

INTERPRETING QUANTUM MECHANICS

LARS-GÖRAN JOHANSSON

INTERPRETING QUANTUM MECHANICS

Presenting a realistic interpretation of quantum mechanics and, in particular, a realistic view of quantum waves, this book defends, with one exception, Schrödinger's views on quantum mechanics. Johansson goes on to defend the view that the collapse of a wave function during a measurement is a real physical collapse of a wave and argues that the collapse is a consequence of quantisation of interaction. Lastly Johansson argues for a revised principle of individuation in the quantum domain and that this principle enables a sort of explanation of non-local phenomena.

This is a fascinating book and I learnt a lot from reading it. Johansson introduces a new, highly interesting definition of realism and carries it through in a discussion of the major issues in the philosophy of quantum mechanics. In doing this he adds important new insights in the subject. I strongly recommend this book to all who are interested in the philosophy of physics.

Sven Ove Hansson, KTH (Royal University of Technology) Sweden

Interpreting Quantum Mechanics argues for a revival of a revised version of Schrödinger's assumption that the quantum physical reality ultimately consists of waves. The book convincingly argues that an ontological program is in fact possible and that it is the first step towards a realistic interpretation of QM. The presentation and discussion is carried out both on a technical and an intuitive level and is therefore accessible to philosophers without extensive background in physics, as well as physicists who do not have a philosophical background.

Martin Edman, Umeå University, Sweden

For Inga, Liv, Hanna and Stina

Interpreting Quantum Mechanics

A Realistic View in Schrödinger's Vein

LARS-GÖRAN JOHANSSON
Uppsala University, Sweden

 **Routledge**
Taylor & Francis Group
LONDON AND NEW YORK

First published 2007 by Ashgate Publishing

Published 2016 by Routledge
2 Park Square, Milton Park, Abingdon, Oxon OX14 4RN
711 Third Avenue, New York, NY 10017, USA

Routledge is an imprint of the Taylor & Francis Group, an informa business

Copyright © Lars-Göran Johansson 2007

Lars-Göran Johansson has asserted his right under the Copyright, Designs and Patents Act, 1988, to be identified as the author of this work.

All rights reserved. No part of this book may be reprinted or reproduced or utilised in any form or by any electronic, mechanical, or other means, now known or hereafter invented, including photocopying and recording, or in any information storage or retrieval system, without permission in writing from the publishers.

Notice:

Product or corporate names may be trademarks or registered trademarks, and are used only for identification and explanation without intent to infringe.

British Library Cataloguing in Publication Data

Johansson, Lars-Göran

Interpreting quantum mechanics : a realistic view in

Schrödinger's vein

1. Quantum theory

I. Title

530.1'2

Library of Congress Cataloging-in-Publication Data

Johansson, Lars-Göran.

Interpreting quantum mechanics : a realistic view in Schrödinger's vein / Lars-Göran Johansson.

p. cm.

Includes bibliographical references and index.

ISBN 978-0-7546-5738-5 (hardcover : alk. paper) 1. Quantum theory. I. Title.

QC174.12.J639 2007

530.12--dc22

2006100194

ISBN 13: 978-0-7546-5738-5 (hbk)

Contents

List of Figures and Tables
Preface

ix
xi

1	Interpretative Problems in Quantum Mechanics	1
1.1	Historical background	1
1.2	What is an interpretation of a theory?	2
1.3	Four core problems of interpretation of quantum mechanics	4
1.3.1	Interpretation of the Born rule	4
1.3.2	Wave particle duality and interpretation of the wave function	6
1.3.3	The measurement problem	7
	Schrödinger's cat	8
	Wigner's friend	9
1.4	Non-local interactions	10
2	Realism and Understanding	13
2.1	Realism	
2.1.1	Realism and quantum mechanics	13
2.1.2	Definition of realism	14
2.2	Interpretation and understanding	18
2.3	Understanding	19
2.4	The semantics of quantum mechanics	24
3	Individuation and Identity in the Quantum World	27
3.1	Introduction	27
3.2	Individuation and Identity	29
	Quine's view	29
3.3	Identity criteria	32
3.4	Individuation and Identity in the quantum domain	33
3.5	Individuation and the Identity Postulate	34
3.6	The Identity Postulate and Identity of Indiscernibles	36
3.7	Individuation by different properties	38
3.8	Cardinal and ordinal numbers	39
3.9	On the individuation of systems	41
3.10	Coupling and decoupling of systems	42
3.11	The state of a system	44
3.12	Observables, operators and properties	45
3.13	Actual and potential properties	47
3.14	Probability in Quantum Mechanics	49

3.15 Identity of states: real change of state versus change of description	50
Wigner's theorem	51
3.16 Schrödinger evolution with changing Hamiltonian	55
4 Quantum Objects are Waves	57
4.1 Introduction	57
4.2 Interpretation of the wave function	58
4.2.1 Schrödinger's idea	58
4.2.2 Rebuttals of these arguments against the wave interpretation	60
4.3 Adiabatic measurements	67
4.4 Matter wave theory	68
4.5 Wave packets	70
4.6 Empirical results supporting the causal relevance of partial waves	71
4.7 Dispersion of wave packets	74
4.8 Doppler effect of partial waves	75
4.9 Interaction between matter waves and electromagnetic waves	78
4.10 Space-time metric and complex phases of wave equations	80
4.11 The path integral picture and the wave interpretation	82
5 Particle Behaviour of Waves	85
5.1 Introduction	85
5.2 What is meant by quantisation?	85
5.3 Quantisation in other processes	87
5.4 Quantisation of interaction between waves explains the particle aspect of waves	88
5.5 Some examples of particle behaviour	90
5.5.1 The Compton effect	90
5.5.2 Particle tracks in cloud chambers	91
5.6 Calculations of the size of the electron in collisions	91
6 The Measurement Problem	93
6.1 Introduction	93
6.2 Measurements	95
6.2.1 Measurements as attributions of numerical values to objects	97
6.2.2 Measurements are energy exchange processes	100
6.3 Collapses	106
6.3.1. The non-linearity of the collapse	107
6.4 Discontinuity	108
6.5 Randomness	109
6.6 Irreversibility	110
6.7 Irreversibility as a result of randomness	112
6.8 Collapse as loss of coherence	113
6.9 Superposition of macroscopic states	113
6.10 Leggett's argument	114

6.11	Interaction-free measurements	117
	Comments	120
6.12	Renninger's negative result experiment	121
6.13	Delayed choice	123
6.14	Quantum Eraser	124
7	Quantum Mechanical Spin	129
7.1	Introduction	129
7.2	A visualisable model of spin-half objects	129
7.3	A vector model of spin-half	130
7.4	The Kochen–Specker Theorem	136
8	Non-locality	139
8.1	Introduction	139
8.2	Derivation of non-locality	141
8.3	Bell's theorem	143
8.4	Bell locality and Einstein locality	145
8.5	Interference of wave packets	146
8.6	Einstein locality and relativity theory	150
8.6.1	The notion of frame of reference and STR	150
8.6.2	Violation of Einstein locality and relativity theory	151
8.6.3	Collapse of the wave function in different frames of reference	152
8.7	Collapse and phase velocity	152
8.8	Violation of locality in interactions between two two-level atoms	154
8.9	The Aharonov–Bohm effect	155
8.10	Conclusion	157
9	Gentle Criticism	159
9.1	Introduction	159
9.2	The Copenhagen interpretation	159
9.3	The measurement problem in the Copenhagen interpretation	162
9.4	The many-worlds interpretation	164
9.4.1	Criticism of the many-worlds interpretation	164
9.5	The de Broglie–Bohm interpretation	166
9.6	The GRW Theory	169
9.6.1	Albert's criticism	170
9.7	The decoherence theory	172
10	Summary and Conclusions	175
10.1	Realism and Quantum Mechanics	175
10.2	The equivalence between matter and energy	175
10.3	Objects in quantum mechanics	177
10.4	Wave–particle duality and complementarity	177
10.5	Historical remarks	178

10.6 From orthodox quantum theory to quantum field theory	182
10.6.1 Individuation in QFT	182
10.6.2 Fields and events – the whole and its parts	183
10.7 An unsolved problem	185
<i>Bibliography</i>	<i>187</i>
<i>Index</i>	<i>193</i>

List of Figures and Tables

Figure 1.1	A delayed choice experiment	6
Figure 1.2	Schrödinger's cat experiment	9
Figure 4.1	Principal arrangement of time-of-flight neutron interferometry. The recombined neutron beam will go to the upper or lower detector depending on the turning of the first phase shifter	72
Figure 4.2	The first curve shows the relative contrast in the entire neutron pulse as a function of the thickness of the inserted Bismuth sample	73
Figure 5.1	Compton scattering	90
Figure 6.1	A measurement on spin-half particles	101
Figure 6.2	Particles with spin up in the x-directions pass two S-G magnets, the first oriented in the y-directions and the second in the x-direction	102
Figure 6.3	Reuniting two separated beams will result in the revival of the original spin state	103
Figure 6.4	Preparation of a definite spin state	105
Figure 6.5	Interaction-free measurement of the position of a 'bomb'	118
Figure 6.6	Renninger's negative result measurement	122
Figure 6.7	A Hong-Ou-Mandel interferometer	124
Figure 7.1	A vector model of spin half	131
Figure 7.2	Vector model of a S^+ eigenstate rotating around a B_z -field	134
Figure 8.1	The function $ 1 + \cos \varphi + \cos \varphi - \cos 2\varphi $. There is a contradiction between Bell's inequality and quantum mechanics when the correlation function is above the line $y = 2$	145
Figure 8.2	Two wave packets travelling in opposite directions	148
Figure 8.3	The function $\phi(k) = \exp\left[-\frac{(k - k_0)^2}{a}\right] 2 \cos kx$ for given $x \neq 0$ This figure represents two wave packets as seen in k -space. The scale is arbitrary	148
Figure 8.4	A principal set-up for testing the AB-effect	156
Table 4.1	Comparison between the motions of a vibrating plate and quantum mechanics of a free particle	64
Table 8.1	The possible values of the function g_n	143



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Preface

We feel the need for interpretation when confronting things we don't understand. Quantum mechanics is one such thing and many people have written books about the interpretation of quantum mechanics. On one hand this is astonishing, since we (or at least physicists) understand the theory well enough in order to use it for making predictions and solving physical problems. In this respect it is well understood and no one doubts that it is very nearly correct; so far as I know there is not the slightest evidence against quantum mechanics. So why, then, does this theory still, 70 years after its inauguration in the 1920s, make us wonder?

The interpretative problems about quantum mechanics have to do with the difficulties in uniting a realistic stance with quantum mechanics. One worry has been the measurement problem: why do quantum systems behave differently when we observe them? Physics is about objects and events existing in space and time and, no matter how strange the micro-world is, the theory talks about real things and these things either have, or have not, the properties ascribed to them by quantum theory. Since measurements are physical interactions, it is disturbing that these events cannot be given the usual quantum mechanical analysis. A realist wants to describe how the world is in itself independent of human interests or perspective, and this has proven difficult – some would say impossible – to achieve. In this book I have tried to do precisely that. Other problems for people with realistic convictions are wave–particle duality and non-local phenomena.

Several physicists and philosophers have proposed stronger or less strong realistic interpretations of quantum mechanics. The usual recipe is to add something to the theory. For example, David Bohm, in the spirit of de Broglie, proposed a kind of hidden variable theory; instead of saying, as is usual, that the wave function gives probabilistic information about discrete corpuscles of matter, it describes a kind of real quantum wave which functions as a 'pilot wave' for the corpuscles. I have always felt uneasy about this theory, basically because it introduces a kind of interaction between the pilot wave and the particles that is left unexplained. This is but one example of the general problem with adding more structure to a theory; there is a high risk that the cure makes the patient sicker. Together with a general adherence to Occam's razor, this has convinced me that a successful interpretation should not postulate any new entities or new fundamental laws.

When reading the founding fathers of quantum theory I felt affinity with Schrödinger's insistence on visualisability ('Anschaulichkeit'). But, unfortunately, all agreed that Bohr and Heisenberg definitely refuted Schrödinger's stance. Reconsidering the arguments I have come to think that he was not incorrect in demanding visualisability and in thinking that quantum objects are a kind of waves, as is argued in this book.

But Schrödinger was, in my view, wrong in not accepting quantisation of interaction. My own view could therefore be described as partly being in Schrödinger's vein, hence the subtitle.

This book is a development and continuation of ideas presented in my doctoral dissertation *Understanding Quantum Mechanics. A Realist Interpretation without Hidden Variables* (Stockholm: Almqvist & Wiksell International, 1992). The two basic ideas – that matter fundamentally is not made up of particles but of waves and that quantisation of interaction is responsible for the corpuscular aspects of matter – are the same in this book as in the earlier one.

Over the years I have discussed my ideas with many colleagues and friends. First and foremost I thank, Paul Needham and Dugald Murdoch, my supervisors when I was a PhD student. During these years I started doing philosophy of quantum mechanics and Paul and Dugald were my friendly critics. Paul later has commented upon and criticised an earlier draft of this book, for which I am most grateful.

When holding a position at Linköping University I had many rewarding discussions with Mattias Severin in the Physics Department and his interest and penetrating questions were most helpful. When arriving at Uppsala University, Kaj Børge Hansen became a friendly critic, who has forced me to clarify my own views at some points. Sören Törnkvist has been my discussion partner over many years. Being a physicist and physics teacher he has forced me to explain myself clearer and to spell out tacit assumptions. Most of all, he has helped me improve my English. In the final preparation of the typescript, George Masterton and Keizo Matsubara have been most helpful in proof reading.

The research behind this book was, to a great extent, made possible by a grant from HSFR, the Swedish Council for Research in Humanities and Social Science, which I hereby gratefully acknowledge.

HarperCollins has kindly permitted me to reprint (in Chapter 10) the discussion between Schrödinger, Bohr and Heisenberg held in Copenhagen 1927 and reported in Heisenberg's *Physics and Beyond*. I thank them for that.

Last but not least I would like to thank Inga, Liv, Hanna and Stina for being what they are, my loving family. Without them this work would have been much more difficult.

Chapter 1

Interpretative Problems in Quantum Mechanics

1.1 Historical background

The interpretation of quantum mechanics has been a source of conflict ever since its emergence in the 1920s. Schrödinger, de Broglie and Einstein were all of a realistic inclination, whereas Bohr, Heisenberg and several others claimed that the then new theory could not be viewed as a ‘literal’ description of nature. As the debate continued, the view held by Bohr and Heisenberg, the Copenhagen interpretation, became established among physicists and philosophers. However, the issue has never been finally settled and opponents have repeatedly voiced their doubts. During the last 20 years, the majority view, i.e. the Copenhagen interpretation, has lost ground to alternative interpretations.

One weak point in the Copenhagen interpretation – among many strong ones – is Bohr’s claim that we must use classical concepts in order to communicate unambiguously the results of experiments because concepts from microphysics have no clear significance. This sharp distinction between the macro and the micro world however seems completely untenable. Furthermore, the Copenhagen interpretation covers a number of views amounting to a kind of idealism, which cannot be ascribed to Bohr. There are reasons to be hostile to such views.

However, none of the most well-known alternatives, i.e. the many-worlds interpretation and the pilot-wave interpretation, is convincing; on the contrary, these two alternatives seem to be as defective as the Copenhagen interpretation. Some new proposals for interpretation, such as the GRW theory, the modal interpretations of van Fraassen and Dieks’ and Healey’s interactive interpretation all of seem to have flaws. The GRW theory has internal problems, as is clearly shown by David Albert (1992). Additionally, the central assumption in this theory, namely that there is a fundamental and irreducible probability of collapse proportional to particle number, seems rather *ad hoc*. The modal interpretations and the interactive interpretation both lack sufficient explanatory force.

The purpose of this treatise is to propose yet another interpretation, or rather reintroduce an old one, based on the idea that quantum mechanics is about real waves, an idea put forward by Schrödinger in the 1920s. However, Schrödinger met severe criticism from, among others, Heisenberg, Planck and Lorentz, and it was generally thought that Schrödinger was wrong. He himself also appeared to come to think so for he did not continue the discussion. However, he was not finally convinced by Bohr and Heisenberg – as can be seen from the recently published

proceedings of the Dublin seminars, held at the beginning of the 1950s, where he reiterated and developed his views from the 1920s (Schrödinger, 1995).

There are two likely reasons why Schrödinger did not succeed in convincing his colleagues about the correctness of his wave interpretation: he held a completely classical conception of what a wave is and he did not think that interaction was a process fundamentally quantised. He believed that quantum mechanics one day would be replaced by a better theory in which one could derive quantisation as an effect of resonance between two waves. He is reported to have exclaimed that ‘If one has to go on with these damned quantum jumps, then I’m sorry that I ever started to work on atomic theory’ (quoted in Rosenthal, 1967, p. 103). As we shall see in Chapter 9, his reliance on a classical wave theory led him into insurmountable difficulties.

The main difference between Schrödinger’s view and the one proposed here is the proposition that quantisation is a fundamental principle rather than something that can be derived within a classical wave theory. The similarity is the belief that objects in the micro world are waves, not particles. Stated in one single sentence: quantum mechanics describes a world of waves, which exchange energy discontinuously, and as quantum mechanics so far has proven true in all tests, it is very reasonable to say that that’s the way the world is.

1.2 What is an interpretation of a theory?

What do we want when we ask for an interpretation of mechanics? What are the criteria to be met by a successful interpretation?

A moment’s reflection over the history of physics tells us that quantum mechanics is alone in triggering such penetrating discussions about interpretation. The reason why we require an interpretation of quantum mechanics whereas we do not, or at least do not to the same extent, when confronted with, for example, relativity theory or classical electromagnetism, is that quantum mechanics, to a high degree, looks like pure mathematics; it is hard to see what kind of things the mathematical entities describe. Quantum mechanics ‘makes contact’ with nature at certain points only, i.e. the expectation values of observables, but these are few and the rest seems rather a lot of quite formal jiggling with formulas. Compare for example the description of the orbit for a planet around the earth on one hand, and the wave function describing the state of the 2p electron in a hydrogen-like atom on the other. The position of a single orbiting planet could be described in a coordinate system with the sun at the origin by $(x(t), y(t))$, where x and y fulfil the relation

$$\left(\frac{x}{a}\right)^2 + \left(\frac{y}{b}\right)^2 = 1$$

and the interpretation of this constitutes no problem: starting with an initial position (x_0, y_0) and velocity $(\dot{x}(0), \dot{y}(0))$ at the time $t_0 = 0$ the planet is to be found in the point $(x(t), y(t))$ for any chosen time t . The functions $(x(t), y(t))$ and $(\dot{x}(t), \dot{y}(t))$ describe the position and velocity, respectively for the planet at any chosen time t .

In comparison, the wave function for a 2p electron in hydrogen is described by the wave function

$$\Psi_{\text{tot}}(r, \theta, \varphi) = \text{const} \times r \exp\left(-\frac{r}{2a_\mu}\right) \sin \theta \exp i\varphi$$

The first and perhaps most obvious fact is that this function is time-independent; the wave function does not change with time. Does that mean that the electron is stationary? Secondly, since the function is complex it cannot be a description of a position (r, θ, φ) of the electron. As we all know, the squared wave function gives, for any chosen point in space, the probability of *finding* the electron at that point. But what, then, does the wave function refer to? And is the probability of finding the 2p electron at a certain place the same as the probability that the electron *is* at that place?

Thus, quantum mechanics has a semantic lacuna in so far as there are some very central concepts that seem to lack a direct interpretation in terms of what is out there in the world. The wave function of a system does not refer to anything at all, some people say.

For a long time this state of affairs was widely accepted as regards all theoretical concepts. In the *received view*, the final view adopted by the logical positivists, no theoretical concept or statement was said to be true or false. Theory only functioned as an instrument for predictions, thus the common label instrumentalism. The only meaning attributed to theoretical concepts was indirect, gained through the logical connection to observational concepts. The orthodox interpretation of quantum mechanics fits in with such a philosophy and that is one reason why it once was almost universally accepted.

However, the tide has turned completely in philosophy, and instrumentalism is nowadays rejected as untenable. Even a philosopher of such strong empiricist inclinations as van Fraassen accepts that central terms refer to real objects and that theoretical statements are true or false (van Fraassen, 1980, p. 10) even if undecidably so. But if so, the naive realist question 'what does the wave function really refer to?' appears legitimate.

Some practising physicists claim that we should not bother about problems of interpretation since quantum mechanics is the most successful theory ever seen; it predicts correctly the statistical outcomes of our experiments and no anomalies have so far appeared. What more could we ask for? But as such a stance is nothing but (thinly disguised) instrumentalism, I do not accept it, as already remarked. The aim of science is not only to predict future events but also to enhance our understanding of nature. In order to understand nature we must understand our theories, which means we should be able to give a semantic interpretation of their concepts and statements, i.e. to tell what kind of objects the terms in the theory refer to.

There is, however, an obvious argument against this position. A semantic interpretation is the association of objects to terms in a theory. To give a semantic interpretation is to select a domain of objects over which the variables in the theory varies. Hence, in order to be able to make a semantic interpretation we need first to identify the domain independently of our grasp of the theoretical framework. But that is surely hard to achieve in the quantum domain. How, for example, are we

to identify such objects and properties as photons, neutrinos, polarity and helicity without recourse to the very theory we want an interpretation of? The conclusion is that our best channel to the microcosmos is our best theory about it, and therefore a demand for a semantic interpretation is inappropriate and becomes completely trivial: ‘electrons’ refer to electrons, etc. Such an approach does not enlighten us.

However an answer is available for the realist. The theoretical terms we use in quantum mechanics signify mathematical objects such as operators, wave functions, S-matrixes and the like. But a realist must certainly hold that at least some of these mathematical objects represent real things, properties or processes. This means that the corresponding terms must be referring to real (i.e. physical) entities. Most hold this to be impossible; to take the most obvious example, the mathematical object of a wave function is usually a complex function and it appears certain that such a function cannot describe any real object, or any property of a real object; for whatever a measurable property a real object has, its measure must surely be a real number. But I will contest the conclusion; I believe that wave functions refer to real things and the issue will be discussed in Chapter 3.

Thus, my conclusion is that an acceptable interpretation of quantum mechanics, and in fact of any theory, entails a semantic interpretation of its terms and statements in terms of real objects and events. Furthermore, in order for such an interpretation to be acceptable the purported real objects and events must meet certain conditions, to be discussed in Chapters 2 and 3, which rightly should be called metaphysical restrictions. Thus, I see the present discussion about interpretation mainly as a metaphysical and semantic discussion.

Besides meeting certain general demands, a successful interpretation of quantum mechanics must solve, or dissolve, or at least elucidate the profound problems inherent in quantum mechanics. These are the interpretation of the *Born rule*, the *wave-particle duality*, the *measurement problem* and the mysteries associated with *EPR-correlations*, *Bell inequality* and *non-local interactions*.

1.3 Four core problems of interpretation of quantum mechanics

1.3.1 Interpretation of the Born rule

The Born interpretation states that the mathematical quantity $\Psi^*\Psi dx dy dz dt$, i.e. the intensity of the wave function in a space-time element denotes the probability to *find* the system in that space-time element $dx dy dz dt$, provided the system is correctly described by that wave function. At first sight this does not seem to be objectionable: if $\Psi^*\Psi dx dy dz dt$ is the probability to *find* the object in that space-time element, this must be equal to the probability that the object *is* in that space-time region. This inference, however, is legitimate only if we accept the principle that an object detected in the space element $dx dy dz dt$, was there before we performed the measurement and that is not a valid principle in quantum mechanics. That fact induces the realist to claim that if we stay content with the Born interpretation we must accept that quantum mechanics is not about nature but about our knowledge about nature. Hence, quantum mechanics is actually a theory about human cognition,

human states of mind. The impossibility to determine the state of a quantum system when not measured upon makes the theory objectionable. Either the theory is incomplete, or we do not yet understand it properly. For my part I think that the theory is complete and we do not yet understand the theory properly. Thus, a new interpretation of quantum mechanics is called upon, an interpretation not being an expansion of the theory, i.e. one which does not yield novel empirical predictions.

