining insight in data we have collected or given as an example. In this paragraph we will explore more of those. To our aid come a number of methods to make the data, we have collected, *visible*. A *graphical plot* of collected data, like the one I presented above, can prove to be very insightful because you can actually *see* in a graph what our data look like. Of course you cannot see the data physically; we only *represent* them with the aid of a graph. What we are doing when we display data graphically, is to show how the data are *distributed* on a scale or in two- or even three dimensional space. We make a distinction between the graphical representation of actual research, discussed in the first sub-paragraph and *special distributions* of data, discussed in the next sub-paragraph.

6.4.1. Distribution of research data

When you prepare your research data for presentation, a lot of computer programs can make graphical representations for you, like spreadsheets and of course, statistical programs. However usually you don't know the way that these programs operate internally, that make these representations for you. For one thing, there are *rules* for statistically correct representation of data. These rules are designed to regulate what you can and can't do in order to be as factual as possible. Take a look at the "Graphical representation of 4 different spreads of data or SD's" in the last paragraph. I have made this graph with the aid of a spreadsheet on the basis of 4 schemes of 10 data. To make things more visible, the computer program *draws lines* in between the points that represent our measured data. *Graphically* this helps us in distinguishing the 4 different plots, but statistically, this is very *incorrect*! There is no way to know whether these drawn lines between the points, correctly display values that lie *in between* the measured values. We call this *unjust interpolation*. Interpolation is a way of making new data points within the range of known data points. So we can only validate these lines as *help-lines*. Only the points, that they are connecting, are displayed *statistically valid*.

A statistically correct way of displaying individual research data, regardless of measurement level, is the bar diagram, like the one below.



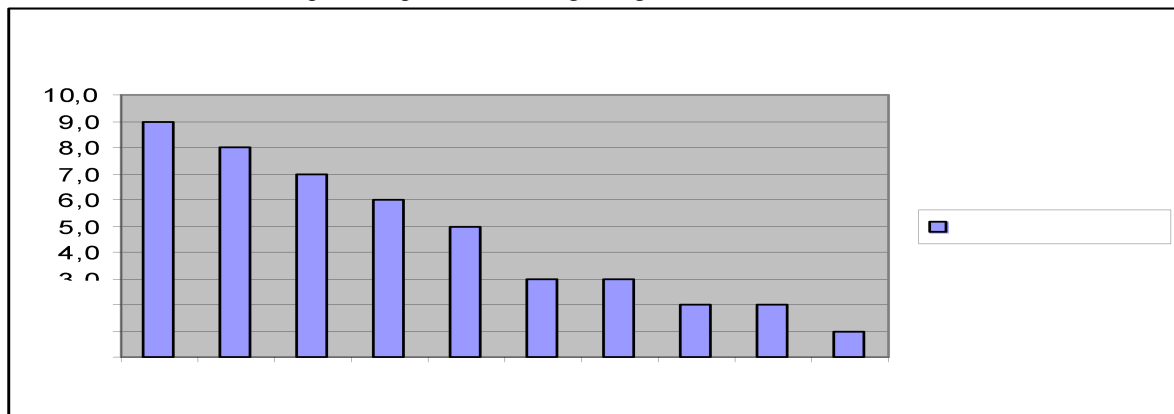
Looking at this graphical representation of our first example of grave goods found in 10 tombs, I can explain some more statistical terms you need to know.

First of all you notice that the *spread* of grave goods found per tomb is very *uneven*. No wonder we calculated a relatively high standard deviation of 2.5. A high standard deviation means that there are large differences in the values regarding the means.

As you can see in the graph, there is one grave with a large number of grave goods found inside, namely 10. Because of the large difference with the others values we have found, we call this number of 10, relative to the other figures, a statistic *outlier*. It is called that way because this point lies far away from the other points in the graph. Interpretation wise, such an outlier could be important. Possibly this was the grave of a noble person, while the other graves belonged to commoners. You should check this. But there are many more possibilities. It could also be that this grave was relatively undisturbed and therefore more grave goods remained inside. Important is, that you take action when you discover such an outlier and try to explain its existence. Yet another possibility is: a research error, meaning the value is incorrect and should be removed from the analysis.

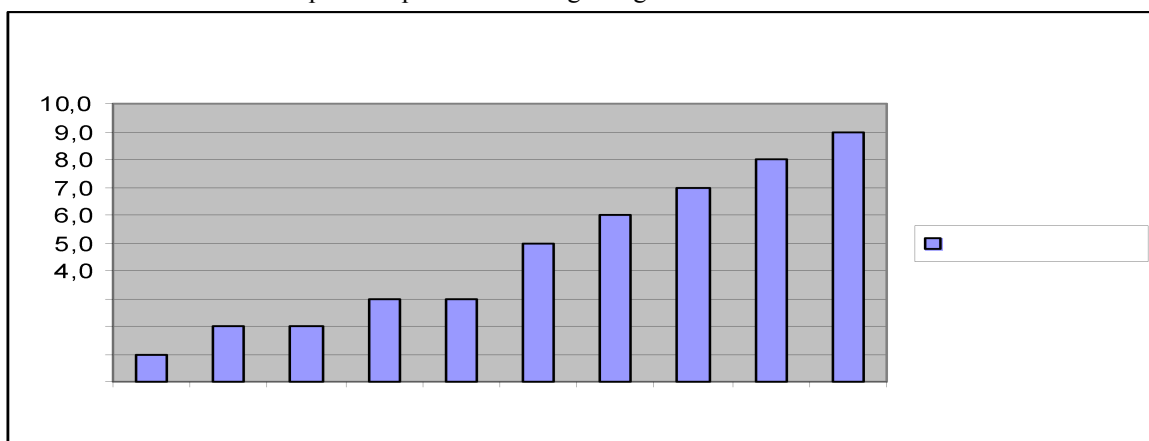
The data plot in our example above is rather scattered or jumbled up. There are more kinds of data plots that show a *tendency*.

Graphical representation of grave goods found in 10 tombs: A



As you can see in the table of Grave goods A, the distribution is *uneven*; *sloping* from left to right. And in the next table of Grave goods B, it is just the other way around. Remember the term *sloping distribution*. It is widely used to indicate uneven distribution. These distributions stand in contrast to the normal distribution or bell-curve, we discuss in the next paragraph.

Graphical representation of grave goods found in 10 tombs: B



I have explained the arithmetic mean in this book and there is another measure with a relation to the middle in statistics. This we call the *median*. Whereas the mean arises as a result of our calculations, the median is a number that is *defined* by mathematicians. The median lies exactly in the middle of our *ordered values*. On either side of the median lie approximately 50% of the observations (i.e. measured values). Like the mean, this is a *calculated* value and need not to exist as an actual value. A median has importance when looking at a graphical distribution of values. This will become clear in the next section. There are more measures like the median. The *quartile* is a likewise measure denoting 25% of the values or a *population*. The *interquartile range* is an alternative measure to calculate spread. The same goes for the *decile* and the *percentile*, covering respectively 10% and 1% of the values.

The *interquartile range* or IQR, belongs in the same range of measures of variability like the *variance* and the *standard deviation*. The IQR divides the population in four equal parts. They function as a sort of milestones and we are *not* interested in the ends but only at the “milestones” in between. These are called Q1, Q2 and Q3 like this:

End | - - - - Q1 - - - - Q2 - - - - Q3 - - - - | End.
 Amount of data: 25% 25% 25% 25%

As you can see Q2 is equal to the *median* which denotes exactly *half* the population. The median splits, the population into two equal segments or 50% of the population.

Definition

The interquartile range IQR is Q3 minus Q1.

Remember the grave goods from table 3 above? Let us calculate with the first column again as an example.

Table: Grave goods from tombs

Tomb	Grave goods 1
1	4
2	7
3	10
4	2
5	5
6	3
7	3
8	2
9	4
10	6

First you have to *order the data points* from least to greatest.

Determine the position of the first quartile using the following formula: $(N+1)/4$ where N is the number of points in the data set. If the first quartile falls between two numbers, take the average.

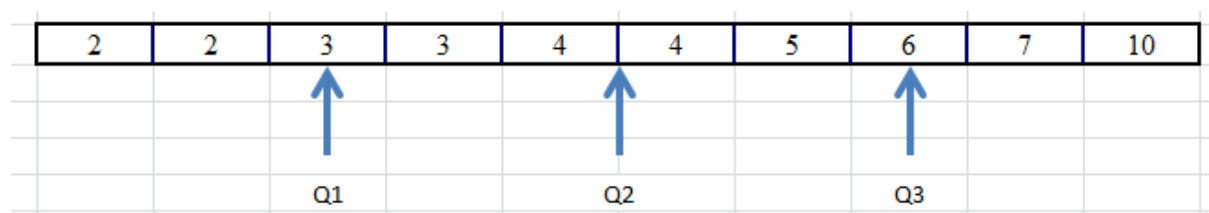
Determine the position of the third quartile using the following formula: $3*(N+1)/4$ where N is the number of points in the data set. If the third quartile falls between two numbers take the average.

Tomb	Grave goods 1
1	2
2	2
3	3
4	3
5	4
6	4
7	5
8	6
9	7
10	10

We have defined the interquartile range as: Q3 minus Q1. Q1 represents the middle of the first half of observations:

End | - - - - Q1 - - - - Q2 - - - - Q3 - - - - | End.

But there is a little bit of a problem here, because we have an *even* number of observations. We hence cannot use the number we want *directly*, because it lies *in-between* two numbers. Therefore, we have to *calculate* Q1 first. Actually, that is not very difficult. We simply take the number *before* our desired value, i.e. 2 and *after* the desired value, i.e. 3, *add* them up and *divide* by 2 because we want the middle of those two values. Q1 is then $2 + 3 = 5$ divided by 2 = **2.5**. We have to calculate Q3 likewise. Again we simply take the number *before* our desired value, i.e. 5 and *after* the desired value, i.e. 6, *add* them up and *divide* by 2 because we want the middle of those two values. Q3 is then $5 + 6 = 11$ divided by 2 = **5.5**. We have defined the interquartile range as Q3 minus Q1, so $IQR = 5.5 - 2.5 = 3.0$.



In some texts, the interquartile range is defined differently. “It is defined as the difference between the largest and smallest values in the middle 50% of a set of data. To compute an interquartile range using this definition, first remove observations from the lower quartile. Then, remove observations from the upper quartile. Then, from the remaining observations, compute the difference between the largest and smallest values. ... When the data set is large, the two definitions usually produce the same (or very close) results. However, when the data set is small, the definitions can produce different results.”⁴³

⁴³ Look at <http://stattrek.com/descriptive-statistics/variability.aspx> for their explanation.

We now compute as follows: remove the observations from the lower quartile: 2 and 2. Then remove the observations from the upper quartile: 7 and 10. Now we are left with the values:

Values
3
3
4
4
5
6

The interquartile range is now defined as the *largest* minus the *smallest* value. The interquartile range (IQR) would be $6 - 3 = 3$.

Because we *do have a small data-set* or set of values (only 10 values is considered to be small) we have indeed a relatively large difference with the value calculated via the first method (i.e. 4.0).

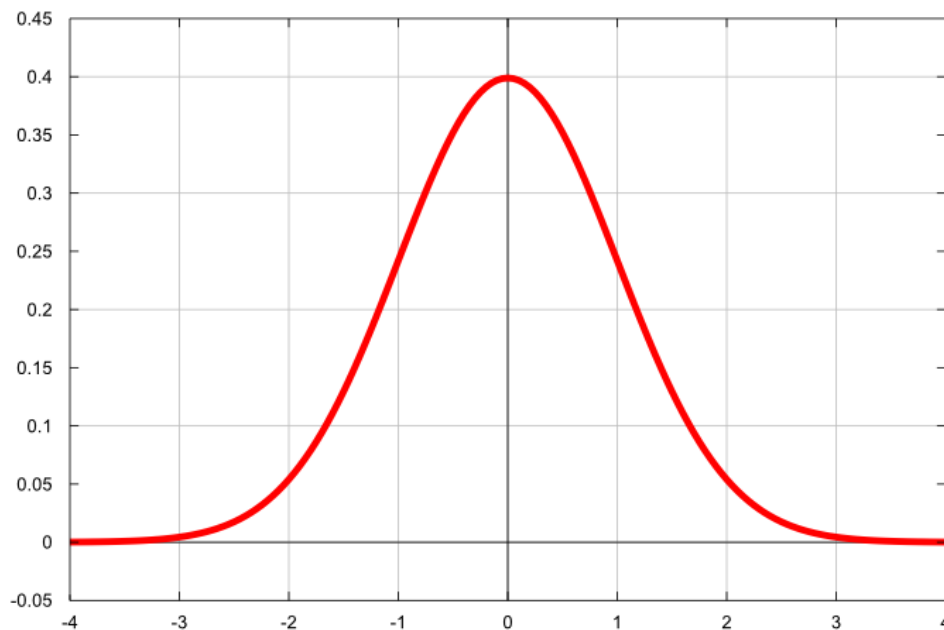
The interquartile range tells us about the spread in our data. A value of 4.0 (via method 1) related to the values we have, is considerably large and a value of 3 (via method 2) as well. Compared to the median these figures all indicate of a *relatively large spread* in our data or values.

A population contains all existing values there are. For instance all the inhabitants of the country you are researching, in a certain era, are a population. The part of data you are researching, the example of the 10 tombs, you may call a *research population*. The term population is important because it denotes the field which you are researching and the strength of your research findings. A fellow student of mine researched a certain kind of statue, and she has found (as far as we know) *everyone in existence*. In this case your research findings could be potentially strong because you are talking about the *whole population*. There is no *margin of error* as could be the case with a *sample*. That is unless another scientist shows you, for instance, that a lunatic emperor destroyed all statues of a certain type from your population. As you can see we have to be careful always.

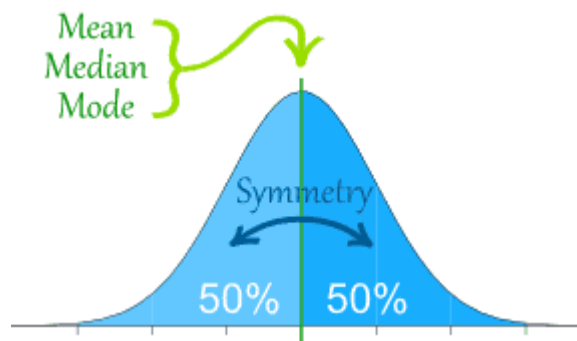
6.4.2. Special distributions: the normal distribution

Take a look at this *special distribution* below of for example the size of potsherds in inches, found in many different locations where an ancient pottery used to be. The form of this distribution has a special symmetrical shape, like that of a bell. This distribution is there called a normal distribution, bell-curve or Gauss-curve.⁴⁴

Normal distribution



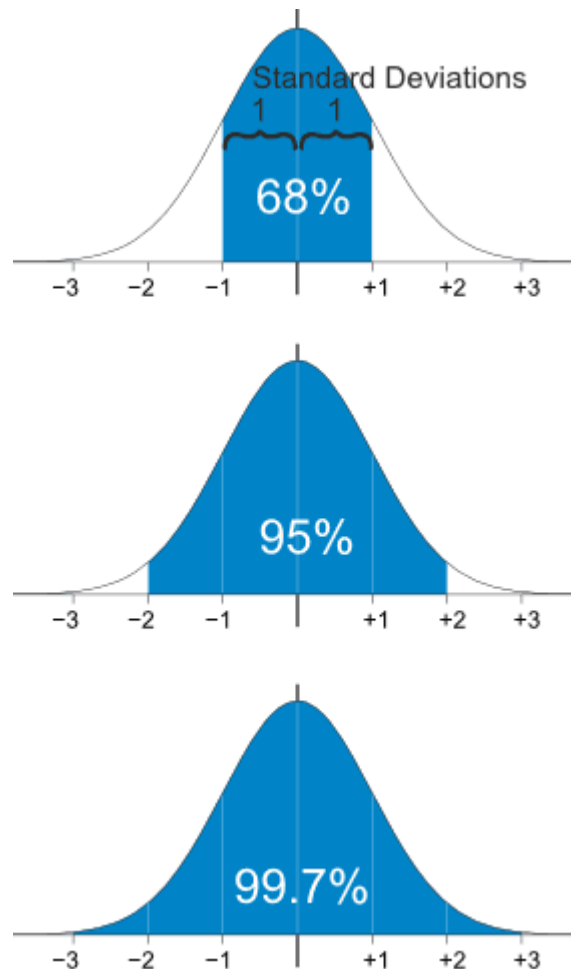
As you can see in the figure above and below, the distribution is *evenly spread*. There are as many cases on the left side as there are on the right side of the diagram. This means that *the median* or 50% of the values is displayed exactly in the middle. In this special case the calculated mean equals the median.⁴⁵



⁴⁴ Picture taken from: <http://www.statlect.com/ucdnrm1.htm>

⁴⁵ Picture and more information from: <http://www.mathsisfun.com/data/standard-normal-distribution.html>

And there is more: one SD or standard deviation, covers 68% of the values in a normal distribution, two standard deviations cover 95% of the values and three standard deviations cover almost the entire distribution with 99.7%. Look at the picture below; taken from the website “Math is fun”.

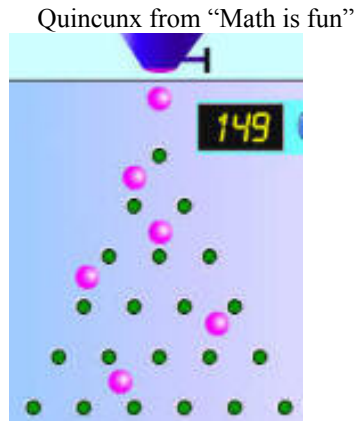


In statistics we assume that, under some given conditions, the mean of a sample population will tend to a normal distribution as the population size increases. almost all populations are distributed according to the normal curve. This is called the *central limit theorem*. Any variables with any distribution having a finite mean and variance tend to the normal distribution.⁴⁶

Examples are as widespread as you can imagine, for instance length of people, not only of their body, but also of their hair or arms, the height of trees, buildings or statues. Even perhaps mechanically fabricated objects, which tend to have a *tolerance*, that is to say that some objects are too small and some too large, but these will not necessarily fall along the normal distribution because of the fabrication process itself. Then again it is not improbable. Most living things, things of nature and most man-made objects tend to have a distribution along the lines of a normal curve. There is a normal distribution even within the scientist community. Some are strong believers of a certain thesis, some do not believe it at all and the majority will have a balanced opinion with mixed arguments for and against.

⁴⁶ Look for more details and explanation at: <http://mathworld.wolfram.com/NormalDistribution.html>

Also very special about this curve is, that it can be formed by dropping little particles or balls on a rack of evenly spread pins in the middle.⁴⁷ The figure that you get when you have dropped a lot of those objects is a bell-curve or normal distribution. This figure of distribution is special also because of the fact that it is derived from pure *chance*. Whether the balls go left or right when they are dropped is not influenced and determined solely by chance.



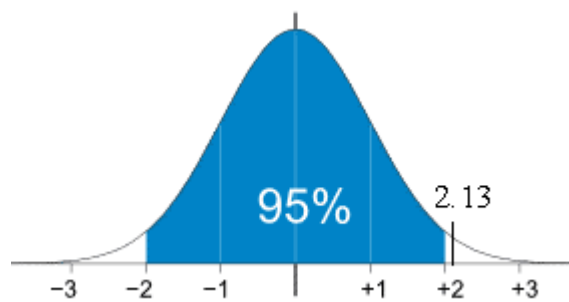
We can adapt the values we have found in our research so that we can *use* a normal distribution *directly* on our research results. To do this we have to *express* the values we have found in SDs standard deviations from the mean. We then arrive at what is called the "Standard Score", or Z-score. We do that by subtracting a score found by the mean and then divide it by the standard deviation, like this:

$$\frac{\text{Score (value)} - \text{Mean}}{\text{Standard deviation}}$$

Example: a normal distribution has a mean of 100 and a standard deviation of 15. What is the Z-score of the value 132 and what does it tell us? First subtract the mean from the value: $132 - 100 = 32$ and then divide by the SD = 15. $Z = 32/15 = 2.13$. This means that the value we have found lies 2.13 standard deviations away from the middle or average of the normal distribution.

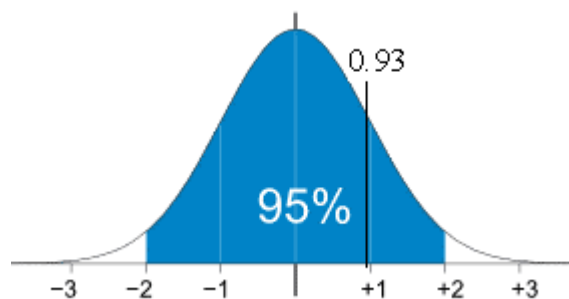
We know that 95% of our distribution is covered by two standard deviations. The value we have found is higher and therefore lies within the 5% range of the right side of the diagram. Usually this 5% area of a normal distribution is considered to be "extreme". The meaning of this is that the value we have found does *not* represent the research population very well. Not at all to be precise! It is an outlier.

⁴⁷ Take a look at the website of "Math is fun" where such a "machine" or Quincunx is animated, here: <http://www.mathsisfun.com/data/quincunx.html>



Now consider another value of 114, with the same mean and standard deviation. Calculate this Z-score. Is this a better score?

First subtract the mean from the value: $114 - 100 = 14$ and then divide by the SD = 15. $Z = 14/15 = 0.93$ This means that the value we have found lies 0.93 standard deviations away from the middle or average of the normal distribution.



This is a 'good' score, because it lies no further than one standard deviation from the middle or mean and therefore represents the research population quite well. As you can see, this is a nice way of telling something about the values we have found in our research. We can report about our research values and show their meaning to our audience with the aid of statistical measures.

6.5. Sampling

6.5.1. Evaluating the population

One of the most valuable tools we have as a researcher is the use of a *sample*. I already discussed a preview in sub paragraph 6.1.2. Take that in mind when reading the next part of this paragraph. Sampling gives us a way of making statements about the population we are researching, *without* having to research the population in its entirety. A sample is a narrowly defined *portion* of the *population* we are researching. As already stated, there are some very strict rules to follow in taking a sample from a population. These rules are there for a very specific reason. The sample that we are about to take from the population we are researching, like tomb paintings, potsherds, grave goods, statues, almost anything you want to study, even the number of terms in a translated text, will have to be *valid*. The validity of a sample depends upon the rules of statistics. In particular, it depends upon the rules of *chance*. If we want to take a *representative* sample (i.e. a sample from which we can deduce what the entire population is like), we must take a *randomly* picked sample. How to do that? For starters, we have to know as much as possible about the *entire* population we are about to research. If we

are researching an ancient civilization, how many persons were there at the period of time we are interested in?

We don't have to know the *exact* number of people there were, but we do need to know the size of the volume by estimate. Are we talking about 300,000 people or perhaps about 3,000? This figure matters, if we want to calculate the number of people in our sample, in order to make statements about the population in its entirety.

Also, we need to be acquainted with peculiarities about the population, as well. For instance, if we know that at the time we are researching, a large army of men was fighting abroad, we can expect a shortage of man versus women. This fact shall have to be accounted for in our research findings and perhaps in the way we are about to draw our sample too. Maybe there are even more peculiarities to know, beforehand. If our research sample shows a lack of man, we know the reason and can explain the shortages to the reader of our research findings.

Example

Let us suppose that you want to research statues of “working people”, laborers on the field and perhaps in the household or even small factories or workshops, of an ancient civilization. You could choose all categories, but you could also choose one of them. The next questions to be answered are: where could we find them and, how many of those statues are there? You have to invest in a little pre-research first, to get these questions answered even before you can decide on what sample to take. Suppose the majority of these statues are in museums, but some remain in the spot where they were found.

First make an estimate of how many statues are in museums and how many there still are in the field. Suppose the ratio museum versus in the field is 90% versus 10%. It seems so that, in order to be representative, our sample has to contain 90% museum and 10% outside statues. But that is not really true. A sample from both, truly selected by chance, could for instance contain less than the 10% outside statues or perhaps even more. That is what could happen if we let chance decide on our sample. If we want to present a good image on both differently situated statues, we best consider them to be *two* different populations and take two representative samples. We call a representative sample chosen by chance, a *randomly* chosen or *a-select* chosen sample. Selecting two different samples *ensures* that we have enough research material on both which will be reflected in our research report. Having both samples introduces great possibilities of testing for differences in both research populations too.

Ok, that is one question solved but now, how are we going to take the samples in practice? First we have to know how many *locations* there are from both the museums and the outside. If we know how many locations there are, we have to know how many *objects* each carries. Now we know exactly how many objects there are, i.e. how large the population is in its entirety. Now we know that number, we have to make another decision. If the number is relatively small, for instance 136 objects or perhaps 310, we can choose to study them all. The research results we will be presenting afterwards will be, statistically speaking, incredibly strong! No margin of error here, because we are talking about the *entire* population. However that is *not* true for all the statues that *originally* must have existed. There are many reasons for certain types of statues either to remain or to get lost over time. If you know not of such reasons, I earlier mentioned a lunatic destroying certain type of statues, you may *assume* that the remaining statues, survived *randomly* and that the *population* you are studying now is a *re-search population* or *sample* of all the statues that once existed.

Since we do not know *how large* the population originally once was, you have to be careful in what you will be presenting your audience. You can be bold in your conclusions if you wish *as long as you present all that you (don't) know*.

Now let us assume that there are not 136 statues but 2136 statues left from antiquity. This population is much too large for us to study, so we do need to take an *a-select* chosen sample. How to do this? First *assign numbers* to each object in the population. You don't have to label them physically, but you can make a list on the basis of museum catalogs with the aid of a sort of *translation table*. Suppose a museum has 50 objects we want to study. If their numbers are neatly in a row (which we do not expect) tie them to your *own numbers*, for instance 201-250, can correspond with museum numbers like those of the British museum BM01001, BM01316 etc. If *your* numbers 211, 233 and 246 are in the sample you have picked, you will only have to research the *corresponding museum numbers*. The British museum, for example, offers a search facility⁴⁸ on their website that present you a way in finding the objects you want. They also facilitate searching by museum number. In the end we have our number series ready from 0001 through 2136. All we have to do now is the *actual* taking of the sample. A *random number generator* can generate *a-select chosen* numbers in the amount we want and within the range we want. In our case, for example, 50 numbers between 0001 and 2136. There are random generators in abundance on the web⁴⁹, like the one from the website random.org⁵⁰. But we do not now yet exactly how many numbers to choose. We discuss the *amount* of objects, i.e. our *sample-size*, and other important criteria, in the next sub-paragraph.

6.5.2. Sampling criteria

“The sampling Guide” from the English National Audit Office⁵¹ declares that sample size will depend on 5 key factors:

1. Margin of Error
2. The amount of variability in the population.
3. The confidence level
4. Population size
5. Proportion of the population with the characteristics you are interested in.

We have already taken care of the last two key factors: 4. “Population size and 5. “Proportion of the population with the characteristics you are interested in” in the sub-paragraph above, and we are going to discuss the three remaining key-factors below.

1. Margin of Error

If we were to research the *entire* population, we wouldn't need to establish a *margin of error* because that is not necessary. We should then be able to present *valid* research findings, on account of the fact that we know the *entire population* and not just a *portion* of it, i.e. the sample. But since we are drawing samples, because we have neither the resources nor the means to research the entire population, we will have to accept the fact that there can be errors. For instance, the sample we take might not reflect the entire population as it should be.

⁴⁸ <http://www.britishmuseum.org/research.aspx>

⁴⁹ and built into Excel from the version 2007 up under the title ASELECT.

⁵⁰ <http://www.random.org/integers/>

⁵¹ Consult: www.nao.org.uk/publications/Samplingguide.pdf

The larger the portion of the population we sample, the less chance we have of making sampling errors. However the size of the sample has to be suited to our means and budget and should not be too big.