Why, then, is it impossible to infer that the system described by the wave function $\Psi(x,y,z,t)$ is in the volume element $dx dy dz$ during the time interval dt , independently of us actually measuring the state of the objects?

Consider a two-slit interference experiment in which a beam of electrons is directed towards a screen with two slits. Behind the screen, the electrons are detected with, for example, a photographic emulsion. It is straightforward to calculate the following probabilities:

1. The probability $p(x)$ of detecting an electron at a certain place x .
2. The conditional probability that the electron is found at x after having been detected immediately behind slit a , i.e. $p(x/a)$, and likewise $p(x/b)$.

As it is impossible to detect an electron at x which has not passed through any one slit, it would seem straightforward to suppose that

$$p(x) = p(x/a) + p(x/b) \quad (1.1)$$

This formula is, however, correct only if $p(x/a)$ is interpreted as '*the probability of detecting an electron at x after it has been detected behind slit a* '. However, if it is interpreted as '*the probability of detecting the electron at x after passing without detection at slit a* ', the formula is not correct, as comparison with the quantum calculation reveals. If we do not perform any measurement behind the slits, the wave function at x can be written

$$\Psi_x = \Psi_a + \Psi_b \quad (1.2)$$

Ψ_a and Ψ_b denotes the partial waves at x coming from slit a and slit b , respectively. The probability for the electron being detected at x is thus

$$|\Psi_x|^2 = |\Psi_a|^2 + |\Psi_b|^2 + \Psi_a^* \Psi_b + \Psi_a \Psi_b^* \quad (1.3)$$

It is easy to identify $p(x) = |\Psi_x|^2$, $p(x/a) = |\Psi_a|^2$, $p(x/b) = |\Psi_b|^2$ and we can conclude that the classical derivation is not correct because equation (1.3) correctly describes the probability distribution and equations (1.1) and (1.3) are not equal.

The assumption that validated the classical derivation was that the probability for an electron passing a slit, i.e. for being there independently of observations, equals the probability for detecting it at that place. Hence, we must conclude that these two probabilities are not equal, and we cannot generally interpret the square of the wave function in terms of probabilities for the electron being at that place. Hence, we must restrict ourselves by saying that the intensity of the wave function is proportional to the probability of *finding the system upon measurement*. This conclusion is not

restricted to position measurements, it applies to all dynamical properties of the system. But what are the values of observables of quantum systems when we do not perform any measurements? Do they not have any properties at all, not even position?

1.3.2 Wave particle duality and interpretation of the wave function

The second problem of interpretation related to quantum objects is that they sometimes behave as particles, show particle properties, and sometimes show wave properties. The problem is not so much that these objects change properties from one moment to another. The core of the problem is that it is impossible to describe a process converting one type of property to the other. We can tell, for any specified experimental situation, whether the quantum system will show particle or wave properties, but we cannot say anything about the system in itself or how the different types of properties transform into each other. ‘Das Ding an sich selbst betrachtet’ – to borrow a phrase from Kant – can be attributed neither wave properties nor particle properties. In fact, it seems as if any quantum system will display wave or particle properties depending on our way of looking at them. In order to be more specific, let us discuss the so-called delayed choice experiment.

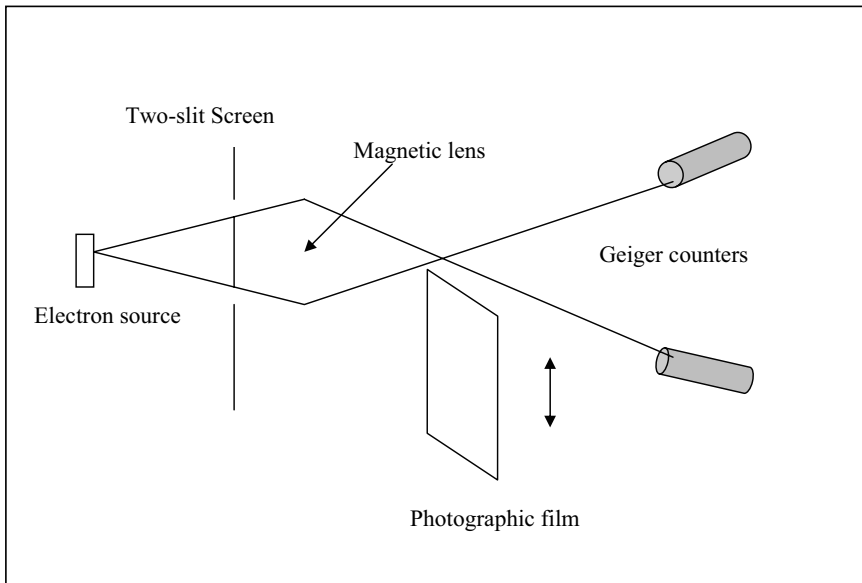


Figure 1.1 A delayed choice experiment

This is a double-slit experiment with *two* detecting set-ups: one photographic film, which can be moved in and out of the electron beam and a pair electron counters, i.e.

Geiger counters. The arrangement allows us to infer through which slit in the screen a particular electron has passed. The principle of the experimental set-up is shown in Figure 1.1 (Wheeler, 1983, pp. 182–184).

The experiment is conducted in the following way: the electron source emits electrons at a rather low rate, low enough to allow for the conclusion that at each moment only one electron at a time is on its way from the source to the target. This restriction is imposed in order to exclude the possibility of the electrons interfering with each other. The photographic film is moved up and down with a frequency such that any electron detected on the film must have passed the double slit before the film was in position. Hence, the ‘decision’ to move the film into the detecting position is made after the electron has passed the screen with the double-slit. If the electron has travelled fast enough to escape the film, it will hit one of the Geiger counters. The net result on the film is an interference pattern signalling that the electrons are waves passing through both slits. But some electrons are counted in the upper detector and some in the lower one and these electrons must have passed one slit only. The following conclusion seems unavoidable: if the upper counter has fired, the electron has passed the lower slit, if the lower counter has fired, the electron has passed the upper slit and if the film is in place the electron must have passed both slits. Consequently, the ‘decision’ whether to pass one or both slits is made after the electron has passed the screen with the two slits! As this conclusion conflicts with both causality and our conception of nature as independent of us, I am not willing to accept it.

1.3.3 The measurement problem

The core of this notorious problem is the elusive concept of measurement as it is used in quantum mechanics. A measurement is most naturally viewed as some sort of interaction between a measuring device and the measured object. It must be a physical interaction and, as quantum mechanics is a fundamental theory describing all kinds of physical objects and all kinds of physical interactions, quantum mechanics has a recipe for this interaction. It goes like this: assume that the measured object before measurement is described by the wave function

$$\Psi_{\text{obj}} = \sum c_i \phi_i \quad (1.4)$$

and that the wave function for the measuring device initially is in a zero state, M_0 . During the interaction M_0 is changed into a superposition corresponding to the superposition for the object

$$\sum c_i \phi_i + M_0 \rightarrow \sum c_i \phi_i M_i \quad (1.5)$$

such that the state M_i corresponds to the state ϕ_i with eigenvalue a_i of the measured object. But the measurement gives a definite result, a_k say. This means that the superposition must collapse into the component k during the measurement;

$$\sum c_i \phi_i M_i \rightarrow \phi_k M_k \quad (1.6)$$

This collapse during measurement cannot be derived from the Schrödinger equation, in fact it conflicts with the Schrödinger evolution. Hence it has to be added to quantum mechanics as a separate postulate, which was done by von Neumann in 1932.

The interpretative problem is due to the fact that the term ‘measurement’ has not been given any physical definition. Von Neumann showed that the interaction between the measurement device and the measured object could consistently be described as a quantum mechanical interaction, if we reckon our sense organs as the measuring device and the entire system object + measuring apparatus as the object we are measuring. In fact, we are completely free to decide what to call ‘object’ and what to call ‘measuring device’; we could even include the observer’s entire nervous system and only reckon his/her mind as the measuring device. But something must be viewed as a measurement apparatus, otherwise no measurement can be done. The problem is that we do not have any physical criteria for something being a measurement interaction. One and the same interaction, for example the interaction between the object and the detector, can be described either as following the continuous Schrödinger evolution or as being a collapse, depending on where we choose to draw the dividing line between object and apparatus. This stance strongly suggests that the collapse is a change of knowledge on the part of the observer and not a physical change of the state of the object.

Schrödinger’s cat

In order to show the absurd consequences of such a view, Schrödinger (1935) construed the famous thought experiment ‘Schrödinger’s cat’ (Figure 1.2). Suppose we have placed a radioactive source inside a detecting device. When the source decays, it emits radiation, which will be detected by the device. The wave function for this radioactive source can be written

$$\Psi = a(t)\Psi_1 + b(t)\Psi_2 \quad (1.7)$$

where the symbols Ψ_1 and Ψ_2 represent the states ‘undecayed’ and ‘decayed’ respectively. The detector equipment is connected to a relay, which acts as a switch in an electric circuit containing a power supply and a cat. When the radioactive source decays the detector will trigger the relay and the cat is given a fatal electric current.

The cat can be in either of two states: dead or alive. These two states make up a one-to-one representation of the states of the source. If the source has decayed, the cat is dead, and if it has not decayed the cat is alive. Let the states ‘cat alive’ and ‘cat dead’ be represented by C_a and C_d . The total wave function for the system radioactive source plus cat can now be written

$$\Psi_{\text{tot}} = a(t)\Psi_1 \otimes C_a + b(t)\Psi_2 \otimes C_d \quad (1.8)$$

Only by opening the box do we know whether the cat is dead or alive. This inspection causes the wave function to collapse into one of its components, dead or alive. With the box closed the state of the cat is a superposition of dead and alive and will not

transform to either dead or alive until the box is opened and the experimenter is aware of the cat's state.

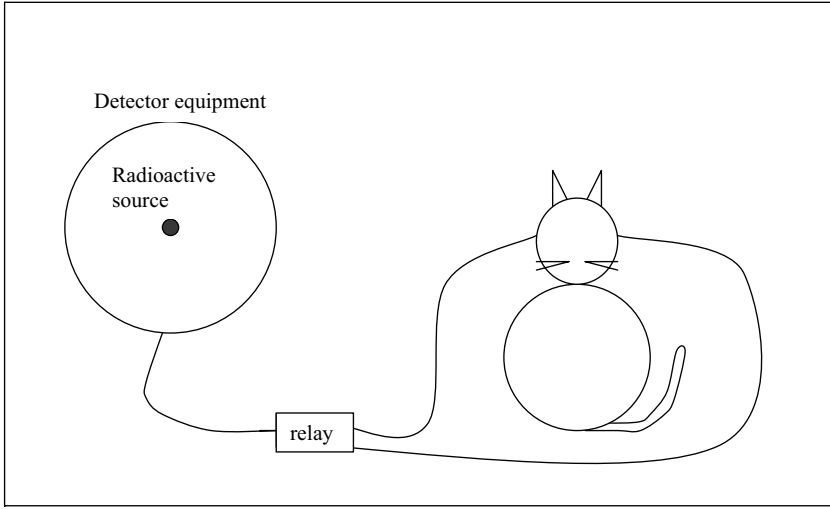


Figure 1.2 Schrödinger's cat experiment

Wigner's friend

In this thought experiment (Wigner, 1961), the point of departure is a quantum system characterised by a superposition of eigenfunctions to some operator

$$\Psi(t) = \sum c_j \phi_j \quad (1.9)$$

The system is connected to a measuring device that measures an observable of the system. As before, let us suppose that the measuring apparatus can be in any of the states M_j . The total wave function after the measurement interaction then becomes

$$\Psi(t)_{\text{tot}} = \sum c_j \phi_j M_j \quad (1.10)$$

It is now supposed that the experimenter, instead of looking himself, asks a friend what the result of the experiment was. From the point of view of the experimenter, the friend is part and parcel of the measuring apparatus. Hence, one must ascribe to the friend a superposition of states in the time interval between his looking at the measuring device and reporting the result to the experimenter. If the possible states of the friend after having observed the measurement result are called F_j , the total wave function for the system object + measuring apparatus + friend can be written

$$\Psi(t) = \sum c_j \phi_j M_j F_j \quad (1.11)$$

According to the quantum theory, the state of the *informed* friend is thus a superposition of all possible values.

By inventing the cat paradox, Schrödinger wanted to demonstrate that orthodox quantum mechanics cannot be accepted as a completely trustworthy theory by anyone having even a minimal realistic conception of the world. Wigner, however, wanted to show that the mind must be the ultimate cause of the collapse of the wave function, his argument being that measurements do not have any common physical characteristic that divorces them from other physical events. The only discernible trait of all measurements, so he assumed, is that human cognition is involved. But as measurement interactions normally change the state of quantum systems, so will human cognition. Thus, the measurement problem can be solved by giving up realism and some have followed Wigner on this route. As already hinted, I do not accept this idealism.

1.4 Non-local interactions

In 1935, Einstein, Podolsky and Rosen published their famous EPR article in which they argued that quantum mechanics is incomplete. Their argument was roughly as follows: the state of a quantum system composed of two particles may be described either as

$$\Psi(x_1, x_2) = \sum \psi_n(x_2)u_n(x_1)$$

or as

$$\Psi(x_1, x_2) = \sum \phi_n(x_2)v_n(x_1)$$

where u_n are eigenfunctions to an operator $\mathbf{A}$ and v_n are eigenfunctions to another operator $\mathbf{B}$, where $\mathbf{B}$ does not commute with $\mathbf{A}$. When measuring the value of the observable A corresponding to the operator $\mathbf{A}$ we will get as a result one of the eigenvalues φ_k , corresponding to the state $\psi_k u_k$, i.e. the wave function collapses. Thus, the other particle is also thrown into a definite state with a definite value of the corresponding observable. But we could instead perform a measurement of the observable B corresponding to $\mathbf{B}$, the result being a collapse to $\varphi_i v_i$ which throws the second particle into the definite state φ_i . But neither one nor the other of these measurements will affect the second particle in any way, so these two wave functions must describe the same real state of the second particle. It is obvious that we can predict with certainty both the value of the observable A and the value of observable B pertaining to the unmeasured particle, without in any way disturbing that particle. This is the criterion of reality adopted by Einstein *et al.* But knowing the wave function is not sufficient for this prediction; hence, the quantum mechanical description of nature is incomplete.

Bohr immediately replied by saying that their notion of completeness was inappropriate and, if the entire experimental situation were taken into account, the description would be complete. Most physicists accepted that argument.

In 1964, John Bell published the now famous article in which he proved the theorem named after him (Bell, 1964). It says that from the following three conditions, predictions contradicting quantum mechanics can be derived:

1. All observables have definite values all the time;
2. The measurement of an observable reveals the pre-existing value of that observable;
3. The value of an observable obtained by a measurement on one particle is not affected by measurements made on other distant particles.

Experiments have been performed to test whether the predictions derivable from these assumptions are correct or whether quantum mechanical predictions are correct. The outcomes are definitely in favour of quantum mechanics. The intense discussion has brought us a number of Bell-like theorems with weaker assumptions, the most important development being that the deterministic assumption (1) can be replaced by weaker assumptions and we still have a contradiction with quantum mechanics. Thus, the focus of interest has in recent years been the locality assumption (3), because it seems increasingly probable that that is the false assumption, and also because of the far-reaching consequences of rejecting locality. The following example highlights the paradoxical situation.

Assume that we have a pair of electrons in a *singlet state*, which means that the sum of their spins is zero in any chosen direction. Now let us assume that the spin of one of the electrons is measured in a randomly chosen direction. The result of such a measurement can only be ‘spin up’ or ‘spin down’. Let us assume that ‘spin down’ was obtained. Then the other electron must have ‘spin up’ in the same direction because the total spin of the pair is zero in any direction. Our first question is: did the electrons have these spin-values right from the beginning or did they acquire them at the moment of measurement? The easiest way of explaining the correlation is of course to assume that the particles had definite and opposite spin directions right from the beginning of the experiment. The correlation can be obtained in a randomly chosen direction, so induction leads us to conclude the correlated particles must have opposite spin in all directions. Hence, both particles would have a definite spin in all directions from the start of the experiment. However, such a generalisation of the assumptions yields Bell’s inequality and this inequality conflicts with quantum mechanics (see Sections 8.2–8.3 for a thorough discussion). Therefore, most physicists and philosophers conclude that the spin in the measured direction must have acquired a definite value at the moment of measurement (except in those cases where the measured direction happens to coincide with the direction determined by the preparation of the particles). As soon as this measured electron has acquired a definite spin value in the chosen direction, the other electron immediately (faster than the speed of light) acquires the opposite value.

The strict and instant correlation between the two particles takes place without the mediation of any known force. Could there be a hitherto unknown force responsible for the distant correlation? That depends on what we mean by a force. The four fundamental forces, i.e. gravitation, the electromagnetic, the weak force and the strong force have three properties in common:

1. They are transmitted by ‘carrier particles’ travelling at most at the speed of

light;

2. Each kind of fundamental force has its own kind of carrier particle; and
3. The force can be described by a potential, which implies that any object of the relevant kind will be sensitive to all other objects of the same kind. For example, a charged particle interacts with all other charged particles.

In all these three respects the spin correlation between a pair of objects in the singlet state differs from an interaction by forces and it should be clear that the spin correlation is not a new fundamental force.

This strict and instant connection is called non-local interaction because it is a violation of the Principle of Locality which roughly says that an object cannot be affected by distant objects.

Alain Aspect and his collaborators (Aspect *et al.*, 1981, 1982a, 1982b) have performed a number of experiments on correlated pairs of particles. The results make it highly plausible that the connection takes place at a speed faster than light and that the predictions of quantum mechanics are correct. However, neither the experiments, nor quantum mechanics by itself indicate *how* the interaction takes place. We cannot escape the conclusion that quantum mechanics is a non-local theory, and so far we have no way of understanding this aspect of nature.

Chapter 2

Realism and Understanding

2.1 Realism

2.1.1 Realism and quantum mechanics

Realism can mean different things, but the common core is the belief that the physical world (or better, that which our beliefs are about) exists independently of the states of our minds. Many philosophers and physicists believe that it is impossible to give an interpretation of quantum mechanics that accords with this realistic attitude. By contrast, I believe them to be wrong in thinking that; quantum mechanics is a theory about the fundamental features of the physical world and the states of our minds can be set aside in the interpretation of this theory.

The contrasting position to realism about scientific theories has changed from time to time, so I will label it simply as anti-realism.

A strongly realistic position holds that an acceptable interpretation of any theory, and in particular quantum mechanics, should tell us more than probabilities for outcomes of possible measurements. For the realist, the main concern is knowledge about the world, not merely predictions about measurements. Furthermore, a realist assumes the principle of bivalence, i.e. that every syntactically well-formed declarative sentence is true or false. In this context this implies that every statement of the form ‘the quantum system S has property P’ is definitely true or false, independent of our knowledge or perspective.

At the opposite end we find the strong anti-realist who rejects such requirements. The only purpose of a scientific theory according to anti-realism is to provide probability distributions for series of experiments. This stance is reinforced by the claim that physics is not interested in the outcome of an individual experiment, which is an individual fact, only in general laws and regularities, and therefore probability distributions are all that matter. The theoretical concepts utilised need not, and cannot, be interpreted in any reasonable sense of the word; they are just conceptual tools. It is then a short step to adopt instrumentalism, which says that theoretical sentences are neither true nor false, but tools for predictions. However, this position is no longer popular in the philosophical debate, due to the insuperable problems of drawing a strict border between theoretical and observational predicates. Such a strict borderline is absolutely necessary because one cannot accept being unclear whether a sentence lacks a truth-value. In spite of this, many physicists take an instrumenalistic attitude to parts of quantum mechanics. For example, many people deny that singular terms containing the expression ‘wave function’ refer to real entities, and therefore statements containing this expression are neither true nor false.

The Copenhagen Interpretation (and the less well-known consistent histories interpretation, a variant of the Copenhagen Interpretation) has by many critics been accused of anti-realism and more precisely of instrumentalism. As the label ‘Copenhagen Interpretation’ covers a family of slightly different views, this accusation is sometimes fair, sometimes not. Roughly, one could say that Bohr himself held the view that quantum mechanics is a theory about nature, not about our knowledge of nature, but, due to complementarity, there are strict limitations on what can be said about nature, and this stance is not instrumentalism or idealism. Heisenberg, on the other hand, held stronger views to the effect that quantum mechanics is not a theory about an independently existing realm but essentially involving human interests and interactions and this is a form of idealism.

The main reason why realism, so natural in everyday life and in classical physics, is often dismissed as a metaphysical ground for the interpretation of quantum mechanics is the impossibility of identifying the probability of *measuring* a value a_k of an observable O with the probability that the object in question *has* the property ‘observable O has the value a_k independently of any measurement’. The realist wants to say how the world is in itself, independent of any measurements or other interventions and quantum mechanics does not seem to allow for that. The problem would be less disturbing for the realist if the concept of measurement could be defined in purely physical terms; for if so, we could at least say that human cognition and human observation has no effect on the measurement outcome, but so far that has seemed impossible. So far, it seems as if a necessary criterion for a process to be a measurement is that it involves a cognitive ingredient, an observer gaining knowledge about the outcome. Consequently, we can only say what we know, but not how the world is in itself, independent of any measurements. A realist cannot be content with such a view because the core of realism is the viewpoint that meaningful statements are about an external world, independent of human minds, points of view or cognitive capacities.¹

2.1.2 Definition of realism

Realism has often been formulated as being that every syntactically well-formed statement is definitely true or false, independently of our possibilities of knowing which is the case (Dummett, 1991a). According to that view one should say that every statement attributing a value to an observable is definitely true or false irrespective of whether we perform any measurement, if we take the concept of measurement to contain a cognitive component. Hence, as Dummett (1991a) has stressed, the realist is committed to the principle of bivalence. However, we immediately confront a peculiar problem in quantum mechanics, namely the uncertainty relations. They tell us that certain predicates expressing dynamical quantities such as momentum, energy, position etc, cannot be used without restrictions. For example, if a system has a definite momentum, the uncertainty in position is infinite and vice versa. This

¹ Realism about social and mental facts must be defined differently, but that is irrelevant in this context.

could be interpreted either ontologically or epistemologically. In the ontological interpretation it means that the system does not *have* any position at all. But still it exists as a physical object in time and space: hence it must be everywhere! Alternatively, the uncertainty relations could be interpreted as stating limits for our knowledge about the system. Here, one holds that objects have definite values on all observables at all times but there are limits on what can be known.

The epistemological interpretation was, for a long time, the common opinion among physicists (Heisenberg, 1930) and philosophers. However, opinions have changed, nowadays the majority view these relations as stating ontological facts. A strong argument for this latter view can be constructed from the two-slit experiment. If we assume that an electron has a well-defined position at every moment of time, it must pass through one of the slits in a two-slit experiment. But then, the resulting probability distribution of many runs of this experiment ought to be the sum of the two probability distributions arising after having passed either one or the other slit. (Individual runs do not interfere with each other.) However, that is wrong, we get an interference pattern, the most natural interpretation of which is that those objects producing the interference pattern pass through both slits. Hence, the most plausible interpretation is that electrons do not have a definite position when passing the screen with these two slits, but are spread out. How that could be understood we shall see in the chapters to follow (and also the ensuing problem how things being dispersed in space can cause individual point-events on detecting screens). Generalising to all observables we arrive at the ontological interpretation of the uncertainty relations, which should better be called the *indeterminacy* relations.