You don't have to calculate a sample-size yourself according to a formula. Nowadays everyone can make those calculations on the internet with the aid of a *sample-size calculator*. However those websites will demand some input from you in order to make a correct calculation.

As the calculators will indicate, the standard margin of error in statistics is 5%. However there can be reasons to increase to 10% because the sample will be too large for us or else to lower the margin of error because of the larger certainty that our research requires (1%).

2. The amount of variability in the population.

In the paragraph 6.4.1 above, I explained some measures of distribution and variability, like the means, standard deviation, variance and the interquartile range. These measures are all used to calculate the amount of variability or spread in your research data. That means that we already know what variability in data is. If we are going to research an unknown population, it is a little difficult to state *beforehand* what the *variability* or *spread* in our population will be. We can only calculate this with the aid of the measures, like variance and interquartile range, if we *already know* the population, which we don't. This problem is usually solved by the use of a default. The figure of 50% is usually taken as a *default spread*. Unless you know a bit more about the population, you are about to research, I advise you to use this number as default. But if you expect larger or less variation, take a correspondingly larger or smaller number or percentage.

3. The confidence level

We use a *confidence interval* to describe the uncertainty with a sample that makes an estimate of a population. In other words: how well suited is the sample we take to represent our population. Earlier on we talked about the normal distribution and we showed that it can be used to display a (research) population. If a normal distribution can be used to display a population, it can also be used to display a population of *samples*. In fact if we take a sample and repeat that process over and over again, our sample results will be distributed like a bell-curve. The more samples we take the more they will tend towards the mean. The general norm in (population) statistics is that we take a margin of 5% of both ends of the curve, to represent parts of the population that are extreme regarding the average of the population. In other words, as long as there is a 95% chance that our sample lies within two SD's or 95% of the normal curve a *confidence level* of 5% is acceptable. A 95% confidence level means that there is a 95% chance that our samples will contain the true values of our population.

6.5.3. The taking of a sample

Assuming a population size of 2136 in our example in subparagraph 6.5.1, we decided to take a sample. Fortunately, we do not need to make difficult calculations anymore. There are a lot of calculators on the web, to do the math's for us.

This one is really simple: <http://www.berrie.dds.nl/calcss.htm>

Let us see what happens if we fill in the form of this calculator, with default values:

- Population: 2136
- Confidence: .95 (i.e. 95% confidence)
- Margin: .05 (i.e. 5% error margin)
- Probability: .50 (i.e. variability of 50%)

Sample Size Calculator for a proportion (absolute margin)

Population	2136
Confidence:	.95
Margin:	.05
probability:	.50
The sample size is:	326

Calculate sample size

The calculator suggests a sample size of 326. Personally, I consider that to be too large a sample. What will happen if we make the population about 10 times bigger, by adding a zero at our population size?

Sample Size Calculator for a proportion (absolute margin)

Population	21360
Confidence:	.95
Margin:	.05
probability:	.50
The sample size is:	378

Calculate sample size

Now the calculator suggests a sample size of 378! While we have increased the population size enormously by about 1000%; our sample only gets larger by $378 - 326$ divided by $326 = 52/326 = 16\%$! Those figures are *not* proportional, and I suggest using different defaults values, however not with this calculator. I suggest a calculator more adapted to our needs. A real nice calculator is the one from raosoft: <http://www.raosoft.com/samplesize.html>. Why? Because it allows us to *change* the numbers of the defaults and see what happens immediately on screen. This calculator also explains a lot and shows different scenario's to choose from. If you have made your self-acquainted with the matter described above, this is the calculator to use. Because, maybe we don't *need* an error margin of 5% and a value of 10% is ok too. This may reduce our sample size considerably and hence all our work associated with the sample and the follow-up analyses. Look this up yourself and play with the different numbers. There is only one *condition*, manipulating the variables: *you have to know what you are doing!* I have tried it for you and came up with the following figures:

- Margin of error: **10%**
- Confidence: 95%
- Population: 2136
- Probability: 50%
- **Sample size: 92**

 Raosoft®		Sample size calculator
What margin of error can you accept? 5% is a common choice	10 %	The margin of error is the amount of error you can tolerate. A larger amount of error than if the margin of error is smaller. Lower margin of error requires a larger sample size.
What confidence level do you need? Typical choices are 90%, 95%, or 99%	95 %	The confidence level is the amount of uncertainty you are willing to accept. If you choose a confidence level of 95%, you would expect that for every 100 times you repeat the survey, you will be wrong 5 times. Higher confidence level requires a larger sample size.
What is the population size? If you don't know, use 20000	2136	How many people are there to choose you from?
What is the response distribution? Leave this as 50%	50 %	For each question, what do you expect the distribution of responses to be? If you don't know, use 50%, which gives the most conservative (largest) sample size.
Your recommended sample size is	92	This is the minimum recommended size of your sample. It is more likely to get a correct answer than a smaller sample size.

By accepting a slightly higher margin of error, i.e. 10% instead of 5%, I have *reduced* my sample size from 326 to 92! Note that 92 is a much more workable number than 326. If I was writing a master thesis, I would choose a margin of error of 10% and *explain* in your thesis *why* you did that. But remember that there are *consequences* regarding this decision. You've *lost* on your research accuracy. For a master thesis, this is no big problem, but if you were working on a new research publication, it probably is. You have to be always cautious of the research decisions you make and you must display them in your research report!

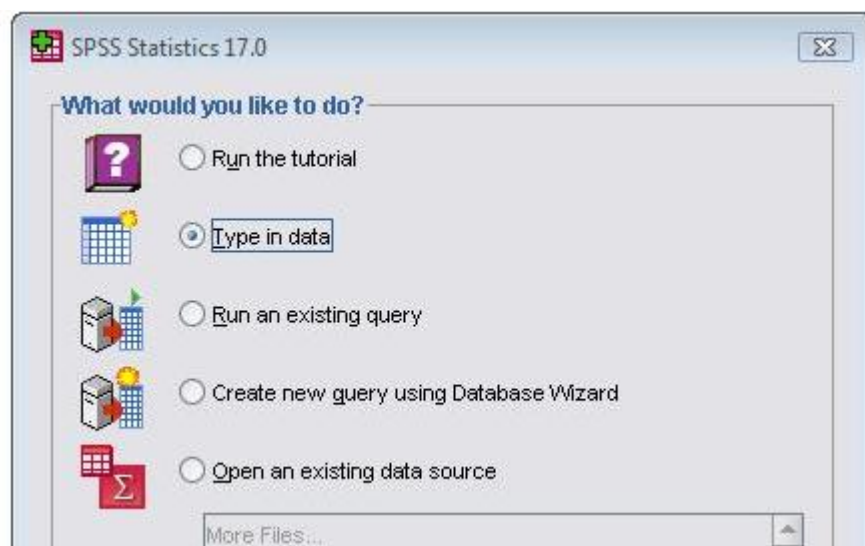
6.6. Using statistical software: SPSS

6.6.1. Introduction.

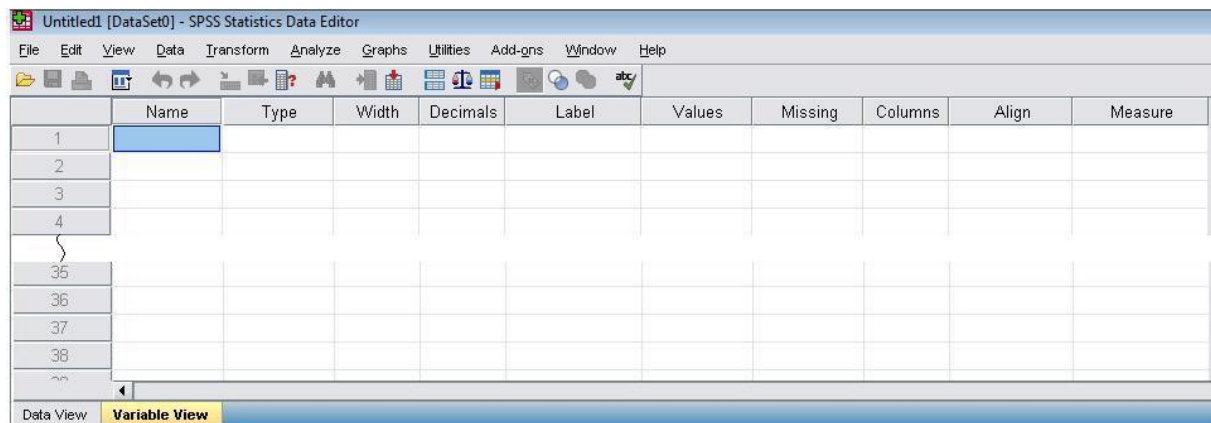
The type of statistical software the reader will be using is largely dependent on the software that his or her university has selected. As the Statistical Package for the Social Sciences or SPSS, is widely used, we will now discuss how to use the SPSS package for the PC here. We will continue our example we have given in chapter 4.2.6. of a simple codebook.

6.6.2. Entering a codebook.

Upon entering the program, SPSS starts up with a screen which asks you what you want to do. Look at the picture below. One of the very first things to do is to let the statistical software know what we have constructed for our research: we have to feed our codebook into the computer. Click “Type in data” and continue.



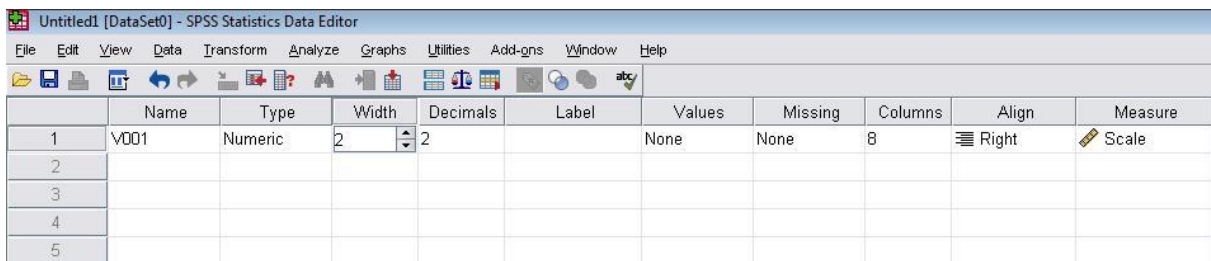
The next screen is called “SPSS Statistics Data Editor” and we now have a kind of spreadsheet-like screen on our PC. You can more or less work with this screen like it is a spreadsheet. However this is the place to enter our codebook. Note the “Variable view” sign in yellow in the left lower corner of the screen. I have shortened the screen below to fit the page.



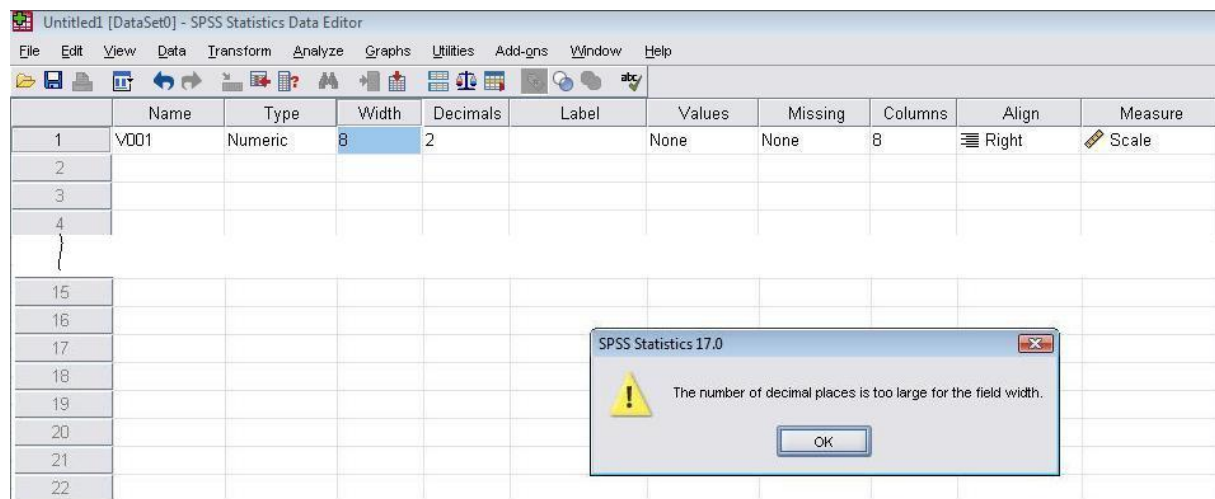
The menu line of this screen, starting with “Name” and ending with “Measure”, represent all the qualities of the variables we have defined in our codebook. We will enter the values of all the menu categories for each variable we have defined in our codebook.

In the first field of the upper left corner under the heading “Name” type in V001. That is the name of our first variable. Refer to chapter 4.2.6. Our first variable is “Casenum”. This variable is used to record the *number* of statues or “cases” that will be fed into the computer. The *variable type* therefore is: numeric. The program *assumes* that our first variable is *numeric*, as can be seen in the second field of this variable. We don’t need to alter that. It is correct.

The next field in the column is “width”. As can be read in the codebook, the values for this variable can be in-between the numbers 01 – 98. Click on the third field and the field splits into two parts. There are arrows in the right part of the field. Select with the aid of the arrows the number “2” for two positions.



Now the program warns that two positions is not enough to account for eight decimals, as the program is assuming.



Go first to the “decimals” field and adjust from “8” to “0” with the aid of the arrows. Now adjust the number of positions of width to “2”.

Next click on the following field in the subsequent column “Label” and enter the label for this variable. Type in “Casenum.”

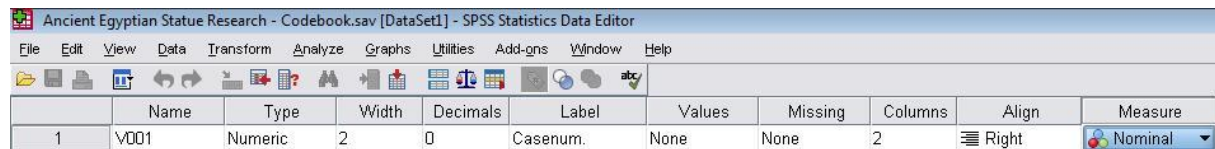
The next column of our variable V001 “Casenum” is Values. The program assumes that *values* for this variable are not necessary to type in, as is often the case with numeric variables. The number that this variable can get equals the value i.e. “67” is case 67. Do not change.

Go over to the next column “Missing”. As previously mentioned, entering a missing value for the variable “Casenum.” is useless because there will be none. Leave this as it is.

Now go over to the next column “Columns”. The width of our variable is already entered as “2”. As no extra digits are needed, the number of columns that are needed for this variable can be changed from the default “8”, into “2” to equal our width of “2”. Click in the field under “Columns” and select with the aid of the arrows the number two for two columns.

The column next to the previous one is called “Align”. The alignment of the variable “Casenum, is not critical”, it refers to the place it has in the columns. There are three options: right, middle and left. The default position “right” is suited to display a numeric field. Leave this as it is because it is only an optical matter and not concerned with the content of the variable.

The last thing the program wants to know is the *measurement level* of the variable. This is explained in chapter 4.2.3 Types of variables. The default option is “Scale” which is suited to display the numeric variable “Casenum.”. However, we will not be manipulating and doing arithmetic, with this variable. This variable is only used to indicate the observations we have, so the measurement level must be changed into “nominal”.



	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure
1	V001	Numeric	2	0	Casenum.	None	None	2	Right	Nominal

By now the file should be saved in order to protect the time and effort invested in our data-entry. Select a proper place and name for the codebook. Go the file “file” option on the menu, select a position on hard disk or USB-drive and type in the name under which, the codebook is to be saved.

Continue with the next variable. This is called V002. Under the heading “Name” type “V002” in the first column of the second line. The next column is “Type” and the standard option is “Numeric”. This is ok, because we have *coded* the different genders Male and female with *numbers* in order to do some statistical analyses i.e. to do some math’s.

The column next to width is “decimals”. Because we have changed the number from “8” to “0”, the default number is now “0”. Check this and correct if this is not the case. “0” is the correct number of decimals. There simply exists no “1.2” type of male.

Moving on to the column “Label” . Type in the label of variable V002, which is: “Gender”. Now move on to Labels and click on the default value “None”. This opens a sub-screen. Type in the values and add a label per value. Type in: Value: 1 Label: Male; and click on “add”. The type in: Value: 2 Label: Female; and click on “add”. That is enough. Do *not* add Value: 9 Label: Missing, because there is a separate column for that in the SPSS Statistics Editor screen.



A slight warning when making labels: If the number of characters is too high, there will be problems with the output of analysis later on because the program will use the space that you enter in the number of characters. Tables and other output could be distorted if the number of characters is set to high. The program will shorten the labels which are too long. That is hardly a problem because you can change the longer names back in the research publication. Remember, this statistical program only performs the statistical analysis and can make some plots, but the makeup of the publication, is not here to be made.

Next click the field with “Missing”. Here we will enter the missing values we have defined in our codebook. Another dialog screen will open. Now select the second option: “Discrete missing values” and type in a “9” in the first of the three fields. Than type in “ok” and we return to the first screen.



In the next option “Columns”, the number of columns must correspond with the width of our variable i.e. “1”. Check if this is the case, change if not ok, and continue.

In order to get a nicer view here on this variable, we will set the alignment to: centre. Click in the field and change the value of the field.

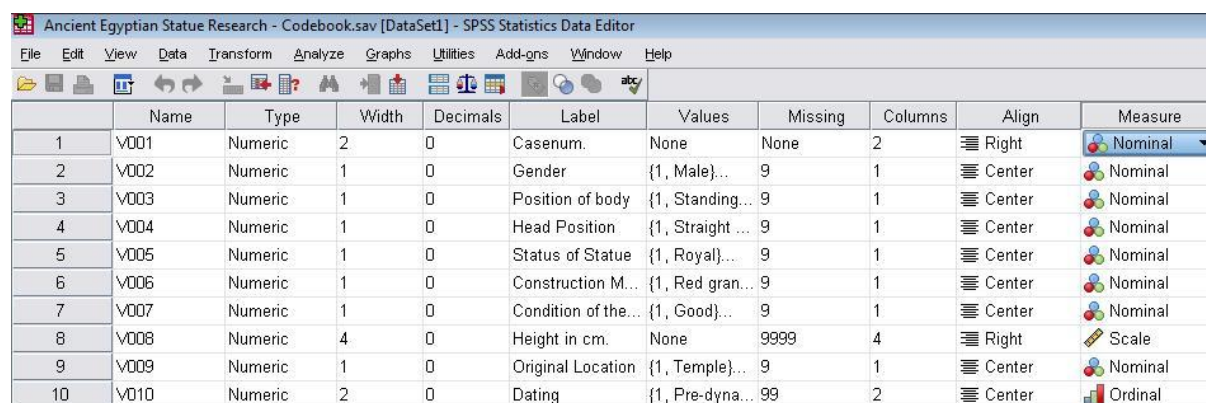
The last column refers to the measurement level. Since we only have coded the values “Male” and “Female” with a number, resp. 1 and 2, the measurement level is only nominal. Change the measurement level by clicking in the field and change the level to “nominal”.

I will now only refer to the differences with the previous variables. The variables V003 trough V007 and V009 can be handled, just like variable V002. Type in all the necessary data for those variables.

As our own codebook states, variable V008 refers to the height in centimeters. We need a width of “4” and no decimals as the height will only be dealt with in cm. whole. Do not forget to type in the missing value of “9999”. 4 columns are needed for this variable. This variable is a true “Scale” variable. Type in all the necessary values.

Variable V010, in our codebook, refers to our dating of the statues. We need a width of “2” and no decimals. Do not forget to type in the missing value of “99”. 2 columns are needed for this variable. This variable is an *ordinal* type variable because, for instance, the third time interval is newer than the second interval and older than the fourth. We do know the exact time distances, but we added variation in them. The chosen intervals varied from 1000 to 500 and even to 250 years; hence the ordinal type. Type in all the necessary values.

If you have correctly filled in all the necessary, below a picture what your screen should look like.

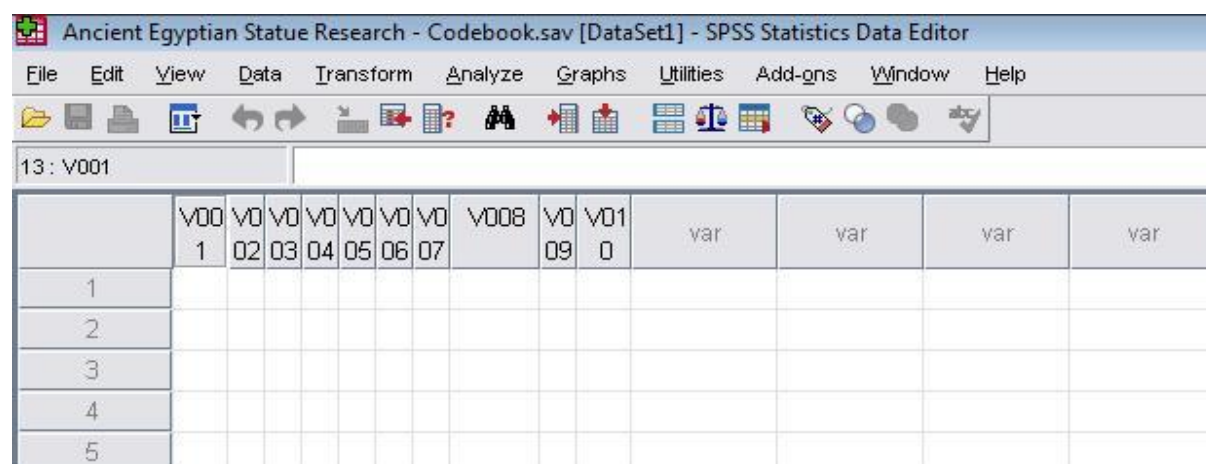


	Name	Type	Width	Decimals	Label	Values	Missing	Columns	Align	Measure
1	V001	Numeric	2	0	Casenum.	None	None	2	Right	Nominal
2	V002	Numeric	1	0	Gender	{1, Male}...	9	1	Center	Nominal
3	V003	Numeric	1	0	Position of body	{1, Standing...	9	1	Center	Nominal
4	V004	Numeric	1	0	Head Position	{1, Straight ...	9	1	Center	Nominal
5	V005	Numeric	1	0	Status of Statue	{1, Royal}...	9	1	Center	Nominal
6	V006	Numeric	1	0	Construction M...	{1, Red gran...	9	1	Center	Nominal
7	V007	Numeric	1	0	Condition of the...	{1, Good}...	9	1	Center	Nominal
8	V008	Numeric	4	0	Height in cm.	None	9999	4	Right	Scale
9	V009	Numeric	1	0	Original Location	{1, Temple}...	9	1	Center	Nominal
10	V010	Numeric	2	0	Dating	{1, Pre-dyna...	99	2	Center	Ordinal

Save your file via the “File” menu option in the “SPSS Statistics Data Editor”. Close SPSS.

6.6.3. Entering research data.

Open the SPSS program. Now click “Open an existing data source” and select your codebook file and click “ok”. At the lower left corner of the “SPSS Statistics Data Editor” there are two options: “Data View” and “Variable View”. Our codebook is written in “Variable View” and this option is now selected as the yellow colour indicates. Now select the “Data View” with the mouse. Once again a spreadsheet like screen appears. The columns represent our variables and do not look very nice.

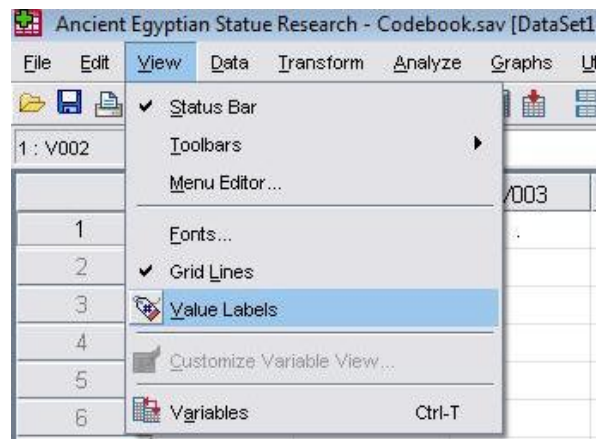


	V001	V002	V003	V004	V005	V006	V007	V008	V009	V010	var	var	var	var
1														
2														
3														
4														
5														

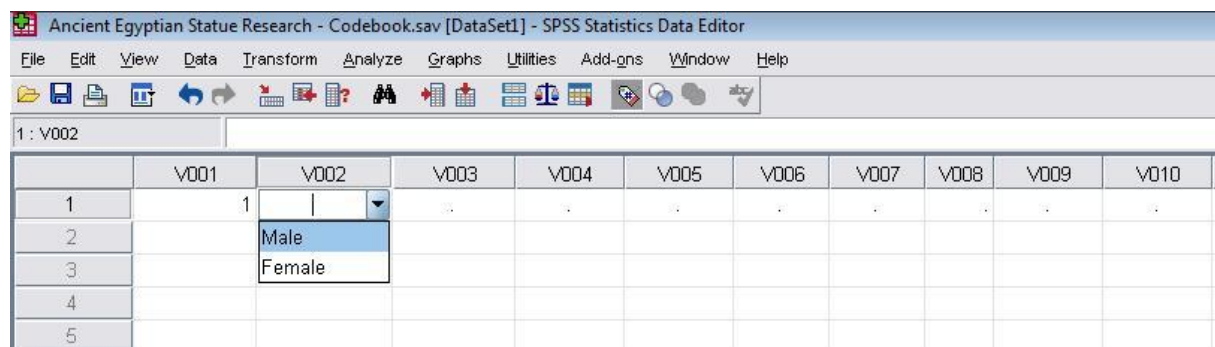
Use your mouse to drag the columns to a wider view and position them nicely on the screen. They must have a good separation among them because we are going to use this screen to enter our research data. There are SPSS versions that have separate data entry parts. If your program has that option, use it! It reduces errors during the data entry process. My version has no data-entry option, so I will continue.

First save your file under a different name, for instance “first research data”.