It follows from these considerations that when a quantum object has a well-defined value on one observable, it has no definite value on any observable incompatible with the well-defined observable. (Observables are incompatible if the corresponding operators do not commute.) Hence, in such a situation, the statement ‘System S has a well-defined value on observable O’ is false and so is consequently all statements of the form ‘System S has the value a_k on observable O’.

Two things must be kept in mind: if it is false that the a system at a certain time has a value a_k on an observable O, this does not entail that an analogous statement attributing another value to the same observable is true; all statements attributing a value to this observable are false. The second thing to keep in mind is that after an interaction with another system, to be discussed in chapters five and six, it might acquire a definite value on this observable O.

The core of realism is the belief that the conditions that determine whether an object can be predicated a certain property or not, must not depend on human minds, human points of view or human knowledge. This is a metaphysical position: it expresses a certain view about the relation between minds, their cognitive acts, and the physical world. Moreover, this view entails a certain condition on interpretations of scientific theories, a condition, which I will call minimal realism:

Minimal realism: any scientific theory should be interpretable as a mind-independent description of the world.

A word of caution is necessary. The world certainly consists of, among other things, human minds and thoughts. Certainly one can be a realist as regards these phenomena. A description of the mind of a certain person is of course dependent

on this person's state of mind, and such a statement does not contradict realism. Minimal realism should rather be understood to mean that the truth of the description should be independent of the state of minds of those who make the descriptions.

In contrast to most realists in this field, I do not presuppose hidden variables. For example, I do not presuppose that if a quantum object in some situations has a definite position it must have that in all circumstances. The proposition that quantum objects have definite albeit unknown values on all observables at all times amounts to *value determinism*. This is an additional assumption, which I, by the way, believe is false.

Value determinism is not the same as general determinism. The latter position holds not only that all observables have definite values at all times, but also that these values are determined by laws and initial conditions in such a way that an omniscient being with complete information can predict with certainty the outcome of those events. Obviously, one can adhere to value determinism and at the same time believe that some laws are statistical, i.e. that there are events in nature, the probability of which cannot be reduced to zero or one even if conditionalised on all antecedent events.

Most researchers in the field lump together determinism and realism. Einstein is one famous example, Barut a less famous one. The latter writes, 'The program of realism, on the other hand is to look for a deterministic description of single events...' (Barut, 1994, p. 10).

Many adherents to the Copenhagen Interpretation think that there is a direct conflict between minimal realism and the fact that it is in general impossible to pass from a statement about the probability for *finding* a certain object in a certain state to the statement that the object *is*, independently of observations, in that same state. However, this need not be so; the conflict arises only if we interpret the concept of measurement (which is involved in talking about *finding* an object at a certain place) as necessarily involving a human act of cognition. If it is possible to state in purely physical terms what is meant by *finding an object in a certain state*, i.e. when an act of measurement is completed, no canons of realism are contradicted. My point of departure is precisely this assumption; I believe it is possible to interpret quantum mechanics realistically and it is precisely this belief which causes most of the interpretational problems. For if we accept that a scientific theory need not fulfil even minimal realism, we have in fact accepted that a physical object has different physical properties depending on our knowledge about that object and that means accepting among other things that:

1. Schrödinger's cat is dead or alive depending on our knowledge about its state;
2. in an double slit experiment the quantum object goes one or two ways from source to target depending on which observations we choose to perform later;
3. the collapse of the wave function occurs as a result of our cognitive act of observation.

I regard all these consequences as absurd.

Minimal realism is, I think, weak enough to be acceptable to most scientists. However, there are some who deny even minimal realism, such as the well-known quantum theorist J. A. Wheeler, as can be seen from the following statement:

The universe does not exist 'out there' independent of us. We are inescapably involved in bringing about that which appears to be happening. We are not only observers, we are participants ... in making [the] past as well as the present and the future. (Goodman, 1984, p. 36)

From the context it is evident that Wheeler does not allude to human matters when he talks about bringing about past, present and future, but the inanimate nature. As far as I know, Wheeler has not given any general philosophical argument for his position. It appears that he has adopted it as one way of solving the interpretational problems of quantum mechanics, or, rather, rejecting the demand for explanation. Wheeler's stance seems to be rather close to that of Heisenberg's. Compare, for example, Wheeler's stance with the following quotation from Heisenberg:

The conception of the objective reality of the elementary particles has thus evaporated in a curious way, not into the fog of some new, obscure, or not yet understood reality concept, but into the transparent clarity of a mathematics that represents no longer the behaviour of the elementary particles but rather our knowledge of this behaviour. (Heisenberg, 1958, p. 100)

This looks rather a lot like the Copenhagen interpretation of quantum mechanics. However, if we should count Wheeler's stance as part of the Copenhagen school, Bohr would certainly not join the party. Dugald Murdoch (1987) argues, convincingly I think, that Bohr must be acknowledged as being a realist; Bohr never said anything to the effect that quantum theory is a theory about our knowledge and not about the external world, nor did he say anything to the effect that the world is mind-dependent. What he did claim repeatedly was that certain concepts have limited applicability and certain seemingly innocuous statements using these concepts are, in some conditions, meaningless.

Many adherents to the Copenhagen interpretation are not as careful as Bohr and more inclined to follow Heisenberg. Although Bohr is the founding father of the Copenhagen interpretation, I think he did not fully endorse Heisenberg's view. Rather he would have been inclined to join Einstein in propounding the following position: 'The belief in an external world independent of the percipient subject is the foundation of all science' (Maxwell, 1931, p. 66).

Thus, the label 'the Copenhagen interpretation' must be used with care when discussing interpretations of quantum mechanics. The Copenhagen interpretation as it is usually understood, violates minimal realism, whereas Bohr's position does not.

In my view, minimal realism is a necessary requirement on an interpretation of quantum mechanics. This is because the goal of interpretation is to achieve insight, to understand nature, and a necessary condition for understanding is that we are told what the world is like irrespective of our thinking. But more is needed.

I want to stress that my aim is not to give arguments for minimal realism; I take it for granted. But since quantum mechanics is often said to prove that realism is impossible, I want to undermine this argument by showing that a realistic interpretation of quantum mechanics, i.e. one satisfying minimal realism, is in fact possible. Hopefully I achieve that goal in this book.

2.2 Interpretation and understanding

The notions of *explanation*, *interpretation* and *understanding* can have various relations depending on how we choose to explicate them. In one perspective, a strongly realistic view, the aim of an interpretation of a theory is to explain certain traits of that theory and that in turn gives us understanding. This in turn is a necessary condition for a theory to be fully believed; it is indeed odd to say; ‘I believe the theory is true, but I do not understand what it says.’

In an opposing perspective, to interpret a theory is to give rules for its application to empirical situations, and explanation and understanding are merely pragmatic virtues of less importance. This is Bas van Fraassen’s (1980) view; he claims, in effect, that explanatory force and *a fortiori* understanding of a theory is context dependent, by which he intends that the explanatory force partly depends on the preconceptions of the audience. Thus, evading the question of whether one could believe a theory to be true without having understood it, he says that the goal of science is not truth but empirical adequacy, i.e. that theories fit all observable phenomena. Since many completely different theories can fit all observable (not only those observed) phenomena, there is never, in his view, reason to demand, or assume, that a theory is true. I’m not convinced by his argument, but I will not pursue this issue here.²²

Being a realist I do not accept such a view. The goal of interpretation is not only to give rules for application, but also express the theory in question in such a way that the described domain is independent of human perspectives, human minds and knowledge. This follows from the stance that the goal of science is not only to make predictions but also to provide explanations of phenomena. This is, in my view, the primary condition for understanding a theory. (In social sciences, though, one must be careful not to exclude the propositional attitudes of the members of the studied society, because these are part of the very facts of interest. In such areas realism means that the description is independent of the state of mind of the *observing and describing subject*. The minds of other people are part of the objects of interest.)

It is highly remarkable that many physicists have been so willing to dismiss this realistic view, at least officially, because in other areas of physics the general aim is to remove, from the description of nature, all traits that have their origin in the chosen perspective. Thus, it is generally demanded of a physical theory that it should be given in a form that is invariant under certain groups of transformations, such as translations, rotations, space inversion, gauge transformation etc. Superficially,

2 In a short paper (Johansson, 1996) I have argued that van Fraassen’s strong distinction between truth and empirical adequacy is rather arbitrary from his own empiricist point of view.

these efforts are motivated by a longing for simplicity and symmetry as a form of intellectual beauty, because invariance under a group of transformations entails that a certain quantity is conserved, according to Noether's theorem. However, a much stronger motivation is available. Translations, rotations, space inversions and time displacements can all be seen as changes of the coordinate system chosen for the description and such changes are not changes in the real world but only changes of quite arbitrary choices of human perspective, i.e. frame of reference. The same can be said of a number of more abstract transformations such as gauge transformations in quantum field theory. Thus, the motivation for these symmetry demands is that these transformations do not represent real changes in the world but only changes of human perspective. Such human ingredients have nothing to do with the real external world, hence the description of that world should be independent of such perspectives, if we want to be realists. Thus, the motivation for these demands on invariance under certain groups of transformations has its origin in a realistic outlook, a clear distinction between human constructs and the real world. It is tempting to ascribe to physicists such a motivation, but I have never seen it explicitly stated.

To be sure, I do not want to suggest that the invariance requirements are less in quantum theory than in other theories. Rather, my point is that the motivation for invariance demands is precisely realism, and some physicists explicitly reject this stance when the interpretation of quantum mechanics is at issue. To reject minimal realism as a requirement of interpretations of quantum mechanics but to require invariance under, for example, space-time transformations, is not fully consistent.

Certainly, minimal realism is only one requirement for an acceptable interpretation of quantum mechanics. More is certainly needed and the most obvious candidate is that the given interpretation should provide understanding. But what precisely is meant by understanding?

2.3 Understanding

Several authors claim that quantum mechanics gives no understanding of nature. Cf. for example the following two quotations, the first from Feynman and the second from Gell-Mann:

There was a time when newspapers said that only twelve men understood the theory of relativity. I do not believe that there ever was such a time. ... On the other hand, I think it is safe to say that no one understands quantum mechanics... Do not keep saying to yourself, if you can possibly avoid it, 'but how can it be like that?' because you will get 'down the drain' into a blind alley from which nobody has yet escaped. Nobody knows how it can be like that. (Feynman, 1967, p. 129)

All of modern physics is governed by that magnificent and thoroughly confusing discipline called quantum mechanics invented more than fifty years ago. It has survived all tests. We suppose it is exactly correct. Nobody understands it, but we all know how to use it and how to apply it... and so we have learned to live with the fact nobody can understand it. (Murray Gell-Mann, quoted in Wolpert, 1992, p.144)

Obviously, both Feynman and Gell-Mann think that quantum mechanics lacks something to be wished for in a physical theory, something necessary for understanding. What could that be? Even though quantum mechanics is difficult to learn and apply to concrete situations, this goal is certainly achieved and the predictions of quantum mechanics are remarkably corroborated in all experiments performed so far. The word ‘understanding’ obviously means something more than understanding how to use the theory.

My guess is that understanding of quantum mechanics can be interpreted as visualisability in space and time. More is perhaps involved but visualisability is at least a necessary condition. James Cushing (1991, p. 351) holds the same view.

Visualisability was once an issue in the debate between, on the one hand, Bohr, Born and Heisenberg and, on the other, Schrödinger and Einstein. Bohr claimed that a spatio-temporal description of events, i.e. a description in terms of trajectories, is in general impossible. Born rejected any physically visualisable interpretation of the wave function and introduced his probabilistic interpretation of the intensity of the wave function. Heisenberg created matrix mechanics, an abstract algebra which defied all attempts at interpretation in terms of objects with motion in space and time. Against these developments Schrödinger protested:

Bohr’s standpoint, that a space-time description is impossible, I reject *a limine*. Physics does not consist only of atomic research, science does not consist only of physics, and life does not consist only of science. The aim of atomic research is to fit our empirical knowledge concerning it into our other thinking. All of this other thinking, so far as it concerns the outer world, is active in space and time. If it cannot be fitted into space and time, then it fails in its whole aim and one does not know what purpose it really serves. (Moore, 1989, p. 226)

Thus, according to Schrödinger, to understand quantum mechanics one has to interpret the theory as being about physical objects and events in space and time. Schrödinger’s term for this demand was *Anschaulichkeit*. The translation of this word into English is not agreed upon; de Regt (1997, p. 471) proposes ‘intelligibility’, some other has translated it ‘possible to intuit’, in accordance with the common translation of the word ‘*Anschauung*’ in Kant’s oeuvre. In my opinion the best translation is ‘visualisability’, although the German word may have a slightly more abstract meaning. De Regt (1997, p. 471) argues that Schrödinger equated intelligibility with visualisability, giving this quotation as evidence:

[W]e cannot really alter our manner of thinking in space and time, and what we cannot comprehend within it we cannot understand at all. There are such things – but I do not believe that atomic structure is one of them. (Schrödinger, 1928, p.27)

One can be sympathetic towards this attitude, but hasn’t Bohr once and for all showed that the goal is unreachable? Mustn’t we accept that a complete description of quantum reality requires complementary pictures, one in terms of particles, one in terms of waves, and these cannot be integrated into one coherent picture of physical objects in space and time? I submit that most physicists, including Feynman and Gell-Mann, would claim precisely this, and so would most philosophers in the field

too. However, I think Schrödinger's view, that the structure of the atomic world is comprehensible in the sense indicated above, is tenable, notwithstanding the strong force of Bohr's detailed arguments.

To begin, I will suggest that we in fact can visualise electron orbitals. In physics and chemistry textbooks we find images of s-, p-, and d-orbitals (f-orbitals being too complex to be of any pedagogical value as pictured). These images are representations of the charge density distribution around the nucleus. An ambitious textbook writer would use colour density to represent the charge density, and we would get a fairly vivid image of an orbital. The picture is not completely faithful in that it does not account for the fact that orbitals extend indefinitely. Very low densities are difficult to portray and neglected as of little relevance. Thus, the picture is an approximation but that is often the case. No picture is a complete representation of all aspects of the object being depicted.

Do not these pictorial representations of orbitals give us understanding? As Feynman and Gell-Mann were (are) great illustrators, they would certainly not deny that these pictures give us some understanding of orbitals. But still they claimed lack of understanding of quantum mechanics! What more could they ask for? Perhaps the simple answer is that they could not integrate interference phenomena into their particle view of quantum objects and that is what they do not understand.

Turning to Schrödinger's views for a moment, I think pictures of electron orbitals give us good examples of what Schrödinger had in mind when talking about visualisability. He extended this to charge densities in general by suggesting that the intensity of the wave function times the unit charge always represents a charge distribution in real space, even in the case where only one electron is present. Such a charge density can be attributed a motion in space and time, the density current, albeit not a particle trajectory. But this is only the first step in Schrödinger's interpretation. The most pressing problem for Schrödinger (not for me!) was to understand 'quantum jumps'. He dismissed this concept and held that interactions could be accounted for in terms of resonance phenomena between two or more waves. The visualisability of such a process is perhaps not evident, but Schrödinger, like all physicists, had, I guess, performed or seen experiments in which mechanical oscillating systems are coupled and their resonances studied. This provided a visualisable model for quantum interactions. In this situation the salient thing for Schrödinger was that we have mathematical models for these interactions and these models are continuous; the function describing the energy distribution between the two coupled systems is continuous and so is its derivative. Schrödinger's requirement for visualisability was thus twofold; A physical theory is visualisable if:

- (1) the theoretical terms can be interpreted as referring to physical objects in space and time,
- (2) the evolution of states in the theory can be interpreted as a continuous state change of these physical objects in space and time.

I think Schrödinger was right in insisting on the first condition but wrong in insisting on the second. My argument for the latter could be given as a question: why is a continuous evolution understandable, whereas a discontinuous one is not?

There are, I think, historical reasons why Schrödinger and many others had problems with discontinuous changes, reasons that go back to a very old metaphysical idea: *Natura non facit saltus*. In early modern times this way of thinking was strongly reinforced by the success of classical mechanics; Newton showed how to describe and analyse the motion of bodies using continuous and differentiable functions and this method became the very paradigm for solving physical problems. We *understand* the physical situation when the differential equation is formulated and its solutions, i.e. continuous functions, are found. Then quantum theory entered the stage and Bohr and Heisenberg claimed that there simply is no differential equation and no continuous function to be found representing the state change from one stationary state to another: the change is discontinuous and irreducibly random. The question ‘what is the cause of a particular state change’ simply has no answer. It is no wonder that Schrödinger thought that this unintelligible, it lacks *Anschaulichkeit* and I suspect most physicists of the 1920s agreed.

As already indicated, I do not agree with Schrödinger, as interpreted by de Regt and myself, on this point. What we claim to understand is always more or less historically determined. To give a simple example; from Aristotle to Galileo, philosophers of nature wondered about the nature of motion and its connection with forces. It seemed obvious for them that as soon as the force on a body terminates the motion also stops. But that view makes many common phenomena hard to understand. If an arrow is shot from a bow it flies its way following a parable-formed trajectory. We do not usually wonder about that, but Aristotle and his followers did; they asked, why does it not fall down to earth as soon as it leaves the bow, i.e. as soon as the force on the arrow has ‘expired’? One can reasonably say that they did not understand its motion. The reason for this perplexity was the implicit presupposition that motion requires a driving force, and no such force was acting on the flying arrow. Nowadays we know better: uniform motion does not require any force, it is a natural state. Once this is accepted the question ‘what is the cause of motion?’ should not be answered but dismissed as presupposing a deficient concept of force.

Analogously, I propose that we reject Schrödinger’s conception of *Anschaulichkeit* with a similar argument. Instead of thinking that only continuous changes are natural, i.e. requiring no further explanation, we might just as well think that discontinuous changes are ultimate facts of nature. Why could it not be said that our understanding of discontinuity is on a par with that of continuity? Using a diagram we can easily give a pictorial representation of a discontinuous change, and such a representation has at least some *Anschaulichkeit*.

Understanding has much to do with what we are conditioned to expect, what we are adapted to. Thomas Kuhn described this well in his (1962) book *The Structure of Scientific Revolutions*. New concepts are formed, new norms are settled, new patterns of reasoning are adopted after a scientific revolution, resulting in a new way of seeing things. This shift of paradigm dissolves old enigmas but creates new ones. It is remarkable that Schrödinger did not fully understand the consequences of the scientific revolution he himself was involved in staging.

But are there no limits beyond which no possible conceptual change can force us? Maybe we must be prepared to reject continuity or determinism as real traits of the world. If so, should we be prepared to give up every conception of nature we

might have when new theories enter the stage? I am inclined to say yes, but with an important proviso. Some features of current physical theory do not reflect traits of nature, but rather conditions for objective descriptions. The most important of these are the requirements of invariance under certain transformations (such as time and space translations and rotations). Objectivity entails that if two observers at different positions in space and time give a description of the structure of our physical world, these two descriptions should be equivalent. This demand entails invariance under space–time translations, which in turn entails conservation of certain magnitudes via Noether’s theorem.³³

Another quantum phenomenon that has aroused an intense debate is nonlocality, perhaps the most perplexing feature of quantum mechanics. Nonlocal correlations appear to contradict a very fundamental principle in our conception of nature, i.e. that influences from one point to another must be transmitted by something, for example a force. So described, nonlocal correlations, which are not transmitted by any force, are impossible to understand. However, there is a way out. I think our perplexity is due to a mistaken ontology, an ontology consisting of particles confined to well-defined places. If this ontology is rejected the contradiction disappears. This is discussed in Chapter 8, the result being, I hope, a way of understanding nonlocal correlations.

My stance is thus that in *some respects* we should be prepared to accept that nature might be completely different from what we expect it to be, and accordingly change our views concerning what is in need of explanation. The respects in which we should *not* be prepared to change our preconceptions are the following:

- (1) invariance principles are valid,
- (2) realism, as explicated above, and
- (3) influences from one object to another must be transmitted by something.

Summarising, to understand a physical theory is to be able to interpret that theory in such a way that it describes objects in space and time and their properties in such a way that the demands (1), (2) and (3) above are met. But the often added requirements of general determinism, value determinism and continuity are not needed for understanding.

I think that when people say that it is impossible to understand quantum mechanics they mostly have a stronger conception of realism than my minimal realism and such stronger realisms often clash with the structure of quantum mechanics.

In recent years, nonlocal phenomena have evoked much interest, but it seems to me that many, perhaps most, participants in the discussion take it for granted that there is a loophole in the argument leading up to nonlocal correlation clashing with principle (3) above. I, on the other hand, believe that there are no such loopholes and consequently I must address the intelligibility of non-local correlations.

3 Noether’s theorem says if a physical system is invariant under a continuous and reversible group of transformations there is a conserved quantity which is a function of the system.

2.4 The semantics of quantum mechanics

In my view, we should add a third demand on an interpretation of quantum mechanics, namely that it provides a systematic semantic account of the theory. The upshots of this demand are several forms of reasoning, purported to give understanding and illumination of the quantum domain, with which I feel dissatisfaction. I will give three examples.

(1) In his reply to the EPR article (where Einstein, Podolsky and Rosen concluded that quantum mechanics was incomplete) Bohr wrote:

From our point of view we now see that the wording of the above-mentioned criterion of physical reality proposed by Einstein, Podolsky and Rosen contains an ambiguity as regards the meaning of the expression ‘without in any way disturbing the system’. Of course, there is in a case like that just considered no question of a mechanical disturbance of the system under investigation during the last critical stage of the measuring procedure. But even at this stage there is essentially the question of an influence on the very conditions which define the possible types of predictions regarding the future behaviour of the system. Since these conditions constitute an inherent element of the description of any phenomenon to which the term ‘physical reality’ can be properly attached, we see that the argumentation of the mentioned authors does not justify their conclusion that a quantum-mechanical description is essentially incomplete. (Bohr, 1935, p. 700)

My problem with this crucial passage is the verb ‘influence’. Bohr clearly rejects a mechanical reading, but physics is more than mechanics. Does he mean that the influence is physical but not mechanical, or does he mean that the influence is conceptual or logical? Perhaps the latter, but then he owes us an explanation how a mere conceptual influence can have physical effects. In short, I do not regard this as a sufficient piece of interpretation of quantum mechanics, although there is, I believe, some truth in Bohr’s position.

(2) Among philosophers and physicists there are those, Paul Teller is an example, who maintain that the two correlated particles in a singlet state in which we can observe nonlocal correlations really are not two different systems but actually one single system. They claim that the relational properties of single quantum particles are not reducible to non-relational properties as they are in the classical domain and that is a holism we can understand. Well, I can’t and Cushing (1991, p. 349) says the same. This holism is clearly not sufficient as an interpretation; we need a general discussion of the concept of system and of criteria for individuality of systems. Cushing, in the article referred to earlier, accords with this criticism of Teller.

(3) Physicists unhesitatingly use particle terms such as ‘electron’, ‘neutron’ or ‘photon’ even when talking about interference phenomena. When accounting for interference experiments physicists usually say things like ‘the wave function travels all possible paths between source and target’ while at the same time maintaining that the ‘corresponding’ particle is very small or point-like. Such talk is doubly problematic. Firstly, a wave function is a mathematical entity that does not exist in space and time, hence you could not literally say that the wave function travels

several paths in space. Perhaps such talk is elliptical; what is really meant is that the physical object represented by the wave function travels several different paths, i.e. that it somehow splits into several spatially distinct parts during its motion. But a point particle, or an object with a well-defined position and held together as a unified piece of matter, cannot be said to travel two paths simultaneously, hence the object represented by the wave function cannot be identical with the particle. Thus, in order to steer clear of an outright contradiction, the word ‘correspond’ is used to describe the relation between what is referred to by the wave function and the particle, respectively. This is the other problematic aspect of this example. In short, by using the expressions ‘the wave function travels both ways’ and ‘correspond’ to describe the relation between object and mathematical description the problem is hidden. (But adherents to the pilot wave interpretation solve this by endorsing a dualistic ontology, claiming that there exist both particles and matter waves. I will discuss their view in Chapter 9.)