To enter values safely and minimize data entry errors, click on the upper menu on “view”. Select “Value Labels”



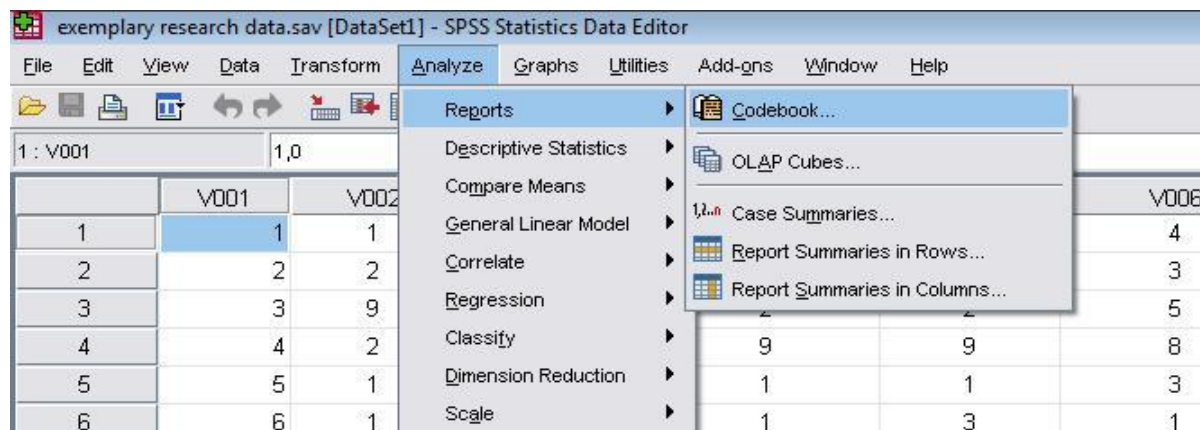
Now you will see your screen change into this:



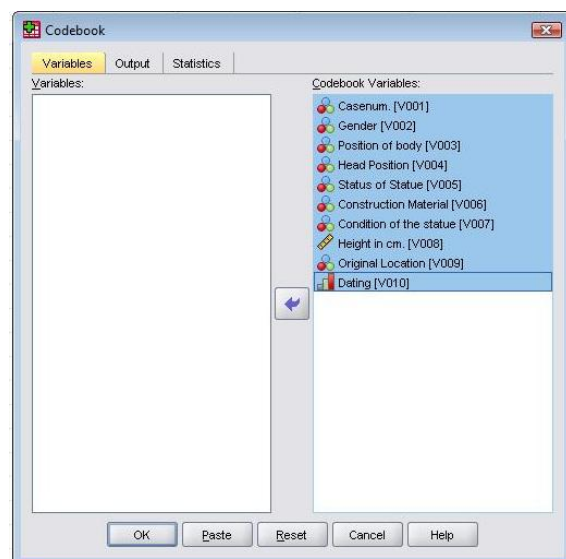
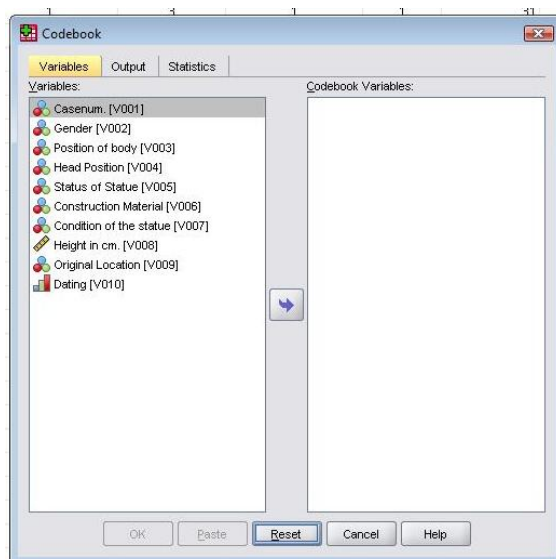
Now “type in” (i.e. select the value labels with your mouse) your research data line after line starting with Casenum. 01. Only V001 and V008 have no predefined value labels and you have to type in these values yourself. Have a printed version of your codebook ready or open up a “Word version” with your codebook for reference of your labels.

Once you have finished, it is best to save your data file in a remote position on your hard disk or on an external USB drive. In this way you can always fall back on the original data file, if you have made some mistakes later on when we can start manipulating data.

At the end of your data-entering process save and close the file and store separately. Now the typed in data are going to be checked for errors. Start SPSS again if it is not already on. Open your data entry file. Either position: “Data View” and “Variable View” in the “SPSS Statistics Data Editor” are ok. In the menu at the upper part of the screen, open “Analyze” and click on “Reports” than select the first option “Codebook”.



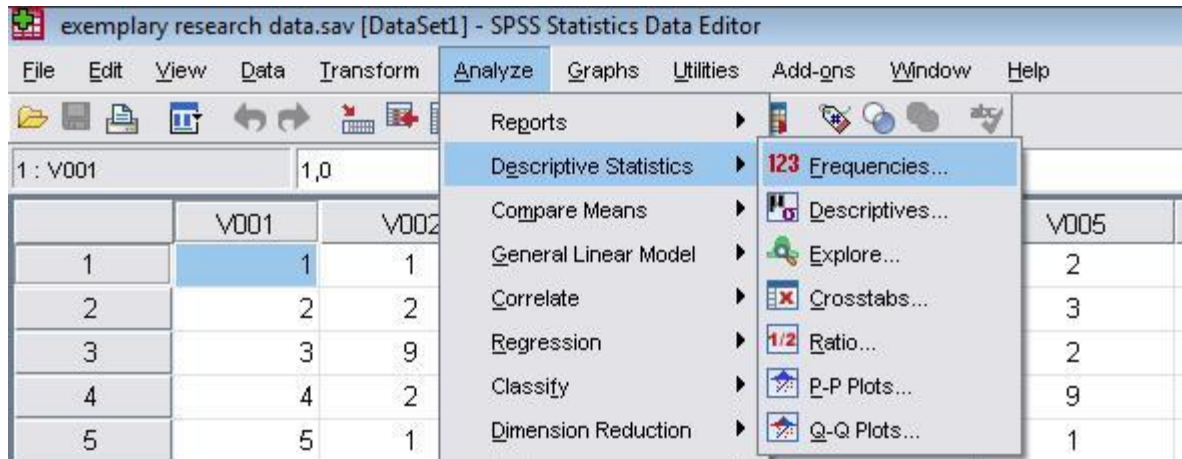
A dialogue box pops up with all variables in the left box and none in the right box. In the right box you can drag from the left box all the variables you want to examine. Select all variables from the left box and, via the arrow, transport everything to the right box and click “ok”.



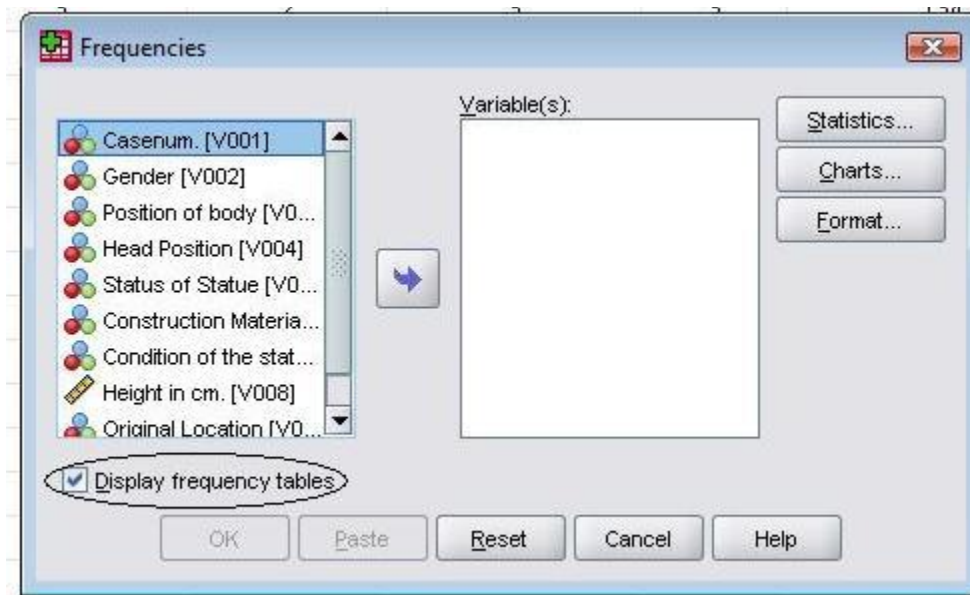
An output file opens and the computer takes some time to display all. Print this output file via the “file” menu at the top left of the screen. After you have done that, close the output screen by click on the “X” at the right top corner of the screen. If the print is ok, you do not have to save this output file. It is very easily produced again, if necessary. If you do not see the output produced by SPSS, click on the “Output1” tab at the very bottom of your screen. Be patient, it will take up some time to get ready.

Now study this paper output file very closely. If you do not know what everything signifies, that is not important. The reason we have made this output is to determine if some errors have been made. If you detect them, open SPSS again and correct them in the “Data view” of the “SPSS Statistics Data Editor” and repeat the procedure above. When you are done, close the program and do not save the output file.

There is an alternative method. Start SPSS again if it is not already on. Open your data entry file. Either position: “Data View” and “Variable View” in the “SPSS Statistics Data Editor” are ok. In the menu at the upper part of the screen, open “Analyze”. Select the second option “Descriptive statistics” and then 123 Frequencies.



A dialogue box pops up with all variables in the left box and none in the right box. Check the box: “display frequencies tables” if it is not on. Otherwise you will have no output. Continue acting like above.



When you are done, close the program and do not save the output file.

Addenda

Addendum 1. Writing Correct translations

Introduction

For their Master thesis students of the ancient cultures often make translations of hitherto unpublished documents or new translations of documents that were translated and published a relatively long time ago. The study of ancient, and usually dead, languages, still gains momentum every year and translations are getting better over time as more of the language is deciphered and learned. In an attempt to facilitate this process of making a translation, and present a clear cut and accepted framework of a scientifically acceptable translation, I contacted Dr. J.G. Dercksen at Leiden University, who is a specialist of Akkadian, and asked him about the prerequisites for making a translation. I consider it a possibility that you know all that is described below. If that is so, than that is great! However, remember, my aim goes further than that. I want to give guidelines to *all* students and strive for international *standardization*.

A brief note on language

This book is mainly intended for students of the *ancient* languages and cultures. Therefore this paragraph will focus only on translations of those languages. They differ from modern and active languages because they are usually *dead* and the exact meaning of certain signs or symbols, words and phrases cannot always be completely known. These languages are still being researched in order to establish the meaning of the unknown and to improve on translations of known words and phrases.

In my contacts with Dr. Dercksen, I discovered that the translation *process* of ancient Akkadian texts does not differ very much from that of ancient Egyptian texts, and I do suspect that the same goes for other ancient and dead languages as well. We both first *transliterate* the signs *before* we start translating the text. The transliterated text then becomes the basis for a correct translation. One important difference I detected at Leiden University was that in the case of Akkadian texts students are required to present a photocopy or digital *image* of the original text before even starting the transliteration process. That means that everyone is able to *verify* your work. Writing my own MA-thesis I discovered that a translator of a papyrus I studied had *omitted* a complete verse. With the procedure for Assyriologists outlined above, this omission will not remain undetected for long. The procedures outlined below will apply to most of the translation processes of the ancient and dead languages.

Prerequisites for making translations

Making a correct translation of a dead language involves having profound knowledge of the culture in which the language is used. We therefore start by stating *demands* on the person performing the translation, i.e. the *translator*.

The following demands are essential:

- Knowledge of the ancient language learned in an *academic training*.
- Solid and integral knowledge of the *culture* of which language you are translating.

One of the most important things in translating a dead language is *knowledge* of the society in which it is used. The social meaning of a text, or the phrasing of it, is very important. It adds to the understanding of the text, because it places it within a *context*. Often words get their true meaning when viewed within a context. When looking at it this way, the better you know the context, the better you are able to make good and adequate translations. But that is only one side of the process. The language to which you are translating is another. Sometimes, to make a certain word understandable, the modern version of the language we are translating to is used. But this can be totally wrong and out of context where the ancient civilization is concerned. A well-known example, and known to most students of the ancient languages and cultures, is the translation of the word in the ancient language that means “grave field” or in Greek “necropolis”. This *is* the correct interpretation. Translating the original word for grave field or necropolis in “Church-yard” would be completely wrong and totally out of context. A church did not exist back then and a grave field did not “belong” or was attached to a ‘temple’. Although the general public would know what you mean, language wise and culture wise it is completely incorrect. Indicating a problem is one thing, solving it another. In the west, words of religion have pervaded our language for a long time. And sometimes we don’t even know that we are using these words or thinking in those terms, which we have learned from our childhood onwards. Our already cited “churchyard” is but one example.

When translating originals, the *purpose* for which you are translating is important. In my own MA-thesis I have used the “Book of the dead” in order to shed light on what the ancient Egyptian thought about “crime” or “deviant behavior”. The ancient 1st Egyptian phrase in the 42 denials of the Book of the dead [*n iri-i isf.t*] is translated by a famous author as “I have not done falsehood”. This is correct grammar and the word falsehood is clear and unambiguous. However, my purpose was to shed light on ancient Egyptian *thinking*. The word [*isf.t*] or Isfet is meant to be the exact opposite of [*Mꜣꜥ.t*] or Maat. The word Maat has a special meaning in ancient Egypt. It is not easily translated. The German Egyptologist Jan Assmann devoted a complete book on what it means⁵² and what it stands for. Simply put, it refers to equality, justice, just treatment of men and gods etc. etc. There is even a goddess in ancient Egypt embodying this concept. Translating the exact opposite Isfet, just by one word “falsehood” is much too simple for my needs. There is a world of ancient Egyptian thoughts behind these two words, and for me to discover what these ancients thought on these subjects, I have to explore them in every detail that I can lay my hands on.

The above shows that a translator has to be familiar with the cultural context as well as the language. This is a *mandatory* prerequisite for the translator.

When you start making a transliteration and translation it is equally mandatory to have available the *most recent* literature on the language from which you are translating, either in books or on the Internet:

- A complete description of the grammar for reference.
- The most recent additions or changes to the rules of grammar.
- A complete list of signs and their most recently established meanings.
- The most recent and complete dictionary.

The study of ancient and dead languages is never fully developed. New discoveries are made regularly and you have to be able to lay your hands on them as soon as they appear in order to make your translation as good and up to date as possible.

⁵² J. Assmann. Ma’at. Gerechtigkeit und Unsterblichkeit im Alten Ägypten. München (1990).

But there are even more prerequisites for translators. It is not only the ancient cultural context that you must be fully aware of, but you must know, as much as possible, about the *background* of the particular text that you are about to translate:

- Where was the text discovered?
- What type of text is this?
- Was it part of an archive? If so which?
- Were there any previous translations? If so which?
- Was the text discussed in scientific publications? If so with what result?

In addition to these points, it is best before you start translating, to pay attention to the object itself. I suggest that you make a brief description of it, paying attention to material, measures, color and perhaps wear and tear on the object. If you have presented an image of the object you can refer to this.

Transliterating, conventions and some tips.

Translating ancient and dead languages is usually preceded by a process we call *transliterating*. In this process we ‘decode’ the ancient signs and put them in a different script especially designed for this occasion. This script shows the *character value* of the signs. The purpose of this is to *add* to the combination of signs. Both together provide information for the translating process. Most transliteration systems are on a ‘one to one’ basis i.e. one symbol represents only one (or more) regular characters at one time. But there are exceptions like in cuneiform writing. In a regular ‘one to one’ based transliteration, in principle you could translate *directly* from the transliteration script, but that may introduce a chance of making errors. The signs remain important despite the transliteration. This should be familiar material for all students, but I have used the previous lines to introduce *how* the transliteration is done.

Cuneiform has a specific format for transliteration and there are multiple conventions for transliterating Sumerian, Akkadian (Babylonian) and Hittite (and Luwian) cuneiform texts. Because the various scripts have many differences between them like characters, forms, meanings, logograms and so on, transliteration requires certain choices of the transliterating scholar, who must decide in the case of each sign which of its several possible meanings is intended in the original document.

Ancient Egyptian seems uniform in the translating process, but in practise, it is not. In Europe alone the Germans, English and the Dutch all have their, albeit slightly, different styles. It is beyond the scope of this book to go into detail about these differences.

Although there might be some form of international standardisation regarding the transliteration process, it is best to consult your university about the standards they employ.

Some tips:

- Pay attention to the style of writing on the original. A certain sign or character may be written in an unusual shape.
- In cuneiform texts, there are conventions to indicate damaged parts. There is a good wiki for that on the web⁵³. Indicate with straight hooks which parts are damaged:

┌ ┐ [A] B
AB

⁵³ http://en.wikipedia.org/wiki/Cuneiform_script.

- Akkadian cuneiform is written in italics. Sumerian can be rendered in interspaced, bold or capitalised letters.
- There may exist variants of your texts. As a rule of thumb, we suggest that you take the *best-preserved* text as the basic one. If the variants add extra text, then discuss that in the comments.
- Motivate all your actions explicitly. There may be other possibilities. Make a clear selection from the alternatives.
- Use a specialized reference library if there is one in your vicinity. In addition, there are several websites that can assist you during your translation process:
 - The CDLI is the online database for a great number of cuneiform texts. Visit: <http://cdli.ucla.edu/>. This is primarily a database and contains no translations. The *Chicago Assyrian Dictionary* can be downloaded from <http://oi.uchicago.edu/research/publications/assyrian-dictionary-oriental-institute-university-chicago-cad>. There are PDF versions of the *Concise Akkadian Dictionary* available on the internet. For Sumerian, see <http://psd.museum.upenn.edu/epsd1/index.html>;
 - Egyptologists can consult the on-line version of *Thesaurus Linguae Aegyptiae*: <http://aaew.bbaw.de/tla/index.html>. This database does contain very recent translations.

The process of translating a text

I would like to recommend the following seven steps in the translating process:

1. First provide a short description of the text itself to put it into perspective. Pay attention to the points we discussed directly above.
2. Indicate the genre of the text and discuss further background information.
3. Who was/were the probable author(s) and what was the author's /their role in history?
4. What was the purpose or intended purpose of the text?
5. Discuss previous translations and / or discussions about the text in scientific publications.
6. Make the translation, using a style-guide if possible, and pay attention to possible *ancient* errors, made by the author or by a copyist, that may turn up in the text:
 - Linguistic errors.
 - Semantics or meaning errors in the text.
 - Omission errors in the text.
 - Style- or form-related errors.
7. Justify your decision to correct what you consider to be such errors:
 - Indicate corrupt use of language in the transliteration.
 - Make a brief analysis of the errors that you found and suggest an explanation for these.
 - Motivate your corrections explicitly in footnotes.
 - Pay attention to translation errors in previous texts.
 - Draw your conclusions about your translation.

Addendum 2. Writing Catalogues of Objects

Introduction

A considerable amount of students of ancient cultures like to make catalogues for their Master thesis. I discovered that more than 10% of the Egyptology students at Leiden University favored this subject. Some students compiled catalogues listing objects from a single museum collection (often after working there as an intern for some time). Others wrote catalogues on particular objects selected from various collections (doing so from behind their own desk). Their approaches were diverse, to say the least. I felt that some guidance was needed and to make this happen I contacted the former curator of the Ancient Egyptian collection in the National Museum of Antiquities (Rijksmuseum van Oudheden) at Leiden, Maarten Raven, who is also a professor of Museology. I wrote this chapter with his help and under his guidance.

Museum policy and method

Museums differ from each other because they all have (slightly) different goals and policies about their collection. The policies cover many areas and will be dealing with acquisition, preservation, cataloguing, storage, ways of addressing the general public, making exhibitions, and so on. In general the student is not required to know them all, though if you are closely working with a specific museum it may be advisable to know at least the *mission statement* of the museum in order to comprehend their way of working and thinking.

Most museums employ their own methods when it comes to their working processes. Fortunately most museums agree on what *main data* there are to be gathered and generally on the process that an acquired object passes through from acquisition to museum display. The method outlined below is based on three main aspects:

1. The acquisition of an object.
2. The various processes that a museum employs on an acquired object.
3. The main data that should be gathered and saved.

Documentation

According to the Small Museums Cataloguing Manual "... Documenting the collection is vital to a museum's active and responsible role in managing its key asset ... enriching the collection's cultural value and enhancing its administration. ... Cataloguing underpins many important museum activities, including *research*, exhibition development, conservation, risk management, *publication* and outreach work, all of which are dependent on detailed and up-to-date collection information."⁵⁴ It is obvious that the significance of the documentation process of museums can hardly be overstated. Hence it is very important that these processes will flow smoothly and comply with international standards in order to facilitate international research. Sadly there are little internationally acclaimed standards. However there is some agreement about the *data* to be gathered and the fact that there should be a clear-cut description and general agreement on *terms* used in the documentation process. A well organized documentation is also of great benefit regarding loss, insurance and retrieval of lost objects.

⁵⁴ P11-12 Small Museums Cataloguing Manual 4th edition published by Museums Australia (Victoria) 2009 Melbourne [my italics].

The acquisition

Museums acquire their objects in different ways and they will have their own policy regarding acquisition. Since the objects that students will be writing about, are already in possession of the museum, the role of the acquisition process is of somewhat less importance for them. But there are a few important remarks to be made.

Sadly in the past many ancient artifacts had a dark background, and regrettably this may still happen today. Although the objects may be genuine, the source might (intentionally) be clouded. I.e. not according to the law regarding the protection of artifacts of the country in which the objects were found, or not compliant with the national legislation of the home country of the museum itself. Next to these *legal* obligations there are *ethical* prescriptions / guidelines regarding the acquisition of objects, as formulated by international acclaimed organizations like UNESCO⁵⁵ and ICOM (The International Council of Museums).⁵⁶ The problem of an object with a dark background is often that there is no knowledge, or has been no standardized and acclaimed *archaeological research* of the place of provenance, and hence much of the scientific data related to the object, are lost. Alternatively, the object may have been stolen from a public or private collection, or from a storeroom of archaeological finds.

Learning the *history* of an object at the museum is of prime importance. Even if there is almost nothing to know about it. Having said that, Professor Maarten Raven has one all-important ground rule: “Doubt everything you will hear and learn about the object.” This is because an object has rarely come to the museum straight from its original provenance, but has followed a circuit involving several dealers and/or collectors. In the process, much information may have been lost or made up. It’s even more than a rule that you have to question the information given. It’s a *basic attitude* that every scientist must have. Even if you do think that you know all the facts, there still is the possibility of new insight that can throw a completely new light on the object at hand. Students usually aren’t involved in these processes, but even so *before* entering your subject, it may be useful to find out as much as you can.

The inventory process

Museums typically will have their own ways of how the intake of new objects should be performed. In this process the objects are *added* to the inventory of the museum. The process in which this is done we can call *registration* or *accession*. Most museums use their own specific *terminology* in the process, and usually there are predefined categories that the museum employs. There is an essential difference between the inventory (the official registration of the object upon its arrival in the museum) and a catalogue (which is any list of objects made for a specific purpose, such as an exhibition, a scientific study, etc.; see next paragraph below).

In a museum, people use *standard worksheets* and act in a very basic and regular working order, when new objects enter the museum. Ideally, they:

Assemble basic data of the object: *assign a museum number and name* to the object.

1. Make a *description* of the object and note its *measurements*.

⁵⁵ There exists a UNESCO Convention on the Means of Prohibiting and Preventing Illicit Import, Export and Transfer of Ownership of Cultural Property of 1970 (“UNESCO Convention”).

You can download these guidelines in your own language from the UNESCO website:

http://portal.unesco.org/en/ev.php-URL_ID=13039&URL_DO=DO_TOPIC&URL_SECTION=201.html

⁵⁶ For the ICOM Code of Ethics, see <http://archives.icom.museum/ethics.html>.

2. Add details about the *provenance* and the *history* of the object, as far as these are known.
3. Mention the available *documents* acquired together with the object in question, or list the *bibliographical sources* which can shed light on its identity.

Most museums agree that a registration should at least involve the following categories⁵⁷:

- Registration date.
- Registration number.
- Object name and description.
- Acquisition method.
- Acquisition date.
- Source's name and address.
- Comments.

Nowadays the registration process is carried out via a computer and a museum will have the most suited software for its needs. Categories will have to be adapted to suit the gathered material as closely as possible. Registering archaeological objects is a special activity. We must adapt the categories or else fill out details in the “Comments” section of the software if adaptation is not possible.

All further data, which may be gathered at a later date, are attached to the assigned number and name for easy reference. Thus the inventory number (accession number) is of prime importance, because it forms the only link between the object and its documentation. Therefore, the number is usually written on the object or attached to it in the shape of a label or bar code. Scientists searching for these or similar objects can usually trace them via the available literature on the subject (which includes museum catalogues; see paragraph below). Catalogues often give their own numbers, not the official accession number! *Scientific articles must preferably give the inventory/accession number, so there can be no mistake about the identity of the object. If this number cannot be found, you may quote a catalogue number instead.* Some catalogues are better than others: try to refer to the most ‘official’ and ‘scientific’ ones and clearly state from which catalogue the number is taken. . If you are writing your own catalogue, you’ll have to present the *original* museum numbers of the objects you are working on, next to your own.

If you will work with museum objects, you will notice that the older acquisitions (say, those of the 19th and early 20th centuries) were not always inventoried in this ideal manner. In those cases, you will have to do much of the inventory work yourself (even though the object already has an official number), before you can think of writing your own catalogue. In other words: first you have to compile the basic data, and then you start arranging them for your own purpose. As stated earlier, every museum has its own practices, so there will not be *one* way of doing it right. To give examples of how the tasks outlined above might be fulfilled, we will discuss a model worksheet in the last paragraph.

⁵⁷ This list is derived from P30 Small Museums Cataloguing Manual 4th edition but is generally accepted among museums around the world.

The cataloguing process

As you might have noticed from what we have been discussing above, the processes regarding inventurisation are pretty much regulated. There is not much free room for your own format. As stated above, there is a distinction between an inventory and a catalogue. A catalogue may be compiled for all kinds of purposes: transports and loans of objects, permanent or temporary displays, lists of masterpieces for the general public or full descriptions for scientists. The inventory has to be complete; a catalogue presents *a selection* of the collection only. For writing a thesis, you will therefore make your own catalogue, as a means of recording detailed information about individual items or groups of related items. The purpose is to provide information to facilitate research and interpretation. However, whatever you do, the detailed information must be directly related to the museum objects by specifying their official museum numbers. If you are working inside a museum, or have a close cooperation with them, it's best to find out first what the museum rules and policies are with the study you're about to perform, and eventually what intentions the museum's staff members have with your study (e.g. display of the results, publication, etc.) are. Make sure that you learn and know everything what there is to know beforehand. Otherwise, you will run the risk of taking a path that is not desired and that eventually might even lead to a repudiation of your work by the museum.

If you are not cataloguing at the request of, and in close cooperation with a specific museum, there might be a little more room for self-expression, but not much more than in the former case. Scientific standards apply in each case. If the objects you are cataloguing for your own purposes belong to different museums, then you should state their museum numbers for easy reference. If your study doesn't involve museum objects, then you will have to come up with an alternative method of uniquely addressing the objects you are about to discuss.

Before you start your work, inquire from the museum staff how exactly the objects that you will be working on are recorded. It is possible that for instance a group of small statues is recorded as a group and not as individual statues.

The cataloguing process primarily exists of adding extra data to the objects already in possession of the museum. Most museums will have one or more *backlogs*. This means that they have acquired objects, but did not find the time to carry out a proper registration and cataloguing process. The professional help of students is therefore appreciated.

We will make a difference in desk studies in which you will not handle the objects and in situations where you do. If you are handling the objects, appropriate gloves are necessary so that the objects are protected from possible acids on your hands. Other objects like measuring-tape, pencil and worksheets are necessary. In most cases the registration process itself will be carried out by computer, so the form only serves the purpose of a medium to prepare the data for computer-entry. Students making their own catalogues need to prepare a data model for collecting data on the objects under description.

Exemplary data gathering form (worksheet)

It is a good practice to make standard worksheets in a text program or database on your computer, to ensure that you don't miss out on details. If you are building your own catalogue you will not know all the details to be gathered at the start of your project, so ensure that you reserve enough space to add categories and data later on.

Before I present a worksheet, some advice:

- Make use of internationally accepted terms and words to describe your object (details) to make sure that no misunderstanding can take place.
- You will find many different worksheets in manuals and on the Internet. Pick one or more that are suited to your work and, if you are working for a museum, from the museum itself.
- Every detail may turn out to be of importance.
- Stay alert on recent publications about similar objects. The aim is that others learn from you and you can learn from others.
- Complete honesty is mandatory. If some data are not known; just state that fact, and make a list of things you don't know on the object.

Worksheet ⁵⁸

*) Mandatory Field

Registration number*:	
Object name*:	
Description*:	
Keywords:	Inscriptions and markings:
Size*:	Dating*:
When/where used:	Probable maker:
Provenance (if known):	
Acquisition details*: How acquired: When acquired: Name and address of source: Comments [known history of the object, as complete as possible]:	
Condition*: _Good _Fair _Poor	
Storage location*:	
Present Location:	
Supplementary file*: Hard files: _	
Digital files: _	
Restrictions*	
Notes:	
Cataloguer*:	

⁵⁸ This worksheet is an adapted version from Appendice 1 of the Small Museums Cataloguing Manual 4th edition.

Addendum 3. Looking into the life of historical persons.