When confronted with these and other examples, one perceives the need for a systematic semantic interpretation of quantum mechanics. The aim should be to answer the following questions:

- (1) which expressions purport to refer to real objects?
- (2) which kind of objects exist?
- (3) how are these objects individuated?
- (4) which properties do they have?
- (5) how are objects identified and re-identified during and after changes of state?

The last item in this list is connected to the problem of understanding Fermi–Dirac and Bose–Einstein statistics; if we conceive of quantum objects as individuals in the usual sense we get Maxwell–Boltzmann statistics, which are not applicable to ensembles of electrons, photons etc. A complete realistic interpretation should give us an understanding of why Fermi–Dirac and Bose–Einstein statistics are the correct statistical descriptions for fermions and bosons, respectively.

If a theory, or an interpretation of a theory, meets these demands I will call it *semantically perspicuous*. Equipped with the answers to these questions we can formulate the theory in question in first order predicate logic with a full semantic interpretation. Nothing less than that is sufficient as interpretation. I fully agree with the following statement by Quine (1994, p. 144): ‘I hesitate to claim that this syntax [of predicate logic] so trim and clear, can accommodate in translation all cognitive discourse. I can say, however, that no theory is fully clear to me unless I can see how this syntax would accommodate it.’

Thus, an interpretation of a scientific theory should meet three demands: it should be realistic, it should give us understanding and it should be semantically perspicuous. Nearly all physicists and philosophers believe that it is beyond human capacity to reach that goal. I do not.



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Chapter 3

Individuation and Identity in the Quantum World

3.1 Introduction

One aim of an interpretation of quantum mechanics is to say which properties are real and which are not. All known interpretations have that goal. According to the *pilot wave theory* both position and momentum of a particle are real and observer-independent properties, although they cannot simultaneously be known. *Modal interpretations*, to take another example, purport to identify the real properties of the *physical state*, which is supposed to be distinct from the *quantum state*, i.e. the state described by the wave function. In the *many worlds interpretation* all the different worlds, as described by the quantum state, are real; we just happen to live in one of them, but the collapse is not real. The Copenhagen interpretation claims that the distinction between real and unreal properties depends on preparation conditions: a property is sometimes real, namely when the conditions for its measurement are satisfied, otherwise not. This stance is easily interpreted as a kind of idealism, and that is perhaps one reason why there is a growing discomfort with the Copenhagen interpretation.

I agree with the common stance that it is important to tell which properties are real, but that is not the whole task. Saying that the goal of an interpretation of quantum mechanics is to tell which properties are real, implies that objects are already individuated and identified. This stance is ‘to put the cart in front of the horse’, since individuation of quantum objects is closely connected to the very theory to be interpreted. We have no direct access, unmediated by any theory, to quantum objects, nor do we have any way of individuating objects independently of theory. Thus, part of the interpretation is to identify individual objects, to give identity criteria.

It is illuminating to compare interpretation of quantum mechanics with interpretation of other highly formal theories such as formal logic. In the latter case, the aim of a given interpretation is to tell which kind of objects the singular terms of a particular formalism refer to. It is taken for granted that those objects can be discerned and coherently talked about without using any part of the theory to be interpreted. It is thus assumed, at least implicitly, that individual objects are given to us directly in perception or, in the case of abstract things like numbers, we grasp them directly with our intellect without recourse to the theory to be interpreted. (I do not think this is correct, but it is a common view.) However, that is not possible in the quantum domain: quantum objects (be they systems, particles, waves or whatever) are not accessible to us independently of the theory to be interpreted. Even though

physicists sometimes talk about observing (e.g. electrons or atoms), this does not mean that observation is in a strict sense theory-neutral. An observation, as this term is used in physics, is only theory-neutral in relation to that or those theories which are tested; in spite of that it is highly dependent on other theories used when construing detecting devices etc. But when it comes to interpretative issues we cannot take a purported observation report such as ‘the size of the electron is less than 10^{-13} m’ as a brute fact beyond interpretation. On the contrary, as we will see, this report is the final conclusion from a number of premises, some of which are far-reaching but unnoticed interpretative assumptions.

The conclusion to be drawn from all this is that questions about individuation and identity in the quantum domain are part of the interpretative issue. An interpretation should, among other things, answer the questions regarding what kind of objects there are and how they are individuated and identified. Although these questions are ontological questions, I prefer to discuss them in a semantic framework; identity among objects is best discussed in terms of criteria for using the identity predicate and the question about individuality is the question about criteria for discerning the referent of a singular term as distinct from the rest of the universe.

Some metaphysicians might object that the semantic questions merely concern the structure of quantum mechanics, i.e. our theory about the world and not the world as it is in itself. Against that I would say that the best we could hope for is a theory, which structurally fits the structure of nature. If it does just that, we will have indications showing this fitness, namely that its theoretical predictions come out true in experiments. Quantum mechanics has been tremendously successful in this respect, thus indicating its fitness. Of course this is no proof of correctness, but such proof is unattainable in any case. But so far we can have good faith in this theory and if it is a correct theory we can discern the ontological structure of the world by analysing the implicit ontology in the theoretical description. The way to do this is to ask the semantic question ‘which objects must we assume as referents of our terms in order for the theory to be true?’.

Some of our problems of understanding quantum mechanics, for example the notions of ‘identical particle’ and ‘entangled state’ are caused by insufficient attention to these questions about individuality and identity. Thus, I think that the common strategy of most attempts to interpret quantum mechanics starts too late, as it were. Before we can ask and answer the question ‘which properties are real and which are not’, we must have a clear conception of how the quantum domain is individuated and what the criteria for identity are. Hence, an interpretation of quantum mechanics cannot simply consist of saying which things take the places as individuals in the theory and what properties these things have. All this shows that giving an interpretation of quantum mechanics is a rather different matter than just telling which objects satisfy the open sentences and which properties are real.

Three things need to be done in this chapter: first we must discuss what kind of objects there are and this task has three components: (i) to tell explicitly how the objects referred to are identified as individuals, i.e. as objects separated from each other and from the rest of universe. In connection with this question we need to explain why it is possible to count the number of a kind of purported objects such as photons, without having identity criteria; (ii) to state explicitly identity criteria

in the quantum domain, and (iii), to discuss re-identification over time for enduring objects. I will also say something about the meaning of some terms often used to refer to individual objects.

The second task is to develop a general theory about quantum properties. The central question is: can a quantum object be attributed a quantitative property in situations where the value of the quantity as a matter of principle cannot be determined? In order to analyse the situation I will join Abner Shimony and reintroduce the Aristotelian distinction between actual and potential properties.

The third task is to discuss identity of states over time. This is an important question whose answer is a prerequisite for the analysis of change and evolution. I will, in this context, introduce a distinction between real and apparent changes, which in due course will shed new light on the measurement problem.

3.2 Individuation and Identity

It is generally assumed that identity and individuation are two sides of the same coin: if we have a principle of identity we have implicitly a principle of individuation and vice versa. If the identity criterion is not fulfilled we have two or more individuals and if it is fulfilled we have one individual. But why must it be so? Couldn't it be the case that we have a principle of individuation but not any identity criterion?

Identity statements have the logical form ' $a = b$ ', where ' a ' and ' b ' are singular terms. Hence it seems that criteria for identity presuppose that we have identified individuals, i.e. those particular things being the referents of the terms ' a ' and ' b '. Therefore, it seems that individuation is more fundamental. This train of thought gets support from the fact that some kinds of quantum particles (photons for example) are identical, they are absolutely similar in all respects. Still, they are not one, we can count them. Doesn't that tell us that we can individuate them but lack criteria of identity? No, on a closer analysis this is wrong; on a very reasonable analysis of the concept of object or thing it is indeed true that individuation and identity are two sides of the same coin. I will here rehearse two influential arguments for this view, namely Quine's and Strawson's.

Quine's view

A principle of individuation is operative when we discern something as being an object. But what does it mean to discern an object? Quine's view, as it is given in (Quine 1981, pp. 1–24; 1990, pp. 23–31; 1995, pp. 27–42) is as follows.

Some utterances, such as 'it's raining' do not involve any assumption of objects on the part of the speaker. The singular term 'it' stands there for mere grammatical reasons, it seems. We could just as well just say 'raining' or 'raining here' or 'raining now' in order to convey the same message. The speaker need only learn that under certain conditions it is correct to utter the sentence, otherwise not. Neither do we add anything by claiming that there must be an object *rain* in order that the statement is true. Raining is simply an observable feature of the outer world. Similarly, nothing is gained by saying that substance words like 'water' or 'sugar' designate anything.

They could just as well be seen as expressions for certain stimuli, or what comes to the same, for observable features in the immediate vicinity of the speaker, if what he or she says is true. One should observe here that the analysis takes as a point of departure that sentences, not single words, are the least linguistic units amenable to analysis, hence truth and falsity (or as Quine expresses it, assent and dissent) are more fundamental concepts than reference and extension. This is almost uncontroversial.

On the face of it, individuation seems to emerge when we introduce individuating words like 'raven' or 'house' into our language. These words are taken to denote many things, each raven and each house respectively. Now, Quine asks us to compare the two sentences 'Fido is a dog' and 'Milk is white'. To claim that milk is white is merely to claim that whenever you recognise milk, you recognise something white. 'Milkhood' and 'whitehood' co-occur, as it were. (That is my formulation; Quine would never allow himself such an expression, due to his hostility towards properties.) But the statement 'Fido is a dog' means more than 'whenever Fido occurs, there also occurs a dog'. 'For whenever you point at Fido's head, you point at a dog, and yet Fido's head does not qualify as a dog' (Quine, 1981, p. 5). Thus, it is *predication*, saying that a certain object belongs to a certain category, which creates the difference between individuating terms and other terms.

But predication is only the first step into full reification, according to Quine. He sees reification, assuming objects, as having intermediaries between totally non-committing occasion sentences ('it's raining', 'milk is white') reporting sensory stimuli, to full-blown reification involving universal quantification with explicit criteria of identity. The most decisive step is the introduction of *essential pronouns*, in contrast to *pronouns of laziness*. He exemplifies this contrast by comparing the two sentences 'I bumped my head and it hurts' with 'whenever there is a raven, it is black.' In the first sentence we can replace the pronoun without changing the meaning: 'I bumped my head and my head hurts'. In contrast, we cannot replace 'it' in 'whenever there is a raven, it is black', for the sentence 'whenever there is a raven, a raven is black.' is a weaker statement. Hence, this pronoun is essential – namely essential for expressing the thought that the black thing is the *same* thing as that described by the word 'raven'. Essential pronouns are used when individual things are assumed to have more than one property (again, talking about properties is not Quine's way of expressing things) and individual things, objects, are introduced into our ontology when we have means for deciding questions of identity.

In our regimented language of predicate logic the pronouns become variables and sentences containing essential pronouns are simply quantified formulas involving two or more predicate symbols. Thus 'Whenever there is a raven, it is black' becomes ' $\forall x(Rx \rightarrow Bx)$ ' and we have a convenient way of stating that if there is such a thing as a particular raven, the very same thing is black. Thus, Quine arrives at his famous slogans 'to be is to be the value of a variable' and 'no entity without identity.'

Thus, an object, in the most general sense of the word, is something that is susceptible to predication, and it can be referred to by a definite singular term, i.e. a name or a definite description. When we use a name to refer to a particular object we must have some means for distinguishing the object from the rest of the world, i.e. we must have criteria for its individuality. If we make these criteria explicit, we

get a singular definite description; hence singular definite descriptions are logically prior to names.

A singular definite description is a singular term constructed out of at least one general term, as we can see in ‘the longest bridge in Europe’, ‘the oldest man in the world’ or ‘the red spot on my skirt’. Therefore, individuation is connected to our rules for the application of general terms when construing singular terms.

In parenthesis, substances such as milk and gold could be considered to be objects if we want, according to Quine. His proposal is that all milk in the world is one single object, *Milk*, and similar for other substances. This is certainly not in accord with how we usually view the world, but Quine’s point is that it is possible to have it that way.

Summarising Quine’s view, we posit objects in the full sense of the word when we describe an entity and make a predication, i.e. when we say something *about* this entity. For a predication to be meaningful we must first describe the object by using a singular term referring to the object and then predicate something of precisely *that* object, i.e. we assert an identity. In Quine’s view, then, individuation and identity are two sides of the same coin.

Quine’s views are controversial on many points, but on this matter I think they are less controversial. I will substantiate this claim by referring to Strawson’s views on this matter. In his *Individuals*, Strawson writes:

To know an individuating fact about a particular is to know that such-and-such a thing is true of that particular and of no other particular whatever. One who could make all his knowledge articulate would satisfy this condition for particular-identification only if he could give a description which applied uniquely to the particular in question and could non-tautologically add that the particular to which this description applied was the same as the particular being currently referred to; but we need not insist that the ability to make one’s knowledge articulate in just this way is a condition of really knowing who, or what a speaker is referring to. (Strawson, 1959, p. 23)

This quotation contains two ideas. The first one says that identification of a particular is attained by using a unique description, i.e. a singular term, and the second one is that we have means to ascertain that at least one other description refers to the same particular. As regards the first idea, we should note that we couldn’t ask whether the particular so described really is unique. The process defines uniqueness, for if it were asked whether we really identified only one object by the description, this question presupposes that we have other means for individuation.

It follows that individuation is implicit in the description. This can be seen also by analysing the structure of descriptions. As already said, we use general terms when construing a singular term in the form of a description. Thus every description contains at least one general term and every description presupposes that we have a principle for discerning different objects falling under that general term.

Why then did Strawson add the second clause? At first sight it seems redundant; however there is, I think, a good argument for adding this second clause, namely that if we posit an object as the bearer of one single property and no more, we do not gain anything. Ontologically, objects differ from mere features precisely in that they admit of identification; *a* is the same as *b*. If we could attribute only one property to

an individual, i.e. giving only one definite description, there is no point in positing an object as bearer of properties, i.e. subject to predication; we could just as well say that a feature occurred. (In such an ontology, properties, features or whatever we call that which is indicated by property words, become in a sense the only objects talked about.) Thus interpreted, Strawson's view is the same as Quine's, namely identity is necessary and sufficient for being an object.

The alternative view concerning individuation is to say that identity and individuation transcends all properties and is independent of predication. This position, labelled *Transcendental Identity* by Post (1963), is too much metaphysics for me.

3.3 Identity criteria

An identity statement conveys the information that two singular terms refer to one and the same object. For our purposes the interesting cases are those in which these terms are descriptions, or names with descriptive content. For enduring physical objects we can discern a subclass of identity statements, namely when two descriptions contain explicit or implicit reference to different times. This is of interest in physics because, as observed already by Aristotle, the concept of change presupposes that there is something that can be re-identified, something that changes. Dynamics presupposes criteria for re-identification over time.

The usual way of expressing identity ('*a* and *b* are identical') is easily misunderstood. We cannot truly say that two things are identical, because if they are two things they are not identical and if they are identical they are not two things but one. A better way of expressing identity is using semantic assent, thus: the identity statement '*a* = *b*' is true if and only if the two terms *a* and *b* refer to one and the same object. Or, which amounts to the same thing, the referent of the term '*a*' and the referent of the term '*b*' are identical.

When identifying and re-identifying an object in the real world we must connect its identity to at least one of its properties. To identify a specific object as the referent of a singular term without at the same time attributing some property to it is impossible; we must use a description for specifying the object.

Typically, but not necessarily, identity is defined in terms of *genidentity*; two observations are observations of the *same* object if the two observations, conceived as events, are connected by a continuous trajectory in space and time. (The fact that we often lack conclusive evidence for the occurrence of this trajectory is beside the point.)

Intuitively, one is inclined to think of identity of material objects as identity of the stuff out of which it is made of. This is however often a mistaken assumption. Quine has made this point, for example in the following passage:

Again there is the stock example of the ship of Theseus, rebuilt bit by bit until no original bit remained. Whether we chose to reckon it still as the same ship is a question not of 'same' but of 'ship'; a question of how we choose to individuate that term over time. ... It shows how empty it would be to ask, out of context, whether a certain glimpse yesterday and a certain glimpse today were glimpses of the same thing. (Quine, 1981, p.12)

In this example genidentity and identity of the stuff give different results: genidentity, i.e. space-time continuity, tells us it is the same ship all along, whereas identity of stuff that it is not. I think common language utilises genidentity in this case, and removal of parts of the ship is irrelevant for identity. But if we replace a major part of the ship at a single occasion, we would perhaps hesitate to say that it still is the same ship. This shows that in common speech the identity criterion is a bit vague. It is not sharper than what is needed in the case at hand.

Some might object by claiming that it is not really the same ship after the removal of all planks. It sounds like a deep metaphysical conflict, but for me, just as for Quine, this is a question of how we use the general term 'ship'. Linguistic practice decides the matter.

For physical objects having a rather well-defined outer boundary it is natural to use genidentity as the identity criterion, the reason being that we discriminate similar-looking objects by attributing to them different places. Thus, for physical objects genidentity is the identity criterion par preference. (Again, this is a statement about linguistic practice!) Physical objects, time and space are a package deal, as it were. We discern physical objects in space and time, thus constituting objects and space-time relations by the very same process. However, in quantum mechanics this is not generally feasible. The reason is that using genidentity as identity criterion presupposes that two objects cannot simultaneously be at the same place; this in turn entails impenetrability and this condition is not fulfilled in the quantum domain.

3.4 Individuation and Identity in the quantum domain

In order to see the significance of the analysis of the previous section in quantum theory, let us look at an example. Suppose we hear someone saying 'the photon, which hit the detector D at time t , had energy E '. In this sentence a particular object is identified as 'the photon hitting the detector D at time t '. This expression is a unique description. Does it describe an individual object? The answer is no. We can predicate a definite energy, momentum or polarisation to this photon, but we cannot give *another* unique description, which identifies the *same* particular. This is so because photons are annihilated in absorption processes. One might be tempted to elucidate the situation by saying that the only way of referring to this particular object is by using this individuating description, but that is to misrepresent the situation, because the photon is not an individual object; no identity between two descriptions can be given. In order to talk meaningfully about an object we must first identify it and then make a predication, i.e. say something about this particular thing. Lacking two definite descriptions, no meaningful statement about an individual object can be made. In general then, quanta of energy are not individual objects. We can compare with our use of mass terms like 'water'. Lacking criteria, we never ask ourselves whether a portion of water put into a lake or into sea is the same portion as is later taken up. We only care about the amount, since any portion of pure water is exactly like any other portion. In specific circumstances it is possible to say of two descriptions of portions of water that they describe the same portion, but generally it makes little sense.

There is one important difference between water and energy, namely that exchange of energy is quantised. Moreover, in microphysics we must take into account that observations of small portions of energy are exchanges of energy. This is an aspect of the measurement problem, to be discussed in Chapter 6.

Of course, all physicists are well aware of the lack of individuality of photons. In a textbook in quantum optics we find this illuminating passage:

Although the word ‘photon’ is ubiquitous in quantum physics, it is commonly used in several different ways. One common use interprets Eqs. (1) and (2) as creating or annihilating a photon:

$$a|n\rangle = \sqrt{n}|n-1\rangle \quad (1)$$

$$a^*|n\rangle = \sqrt{n+1}|n+1\rangle \quad (2)$$

According to this interpretation, the photon is a quantum of a single mode of the quantum field. As such it fills the cavity of quantisation. One might argue that although precise, such a definition of the photon is not particularly useful from an experimental point of view. In this context ‘photon’ is often used to describe a relatively localised (and therefore multimode) packet of radiation with an average energy of $\hbar\omega$. In such situations, photons are always multimode objects since Fourier analysis tells us that the localisation of any wave packet requires a superposition of modes. This ‘particle’ interpretation has intuitive appeal, but presents the drawback that each source emits its own kind of wave packet and hence such ‘photons’ have a wide variety of analytic forms or worse, no analytic form at all. ... The most common use of the word photon is as a unit of electromagnetic radiation with the energy $\hbar\omega$. This use is found in chemistry, in some parts of atomic physics dealing with ‘multiphoton processes’, and often in semiconductor physics and in engineering optics. In these and similar cases, the field is classical, with absolutely no quantum character. ‘Photon’ is used as a catchy synonym for ‘light’ as in ‘photon echo’ and ‘photonics’. Because the photon is used in so many ways it is a source of much confusion. The reader always has to figure out what the writer has in mind. (Meystre & Sargent III, 1990, pp. 325–326)

A photon is thus merely a quantity of electromagnetic energy and it has in general no well-defined position, except in interaction processes, i.e. in the moments of its creation and destruction.

What then are the individual objects? My suggestion is that the only entities we need quantify over and accept as individual entities are those things referred to by the word ‘system’. More about this issue below.

3.5 Individuation and the Identity Postulate

Several classes of objects are said to be classes of identical particles, for example electrons and photons. Such a locution has always made me wonder; if two objects are identical, they are not two objects but one; how, then, is it possible that they are two? The expression ‘ x and y are identical particles’ must hence mean something else than that x and y are the same object. And so it is: the said identity among

particles means that permutation of particle indices in multi-particle states has no empirical consequences. This is expressed by *the Identity Postulate*:¹

IP: A permutation of particle labels in a multi-particle state gives the same expectation values of all observables.

Another version is:

IP': A permutation of particle labels in a multi-particle state gives the same state.

It is possible to claim that the second version is logically stronger, which amounts to assuming that two different states can have exactly the same expectation values for all observables. However, as far as I know, all agree that these versions express the same fact of the matter. In any case I will take it for granted that they are equivalent.

A redistribution of the particles in a multi-particle state does not change anything (except the sign of the wave function, but that is a change of phase factor and has no effect on anything observable) and we have lost any possibility for even asking whether an individual particle is in this or that particular state. Consider a two-particle state comprising two fermions, labelled 1 and 2, both of which can be in two different states *a* and *b*. As no two fermions can be in exactly the same state we have only two possible states of the system. The wave functions for these states can be written

$$\Psi_1 = \phi_1^a \otimes \phi_2^b \quad (3.1)$$

$$\Psi_2 = \phi_2^a \otimes \phi_1^b \quad (3.2)$$

According to the IP, a permutation of particle indices does not change any observable feature of this system, so these two *state descriptions* refer to the same physical state. In other words, a permutation of indices does not represent any real change. This shows that particle indices do not confer individuality on particles. That in turn implies that there are no intrinsic properties connected to particle indices. Or stated otherwise, if a particle label is to have any significance, it must be connected to some property, which is denoted by a general term used in construing a singular term, i.e. the unique description referring to that particular particle.

The two fermions have exactly the same intrinsic and permanent properties. The difference between them is that one is in state *a* and the other in state *b*. *This* gives individuality to them. Hence, the locution ‘particle 1 may change from state *a* to state *b*’ makes no sense because such an expression presupposes that individuality is carried by the particle label and not by the state, which is wrong. This is why equations (3.1) and (3.2) refer to the same state.