Introduction

Many students of ancient cultures are interested in the life of historical persons and study them for their Master thesis. I have found this to be one of the most favourite subjects next to “art” and “religion”. The descriptions, which I considered, were constructed from ancient documents and the students usually made a re-interpretation or retranslation of them. Their approaches had much diversity and I considered this a reason to pay attention to this topic. I contacted Prof. Dr. J. van der Vliet, who teaches Egyptology, Coptic and early Christianity at Leiden and Nijmegen University in the Netherlands. I wrote this chapter with his help and under his guidance.

Writing History

Making a description of the life of persons in history, is actually *writing history*. In addition, if that is true, then we have to abide the rules that are associated with that. The most important rules follow below.

Sources

The most important things in writing correct history are our sources. Finding the accurate sources to facilitate an analysis is not easy. I already discussed that in the first chapter of this book. In early Egypt, we had the good fortune of discovering a large amount of ostraca, papyri and inscriptions. Although the ostraca (messages contained on a flake of stone or potsherd) are relatively short, they provided detailed information.

Regarding our subject, we suggest a division into two types of sources:

- A. The literary source.
- B. The non-literary source.

Ad. A. The literary source.

Studying literary sources is a discipline in its own right. We have to consider ourselves with philology as well as literature study. Philology occupies itself with the explanation of language and scriptures of a people; often within the broad framework of culture and culture history. Literature study⁵⁹ is focussed on form, structure and construction of the source document we are considering.

The aim is to get as close as possible to the original (intended) *meaning* of a document. When studying Egyptology, there are for instance 4 different versions of spell 18 in the “Book of the Dead” (A recitation for the day of burial). The object is to determine what is *originally* intended. A text out of a literary source had a special purpose and it was meant to be “passed on”. The elements from the text can be traced and analysed.

⁵⁹ A well-known author on the subject is Mieke Bal. She wrote a much-acclaimed book entitled: “Narratology, introduction to the theory of narrative.” [Transl. from the Dutch by Christine van Boheemen], Toronto, 1997.

To make matters a bit more complicated, there is a difference nowadays between *traditional* and *new* philology. The new philology is a school of thought within (traditional) philology, trying to recreate the history of a certain period in time, making use of the written documents from that culture and reconstructing the way in which people viewed their own society. An important historian within this tradition, who is actually its founder, is J.M. Lockhart.⁶⁰ Up to then, the traditional philology, only studied *sources* in order to understand certain *events*. The new philology goes further than that and tries to understand the *culture* that produces these texts, itself. This type of thinking is important for us because we can now try to reconstruct how the authors tried to reach their goals and what they originally intended to do (how they wanted to influence their public). In other words we may acquire insight in the processes in (certain circles of) society that were involved at the time. Texts are not documents in its own right. They are now viewed as *means* in societal *processes* to reach certain goals.

Genre

A description about the life and work of Christian historic people is actually a literary *genre* called *hagiography*. According to an article in the Encyclopaedia Britannica,⁶¹ “hagiography, the body of literature describing the lives and veneration of the Christian saints. The literature of hagiography embraces acts of the martyrs (i.e., accounts of their trials and deaths); biographies of saintly monks, bishops, princes, or virgins; and accounts of miracles connected with saints’ tombs, relics, icons, or statues. Hagiographies have been written from the 2nd century AD to instruct and edify readers and glorify the saints. ... The importance of hagiography derives from the vital role that the veneration of the saints played throughout medieval civilization in both eastern and western Christendom.”

This literature preserves much valuable information not only about religious beliefs and customs but also about daily life, institutions, and events in historical periods for which other evidence is either imprecise or non-existent.

The Encyclopaedia continues “... The hagiographer has a threefold task: to *collect* all the material relevant to each particular saint, to *edit* the documents according to the best methods of textual criticism, and to *interpret* the evidence by using literary, historical, and any other pertinent criteria.”⁶²

As we observe this text a little closer, we can judge this definition to be limited because it mentions only people associated with Christianity, and there were leaders of *other* religions that had a profound influence on their environment as well. According to Felix⁶³, this literary genre was often used as ecclesiastic and political propaganda. Nowadays these works about the lives of Christian leaders give us a valuable insight in the culture in which they operated. They were inspirational stories and told legends about these people. These hagiographies were also used to ‘sanctify’ people.

Nowadays hagiographies are not written anymore. However studying them, is still very fashionable and is actually done more often in the last 20 years or so. If properly analysed, hagiographies can shed light on the ongoing civilisation back then and offers a great opportunity to get a grip on important values that governed life in the period. A warning is in place here. Studying Hagiographies is studying literature (e.g. narratology) and one has to be *very critical*

⁶⁰ He is an expert in the study of historical sources in the Nahuatl language and the postcolonial Nahuatl people. He is now a professor emeritus at UCLA.

⁶¹ <http://www.britannica.com/EBchecked/topic/251586/hagiography> consulted at 4-3-2016.

⁶² My italics.

⁶³ Propaganda in Hagiography By Theodore Felix, 20 August 2008; Revised - All Empires History

about the information presented and value the content with great care. You are reading ‘filtered’ information. To present an example: one hagiography claims that the person described was “a shepherd” in early life. Another hagiography of someone else 200 years later, claims the same. Do not take that literally! The item can be traced back to a biblical king ‘David’ who was a shepherd. These kinds of statements are called a *Topos*: regular habitual description items. As far as the hagiographers are concerned, *every* saint once was “a shepherd”. This symbolises their ability to ‘lead a flock’.

Another important genre is the *autobiography*. They are to be found in many ancient cultures. An important function of this is to *demonstrate the outstanding life* that the person in question, has lived. This in turn, is an important factor according to the social norms at the time the biography was written. This presupposes an excellent knowledge of the ancient society, religion and culture in particular and is prerequisite to the task of learning more about these people.

Ad. B. The non-literary source

These types of sources contain *uncategorised* texts of a society that come to us in a random way. The selection of these sources is based on *coincidence*. These texts are fragments or reflections of a society gone by. Almost every kind of written material is suited to be analysed and even the most trivial messages may shed light on what we are looking for: *information* about our subject that could be worth something. For instance, we know of an ostrakon in ancient Egypt, in which a person asks to lend money from a bishop. We start analysing immediately: normally such a question would be outrageous for a commoner to ask a bishop. We have to assume that the person asking this question was at least in a client – patron relationship to that bishop. As you can see, even this small piece of “conversation” conveys quite a bit of information. These texts are mostly *unique*, and you can find other similar but not identical texts as would be the case where literary sources are concerned. This research could benefit much from *social network analysis*, a type of analysis, which we will discuss below. In Egyptology, we think of sources from Deir el-Medina and in Assyriology of randomly found tablets not belonging to a certain kind of archive. Especially the kinds that convey short messages about daily life are the most interesting in this category.

Method

As outlined above, we now consider the description of only ‘the acts and the facts’ of historic persons to be old fashioned and outdated. This is no longer the way to write history anymore. The *significance* of their acts must be traced against the structure, rules, customs and constellations of their ancient society. In addition, of course, we are not just summing up *facts*. We are actively exploring a *research question*, which we have thought over and stated beforehand. Your research therefore must have a *direction* and your task should be to actively *explain*, why events are represented as they are. As a student, you must operate from within a scientist’s view, and always ask the right questions about the subjects that become known in your research. Your method of working should also be *dynamic*. There are connections between the research question, your sources, your data collection and the results. These are all intricately connected. Change in one element implies changes in the others. You are allowed to make modifications in all as long as you state your reasons for that. However, as I explained earlier in this book, *not* for adapting the results to your liking. Be careful with this be-

cause your credibility as a scientific researcher is at stake. You have to accept your research outcome, but are allowed and even encouraged to explain “strange” or unexpected phenomena.

Research

Before you start analysing, you should perform a research on sources first:

- Indicate the *kind* of source that you will be using and the associated types of analysis.
- Discuss any *peculiarities* about the source.
- Indicate why you are using this source and what you expect to find. When analysing a *literary* source like a hagiography, you should always pose at least these questions:
 - By whom?
 - Intended for whom? And,
 - What was the (original) purpose?

These texts were *always* written with a certain *purpose* in mind. Important to discover is what that was and what it meant or, what rationale it served, in the society back then. The most essential fact about these texts is their *function*. There are quite a few of functions. One main function, in case of the bishops is, “to bring piety”. Not the birthplace of a bishop is important but his *actions* are. The question is “What actions did the bishops carry out and why?” “What picture of a bishop emerges”. “Can we reconstruct a ‘bishop’ historically?”

When you start analysing non-literary sources, using *your own knowledge* of the ancient society, you should be able to reconstruct several of the *social networks* in which the subject person was involved. A social network analysis can prove to be fruitful. A large number of available materials may paint a vivid picture of all the connections a certain person had. Not only the *contacts* spring to mind but there is an indication of the *actions* a person performed as well. Prerequisite of making such an analysis is an in depth knowledge of the ancient society and of the background of the original sources. Precondition is that there must be sufficient related material available.

The main object of our research must be to interpret the events we are analysing with the view of the people of society back then. As indicated before in chapter one, this is not easy because, as a modern scientist, we are inclined to look at matters *our* way.

Social Networks

Numerous people in history were *known*. In order to become known, many people will have to be acquainted with these persons and report on them. This indicates that these people were probably a member of many different groups in society. We now call them “social networks”. A social network puts all their members in a certain *association* towards each other. One way of representing this, is that it looks like the circles originated from a stone thrown into the water. The waves are concentric, small and high nearby, large and dying out, the further away.

The point of origin (the stone) is the *person* and the waves represent, for example:

- Their living quarters.
- The organisation from which he is a part.
- Society and perhaps ...
- The entire world.

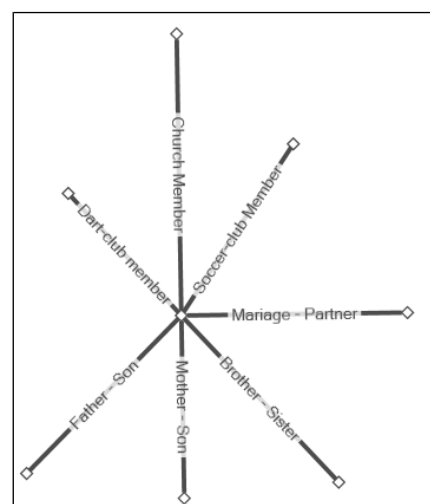
The networks, of which a person was part of, are very important in understanding and interpreting their behaviour. Social Network Analysis or SNA, can make a difference when evaluating a person's actions, and is therefore highly recommendable. When talking about important historical persons, you have to know, what their networks looked like and how the communities in which he worked, functioned.

Social Network Analysis

The study of social networks can serve as a valuable extra method of discovering the relationships of a person that can bring new information to the surface. This type of analysis can draw our attention to streams of information and dependencies thus providing innovative insights about our subject of analysis.

There are many kinds of networks, like neural networks, computer networks, health networks (for the analysis of spread of a disease) and so on. Social networks belong to the group of networks too. A sociologist, J.A. Barnes, started the term in 1954, to denote patterns of ties within groups like families and social identifying categories like gender. The theory of social networks is based on a *general* network theory and because of that, the terms that are used seem to belong more to the computer world, than that of real humans. There are not so many terms to explain, and it is not that complicated but the analysis later on could be. Unfortunately, there may be variations in the names given to analytic categories, depending on the system / software you can apply / use. I will explain using one method.

A social network could exist between individuals, groups or organisations. We will call all these elements *nodes*. Since we are interested in persons, the individual will be our smallest unit or node. That does not mean, that individuals could not be related to groups or larger units. Of course they are! The relations that exist between the individual and a group or between a group and an organisation we call a *tie*. Therefore, a (social) network is made up of nodes and ties. Ties are also called: links, connections or edges. The nodes represent the *actors*, i.e. the individuals, inside the network and the ties represent the *relationships* between these actors. Lines in a diagram represent all these relationships, the ties. These are the fundamental terms or building blocks of a social network. In its basic form, a picture of ties, which are to be specified, represents a social network.



Example Social Network

An example: John (our node) is tied to his father (father-son relationship) to his mother (mother-son relationship) and to his sister (brother-sister relationship). They are all within the same group: the *family*. He is also tied to the *church* of which he is a member (church-member relationship), to the *soccer-club* for which he plays (soccer-club-member relationship) and to his friends at he *pub* in which he plays darts (friends-individual relationship). As we can observe, there are many *different* associations, with diverse *intensity* and/or *importance*.

Later in life, John's networks will change. Colleagues will join in the relationships and the relationship with his boss will gain in importance; all within his *work*. As John marries, his *nodes* will change even further. He will get his own family and his relationships will be extended with a mother, father, brothers, and sisters in law.

What I mention in my example, is actually quite common and if you make a graphical representation of all these nodes and ties (each type of tie represented by another colour), the resulting picture can become quite complicated. You can use a different colour in the diagram to represent each *type* of relation, but you can also assign a value to the *importance*, i.e. "weigh", the relationship in respect to the object you are researching, instead. You could, for instance, measure the *contribution* of the relationship to the political or spiritual power that the person you are analysing, has in possession. If you assigned a *value* to this contribution, you would have to explain how you did it of course, you could drop all relationships below a certain level, to simplify your diagram and hence analysis.

The importance of a social network analysis lies within the basics. It cuts right through traditional associations of which we once thought that they are important, like family, class, church and soccer club. Because of the technocratic way the analysis is executed, quite another picture emerges from, what seems to be, a different world. It places matters into a very different perspective that can be very refreshing and useful for our analysis. A 'normal' person could be just four or five levels away from the president of your country! The 'traditional' view on society loses its importance. However, keep in mind that we have to have sufficient material from our sources. Otherwise performing this kind of analysis can be very hard to do.

Many different programs on the computer support social network analysis. One of the easiest and cheapest ways of attaining such a program is by inserting a so-called "plug-in" in the spreadsheet program Excel. Attainable for the version 2007 and up⁶⁴. However, you will have to invest some time in learning how to work with the program.