The state labels ‘*a*’ and ‘*b*’ are short for sets of quantum numbers. If the particles are not in bound states, i.e. if no quantum numbers can be attributed to them, then

1 When identical particles are talked about in quantum mechanical textbooks one does not mean that identical particles are one and the same particle, as is the common meaning of ‘identity’ in philosophy, but rather that they are indistinguishable by their state properties.

they lose all individuality. Thus, they are rather to be considered as two completely similar portions of a common substance. However, still they are called *two* portions, e.g. two electrons. How come? We cannot give any individuating property, not even position, because these objects are not impenetrable, they can simultaneously be in the same place. Additionally, a definite position is usually not attributable to them.

The reason why it still makes *some* sense to talk about two things is quantisation of charge. The term ‘electron’ could be seen as short for ‘a portion of electric charge equal to -1.6×10^{-19} C’. That is the reason why we can intelligibly use count terms in spite of lacking generally applicable identity criteria, since this fact enables us to tell how many electrons there are in a situation in which we only know the total negative charge. The lesson to be learned is that identity is not a necessary requirement for using count terms, contrary to what one might suspect. Calculations of the number of, for example, electrons are hence not really a counting of individuals; it is a calculation of the amount of charge. We have here an interesting case where numbers can be interpreted only as cardinal numbers and the ordinal interpretation is excluded.

3.6 The Identity Postulate and Identity of Indiscernibles

One might reasonably ask how it is possible that we can talk about several electrons, several photons, etc, when it does not matter how we distribute them in a common state. How can we know that there are several such things when they are completely similar and nothing observable follows from exchanging two of them? I am thus rehearsing the intuition behind Leibniz’s principle of identity of indiscernibles: if two things have all their properties in common, then they are identical and there is in reality only one object. Leibniz’s principle can be stated formally in second-order predicate logic as

$$\text{LP: } \forall F(Fa \leftrightarrow Fb) \Rightarrow a = b$$

where F is a predicate variable which, as values, takes the possible attributes of the individuals a and b . The question is thus whether Leibniz’s principle is valid for quantum objects?

Stephen French and Michael Redhead (1989) have argued that Leibniz’s principle is actually false. Their argument applied to fermions goes as follows:

Assume that two fermions may be in one of two states $|r\rangle$ and $|s\rangle$. There seems to be four possible combinations of states; (1) both particles are in state $|r\rangle$; (2) both particles are in state $|s\rangle$; (3) particle 1 is in $|r\rangle$ and particle 2 is in $|s\rangle$; and (4) the reversed state. However, this is not the quantum way of describing the possible states. Instead the four possible combinations are

$$(1) \quad |r\rangle \otimes |r\rangle$$

$$(2) \quad |s\rangle \otimes |s\rangle$$

$$(3) \quad \frac{1}{\sqrt{2}}(|r\rangle \otimes |s\rangle + |s\rangle \otimes |r\rangle)$$

$$(4) \quad \frac{1}{\sqrt{2}}(|r\rangle \otimes |s\rangle - |s\rangle \otimes |r\rangle)$$

where the first component of the tensor products refers to the state of particle 1. As fermion wave functions are antisymmetric in particle labels, the states (1) (2) and (3) are not accessible to this pair because these are symmetric upon particle permutation. Therefore, the fermion pair can only be in state (4).

All properties, both monadic and relational, can be expressed as probabilities. Monadic properties of the first particle can be given on the form $\text{Prob}^w(\text{quantity } Q_1 \text{ has the value } q^a)=p$ and its relational properties as conditional probabilities of the form $\text{Prob}^w(\text{quantity } Q_1 \text{ has the value } q^a/Q_2 \text{ has the value } q^b)=p$. It is then not difficult to show that both particles have exactly the same properties *as long as they belong to the state described by (4)*! If, on the other hand, as French and Redhead noted, we identify the relevant properties of the two particles with the pure states $|r\rangle$ and $|s\rangle$ they are in, we would have to conclude that there is no answer to the question ‘Do they have the same properties?’, because the two particles cannot individually be ascribed pure states. An analogous argument concerning bosons gives the same conclusion, i.e. that two bosons may have all properties, interpreted as probabilities for observable values, equal.

French & Redhead concluded that Leibniz’s principle is false. I think they are wrong; in a state described by (4) we do not have two objects, two particles, with exactly the same properties. We have *one object*, i.e. one quantum system, because individuation must be associated by some property and if there is no property that can be used to individuate between these particles, there are not two objects but one! This object is not made up of two individual things. We could say that it is constituted of two portions of stuff, i.e. charge, but not that it is constituted of two individuals. But if we have only one object, Leibniz’s Principle of the Identity of Indiscernibles does not apply, and so it cannot be false.

This analysis provides an explanation of Fermi–Dirac and Bose–Einstein Statistics, both of which appear a bit astonishing to the preconceptions of the layman. When calculating probability distributions we do as follows: the first step is to count the number of objects and the number of possible states. Let us take as example throwing two coins. We have two objects and each object can be in two states, *head* and *tail*. That gives us four possible outcomes: (head, head), (head, tail) (tail, head) and (tail, tail). (It should be observed that the two coins can be distinguished!) Hence the probability for one head and one tail is 1/2. Now if these coins were bosons following Bose–Einstein statistics we would get the result that the probability for one head and one tail would instead be 1/3. The usual explanation is that bosons are ‘identical’ which is supposed to mean that the two outcomes (head, tail) and (tail, head) are physically identical. This explanation is in my opinion too vague because it is left open whether the two states really are physically the same state, or if it merely *counts* as the same state. But what could it mean to say that the two states count as one state? Either it is one state or two and we must count correctly in order to perform a correct calculation. The only coherent interpretation is that the two state descriptions ‘(head, tail)’ and ‘(tail, head)’ refer to one and the same state. But that implies that there is no individuating property for these objects, which in turn implies that they are not two different objects. We have only one object, usually said to be made up of two particles, and this object can be in three different states, *hh*, *ht*

= *th*, *tt*. In my view Bose–Einstein statistics confirms the view that bosons are not individual objects.

The example with two coins falling head or tail up is not suitable for discussing Fermi–Dirac statistics because the exclusion principle forbids the states (head, head) and (tail, tail) leaving only one state left, and the probability distribution collapses to unity. However, in more complicated cases we should count all outcomes transformable into each other by a permutation of particle indices as one and the same state. Just as in the boson case, this can only be understood as the so-called ‘particles’ not being individual objects that can be quantified over.

Nevertheless, we talk about two or several particles. Is that simply a mistake? It need not be; it depends on which ontological commitments we make. It is not wrong to use the individuating term ‘particle’ if this term only refers to a portion of some quantity and nothing more. We should keep in mind that the individuating term ‘particle’ can be used only because interaction is quantised. If an interaction could consist in an exchange of any amount of energy, big or small, we would have no use for the concept of a particle; rather, we would talk about the amount of energy, portions of charge etc. But as interactions between charges are quantised, it is in some contexts possible to talk about particles, and likewise electromagnetic energy is quantised, giving us a reason to use the word ‘photon’. Two or more such portions can be united to form a multi-particle system, which behaves as one object. Hence, the individuation associated with properties such as spin, energy and momentum is such that the state description (4) above refers to *one* object.

This does not preclude the possibility of talking about one specific electron passing from one place to another in cases where only one portion of negative charge is present. It only shows that in multiparticle systems we have not several objects, but only one, namely the entire system.

The state description (4) does not tell us anything about the places of interaction; therefore it is not contradictory to say that this system is one object as regards all quantitative properties except charge, but two objects as regards interaction. (I will in Chapter 5 show that all interactions take place at well-defined places and are basically an exchange of energy.)

In my view, the Leibniz principle is a metaphysical principle, which cannot be falsified by empirical facts. Its real power is to provide a second-order principle regarding individuation, a principle that says that individuation is connected to at least one property. Seen in this light it is a precursor to Quine’s and Strawson’s views on individuation.

In conclusion, to individuate is to construct singular terms out of general terms. Thus, individuation presupposes a principle connected to the general term used for singular descriptions.

3.7 Individuation by different properties

I guess that most readers find it paradoxical to claim, as I did in the preceding section, that we have one object being made up of two portions of charge. I admit that it seems paradoxical, but it is not self-contradictory. I do claim, though, that as regards

charge or exchange of energy, a system made up of two fermions is to be regarded as two separate objects, whereas when we consider other observables of this system, it is one object. There is no reason to suppose that if we apply two different predicates, i.e. observe two different properties of a system, these two properties individuate the domain of discourse in the same way; that they, as it were, cut the cake similarly. When we describe a quantum system by giving its wave function we implicitly describe it by means of its observables, as that term is used in quantum theory. When we describe the system by saying that it is made up of two electrons we describe it using the property that interaction is quantised; to say that we have two electrons in a system implies among other things that the system may interact at two different places simultaneously. I will expand on this in the next section.

Thus, the question how many objects we have in a certain region of space-time makes no sense out of context, just as sameness simpliciter makes no sense. We can only ask whether *a* and *b* are the same *person*, the same *number*, the same *state* etc, and similarly, we cannot ask how many things in general we have, only how many persons, how many numbers, how many states etc.

3.8 Cardinal and ordinal numbers

It also sounds paradoxical that we can count, e.g. photons without having any criterion for identity, because identity is involved in the definition of number. According to the standard Frege–Russell analysis, numbers are conceived of as sets of sets of objects. The number two, for example, is the set of all sets which are such that there is in the set an object *x*, an object *y*, *x* and *y* are not identical and if there is an object *z* in the set, it is either the case that $z = x$ or $z = y$. Thus, counting presupposes identity. How then is it possible to count the number of, for example, electrons in a container or the number of photons in a cavity filled with electromagnetic energy?

The answer is that numbers have two properties and can be used in two different ways, namely to decide the cardinality or to decide the ordinality of something. A number can either be used for stating the *number of things* in a set or to state the *quantity*, or the *intensity*, of some measurable attribute. When using a cardinal number you answer the question ‘how much’. Example: ‘How much energy is required to melt one gram of ice at 0°C? Answer: 333 kJ.

Numbers usually also have the ordinal property, i.e. each ordinal number can be described as an ordered set containing all inferior numbers as subsets. Thus, when we attribute an ordinal number to a set of real physical objects we thereby attribute to these objects an ordering. Ordinal number can thus be used for counting the number of objects in a set, for example: ‘How many were there before you in the queue?’ This analysis shows that it is possible to attribute a number to a quantity, without it being possible to subdivide this quantity into smaller parts; in such a case we only use the cardinal property of numbers.

Most often we need not bother about the difference between the cardinal and ordinal property of numbers, but in quantum mechanics we must do that, just because electrons and other quantum objects lack individuality.

The important thing is that we can use numbers in the cardinal sense when talking about quantitative properties of certain objects without presupposing that these objects can be looked upon as sets of constituent objects. If only we have a well-defined procedure of measuring a quantity, i.e. of construing a homomorphic mapping of (a property of) objects onto the real numbers, the attribution of quantitative properties is fine.

When we talk about multiparticle states we do not really give the number of individual things when we say that it is made up of x electrons, but that the total amount of charge is x times that of the charge unit. When talking about the number of electrons, photons and the like, we apply cardinal numbers, but this does not logically imply that these statements should be interpreted as statements involving ordinal numbers, i.e. about the number of individual things.

One might ask why it is possible to express the amount of charge in a confinement of some sort by saying that there are x electrons without being able to count the number of individuals and the answer is of course that the quantity of relevance, i.e. charge, is quantised. It is a brute fact that charge cannot be infinitely divided into smaller portions in interactions. (The fact that quarks have non-integer charges does not contradict this statement since quarks do not interact individually.) Essentially the same argument applies to photons, just by replacing the word 'charge' by 'electromagnetic energy'. In the case of photons we count the number of interaction processes in which energy is exchanged. Suppose we produce an electromagnetic field with a reasonably well-determined energy and frequency. Can we then say that it is made up of x individual photons of this frequency? I am prone to say no, for two reasons. The first is that we cannot count photons without destroying them and hence we cannot satisfy Strawson–Quine's criterion for something being an individual object. The other reason is that this radiation field interacts with its surroundings as a unit, not as a collection of individual things. We find ourselves in the peculiar situation that we can say that we have produced a definite number of photons, thus getting an electromagnetic field of a certain strength, and we can say that this field can produce exactly the same number of excitations in the surroundings when these photons are destroyed, but we cannot say that in the time interval between production and destruction we have a definite number of individual photons, because counting the number of photons is, in that situation, impossible. The question of identity of photons makes no sense in the case of the isolated ensemble. There is even empirical evidence for this. Phlegor & Mandel (1968) performed a series of interference experiments with two independent laser beams directed to a stack of thin glass plates. In the plates, the two beams overlapped, thus making interference possible. The outgoing beam was split into two, directed to two photo-detectors. If the two incoming beams interfered this should be observed as a correlation between the counts of the detectors and the phase angle between the incoming beams because the phase angle between the two incoming beams determines the probability distribution for detection in the two photo-detectors. (At a certain phase there should occur clicks only in the left counter, and vice versa.) This is itself not completely refuting the assumption that photons are individual objects with identity criteria. However, if, as was the case, the intensities in the beams are so low that at most only one photon was present at each moment of time the conclusion is unavoidable; there are no

distinguishable objects in the beam. The result was in excellent agreement with the predictions.

The remarkable thing is that it proved possible to produce interference effects with two independent light sources; Dirac, among others, had predicted that interference would not be observed with independent sources. But with the long coherence length of laser light this is possible.

This experiment is a telling argument for the view that photons are not individual particles. The *electromagnetic field*, suitably specified in a specific context, is the individual object and that field is able to interact with one of the detectors at each point of time. A photon is simply a portion of energy given away from, or taken up by, the field.

In passing, it could be mentioned that adherents to the pilot wave theory have a different analysis of this experiment; they say that it proves the existence of empty waves which guides the particles. I will discuss their views in Chapter 9.

In conclusion, the fact that we in some circumstances can count the number of electrons, photons and other quantum particles without being able to confer individuality to them is explained as a consequence of the quantisation of charge and quantisation of interaction, which implies that using cardinal numbers is meaningful. It does not, however, imply that the conditions for using ordinal numbers are fulfilled.

3.9 On the individuation of systems

A common noun in quantum theory is 'system'. Its generality enables us to use it without presupposing anything at all about the properties of the things denoted. But what is a system? How is it identified and how is the domain individuated?

It seems to me that most physicists take it that systems are identified through their wave functions, or, more generally, state vectors, although this view is seldom explicitly stated. I will follow this terminology. Hence, the identity question for systems is a question about identity for wave functions.

Identity among wave functions is given by the condition that two wave functions are identical if they return the same value when the same argument is inserted. Furthermore, as the total phases of wave functions are irrelevant for expectation values, two wave functions differing only by a constant factor refer to the same system. (Identity over time will be considered in Section 3.15.)

As the product of two wave functions is a new wave function this mathematical operation corresponds to the fusion of two systems into one system. This mathematical operation has physical significance in those cases where some properties (i.e. probability distributions of values of observables) cannot any longer be attributed to any of the formerly individual systems independently of the properties of the other. The criterion for having one individual system is thus: A and B are two individual systems if and only if all probability distributions of values of observables attributable to A and to B are statistically independent of each other. It follows from this criterion that an entangled state must be conceived of as one system, because an entangled state is a state for which it is impossible to ascribe definite probability distributions for each part separately.

This principle of individuation is often violated in discussions about quantum mechanics. For example: rather than saying that two systems have undergone fusion into one, it is often said that the two systems are in an entangled state. This way of expressing things causes conceptual confusion.

The term ‘quantum system’ is seldom used rigorously as regards individuation. For example, in some discussions about the measurement problem, the object to be measured and the measuring device are referred to as two interacting systems, whereas in some other discussions the two are regarded as one single system, in both cases without reference to any explicit criterion. Similarly, singlet states made up of two fermions are sometimes considered as two interacting systems, sometimes as one system made up of two particles. The usual way of conceiving the situation is to regard them as two systems without referring to any explicit criterion.

Paul Teller has argued that a holistic conception of a certain state of affairs, often dismissed as mysticism, can be made precise in terms of relational properties that do not supervene upon non-relational properties. Using this ‘relational holism’ as he calls it, provides a kind of explanation of nonlocal phenomena (Teller, 1986). This is certainly one way to take it, but I prefer to individuate differently, in order to avoid irreducible relational properties.

Some authors claim that the wave function does not refer to an individual physical system, but to ensembles of identically prepared systems, and that the interpretative problems occur when we try to say something about individual systems. It is claimed that using the formalism as a description of individual systems is not justified. My response is that there is nothing wrong in restricting one’s talk to ensembles only, but the conclusion that talk about individual systems is unjustified and illogical is not well motivated. I agree that most problems of interpretation are avoided if we talk only about ensembles, but that does not mean that talk about individual systems is incoherent. In my view, to impose the restriction only to allow talk about ensembles begs the question of interpretation.

3.10 Coupling and decoupling of systems

In many situations, for example when considering a collision or a measurement process, we start by describing two systems, using two independent wave functions. Then these systems are brought into contact, thus forming one system, according to the criterion proposed above. The mathematical representation of this fusion is the multiplication of the two wave functions, with the effect that the time evolution of the entire system will be governed by one common operator $\exp(-i\mathbf{H}t)$, where $\mathbf{H}$, the Hamiltonian, represents the energy. This physical operation forces the two original systems into resonance with each other, which in turn implies that an individual probability distribution for each subsystem cannot be defined any longer, only a joint distribution is well defined. One might now ask about the necessary physical conditions for this fusion of two individual systems into one; what do we mean by ‘contact’? What are the physical conditions for this to occur? And similarly, what would cause the joint system to break up into two separate systems again? These questions are, it seems, closely connected to the measurement problem, for when we

measure a dynamical variable attributable to one of the subsystems we break up this joint distribution. More about this later.

Returning to our original question, what more precisely do we mean by saying that two systems are brought into contact and under what conditions will this result in one system?

If we adopt the field theoretical viewpoint for a moment we see that the process usually called ‘bringing two systems into contact’ is not a change from one state of two completely isolated systems to a state of these two being in perfect contact. In the field theoretical perspective these two systems are two different fields and they are never completely isolated. Bringing two fields into contact means increasing the coupling between them and that in turn means that we increase the probability for interaction. The contrast ‘not-contact/contact’ is thus not a sharp dichotomy. Bringing two systems into contact is better described as a continuous change from a very weak to a stronger coupling. Assuming that two distant systems are not in contact means neglecting a minute probability for interaction. This omission is of course completely rational when calculating probability distributions of observables, but we should not confound a pragmatic distinction for a fundamental one.

That the assumption of locally isolated fields is not a fundamental feature is evident from the motivation of the so-called *cluster decomposition principle*. Here is how Stephen Weinberg introduces this principle:

It is one of the fundamental principles of physics (indeed of all science) that experiments that are sufficiently separated in space have unrelated results. The probabilities for various collisions measured at Fermilab should not depend on what sort of experiments are done at CERN at the same time. If this principle were not valid, then we could never make any predictions about any experiment without knowing everything about the universe.

In S-matrix theory, the cluster decomposition principle states that if multi-particle processes $\alpha_1 \rightarrow \beta_1$, $\alpha_2 \rightarrow \beta_2$... $\alpha_N \rightarrow \beta_N$ are studied in N distant laboratories, then the S-matrix for the overall process factorizes. (Weinberg, 1995, p. 177)

That the S-matrix factorises means of course that the outcomes of the N different processes are statistically independent, which in turn means that there are no interference terms between the elements representing different experiments. The cluster decomposition principle is thus a locality principle. Two things are interesting here. The first is that the cluster decomposition principle is not derived theoretically, but motivated empirically. It is stated as a first principle and is extensively used for solving calculation problems in quantum field theory; it enables the physicist to set many terms in the matrix equal to zero, which is to say that the probability for the transition represented by that term can be neglected. It is not *proved* to be zero!

The second observation to be made is that Weinberg does not discuss the exception from this principle, namely the nonlocal correlations occurring in EPR experiments. Perhaps he thinks that this exception is not relevant for an exposition aimed mainly at describing the fundamental structure of quantum field theory relevant for the working physicist. It could reasonably be claimed that the nonlocal correlations as a matter of fact can almost always be ignored by the particle physicist. However, in the context of discussing EPR phenomena it cannot be ignored.

The occurrence of nonlocal correlations shows that the criterion for application of the cluster decomposition principle, events being sufficiently separated in space, is not always a sufficient criterion. I will leave this discussion for now, returning to the subject in Chapter 8.

Returning to the main line of exposition, i.e. how to individuate systems, we can thus say that systems are individuated provided that weak enough couplings are neglected; if so, each system can be attributed properties independently of other systems; and if the couplings to other systems in fact are very weak, the resulting probabilities are very nearly correct.

How weak is weak enough? It all depends on the circumstances. Invariably all systems are more or less coupled to all other systems. If the strength of one coupling is 100 times stronger than that of any other coupling present, we can safely neglect these latter as long as we do not use very sensitive or low-noise detectors. This situation is well known to all physicists working with calculating real problems.

Neglecting weak couplings thus enables us to describe a chosen system as independent of other systems and to do so uniquely without reference to the properties of other systems. This also means that the spatial extension of a system is assumed to be confined within a certain volume and minute ‘tails’ of the wave functions are neglected. As we will see in Chapter 4, real systems are spatially extended entities, which mostly behave as waves and the extension in space of these waves is, as a matter of principle, unlimited.

Omitting the effects of the non-confinement of quantum systems as empirically irrelevant is however not always legitimate as we will see in Chapter Eight.

The concept of a system has the same place in quantum mechanics as the concept of a material thing or a body in classical physics. In experimental physics, where individual systems are studied the identity criterion at work is that of classical physics, namely *genidentity*: two descriptions of a quantum system are descriptions of the *same* system if there is a continuous space–time trajectory connecting the referents of these two descriptions and no other trajectory for a system of the same kind is close enough to cause trouble. This identity criterion is applied in many experimental situations in which systems are produced in some sort of generating equipment and then injected into an experimental device. Effective isolation is crucial for the validity of the empirical conclusions, for we want to be able to say such things as ‘if a system is produced in this type of generator, it behaves such and so.’ Observe the essential pronoun ‘it’ here. It signifies the use of an identity criterion.

3.11 The state of a system

Chris Isham (1995, p. 70) has discerned five different interpretations of the expression ‘the state of a system’. They are:

- An individual system,
- Our knowledge of the properties of such a system,
- The result of any measurement that could be made on such a system,

- A collection (real or hypothetical) of identically prepared copies of the system on which repeated measurements are to be made,
- The results of repeated measurements that could be made on such a collection.

From the context it is obvious that the first alternative should be understood as that the state is a general attribute of an individual system, an attribute that is independent of our knowledge. Thus, the statement ‘the system is in state *S*’ is true or false independently of our knowledge or state of mind. I have already indicated that this is a realistic view. The second alternative is obviously a kind of idealism. The third alternative is expressed in counterfactual terms and invites a host of problems concerning the meaning of counterfactuals. The fourth and fifth alternatives are different ways of giving a statistical interpretation of quantum mechanics. They are both acceptable as far as they go, but they do not answer the realist’s question about the state of an individual system; for as already discussed in Chapter 2, when we try to infer the state of an individual system from the collective state we confront problems that cannot be avoided without introducing stronger assumptions.

3.12 Observables, operators and properties

I have referred to Quine’s slogan ‘no entity without identity’ with approval and used his arguments for this doctrine against viewing electrons, photons etc as individuals. Clearly, then, I should be consistent and apply the same principle when discussing properties. Hence, we need a principle of identity when quantifying over properties, and if such a principle cannot be found we should dismiss talk about properties as incoherent.

To begin, it is quite clear that physicists quantify over properties. We can hear expressions such as ‘some properties are conserved’ or ‘some properties are invariant under a class of transformations’. Such talk indicates a commitment to properties as real entities and we might reasonably ask how properties are identified.

As a preliminary remark we should observe that physicists and philosophers attribute different meanings to the word ‘property’. Some philosophers use the word ‘property’ indiscriminately as the referent of every general term. For example, the general term ‘...being 100,000 km above the surface of the moon’ denotes a property according to this usage. Other philosophers with more nominalistic tendencies tend to think there are no properties at all, i.e. that general terms are non-referring. Quine is a good example: he dismisses all talk about properties and restricts himself to mere talk about predicates, his main argument being that properties are intentional and it is impossible to give clear identity criteria for intentional objects. Predicates, on the other hand, are linguistic entities and easily identified and individuated.

Physicists usually think of a rather limited class of quantitative magnitudes, such as mass, charge, momentum etc, when they talk about properties. But the word ‘property’ is not often heard; physicists instead use the expression ‘physical quantity’. Are these physical quantities properties and is it possible to individuate them in any clear way?

Let us first observe a distinction. The expression ‘physical quantity’ is used in two slightly different senses, a generic one (referring to a universal or property if such things are accepted) such as mass or length, and a specific example of that quantity (a trope, as it is called in metaphysics) such as *the length of a particular object*. The following definition of physical quantity in the latter sense is arrived at by *Union Internationale de Physique Pure et Appliquée*. In the document IUPAP-25² the following definition is given:

A physical quantity is expressed as the product of a numerical value (i.e., a pure number) and a unit:

physical quantity = numerical value * unit.

The first, generic sense of ‘physical quantity’ is not defined in this document, but it is easy to supply one. A physical quantity in the first sense can simply be defined as the set of all physical quantities in the second sense having the same unit (or a multiple thereof). To give an example, the universal *length* can be defined as the set of all attributes having the unit metre or a multiple thereof. Thus, we have a purely extensional definition of a physical quantity in the first sense and its identity criterion can be given in set theoretical terms; two quantities are identical if and only if they have the same members. Thus, when I talk about properties, this should be understood as physical quantities in the first sense. I hope this is clear enough even for the most austere nominalist. We need not take any stance, in this context, as regards the difficult question of whether there exist properties that cannot be defined extensionally.

Quantum objects have many properties, some of which are observables. Observables are the quantum analogue to classical dynamical state functions and they can be identified as those properties of a system that are represented by Hermitian operators.

Some properties attributable to a system are not observables represented by operators. Two obvious examples are mass and charge. These properties cannot be used to define the dynamical state of the system.

One might be tempted to think that observables are rather like dynamical classical functions such as velocity or momentum, except that observables do not have well-defined values at all times. However, there is a further difference. There is a class of *maximal observables* such that for each such maximal observable, if its value a_i is known, a_i determines the wave function for the system and, given the realistic view presented earlier, a_i constitutes the complete information about the system, which is to say that the value of one single observable *determines the state*. This is remarkably different from classical mechanics, where we need to know the values of at least two functions such as position and momentum in order to determine the state completely.

How is the correspondence between an observable and an operator found? In elementary textbooks in quantum mechanics the correspondence is normally

2 Printed and distributed in *Physica 146A* (1987, pp. 1–68).

introduced as two postulates introducing position and momentum operators to be used on the wave function:

$$x \leftrightarrow x_{\text{op}}$$

$$p_x \leftrightarrow i\hbar \frac{\partial}{\partial x}$$

From these two, other operators can be derived. But why use these operators in the first place? Ultimately, of course, the reason is that when used they give predictions in agreement with experiment. But the philosophical quest for reason is not fully set to rest with this answer. In the context of this book deeper reasons must be sought for the use of a derivative as mathematical representation of momentum.

An illuminating approach can be found in the path integral formalism. I will rehearse a derivation of the operator-observable correspondence given by Michio Kaku (1993) in the next chapter. The net result is that, starting with two postulates about transition probabilities, we get transition functions of the form $\exp(ikx)$. (The same argument was once given by Landé, 1965.) The derivation $\partial \exp(ikx)/\partial x = ik \exp(ikx)$ brings down ik which is proportional to the momentum eigenvalue $p = \hbar k$. The operator $i\hbar \partial/\partial x$ thus gives precisely the momentum p and we have an explanation of why this operator is a representation of momentum.

3.13 Actual and potential properties

What does a measurement on a quantum system tell us? Does it give us information about the state of affairs before the measurement, or does the measurement itself produce a new state of affairs, for which it gives us ‘updated’ information? The first alternative means that the measured object has some property, which can be expressed as ‘the observable O has the value v_0 ’ and that this property can be attributed to that object independently of the measurement process. This stance is classical realism in the form of value determinism and anyone adopting it in the quantum domain will run into serious trouble. I believe such a stance is impossible.

The alternative view is that the value of the observable is indeterminate before the measurement. I interpret this as that the measurement process brings about a new state of affairs. Thus, for example, a quantum system not being in a momentum eigenstate does not have a well-defined momentum before the measurement interaction. This view invites us to adopt an Aristotelian distinction between *actual* and *potential properties*. It can be said that a measurement is an *actualisation* of a *potentiality*.

Heisenberg (1962, p. 185) and Shimony (1986, p.184) have used the same concepts. Shimony holds more realistic attitudes than Heisenberg and thinks that actualisations need not be caused by our observations. Thus, he admits that actualisations might not be uniquely occurring in measurement situations. The question is what this might mean; what are the truth conditions for an ascription of a potential property? How do we know that an object has a potential property if it never is actualised?

Essentially the same problem arises in connection with concepts like dispositions, abilities, capacities, and causal powers. In all such cases we ascribe to an object a permanent property which may or may not manifest itself. And consequently we face the problem of giving the truth conditions for an ascription of a property not yet manifested.

According to a common view, dispositions are scientifically respectable only if they are reducible to structural properties. For example, we explain the dispositional property of fragility by saying that the chemical bonding in the object being fragile is weaker than what is necessary for the constituents to keep together when moderate external forces are acting on it. Thus, we are justified in saying that the dispositional property of fragility can permanently be predicated of the object because it can be translated to categorical form. If this analysis in terms of structural properties is not feasible, dispositions are not acceptable because we can give no clear conditions for their application.

Can we give an analysis of potential properties in terms of one or several categorical properties attributable to the objects in question? The answer is yes; the wave function for a specific system at a specific time can be seen as such a categorical property or categorical description, because the wave function is an eigenfunction to a set of operators and the observables corresponding to those operators that have definite values, and these are the actual properties. If we know the total wave function for the system we can calculate probability distributions for all other observables which do not have precise values, i.e. the potential properties. When measuring such an observable one of the possible values will be the outcome; therefore, we can say that a certain object has a *disposition to actualise* a certain property. We can give the probability for this actualisation, and the truth conditions for this statement can be given independently of whether this property really is actualised or not.

Let us look at a simple example. Suppose a spin half particle has been prepared to have spin = $\hbar/2$ in the z -direction; this is an actual property, whereas the spin in the x - and y -directions are potential properties. It has potentially both spin up and spin down in both the x - and y -directions. A spin measurement in the y -direction transforms one of the potential spin values to an actual one, namely that corresponding to the measured result. At the same time the actual property of having definite spin in the z -direction is destroyed.

The discussion above rests on the assumption that the probability for actualising a potential property is higher than zero; it goes without saying that an object cannot reasonably be attributed a potential property if this property as a matter of principle can never be actualised.

The distinction between actual and potential properties is closely related to von Neumann's (1932) distinction between measurement of the first and of the second kind. Potential properties are actualised through measurements of the first kind, whereas measurements of the second kind concern actual properties; they merely reveal what we already know, if we know the wave function.

The mathematical entities corresponding to both actual and potential properties are self-adjoint operators. This is, in fact, a further good reason to introduce the concept of potential property, for otherwise we would be forced to say that those operators correspond to real properties in some situations and in some others not,

thus making a coherent interpretation of operators impossible. Some physicists, for example, Asher Peres (1998), claim that we should give up talking about properties as values of observables altogether and replace this concept by the mathematical well-defined concept of a self-adjoint operator. In my view this is to reject the need for a physical interpretation: operators are mathematical objects, not physical objects, events or actions. The crucial question is what these mathematical objects correspond to in nature and my answer is simple: properties, actual or potential, as the case may be.

3.14 Probability in Quantum Mechanics

The interpretation of the concept of probability is a controversial topic in the philosophy of science. The main dispute is whether probability is to be understood as subjective expectation, given a specified amount of information, or whether it is an objective and observer-independent attribute. And if the latter is the case, attribute of what: to singular events, singular objects, sequences of events, sequences of experimental situations, or something else? A thorough discussion of these general questions is beyond the scope of this book, but it is necessary to say something about the use of probability concepts in quantum mechanics.

The remarkable thing with quantum probabilities is that they can be defined independently of both subjective expectations, relative frequencies and assumptions about indifference. Once the complete wave function for a system is known, we can calculate the probability distribution for every observable by applying its corresponding operator to the wave function Ψ . That procedure returns the probability distribution without any use of measured frequencies or any a priori assumptions about equal probabilities. If, moreover, the (maximal) wave function really is a complete description of the system, and nothing whatsoever testifies so far against that assumption, we can conclude that the probability distribution for this observable is an objective and irreducible property of each single object described by that wave function. It is not a measure of our ignorance. Observe the conditional: *if* the wave function is a complete description, *then* probability attributions are not measures of our ignorance. Hidden variable theorists assume that even a maximal wave function is an incomplete description of the system, but I cannot see that they have any good arguments for their view.

Several authors, for example Popper, have taken the same stance regarding probabilities and, in order to distinguish the ontological character of quantum probabilities from the epistemological character of the probability concept in many other applications, he introduced the concept of propensity as denoting the objective and irreducible statistical features of quantum systems. Some have criticised his introduction of this concept as lacking explanatory value. This criticism can be rebutted by the simple argument that the concept of propensity is used for distinguishing between two ways of using probabilities. When we say that the probability of an event is p , this might either be an expression of incomplete information, or it might express an irreducible randomness of the event. In the latter case the probability of the event is not conditional on our amount of information and

it is precisely this property, unconditional irreducible probability, which is called propensity by Popper.

In this discussion, I have taken for granted that the state of the system in question is a pure state and not a mixture. Like most participants in the discussion I interpret mixture states as incomplete descriptions giving probability distributions over ensembles of systems governed by different wave functions.

3.15 Identity of states: real change of state versus change of description

Physical states are described by wave functions, or equivalently, by state vectors in Hilbert space. If we believe that quantum mechanics is a complete theory, i.e. believing there are no hidden variables determining what is undetermined in orthodox theory, any change of the state of a system must result in the change of the wave function (or state vector). The converse need not be true, however; different state vectors may describe one and the same state.

In order to decide when two wave functions describe one and the same physical state, we need a *principle of identity of states*. The following proposal seems reasonable.

Principle for identity of states: Two state descriptions ' Ψ_α ' and ' Ψ_β ' describe the same physical state if and only if the probability distributions of the possible values of any observable O are the same when applied to Ψ_α and Ψ_β .

Some, perhaps most, physicists do not view things in this way, but claim that there is a class of state changes that do not result in changes of probability distributions. I will here rehearse and comment upon Stephen Weinberg's (1995) view.

First, he states that two state descriptions, i.e. normalised state vectors, which differ only by a complex number belong to the same *ray*, and such a ray, i.e. set of state vectors, represent a physical state. Then he proceeds by considering the situation where two observers look upon the same system. He writes: 'if an observer O sees a system in a state represented by a ray R or R_1 or R_2 ... then an equivalent observer O' who looks at the same system will observe it in a *different state* [my italics] represented by a ray R' or R'_1 or R'_2 ... respectively, but the two observers must find the same probabilities' (Weinberg, 1995, p. 50).

I take it that Weinberg here is considering the fact that different observers may use different coordinate systems when observing a system and determining its properties. Since the system might change its state with time, it is obvious that Weinberg implicitly presupposes that the two observers look at the same system at the same time, otherwise he couldn't say that they *must* find the same probabilities; the system might change its state over time. Hence it is quite clear that Weinberg takes states of a system to be relative to an observer, or more precisely, relative to a coordinate system. This is all right as long as no one mistakes this notion for one that is independent of the observer.

If we instead insist on the notion that states are objective and not relative to an observer we must reject the notion that two observers might attribute different states to the same system at the same time. The intuition is that the expression 'the same state' should be interpreted as 'the same properties', i.e. the same values of the

quantitative variables. However, this condition must be further elaborated since we are discussing quantities expressed in different coordinate systems. The following formulation seems appropriate:

Two state descriptions Ψ_a and Ψ_b describe one and the same state if (1) there is a coordinate transformation T such that $\Psi_a = T\Psi_b$, and (2) for all operators O_a and O_b defined on the relevant Hilbert spaces, $O_a = TO_bT^\dagger$.

This latter condition is clearly a sufficient criterion for obtaining the same probability distributions in the two coordinate systems.

The question about identity of state is intimately connected with the interpretation of unitary/anti-unitary evolution of quantum states. These transformations are by all physicists regarded as state changes. This is not my view and I will here give my argument, which essentially is built upon a theorem proved by Wigner (1931).³

Wigner's theorem

If the probability distributions of observables on a system are invariant under a transformation, we can define an operator U representing that transformation on the relevant Hilbert space such that if Ψ is in a ray $\mathfrak{R}$, then $U\Psi$ is in the ray $\mathfrak{R}'$ and U fulfils either

$$\langle U\Psi, U\Phi \rangle = \langle \Psi, \Phi \rangle \quad (3.3)$$

$$U(\xi\Phi + \eta\Psi) = \xi U\Phi + \eta U\Psi \quad (3.4)$$

or else

$$\langle U\Psi, U\Phi \rangle = \langle \Psi, \Phi \rangle^* \quad (3.5)$$

$$U(\xi\Phi + \eta\Psi) = \xi^* U\Phi + \eta^* U\Psi \quad (3.6)$$

Equations (3.3) and (3.5) are the unitary and anti-unitary conditions respectively. The well-known unitary condition, $U^\dagger = U^{-1}$, follows from equation (3.3). Because the adjoint operator $A^\dagger$ to an anti-linear operator A is defined (Weinberg 1995, p. 51) by

$$\langle \Phi, A^\dagger \Psi \rangle \equiv \langle A\Phi, \Psi \rangle^* = \langle \Psi, A\Phi \rangle \quad (3.7)$$

it follows that the conditions both for unitary and anti-unitary transformations take the form $U^\dagger = U^{-1}$.

Together with the principle stated above we have that any transformations between frames of reference are represented by either unitary or anti-unitary operators. Unitary transformations are all continuous transformations as can be seen from the following argument.

3 For a modern proof see Weinberg (1995, pp. 91–96).

The identity transformation $\mathfrak{R} \rightarrow \mathfrak{R}$ represented by the operator $U=1$ is, of course, unitary and linear. It follows that any transformation that can be changed into the identity transformation by a continuous change of some parameter must also be unitary and linear rather than anti-unitary and anti-linear. Hence, rotations, translations in space, translations in time and boosts, which all have the identity transformation as the limit when the relevant parameter approaches zero, are all unitary and linear.

There are two other transformations that properly must be seen as changes of frames of reference. One is space inversion, i.e. change of relative orientation of the coordinate axes x, y, z , the other is the change of time direction. This last change has nothing to do with the real time ordering of events; it is rather the result of a choice between representing the temporal relation ‘...after ...’ by the mathematical relation ‘...’ or by its converse ‘...’, when events are attributed time coordinates. In practice, of course, we always choose the first alternative of representing later times by bigger numbers. Since this is a mere convention, time reversal is a mere change of convention, of choice of parameterisation, which of course cannot change any probability distributions.

Space inversion, represented by the parity operator in quantum mechanics, is different. Although parity operations do not in general change probability distributions, there are exceptions, namely in weak interactions. Hence, space inversion cannot be seen as a mere change of convention of description; nature has a right–left-hand asymmetry. This appears astonishing, but there is an explanation. If we view the world as being five-dimensional, the fourth dimension being time and the fifth charge, the asymmetry disappears. Then the true symmetry is space inversion + time reversal + charge conjugation, i.e. CPT symmetry and such a complete set of transformations is a mere change of conventions.

It can be shown that space inversion must be unitary and linear, whereas time reversal is represented by an anti-unitary and anti-linear operator (Weinberg, 1995, p. 76). An important question is whether the converse to Wigner’s theorem is true, i.e. does it hold that for every unitary or anti-unitary transformation the probability distribution of all observables are left unchanged? That this is the case is rather trivial. Consider a wave function Ψ and an operator $\mathbf{O}$ having the set of eigenfunctions $\{\phi_i\}$ $i=1,2,\dots,n$. The eigenvalue corresponding to ϕ_i is c_i and the probability for the value c_i is $|\langle \phi_i, \Psi \rangle|^2$. Then let us apply a unitary transformation U on Ψ . The result is

$$U\Psi = U\sum c_i\phi_i = \sum c_i U\phi_i \quad (3.8)$$

because unitary operators are linear. In order to have the same probability as before we must require

$$|\langle U\phi_i, U\Psi \rangle| = |\langle \phi_i, \Psi \rangle| \text{ for all } i \quad (3.9)$$

But this condition is easily seen to be true since it is a straightforward consequence of the unitary condition (3.3). A similar argument can be given for anti-unitary operators. Hence every continuous unitary or anti-unitary transformation preserves the probability distributions (except in weak interactions) and conversely every transformation that preserves probabilities is either unitary or anti-unitary.

We now have completed the argument; all transformations that belong to the class of continuous changes of frames of reference, namely space translations, rotations, and boosts, time translations + time reversals, are represented mathematically by operators that fulfil the condition $U^\dagger = U^{-1}$. The converse is not exactly true because of parity violation in weak interactions. With this exception we have that all transformations for which $U^\dagger = U^{-1}$ holds are not real changes of the *state* of the system but only changes of its *description*.

What has been said is, I hope, easy to accept, except perhaps the non-reality of time translations. But just as translations along the spatial axes can be interpreted as mere change of position of the observer in space, a mere change of the observer in time, or what amounts to the same, a mere change of the zero point of the time axis, should be irrelevant for the physical state. The time translation operator is $\exp(i\mathbf{H}t/\hbar)$, hence the transformation

$$\Psi \rightarrow \exp(-i\mathbf{H}t/\hbar) \Psi \quad (3.10)$$

which is a unitary transformation, represents a mere change of position along the time axis. This translation does not change the probability distributions for observables, provided the Hamiltonian is a constant of the motion. Could it reasonably be said that this is no change of the physical state? The answer is yes, and the analogy with uniform motion in classical mechanics is obvious. A classical point particle moving uniformly is in one and the same *state of motion* or, equivalently, *the same dynamical state*, as long as no external forces act on it, even if it changes position relative to some chosen external set of objects. There is no reason not to say the same in the quantum case.

Many would perhaps voice a protest against this last statement, arguing that the state of a classical particle is given by its position in configuration space and as a particle in uniform motion constantly changes its position in real space, it must also change its position in the configuration space, thus changing its state.

The conflict here is really a conflict about the proper use of the term ‘the same state’. Both in classical and quantum physics the notion of the *state of a system* is used both in kinematics and dynamics. The kinematical problem is essentially the problem of keeping track of different objects and, in order to do that in the classical case, we attribute trajectories in configuration space. The dynamical problem is the proper description of *interactions* between physical objects. A physical object that has not interacted with other objects is still in the same dynamical state, i.e. it has the same velocity. Newton formulated this in his first law: a body remains in its *state of motion* if no forces act on it.

Going to the quantum case we observe one important difference, namely that the kinematical state cannot be given as a point in configuration space, but as a point in Hilbert space. Thus, in the quantum case, uniform motion results in changes of the state vector in Hilbert space instead of changes in configuration space.

The most important difference between the classical and the quantum case is the effect of interactions. In the classical case, an interaction is a continuous change of velocity of the interacting particles and the dynamical evolution can be accounted for using continuous and differentiable functions. In the quantum case by contrast,

interactions are discretised and hence no such functions can be used. (This will be further discussed in Chapter 5, especially Sections 5.2–5.3.) Instead we use operators with discrete spectra and the result of such an interaction is a discontinuous change of dynamical state. Usually this is held to be true only of measurement interactions, but in Chapters 5 and 6 I will give arguments for the view that that is a mistake; interactions which change the dynamical states of the interacting objects occur frequently.

Summarising, a *change of dynamical state* of a quantum system occurs if and only if an interaction with something else, ‘the environment’, has taken place.

The notion of interaction with the environment is, however, not as clear as could be wished. Consider, for example, an electron going through an insulator in which the energy spacing between the valence band and the conduction band is larger than the kinetic energy given to the particle. Then the particle will polarise the insulator as it passes and when it leaves the insulator the polarisation disappears. But if the energy gap is less than the available kinetic energy (which means that it is not a perfect insulator), there is a non-negligible chance that the insulator takes up a portion of energy from the passing particle. According to my suggestion, then, the first situation is not to be considered as an interaction with the environment (i.e. the insulator) and no state change occurs, whereas in the second case, when energy is exchanged, we have an interaction and thus a state change. The important distinction is that in the second case the environment has undergone an irreversible change. Observations of its state before and after the passage of the particle will reveal this.

If this is accepted we must say that the time dependent Schrödinger equation does not describe how quantum systems change a dynamical state over time. Rather, Schrödinger’s time-dependent equation gives a rule telling us how to perform transformations between observations at different times of one and the same system in one and the same dynamical state. This is exactly the same as performing a transformation of the state description between two coordinate systems using clocks started at different times. We could now introduce the notion of apparent change: an apparent change is a transformation of the description only. These transformations are those represented by unitary or anti-unitary operators. In contradistinction, real dynamical changes are represented by projection operators, because a projection operator changes the probability distribution of one or several observables.

Changes of dynamical state are represented by projection operators and are hence in-deterministic. One is then tempted to ask if the converse is true, i.e. whether all in-deterministic state changes are changes of the dynamical state. The simple answer is yes but I have one proviso: since indeterminism enters quantum theory with the use of projection operators, and these are used to represent measurements of the first kind, it seems obvious that indeterminism is the effect of measurements only. This is not my view. As will be argued in Chapters 5 and 6, a measurement is only a sufficient criterion for a collapse of the wave function, not a necessary one. And since projection operators represent collapses, they could represent a broader category of events than measurements of the first kind. In short, I think we should decouple the strict one-to-one correlation between measurement of the first kind and indeterministic change. We can summarize the argument:

1. All indeterministic state changes are represented by projection operators on the Hilbert space.
2. All projection operators represent change of probability distributions of the associated observables.
3. All changes of probability distributions are real changes of dynamical state.

Hence we have the following result:

A change of state of a quantum system is a real dynamical change if and only if it is indeterministic.

3.16 Schrödinger evolution with changing Hamiltonian

The existence of continuous and deterministic evolution of a system governed by a non-constant Hamiltonian (a Hamiltonian for which the time derivative of the expectation value is non-zero) provides a *prima facie* objection to my conclusion that all real changes are indeterministic. For if the Hamiltonian is a deterministic function of time, the probability distributions for observables compatible with the total energy for the system will change accordingly.

A non-constant Hamiltonian represents the physical state of affairs that energy is exchanged between the system and its surroundings. This exchange of energy is, in quantum mechanics, an exchange of electromagnetic energy, i.e. photons. (It could also be other interactions and other exchange particles, such as vector bosons, but that does not matter in this context.) If this process is ongoing for some time it means that an electromagnetic field is constantly present where the system is situated and thus, if the energy exchange is an ongoing process, the phase correlation between the system and the electromagnetic field is not broken. In other words, the quantum system is not a separate object that can be attributed properties of its own independent of the rest of the environment. Rather, it is a part of a bigger system, which includes the sources of the field. Hence, if we stick to our principle of individuation for systems, there is simply no such thing as an individual system with a changing Hamiltonian. If the Hamiltonian changes, the system must be connected to its surroundings and thus it is not a well-defined individual system, it is only part of a bigger system.