Briefly, this is what you need to know, about social network analysis in order to construct your own description of the life of historical persons⁶⁵. The value of the process of making an analysis is, that it gives you *clues* and material to think about and to elaborate on, in your description and analysis of the person you are interested in. You can make it as complicated or as simple if you want. The most important thing is that it has to bring you a contribution to the analysis you are performing.

⁶⁴ You can download it free at <http://nodexl.codeplex.com/>. There are several guides on the web too, to help you learn to use this plug-in.

⁶⁵ Of course, there is much more to be discussed about this topic. If you are interested and want to learn more, the following book is considered among many authors to be of prime importance: Wasserman, Stanley / Faust, Katherine: *Social Network Analysis. Methods and Applications*, Cambridge, University Press (2008).

An example of what is intended is to be found in a recently published book by an American Associate Professor of History and Classical Studies at Fairfield University and papyrologist, Giovanni R. Ruffini. He performed social network analysis on ancient Egyptian papyri (in Greek).⁶⁶

Arrangement of your description

In an earlier addendum, I suggested the outline of a general research report:

1. Introduction.
2. Research question.
3. Delineation.
4. Method & Operation.
5. Data collection.
6. Organizing and analysing data.
7. Drawing conclusions: answering the research question.
8. Writing the report / master thesis

Below you will find a *suggested* arrangement of your report or analysis based on the general outline. The format is 'free' but I recommend that you pay attention to all the elements noted below and let them return in a *logical* place within your report.

1. Introduction, research question & delineation

- Explanation of your choice of subject.
- Earlier research into the subject and a brief evaluation.
- State your research question; what are you after or hope to find or display?
- What exactly will you be studying? (set the boundaries of your research field)

2. Method & operation

- State your working method.

3. Data collection

- Sources
 - What kind of sources will you study and why?
 - Literary sources (hagiographies).
 - Non-literary sources.
 - How much and what sort of texts have been written about the person of your choice?
 - Evaluation of your sources. (provenance, quality, other attributes)
 - If necessary (re-)construction or retranslation of the text(s).

⁶⁶ In his book: *Social networks in Byzantine Egypt*, Cambridge/New York: Cambridge University Press, 2008.

- Society; roughly, how does society function around the chosen figure on three levels:
 - Macro-level (society).
 - Meso-level (organisations or other, people related, structures).
 - Micro-level (the level of the individual).
 - Pay attention to the position or place of your chosen subject within these spheres of life.
 - Describe and analyse the social network of the chosen person.

4. Organize and analyse your data.

- Fashion your data in such a way that you can infer conclusions from them.
- Make a final judgment about the quality of your data.

5. Conclusion

- Give an answer to the research question.
 - Are you able to do that (else: why not?)
 - Draw your final conclusion
 - What is needed in the future

6. Writing the report / master thesis

Addendum 4. Research practice, reporting and follow-up

Research practice

Carrying out a research, usually on your own, is not something to be taken lightly. In the paragraph above, you can see what steps there are to be taken. Most students do not plan ahead and see for themselves how much time they need, to read, research and write the final report or master thesis. It is possible to plan out everything in detail, but chances are that you'll not be able to keep your time planning. Therefore generally most research is cut into a few phases and these are roughly planned. Usually by assigning time scheme to the phases.

Consider the steps outline above, to which I added step 8. writing the report:

1. Introduction.
2. Research question.
3. Delineation.
4. Method & Operation.
5. Data collection.
6. Organizing and analysing data.
7. Drawing conclusions: answering the research question.
8. Writing the report / master thesis

We could roughly divide those steps into three phases

- I. Orientation phase: step 1 thru 4.
- II. Data collection phase: step 5
- III. Analyses and report phase step 6 thru 8.

Alternatively, you could add an extra phase:

- I. Orientation phase: step 1 thru 4.
- II. Data collection phase: step 5
- III. Analyses phase: step 6 thru 7.
- IV. Report phase: step 8

Unfortunately there are only a few guidelines that help you make up a planning. The only thing that you can be sure is that the data collection phase is by far the longest phase in your research. The advantage of making a planning is that you can keep track of your own advancements. Keeping track of your progress, can offer you *incentives* at times that you are not working as hard as you may wish. Nowadays, studying becomes increasingly expensive and making sure that you can write you master thesis in one year, saves the expenses of an additional year. Planning comes to the rescue! Setting time frames for your research phases may limit you fully justified. Most of the students in their master phase start writing their master thesis at the beginning of the new year. Depending on your university your work has to be handed in at the very end of the academic year, i.e. end of august. That gives you about eight months of planning time. Make sure of the exact dates of your university, before you start planning.

Because of the difficulty of making an estimate for the data collection and analyses phase, I will base my planning on the two phases that I *can* manage. The first and the last phase, and I will explain why. Considering the first, orientation phase, the most important part is the *research question*. You should spend the most time in this phase on that, because it determines all that happens afterwards. I planned six weeks for the entire phase, starting at the very beginning of that year. Need extra time? Earn that by starting early in the last months of the previous year, work harder in the next few weeks or take a bit out of your data collecting phase II. See to it that you start on time with phase III Analysis. I cannot determine exactly how much time you need for the phase II. But I do know that your report has to be handed in at the end of august. I will plan 8 weeks for writing and finishing the report. That has to be more than enough for you to finish it, provided that you do wrote already down what you did during all the other phases. And of course, there is time planned for you to consult your supervisor. The exact planning for the middle phase(s) is not very interesting, provided that you end the last middle phase on time. I recommend that you drop your last detail analysis, if you are still busy in the last week, and start writing your report at planning week 26. Maybe if you are quick at this, you are able to earn some extra time, so that you can perform your “lost” detail analysis. Otherwise you have to settle for what you’ve got. May be I should let you in on a tiny secret; unless you are an A-plus, plus, plus student, the grade you’ve earned isn’t important anymore on the job market.

The planning scheme below is only *my suggestion*. Make it as you wish, it’s your responsibility. One final word about this planning thing. Get used to it! In your job to come you will be confronted with it all the time. Better start early or ASAP, as soon as possible! You should also start writing almost immediately so as to get a grip on your subject quickly.

Planning scheme

Phase	Planning week nr.
I. Orientation phase	01 thru 06
II. Data collection phase	07 thru 21
III. Analyses phase	22 thru 25
IV. Report phase	26 thru 32

Reporting your research

Your report has to offer an account of all the steps mentioned in the previous sub-paragraph. The decisions you’ve made, the data you’ve gathered, the analysis you undertook and the conclusions you’ve drawn. All within a logically stated and readable text, that complies with all academic rules about the use of theories, remarks from other authors, research methods, the use of books, notes and the like.

The best thing to do is make an *outline* at the very start of your research, during the orientation phase. An outline is a sort of design of the full report, a framework that will contain the texts that shall be written. You can change that as much as you like during your research, but you will have to have a plan in advance of what you are striving for. Given the subjects of your report or master thesis, most of the chapter titles are already known, and can be drawn from the list above.

The complete orientation phase can be written at the start of your research project. Make sure that you record previous *editions* of all your texts with the aid of *version numbers*. In this way you can always fall back on older texts if you happened to cross out a part that you later may want to reinstall. Make provisions for the dynamic nature of your research, so that you can easily change texts. Nothing is permanent unless you feel comfortable about your end-design of your chapters and paragraphs. Only at that time can they be made permanent.

Follow-up on your research

There are quite a number of reasons why your research, may not stay static once you have published this in the form of a research report or a master thesis. They are:

1. Replication.
2. Cohort study.
3. As a bases for additional study.
4. Need for updating / improvements.

Ad 1. Replication

Replication research, in which another researcher *repeats* your research with all the same data, is not very well known or used within the sciences of the ancient languages and cultures. There are two main reasons for this, which lie in partly line with each other:

- Limited research budgets lead to different priorities. Replication usually is of low priority.
- The research field still contains many white spots. Usually the priority is focused on removing as many white spots as possible.

Nevertheless if the science keeps on developing, we might expect to see some replication research soon, particularly when the subject of the research has or gains momentum in scientific circles. This means that you must keep your research data available for this purpose even after you have closed your research and written your final report.

Ad 2. Cohort study

In a literary sense a cohort study is a study whereby people from a certain age (group) are followed over the years. If I do not take that quite literary, and apply it on research over time, I can imagine that you carry out the same study after, say, ten years to see whether the new and improved state of the art of your science, allows you to come up with a different result. I don't think that this will be popular but, certainly if the importance is considered high enough, I do not rule out that this could be happening. A lot of research is based on the research of others and in the course of time, the science improves itself in recent literature.

Ad 3. A base for additional study / studies

It might even be so that your research forms the basis of one or many additional research (-es) to come. It might even be you, yourself that triggers that of. May be your master thesis was sufficiently valid of laying the groundwork for additional study. If you researched a basic category of objects, the objects themselves might be the aim of further exploration.

Ad 4. Need for updating / improvements.

A science that stays in the interest of scientists and the general public never stays static. There will be improvements continuously. The reasons for that are:

- new research methods;
- new theories;
- new discoveries.

It is obvious that those reasons might motivate you, yourself, or others to look at your research results anew and try to find out if these remain valid or must be improved upon. A good master thesis can be improved by turning it into a PhD thesis.

Consulted literature & Digital sources

Consulted literature

ENGLISH SECTION

Babbie, Earl, 'The Practice of Social Research', 10th edition, Wadsworth, Thomson Learning Inc.

Blumer, H. *Symbolic Interactionism; Perspective and Method*. (1969) Englewood Cliffs, NJ: Prentice-Hall.

Denzin, N., *Sociological Methods: A Sourcebook*. Aldine Transaction. (2006).

Elias, Norbert., *Problems of Involvement and Detachment*. The British Journal of Sociology, Volume 7 Issue 3 (1956).

Glaser B.G., Strauss A.L. *The discovery of grounded theory; strategies for qualitative research*. New York (1967).

Hempel, C.G. and Oppenheim, P, *Studies in the Logic of Explanation*, Philosophy of Science, Volume 15. (1948).

Hindess, Barry, *The Use of Official Statistics in Sociology*. London (1973).

Holsti, Ole R.: Content Analysis for the Social Sciences and Humanities. Reading, MA: Addison-Wesley 1969.

Kaplan, Jim The Auditnet Monograph series: How to use statistical sampling. AuditNet (2003).

Levin, Michael., *What Kind of Explanation is Truth?* In Jarrett Leplin. , Van Fraassen, Bas., (1980). *The Scientific Image*. Oxford (1980).

McCall, Robert B, *Fundamental Statistics for Behavioral Sciences 7th Ed.* 1998 Pacific Grove CA USA.

Popper, K., *The Logic of scientific discovery*, Taylor & Francis e-Library, (2005).

Russell, Bertrand *The History of Western Philosophy*. New York (1945).

The journal *Literature of Liberty: A Review of Contemporary Liberal Thought* , vol. II, no. 1, January/March (1979).

The Oxford Companion to the History of Modern Science. New York (2003) Oxford University Press.

Wilson, Edward O., *Consilience: The Unity of Knowledge* (1st ed.). New York, (1998).

DUTCH SECTION

Berselaar V., *Wetenschapsfilosofie in veelvoud*. Bussum (1997)

Boer D.J., Bouwman H., Frissen V. & Houben M., *Methodologie en statistiek voor communicatie-onderzoek*, Houten/Diegem 1994.

Colder, J.C. m.m.v. E. Nuyten-Edelbroek, *Het winkelcentra project*, WODC - SDU, Den Haag 1988

Ende, H.W. v/d, *Beschrijvende statistiek voor gedragswetenschappen*. Amsterdam/Brussel 1973.

Ende, H. v/d / Verhoef M., *Inductieve statistiek voor gedragswetenschappen*. Amsterdam/Brussel 1973.

Segers, J.H.G., *Sociologische onderzoeksmethoden*. Assen/Amsterdam (1997).

Verschuren P.J.M.. *De probleemstelling voor een onderzoek*. Utrecht 1992.

GERMAN SECTION

Marx, K., *Das Kapital*, Hamburg (1867).

Digital sources

The following digital sources on the net are all viewed, checked and tested on: 8-8-2014.

<http://www.acastat.com/Statbook/normaldist.htm>.

<http://www.berrie.dds.nl/calcss.htm>.

<http://www.bjmath.com/bjmath/Stats/sd.htm>.

<http://www.britannica.com/EBchecked/topic/528804/philosophy-of-science/271823/Scientific-change>.

<http://www.britishmuseum.org/research.aspx?ref=header>.

<https://www.boundless.com/statistics/measures-of-variation/describing-variability/standard-deviation-definition-and-calculation/>

<http://en.wikipedia.org/wiki/Science>

http://en.wikipedia.org/wiki/Standard_deviation.

<http://www.iep.utm.edu/>

<http://www.iep.utm.edu/hempel/>

<http://www.marxists.org/>

<http://www.mathsisfun.com/data/standard-normal-distribution.html>

<http://www.mathsisfun.com/data/quincunx.html>

<http://math.ucr.edu/home/baez/physics/General/occam.html>

<http://mathworld.wolfram.com/NormalDistribution.html>

www.nao.org.uk/publications/Samplingguide.pdf

<http://nodex1.codeplex.com/>

<http://oxforddictionaries.com/definition/english/science>

http://www-personal.umich.edu/~johnpi/stat/module_mean.htm

<http://philosophy.hku.hk/>

<http://philosophy.hku.hk/think/>

<http://philosophy.hku.hk/think/m/quiz.php>

<http://plato.stanford.edu/>

<http://www.qualitative-research.net/index.php/fqs/article/view/466/996Volume>

<http://www.random.org/integers/>

<http://www.raosoft.com/samplesize.html>

<http://stanford.library.usyd.edu.au/archives/spr2009/entries/scientific-realism/>

<http://strangebeautiful.com/other-texts/popper-logic-scientific-discovery.pdf>

<http://stattrek.com/descriptive-statistics/variability.aspx>

<http://stattrek.com/statistics/dictionary.aspx?definition=Variance>

<http://www.techrepublic.com/blog/10things/10-tips-for-sharpening-your-logical-thinking/2590>

http://wiki.uva.nl/kwamcowiki/index.php/Gefundeerde_theoriebenadering