But when is this connection broken, when does the continuous exchange of energy between the coupled systems come to an end? Answer: when a collapse, an irreversible state change, occurs. But what brings about this collapse? I have no answer. We know that it does occur, and roughly we know how to bring about a collapse. It has to do with decoherence, the disappearance of interference terms. Under what conditions will the interference terms become zero? Suppose the interference terms have diminished to the order of magnitude, say, 10^{-6} ; is that sufficient for making it impossible to arrange the situation so that they start to increase again? I don't know. The question is intimately connected with the explanation of quantisation, since quantisation entails collapse, as will be proved in Section 6.7. Present quantum mechanics has nothing more to say, since quantisation is a basic postulate. Since this book is an interpretation of quantum mechanics, not a proposal for an improved theory, I must leave this problem unsolved.

Summarising the core idea: if we are strict as regards what to count as an individual object we should say that objects, namely quantum systems, should be individuated by intrinsic properties, i.e. properties that can be attributed without reference to other objects. Then taking into account the core idea of quantum theory, namely that all interactions are quantised, we arrive at the conclusion that, although it sometimes can be convenient to treat a system as not isolated and thus governed by a non-constant Hamiltonian, such a system is not a well-defined individual object.

Chapter 4

Quantum Objects are Waves

4.1 Introduction

Many, if not all our problems of understanding quantum mechanics seem to have one common source, namely a mismatch between the implicit ontology of quantum theory and the ontology we take for granted when we talk about micro-phenomena. The foundation of quantum theory is wave mechanics. But when we talk about the quantum world we usually take a particle ontology for granted: we talk about electrons, neutrons, photons etc as if they were corpuscles, objects with well-defined positions and confined to certain volumes of space. The fact that, for example, electrons cannot in general be attributed a definite position is thus often interpreted epistemologically: we cannot *know* an electron's exact position, although we take for granted that it *has* a well-defined one. As waves and corpuscles have different properties, we meet difficulties when we use the superposition principle (a principle valid for waves) for the description of the motion of quantum objects conceived as corpuscles. By its very nature, a corpuscle cannot be subject to the superposition principle.

However, many readers may wonder why we bother so much about wave-particle duality. Is not this issue once and for all settled by Bohr's complementarity principle? Why not proceed to more recent issues such as the measurement problem and non-local correlations? My response is that these very problems are, in fact, intimately connected to wave-particle duality. Our way of looking at the measurement problem depends heavily on our ontological stance. This is quite obvious when adopting the pilot wave interpretation, a stance that at the same time explains both the collapse of the wave function and the wave-particle duality. Alternatively, if we adopt Schrödinger's view that quantum objects are real waves (of a kind yet to be described), the collapse must be a collapse of a real wave and that proposal immediately triggers the question concerning the cause of the collapse. Or, if we adopt a particle ontology by saying that quantum objects really are kinds of particles, the most natural stance towards the collapse of the wave function is to view it as an ordinary statistical collapse of a probability distribution when the scientist obtains new knowledge. According to such a view, the only difference between quantum mechanics and classical mechanics is that probabilities evolve according to different laws in the two contexts. In short, our choice of ontology decides our general view on the measurement problem.

The problem of understanding non-local correlation is also connected to our opinion regarding what kinds of things exist. There are two problematic aspects of non-locality, namely (1) that distant correlations may occur instantaneously, and (2) that these distant correlations are produced without any mediating force. Both aspects

will conflict with deep metaphysical convictions if we adopt a particle ontology. However, if the real objects are some sort of waves, the metaphysical conundrum is less disturbing. Waves might be more or less overlapping although their centres are widely separated; if that is the case, non-local correlation is less incomprehensible, because then they are not completely separated from each other and can interact without the mediation of anything else. Similarly, if we accept the pilot wave solution we will have a sort of solution to non-local interactions, for, as is well known in this interpretation, a quantum potential that guides the particles is assumed and this can change instantaneously all over its entire space, thus bringing about non-local correlations between the correlated particles. This pilot wave theory has however never aroused my sympathy, in spite of its being a realistic interpretation. The basic problem is that it assumes both waves and particles and this dualistic ontology leads to difficulties, as I will argue in Chapter 9.

I believe that Schrödinger was on the right track when he assumed that those objects we ordinarily regard as particles, namely electrons, neutrons etc, are some sort of waves. In this chapter I will develop this idea and try to counter all hitherto given arguments against viewing quantum objects as waves. In the next chapter particle behaviour will be accounted for in terms of these waves.

4.2 Interpretation of the wave function

4.2.1 Schrödinger's idea

It is well known that Schrödinger, contrary to Born and Heisenberg, struggled with a realistic interpretation of the wave function. More specifically, for the electron he proposed that the quantity $e\Psi^*\Psi$ should be interpreted as a charge density, i.e. the electron is not a particle at all but a charge density travelling in space according to the wave equation. In a letter to Lorentz he expressed this idea in the following way:

Now I want to choose more simply $\Psi\Psi^*$, that is, the square of the absolute value of the quantity Ψ (as representing the real stuff of the world). I now have to deal with N particles, then $\Psi\Psi^*$ (just as Ψ itself) is a function of $3N$ variables, or as I want to say, of N three dimensional spaces, $R_1, R_2, \dots, R_N$. Now first let R_1 be identified with the real space and integrate $\Psi\Psi^*$ over $R_2 \dots R_N$; second, identify R_2 with the real space and integrate over R_1, R_3, R_N and so on. The N individual results are to be added after they have been multiplied by certain constants which characterize the particles (their charges according to the former theory). I consider the result to be the charge density in real space.¹

This interpretation immediately explains all interference phenomena satisfactorily. However, several problems were encountered by Schrödinger, some of which he recognised from the very beginning, some of which he met as objections from other physicists and failed to rebut in a convincing way. Did Schrödinger give up his realistic attitude towards the wave function when confronted with these difficulties? According to Victor Weisskopf who knew Schrödinger well, he did not; Weisskopf

1 Letter to Lorentz, June 6, 1926. Reprinted in Przibram (1967, pp. 55–66).

states that ‘Schrödinger persisted in his belief in the real wave function until his last days.’² Michel Bitbol has studied Schrödinger’s published and unpublished papers and draws a more complex picture of Schrödinger’s intellectual evolution. He identifies three stages in Schrödinger’s thinking about quantum mechanics, first a period of strong belief in a real wave function, then a sceptical attitude during at least the 1930s and then a third, more constructive period during the 1950s (Schrödinger, 1995; Bitbol, 1996). Since this book is not an exposition of Schrödinger’s evolving thoughts, I won’t discuss the matter.

I think Schrödinger was fundamentally right in believing that matter fundamentally consists of some sort of waves and that the objections he encountered can be rebutted. So let us start by listing the difficulties encountered by Schrödinger and anyone else proposing a wave interpretation. The following seven problems make up a complete list, I think:

- (1) In quite a number of situations, electrons and other quantum objects show particle behaviour. How is that possible if every single one of these objects is a spatial charge distribution (or distribution of some other quantity) moving according to a wave equation? For example, does not the collision between an electron and a photon in the Compton effect prove that electrons and photons are particles? (The Compton effect is called by Jammer, 1974, p. 29, ‘the paradigm of particle physics’.)
- (2) An electron, conceived as a charge distribution in the form of a wave packet, will inevitably be dispersed during motion. Even if the medium is truly non-dispersive (which hardly ever is to be expected) the wave packet will disperse in the transverse as well as longitudinal direction. How can this possibly be reconciled with the individuality of the electron, manifested in interactions?
- (3) If the process of emission of electromagnetic energy from an atom is a resonance phenomenon (an assumption most physicists in the 1920s found reasonable), we should see harmonic multiplets of the fundamental frequencies in the atomic spectra, which we don’t (Heisenberg’s argument, cf. Jammer, 1974, p. 32).
- (4) When solving scattering problems, the incident particle is often described as a wave packet scattered against a heavier particle, which is described by a potential, the source of which is a point particle. But assuming a point particle as being the source of the scattering potential while attributing wave properties to the incoming electron is inconsistent.
- (5) The wave function is usually a complex function. How is it possible to view this function as representing something real?
- (6) Schrödinger had a purely classical conception of a wave as being an infinite train of troughs and peaks. Such a concept cannot be used for an interpretation of quantised electromagnetic waves: if each quantum with frequency f has an energy hf , it cannot be an infinitely extended wave because, as Heisenberg pointed out, the energy in a classical infinite wave is infinite. Schrödinger is reported to have exclaimed: ‘If one has to go on with these damned quantum

2 Letter to Robert Chen from Viktor Weisskopf, quoted in R. Chen (1990, p. 183).

jumps, then I'm sorry that I ever started to work on atomic theory' (Rosenthal, 1967, p. 103).

Furthermore, there is one very important difficulty which, it seems, was not fully recognised by Schrödinger:

- (7) Not all wave functions for multi-particle states are factorable into one-particle-components. In these cases the $3N$ -dimensional configuration space cannot be mapped into three-dimensional physical space containing N individual particles. Hence, the procedure Schrödinger proposed in the letter to Lorentz quoted above will not always work.

4.2.2 Rebuttals of these arguments against the wave interpretation

Of these seven difficulties some are, from a modern point of view, less troublesome, namely (3), (4) and (6).

The third difficulty is no longer considered a real problem because it arises only if one conceives of photon emission from atoms as a classical resonance phenomenon, thus assuming that the frequency of the emitted photon is an integer multiple of the frequency of the emitter. However, this assumption is not warranted and it is in fact patently false. But that is no argument against viewing material objects as a kind of wave because something could be a wave phenomenon even if it does not comply with all classical laws of wave mechanics. Exchange of portions of energy follows its own laws.

A similar comment can be made about the problem discussed in point (6) above. A classical infinite wave cannot have a finite energy, but we need not restrict ourselves to this purely classical concept. A slightly more generalised wave concept, to be discussed in the following section, will fit.

The problem in point (4) is different because it is indeed inconsistent to assume that waves are scattered from point particles if waves are all there are. However, it should be observed that the assumption about a point scatterer could be seen as a model for calculating scattering probabilities only, and it is not implied that such a target in reality is a particle, it could just as well be a suitable approximation. In a discussion of fundamentals the crucial question is not whether the assumption of a point scatterer is a good approximation, but whether similar or better predictions are possible under the assumption of a wave scatterer. In fact, Schrödinger himself showed that scattering problems can be solved by assuming that the interacting objects are two plane waves. His argument will be rehearsed in the next chapter.

What remains to be discussed are points (1), (2), (5) and (7) above.

(1) In order to analyse the first problem (how waves can show particle behaviour), let us consider how electrons are registered in a film emulsion. We know that a grain of the silver bromide in the film emulsion will change state upon acquiring kinetic energy from the electron. Now, suppose Schrödinger is right, electrons are waves, which do not hit only one single grain but many. That means that the contact area on impact will cover a huge number of silver bromide grains. Why, then, does only one grain take up energy? My answer is: because energy exchange is quantised, an

argument which is in line with the interpretation of measurement itself as collapse of the wave function. I will here give only a brief outline of the complete argument to be elaborated on in Chapter 6.

The process described as ‘the electron impinging on the film with silver bromide grains’ is a position measurement. (Provided we use the word ‘measurement’ as referring to a purely physical process, the result need not necessarily be observed.) Before the interaction the wave function is

$$\Psi_{\text{before}} = \sum a_i \phi_i \quad (4.1)$$

where $\phi_i = \delta(x - x_i)$, x_i being the position of the i th silver bromide grain. Prior to the measurement the state of the electron is thus a superposition of a great number of different positions. (This description is highly idealised. It is implicitly assumed both that the electron cannot pass beside the screen, which it can, and that the screen is a perfect detector reacting on all impinging electrons, which it is not. Furthermore, silver bromide grains are not points.)

When the interaction has taken place the wave function is

$$\Psi_{\text{after}} = \phi_k \quad (4.2)$$

where k denotes the position of the silver bromide grain that actually changed its state during the interaction. The collision between the electron and the film is thus a collapse of the wave function. This collapse is postulated by von Neumann to occur whenever we perform a measurement of the first kind (i.e. when the wave function is not an eigenfunction to the applied operator). If we adopt the view that a measurement must be understood as a purely physical interaction, the ensuing question is ‘what is the (physical) cause of the collapse of the wave function during measurement?’ and my answer is ‘exchange of energy’. The argument goes as follows. The fundamental feature of micro physics is that energy exchange is quantised. In the mathematical formalism this is represented by the replacing of continuous functions by operators with discrete spectra. The physical fact represented is that exchange of energy (and other conserved quantities) occur in discrete portions that cannot be further subdivided. Hence, an object cannot interact with more than one other object at each moment of time. It follows that the wave function must collapse to one particular place, namely the position of the object taking up energy. (The argument will be elaborated in the next chapter.) But then, the wave behaves as a particle at the moment of interaction.

The particle behaviour of waves in interactions was a great obstacle for Schrödinger, because he did not view energy exchange in this way. As is well known, he strongly opposed the notion of ‘quantum jumps’, i.e. discontinuous state changes. His classical conception of waves entailed that energy exchange is a resonance phenomenon which is analysed as a continuous exchange of energy, both in interactions between material objects and between matter and radiation. From Schrödinger’s point of view, the questions of measurement and of particle behaviour respectively, were two different problems, neither of which he succeeded in solving. This was a major reason why Schrödinger’s views were rejected by other physicists.

(2) The second problem, the dispersion of wave packets, can be handled in two completely different ways depending on whether we accept the premises of the argument or not. The point was raised by Lorentz in a letter to Schrödinger.³ If Lorentz's assumption that wave packets inevitably undergo dispersion during motion is accepted, we face the problem of explaining how a widely dispersed object can interact at a rather narrow location. In other words, how can a dispersed wave packet collapse into a small volume? But again that is the core of the measurement problem, namely to explain the collapse, touched upon above and discussed at length in Chapter 6.

We could alternatively reject the premise that wave packets must be dispersed during motion. By adding a non-linear term to the Schrödinger equation it can be shown that solutions to the relevant wave equation are so called solitons, i.e. localised wave packets in the form of hyperbolic secants, and such objects do not undergo dispersion. The crucial question is of course whether the addition of non-linear terms represent real states of affairs or not. That is fundamentally a physical question, not a philosophical one. My stance is that as a philosopher trying to interpret a physical theory, one should take the physical theory, i.e. the formalism, for granted. Quantum mechanics, as it is presented in all textbooks, takes for granted that the Schrödinger equation without non-linear terms is the correct equation of motion (neglecting relativistic considerations). Consequently, I think that non-linear corrections to the Schrödinger equation should be rejected in this context. This does not mean that I reject the possibility that such corrections might be legitimate. It is a methodological choice to take the present form of quantum mechanics as a point of departure.

(5) The fifth problem centres on the question of why we must use a complex function to describe a real object. Schrödinger assumed that nothing prevents us from using a complex function to describe the motion of real waves. He thought it must somehow be possible to express the imaginary part of the wave function as a function of the real part, thus assuming that the use of a complex wave function is only a matter of convenience.

But is it only a matter of convenience? Schrödinger struggled with this problem, as can be seen from his letter to Lorentz of June, 6, 1926:

What is unpleasant here, and indeed directly to be objected to, is the use of complex numbers. Ψ is surely fundamentally a real function and therefore in Eq. (35) of my third paper

$$\Psi = \sum c_k u_k(x) \exp\left(\frac{2\pi i E_k t}{h}\right) \quad (4.3)$$

I should be good and write a cosine instead of the exponential and ask myself: is it possible in addition to define the imaginary part unambiguously *without reference to the whole behaviour of the quantity in time*, but rather referring only to the real quantity itself and its time and space derivatives at the point in question. This actually can be done, at least for Ψ . I write the 'wave equation' for the sake of brevity in the form

$$L[u] + Eu = 0 \quad (4.4)$$

3 Letter to Schrödinger of May 27, 1926, reprinted in Przibram (1967, pp. 43–54).

where $L[]$ means a certain differential operator. Furthermore let Ψ_r be the *real* wave function, the only one known originally, therefore the real part of Ψ , whose imaginary part is to be defined in addition. That can be done this way:

$$\dot{\Psi} = \dot{\Psi}_r - \frac{2\pi i}{h} L[\Psi_r] \quad (4.5)$$

By this method therefore, the magnitude $\dot{\Psi}$ is in any case represented by the space and time derivatives of the real quantity Ψ_r , independent of the complex representation; so that one does not get into difficulty even if there is a Ψ_r that does not correspond to a stationary superposition of proper oscillations. Now to be sure one has only $\dot{\Psi}$ and integration with respect to time would involve an underdetermined additive purely imaginary function of the coordinates. I do not know whether this can be fixed in a rational way.⁴

Schrödinger was not able to solve this problem. It has, however, been studied by Robert Chen (1990, 1993) who maintains, correctly I think, that there is a natural solution to the problem, as well as an answer to the question of why Schrödinger was unable to see that. According to Chen, Schrödinger failed because when he used the vibrating plate as an analogy (as implied in the letter to Lorentz) he did that incorrectly; Schrödinger assumed that the lateral displacement of a vibrating plate is the analogy to the real wave function and in so doing there is nothing in the vibrating plate corresponding to the imaginary part of the wave function. But if he had taken the velocity field as the sought analogy he could have solved his problem. Here is Chen's argument.

Assume that $\eta(x,t)$ is the lateral displacement of a vibrating plate. (x must here be taken to be a two-dimensional vector.) The local deformation $u(x,t)$, i.e. the bending of the plate at x at time t , is then given by

$$u(x,t) = a^2 \nabla^2 \eta(x,t) \quad (4.6)$$

where a is a material constant. The lateral velocity v is then given by

$$v = b^2 \partial \eta / \partial t \quad (4.7)$$

With suitably chosen constants a and b , u^2 and v^2 give respectively the potential and kinetic energies of the plate. η obeys the wave equation

$$\partial^2 \eta / \partial t^2 = \text{const.} \cdot \nabla^2 \eta \quad (4.8)$$

The function $\Psi_r = A \cos(kx - \omega t)$, which is the real part of $\Psi = A \exp[i(kx - \omega t)]$ – the wave function for a plane travelling wave – is a solution to this equation. Thus, there is a formal analogy between the lateral displacement and a travelling wave, which misled Schrödinger. Chen explains:

In his [Schrödinger's] version of the analogy, he identified Ψ_r with the lateral displacement [i.e. η] as its counterpart (not with the deformation u). It so happens that u obeys the same equation as η ...

4 Letter to Lorentz June 6, 1926, reprinted in Przibram (1967, pp. 56–57).

$$\partial^2 u / \partial t^2 = -L^2 u \quad (4.9)$$

[‘ L ’ is Schrödinger’s expression for ∇]

One might therefore put Ψ_r equal to either u or η . With the latter, however, one could not find a physically meaningful quantity (in the vibration of the plate) that relates to η as Ψ_i does to Ψ_r – nothing like the velocity field mentioned above. Thus, in his mind, Schrödinger could see only a limited analogy – rather than the complete analogy shown in table 1 – which was not helpful at all to his task. (Chen, 1993, p. 257)

The table referred to in the quotation above is Table 4.1 here.

Table 4.1 Comparison between the motions of a vibrating plate and quantum mechanics of a free particle

Quantum mechanics of a free particle	Classical motion of a vibrating plate
Ψ_r and Ψ_i	Deformation and velocity field
Ψ_r^2 and Ψ_i^2	Potential and kinetic energy densities
Probability density $\Psi^* \Psi$	Total energy density

Using this more complete analogy, it is easy to solve Schrödinger’s problem, namely to fix the integration constant when deriving the imaginary part of the complex wave function from the real part. As the deformation and velocity fields for a vibrating plate will differ by a constant phase of $\pi/2$, the deformation is zero when the velocity reaches its maximum and vice versa. Thus, the potential energy is zero when the kinetic energy reaches its maximum and vice versa.

But the main question remains: why do we need complex phases to describe the evolution of these waves? No realistic interpretation can be considered satisfying unless a reasonable answer to this question has been given.

Generally, complex functions can be mapped one-to-one to pairs of real functions which are coupled to each other by a fixed phase $\pi/2$. It is the coupling that is important; without it there would be no mapping between complex functions and pairs of real functions. Hence, the use of complex phases shows that there are two components in the physical wave, that these components are coupled and the description of the coupling between these components requires imaginary numbers, as can be seen from equation (4.5). The two components in the real wave are those governing the evolution of the potential and the kinetic energy density, respectively. But why are imaginary numbers required to describe mathematically the coupling between two real wave phenomena? The answer has to do with the metric of spacetime itself. I will return to this question in Section 4.10.

Let me stress that the issue is not whether it is mathematically possible to reduce the complex wave function to a pair of real functions in a uniform way. The issue is

whether we can find a *physical reason* for using two coupled functions for describing real events, i.e. whether both the real and imaginary parts can be interpreted as representing something real. Chen has, in his paper on Schrödinger's views on these matters, shown that such an interpretation – as presented in the table above – is possible. The discussion about Schrödinger's views is merely of historical interest.

(7) The wave function for an N -particle state is a function in an abstract $3N$ -dimensional Hilbert space, which in general cannot be mapped onto N 3-dimensional spaces containing one-particle wave functions. Therefore a realistic interpretation of the Ψ -function meets a difficulty: where is this $3N$ -dimensional real wave?

The same problem occurs in the case of spin coupling in a singlet state. Assume that the prepared spin state is orthogonal to the direction in which the states $|\text{up}\rangle$ and $|\text{down}\rangle$ are defined. Thus each electron in the spin pair could be represented by

$$\Psi = \frac{1}{\sqrt{2}}(|\text{up}\rangle - |\text{down}\rangle) \quad (4.10)$$

The wave function for a two-electron state prior to contact is simply the sum of two such functions. Such a function causes no interpretative problem as long as the two particles are independent of each other, because such a sum *can* be separated into two independent wave functions, one for each particle; all terms in the sum contain argument places taking values referring to only one of the particles. However, when they interact the result could be a singlet state, i.e. a state with zero total spin. The wave function for such a state is given by

$$\Psi = \frac{1}{\sqrt{2}}(|\text{up}\rangle \otimes |\text{down}\rangle - |\text{down}\rangle \otimes |\text{up}\rangle) \quad (4.11)$$

where the first factor of the tensor products refers to particle 1 and the second factor to particle 2. (We should remember that equation (4.11) is not the complete state description; all other degrees of freedom but spin are suppressed.) In this state both terms contain information referring to both particles and it cannot be written as a sum of two one-particle states. The interaction between the particles introduces a correlation between them: whenever particle 1 has spin up, particle 2 has spin down and vice versa and this correlation obtains in any randomly chosen direction.

This is but another way of expressing what was argued in Chapter 3, namely that in multi-particle quantum systems the component parts in the system cannot be attributed any properties of their own, independently of the other parts of the system. The Hilbert space for a two-particle singlet state is of the same dimensionality as the Hilbert space for one particle spin states.

All this shows that we cannot correctly say that we have several different objects in the system; rather the entire system is one quantum mechanical object. Hence, in a correct individuation, no problems arise. In short, the problem arises because of a mistaken individuation of quantum objects; quantum 'particles' cannot generally be considered as individual objects.

It is worth pointing out that this discussion is relevant also to non-locality and EPR correlations, because these phenomena occur precisely when we have two-particle states with correlations. In fact, one aspect of the locality problem is the impossibility

to factorise the joint probability distribution by conditionalising on some hidden parameter, which is a consequence of these particles not being individuals.

If it were possible to factorise equation (4.11), we would have no problem with non-locality, because then we could describe the state of one of the particles independently of the state of the other one, in which case there would be no non-locality. This would also enable us to describe the spatial part of each wave function as a function in real space, independently of the other wave, making a mapping of the total state into a sum of two independent states possible and the multi-dimensionality would be no more a problem than the multi-dimensionality of multi-particle states in classical mechanics.

A realistic interpretation of multi-particle wave functions, on the one hand, and understanding non-local interactions, on the other, turn out to be basically the same problem. This latter problem, which has evolved into the most discussed topic in the philosophy of quantum mechanics for the last three decades, will be thoroughly discussed in Chapter 8.

Let us sum up.

Problems (3) and (6), which confronted Schrödinger, are problems only because Schrödinger had too narrow a conception of real waves and their interactions. They disappear altogether once we relax our conditions for something to be a wave.

Problems (1) and (2) are in fact aspects of the measurement problem; thus far, they are by no means solved, but the point is that we need not view these as devastating objections against a *wave* interpretation; the measurement problem has to be dealt with no matter which interpretation we prefer. When adopting a wave interpretation together with the postulate of quantisation of interaction, the objections (1) and (2) do not give us additional problems; we are again left with the measurement problem (or part thereof). By the same token, a solution to the measurement problem also provides answers to objections 1 and 2.

The use of complex functions in wave mechanics is convenient when doing mathematics on two coupled dynamical functions. An exposition in terms only of real functions is possible; only the coupling between these requires an imaginary number. The magnitudes $|\Psi_r|^2$ and $|\Psi_i|^2$ can be given an interpretation in terms of classical quantities, i.e. as potential and kinetic energy respectively.

The failure of representing a multi-particle system with a wave function written as a sum of independent one-particle wave functions indicates that waves may interact and become connected, which is to say that we have a new object whose parts cannot any longer be attributed properties of their own. For objects far apart this connectedness manifests itself as non-local correlations, a difficulty no matter which interpretation we adopt.

Viewing an entity described by a wave function as a sort of wave thus poses no *additional* problems in our understanding of nature. Hence when counting pros and cons for this interpretation we have, in fact, no serious cons.

However, if the wave function really represents some kind of real wave, would it not be reasonable to suppose that we could measure, i.e. directly determine the form of this wave? Orthodox interpretation of quantum mechanics does not seem to allow for that. As soon as we measure some observable of the system we exchange energy with it and this process changes the wave function and thus the real wave abruptly.

However, there is a way of determining the wave, i.e. the wave function; the idea is to make the energy exchange so small so one can neglect the state change. This idea is proposed by Aharonov & Vaidman and the details follow.

4.3 Adiabatic measurements

Aharonov & Vaidman (1993, pp. 38–42) claim that it is possible to determine the wave function of a quantum system without changing its energy state. Their idea is to perform an interaction slow enough to make the expectation value of the interaction Hamiltonian approach zero as the time of interaction approaches infinity. Such a measurement is called a *protective measurement*, because the state of the object is protected during the measurement. Here is their argument.

The interaction Hamiltonian has the form $H = g(t)pA$, where p is the conjugate to the pointer observable q , A is the measured observable and $g(t)$ the time-dependent coupling factor. For any value of p the energy of the eigenstate shifts by an infinitesimal amount given by first-order perturbation theory (T is the absolute temperature and they put $\hbar = 1$):

$$\delta E = \langle H_{\text{int}} \rangle = \langle A \rangle p / T \quad (4.12)$$

Now the pointer will be shifted by the average value $\langle A \rangle$ by the time evolution $\exp(-ip\langle A \rangle)$, in contrast to the usual (strong) measurements, which yield one of the eigenvalues of A . By measuring the average of a sufficient number of observables A_n , for example the observables ‘position between x and $x+dx$ ’, where x varies along a sufficiently long distance, we can reconstruct the probability density, i.e. the squared wave function $|\Psi(x)|^2$. The Schrödinger wave can then be reconstructed by changing sign at the nodal surfaces.

In the more general case we also want to know the phase of the wave, which can be determined by measuring the current density;

$$\mathbf{j} = \frac{1}{2im} (\Psi^* \nabla \Psi - \Psi \nabla \Psi^*) \quad (4.13)$$

Writing $\Psi(x) = r(x)\exp\{i\theta(x)\}$, where $r(x) = \sqrt{\rho(x)}$ we find that

$$m\mathbf{j}(x)/\rho(x) = \nabla\theta \quad (4.14)$$

and the phase $\theta(x)$ can be found by integrating $\mathbf{j}/\rho$.

For degenerate energy eigenstates this method is not valid, because the vector potential must be taken into account when calculating the density current.

There is a further condition, not mentioned by Aharonov & Vaidman, for this method to work, namely that the energy spectrum of the measured object must be continuous. This is so because every measurement is an exchange of energy (as will be argued in Chapter 7), in other words, no measurement can be performed without changing the energy state. But if an extremely small amount of energy is exchanged (using very soft photons), the change of state of the measured object will

be very small and the described method might be applied in order to decide, with good approximation, the state (i.e. the states before and after the measurement are very nearly the same) and thus the wave function. Thus, this protective measurement is not absolutely adiabatic, only so to a certain approximation.

Thus, although the limitations for this procedure are rather strict, it is in some cases possible to determine the wave function. This possibility is one further argument for adopting the view that quantum waves are real.

4.4 Matter wave theory

If we start with Einstein's postulate of the equivalence between matter and energy and de Broglie's notion that to any material corpuscle corresponds (de Broglie adopted a dualistic ontology, assuming particles and waves as two distinct entities, which is why he used the word 'correspond') a vibration of frequency

$$f_0 = \frac{E}{h}$$

we get

$$f_0 = \frac{m_0 c^2}{h} \quad (4.15)$$

where m_0 is the rest mass of 'the corpuscle' and the vibration frequency is measured in the rest frame of the corpuscle. If we generalise to a frame of reference moving with velocity v relative to the rest frame of the corpuscle we get

$$f = \frac{mc^2}{h} \quad (4.16)$$

where f is the frequency as measured in the moving frame, and m the relativistic mass according to

$$m = \gamma \cdot m_0 \quad (4.17)$$

where

$$\gamma = \frac{1}{\sqrt{1 - \frac{v^2}{c^2}}}$$

De Broglie assumed that the corpuscle was a real particle and the vibration was a vibration of something different, some sort of field associated with the corpuscle. This idea later was developed to the pilot-wave theory by Bohm. However, the assumption of two different kinds of entities is unnecessary. Assuming that we do not have two entities, corpuscle and wave, but one having both particle and wave properties, albeit not simultaneously, we can still use de Broglie's formalism. The vibration is a kind of standing wave in the object in question. Then the vibrations are not mere artefacts of the theory but just as real as, for example, the mass of the

corpuscle. If this is correct, the transformations between different frames must give consistent results, which means that

$$f_1 = f_0 \gamma \quad (4.18)$$

because the vibration frequency is proportional to the energy. What are the consequences?

In the moving frame the corpuscle has momentum $p_1 = mv_1$ and according to de Broglie's formula the wavelength of the matter wave is

$$\lambda_1 = \frac{h}{p_1} = \frac{h}{m_0 v \gamma} \quad (4.19)$$

In order to determine the frequency in this frame we utilise the wave relation

$$\lambda_1 = \frac{v_p}{f_1} \quad (4.20)$$

where v_p is the phase velocity of the propagating wave. Together with equation (4.19) this yields

$$\frac{v_p}{f_1} = \frac{h}{m_0 v \gamma} \quad (4.21)$$

from which we infer that

$$f_1 = \frac{m_0 v \gamma v_p}{h} = f_0 \gamma \frac{v_p v}{c^2} \quad (4.22)$$

Comparing with equation (4.18) we see that the consistency of the wave picture is fulfilled if the phase velocity v_p fulfils

$$v_p = \frac{c^2}{v} \quad (4.23)$$

The consequence of this is that in all frames of reference (which necessarily move with a velocity $v < c$) the phase velocity is higher than the velocity of light. This looks like an outright conflict with relativity theory; however, that is not the case. Relativity tells us that no signal can be sent with a velocity exceeding that of light. Hence, we would have a conflict with relativity theory if it would be possible to send a signal using the phase information. But that is impossible: if the phase is well-determined so is the momentum, i.e. $\Delta p = 0$. Consequently the extension in space of the wave is indefinite, $\Delta x = \infty$. Thus the concept of signalling cannot be applied, because the object is everywhere. The same argument could be given in quantum field terms: if the phase is precisely determined, the occupation number of the field mode is completely indeterminate and that means that we cannot use the field for sending any message.

Some people, those with strong realistic attitudes, might say that the matter wave description cannot be a literal description of the real world, it must reflect mere mathematical traits of the theory. But why should reality conform to our present comprehension about it, why should common sense be a trustworthy guide? The correct way of finding out what is real is to look at the structure of our best theory in

the field, in this case quantum mechanics, and apply the criterion that only objects which can be substituted for variables, when quantum mechanics is expressed in the logical notation of quantification theory, are real. (This is Quine's famous dictum 'to be is to be the value of a variable'.) This is not a guarantee for success, of course, but our best option.

Another argument to the effect that a superluminal phase velocity does not conflict with relativity theory is the following: a necessary condition for deriving a conflict is that a frame of reference in which the partial wave is at rest is definable. But that is not possible, because using a frame of reference presupposes that we have at least one macroscopic material object with well-defined position and internal periodic state change that can be used to define a unit length and a unit time interval. But this macroscopic object cannot be identified with a partial wave, it must be a wave packet. So no frame of reference can be attributed a superluminal velocity, even if the phase velocity is superluminal.

We cannot, however, exclude that non-local correlations might be an effect of this superluminal velocity of the phases of partial waves. This will be discussed in detail in Chapter 8.

4.5 Wave packets

From Heisenberg's uncertainty (or as I prefer to say, 'indeterminacy') relations it follows that if the momentum of an object is determined with infinite precision, its position is left completely indeterminate. This means that it is spread out in all space and the velocity of the object loses all meaning. In order to get a more localised object we have to allow for a spread in momentum. The mathematical object will then be a wave packet comprised of partial waves with slightly different frequencies. This could be viewed as a roughly correct representation of the real object; the 'stuff' of this object is spread out, but it is not a rigid body.

A second argument for real objects being wave packets is that a homogeneous wave $\Psi(x,t) = A \cdot \exp(i(kx - \omega t))$ cannot convey any information; the wave must be deformed in order to be able to do so. But a wave packet can transmit information. The mathematically simplest form to handle is the minimum uncertainty wave packet. The signal velocity is the group velocity of the wave packet and it is defined as

$$v_g = 2\pi \left(\frac{df}{dk} \right)_{k=k_0} \quad (4.24)$$

where k_0 refers to the average wave number in the wave packet (McGervey, 1971, p. 116). For a matter wave packet this becomes

$$v_g = \left(\frac{dE}{dp} \right)_{p=p_0} = \frac{p_0 c^2}{E_0} = v_1 \quad (4.25)$$

which is just as it should be.

The phase velocity of equations (4.20) – (4.23) is the average phase velocity in the wave packet and it must be assumed that the spread around this value is small in order to get a wave packet.

The frequencies f_0 and f_1 in equations (4.18) – (4.22) are the average frequencies, in different frames, of the wave packet. The picture is that of rather ‘amorphous’ object made up of an indefinite number of vibrations of slightly different frequencies. The reason why this object is usually referred to as a particle is that it interacts as one inseparable whole at a rather well-localised places in space. The wave description is appropriate when and only when discussing its motion, whereas in interactions we use the particle picture.

As the wave packet is composed of many plane waves with infinite extension the one-dimensional case can be expressed as

$$\Psi(x, t) = \int_{-\infty}^{\infty} G(k) \exp[-ik(x - v_p t)] dk \quad (4.26)$$

where $G(k)$ is the envelope function. If, for the sake of simplicity, we chose a Gaussian function for $G(k)$ we will get

$$\Psi(x, t) = \int_{-\infty}^{\infty} \exp\left[-\frac{(k - k_0)^2}{a}\right] \exp[-ik(x - v_p t)] dk \quad (4.26)$$

where k_0 is the mean wave number in the wave packet. The rms-width of this wave packet is proportional to $\sqrt{a}$. When this wave packet travels in the positive x -direction with a speed given by the group velocity

$$v_g = \frac{k_0 \hbar}{m}$$

the wave packet can be expressed as

$$\Psi(x, t) = \exp(i(k_0 x - v_p t)) \exp\left[-\frac{a(1 - ibt)}{4(1 + b^2 t^2)}(x - v_g t)^2\right] \quad (4.27)$$

where $b = \frac{\hbar}{2m}$.

One salient feature of this wave packet is that every plane wave with wave number k is infinitely extended in space, but outside the bulk of the wave packet (with a dimension of the order of magnitude of $\sqrt{a}$) the total amplitude is nearly zero. But still, each component wave exists outside the wave packet, at least in the mathematical model. Hence the object is not completely confined within the limits of the bulk of the wave packet, it has ‘tails’. Are these tails real?

4.6 Empirical results supporting the causal relevance of partial waves

The question is thus whether the decomposition of wave packets into infinitely extended plane waves is an artefact of the mathematical model or a trait of nature? Empirical evidence suggests that most plausibly it is a trait of nature. Rauch *et al.*

(1992, pp. 45–57) have performed a number of neutron interference experiments, and one of these corroborates the existence of partial waves outside the wave packet. Their experiment is in outline as follows.

A beam of neutrons is incident on a perfect Si-crystal neutron interferometer. In this device the neutron beam is split into two coherent parts. These partial beams are later recombined at a plate at which the recombined neutron beam is reflected in one of two possible directions depending on the phase relation between the incoming partial beams. Thus, the intensity distribution in the two detectors is a function of the phase differences. By inserting a phase-shifter in one of the paths, the intensity can be modulated and the resulting pattern shows typical interference effects. This confirms the wave nature of the neutrons. By keeping the intensity in the neutron beam so low that at most one neutron is present in the interferometer at any time, the possibility of interference between different neutrons can be excluded. Hence, the interference pattern shows that each individual neutron has travelled both ways. (Adherents to the de Broglie–Bohm theory would not admit that, but I will ignore their interpretation for the moment. The purpose in this context is to provide an explanation without introducing extra assumptions outside quantum theory *per se*.) So far this experiment is thus very much like a double-slit set-up.

In the next step, an additional device, a Bismuth-sample which interacts with the

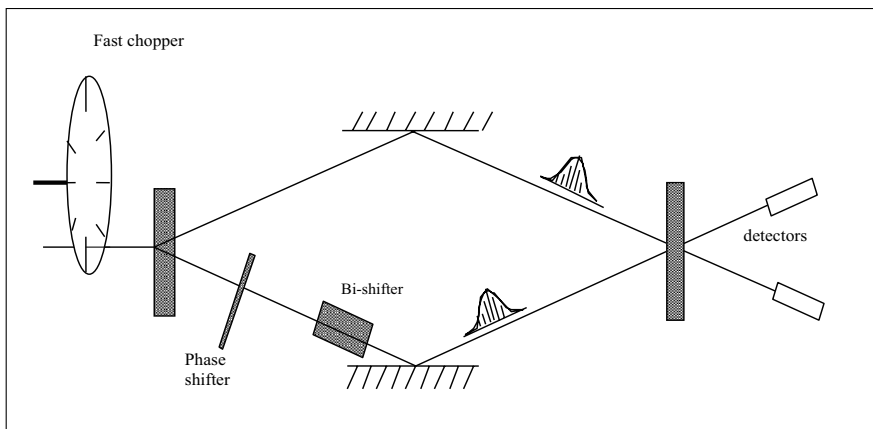


Figure 4.1 Principal arrangement of time-of-flight neutron interferometry. The recombined neutron beam will go to the upper or lower detector depending on the turning of the first phase shifter

neutrons, is placed in one of the beams. The velocity of the neutrons in the Bi-sample is lower than that in vacuum. The result is that this partial beam will arrive later than the other part to the plate. Thus, there is no place where the two partial beams overlap, thus making interference impossible. Hence, no intensity variations can be observed even if the intensity is modulated with the phase-shifter.

In the third step of the experiment, the pulses are chopped in short time ‘slices’. The experimental set-up is shown in Figure 4.1.

This will increase the definiteness of the momentum of the neutron pulses, which can be seen from the following calculation. We denote the distance between the chopper and the detector by L , the time of flight t and momentum of the neutron pulse $p = mv$. We thus have

$$p = \frac{mL}{t} \quad (4.28)$$

and by differentiating

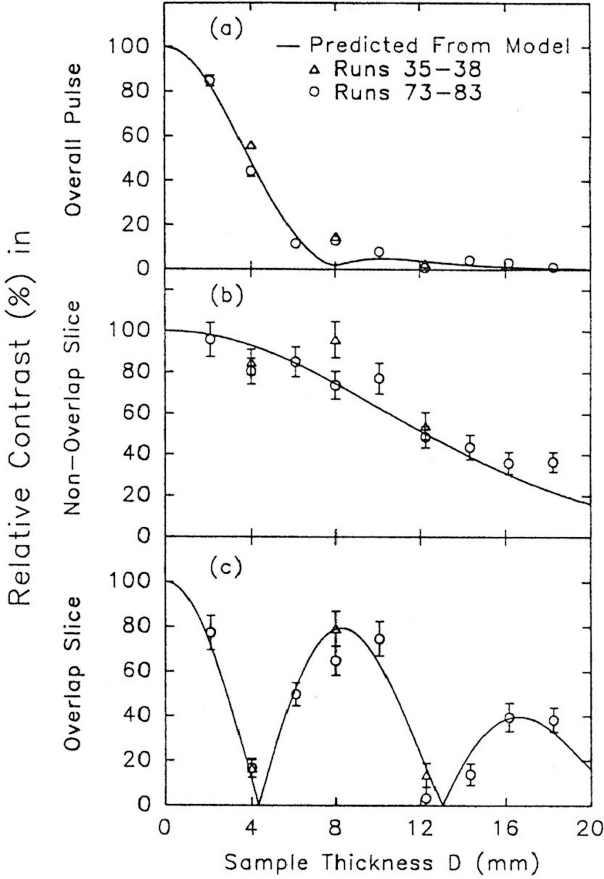


Figure 4.2 The first curve shows the relative contrast in the entire neutron pulse as a function of the thickness of the inserted Bismuth sample. The interference is effectively destroyed when the sample is thicker than 8 mm. The curves (b) and (c) show the relative contrast in two time slices of 2 ms each. In (b) minimum overlap is expected, in (c) the expected overlap is at its maximum. Interference, i.e. maximum contrast, occurs with $D \approx 8$ mm. The solid lines represent theoretical predictions and data points observed results. Adapted from Fig. 9 in Rauch *et al.* (1992).

$$\delta p = -\frac{mL\delta t}{t^2} \quad (4.29)$$

The momentum indeterminacy δp is thus proportional to the time indeterminacy δt . The time indeterminacy is composed of two components, the time window of the chopper and the time resolution of the detector. It is evident that for a short time window, the momentum spread in each of these time slices will be less than that in the total neutron pulse.

It should be pointed out that δt must be much longer than the coherence time $\Delta t^c \approx h/v \cdot \delta p^c$, where δp^c is the momentum indeterminacy in the minimum uncertainty wave packet of the neutron pulse. (The coherence time is thus the time it takes for the wave packet with coherence length $h/\delta p^c$ and velocity v to pass a point.) We observe that if δt is shorter than Δt^c , there will be an increase in the momentum indeterminacy.

However, narrowing the momentum distribution will produce a spatially more extended wave. The time slices of two partial waves, which previously arrived one after the other will now overlap, due to more well-defined k -vectors in each time slice. This makes interference possible again, resulting in intensity variations in each time channel. Hence interference can be, and indeed is, observed again, as can be seen from Figure 4.2. The empirical results are in very good agreement with theory. This is true for the destruction of interference when the Bi-sample is inserted in one of the paths, as well as for the recovery of the intensity variations when the wave packet is chopped into short time intervals.

My interpretation of this experiment is that the conception of a wave packet as a sum of many widely extended waves with different wave numbers is not only a mathematical fiction but a real trait of the world. This is of course necessary if the mathematical model shall be of any value as an explanation. Using a model in an explanation always triggers the question: do the properties of the model have a counterpart in nature or do they not? If not, the purported explanation lacks explanatory force.

4.7 Dispersion of wave packets

From equation (4.24), the definition of group velocity of a wave packet, and utilising the general wave equation applied to phase velocity,

$$v_p = \frac{\omega}{k}$$

we can infer that

$$v_p = v_g - k \left(\frac{dv_p}{dk} \right)_{k=k_0} \quad (4.30)$$

Thus we see that since the phase velocity depends on the wave number, the group velocity and phase velocity will differ. In other words, the wave packet is subject to a longitudinal dispersion with the effect that a ‘tightly bundled’ wave packet will be smeared out as it goes along. As this effect is derivable without any specific

assumptions regarding the nature of the waves described, it also applies to quantum waves. This can also be seen in equation (4.27) where the factor t^2 in the denominator of the envelope function

$$\exp\left[\frac{a(1 - ibt)}{4(1 + b^2t^2)}(x - v_g t)^2\right] \quad (4.31)$$

accounts for the time-dependence of the width of the wave packet.

Dispersion also occurs in the transversal directions, although the mechanism is different. Hence, considering a wave packet made up of one single electron, this single electron can easily be distributed over several centimetres in a laboratory experiment. Still, it interacts at one particular place at each point of time. Hence, the wave packet must collapse during its interaction with a macroscopic object. What this strongly suggests is that the collapse of the wave function during measurement is a real collapse of a real wave. It also suggests that collapses occur not only in measurement situations, for interactions occur all the time. This will be thoroughly discussed in Chapter 7.

4.8 Doppler effect of partial waves

As all known wave phenomena exhibit the Doppler effect, this is likely to be true also for matter waves, and Christian Cormier-Delanoue (1996) has derived a formula for the Doppler effect for matter waves. This formula is, however, inconsistent with the usual relativistic Doppler formula relating frequencies of any sort of wave in two frames of reference

$$f_1 = f_0 \frac{\sqrt{1 - \beta^2}}{1 - \beta \cos \theta} \quad (4.32)$$

where $\beta = v/c$ and v is the relative velocity between the frames of reference in which the frequencies f_0 and f_1 are measured and θ is the angle between the direction of the wave vector and the direction of motion of one frame S1 as seen in the other frame S0.

Cormier-Delanoue concludes that matter waves cannot be real waves, for if they were, they would conform to equation (4.32). More specifically, his argument goes as follows. In order to be objective an entity must exist independently of an observer's consciousness. Therefore, 'an objective science is concerned with real objects, i.e. objects which exist independently of their actual observation.' From this he concludes that, for example, special relativity is an objective theory because 'it describes different aspects or values of the same physical quantity when measured by observers in different frames of reference. Quite evidently, this entity has to exist objectively in all inertial frames in order that its various aspects may be described by the theory' (Cormier-Delanoue, 1996, p. 96). This is Cormier-Delanoue's necessary and sufficient criterion for something being real and objectively existing. As this criterion is not fulfilled by matter waves, he concludes that they are not real.

As an argument against matter waves as something separate and distinct from the corpuscles, the argument has considerable force and it is similar to the argument

often directed against Bohm's pilot wave theory, namely that pilot waves are not Lorentz-invariant.

However, if we assume that there is only one type of thing, objects showing wave behaviour during propagation while showing particle behaviour during interaction, the situation becomes different. In this perspective it is important to observe the distinction between wave packets and partial waves within the wave packet.

Cormier-Delanoue's argument is not applicable to wave packets. The reason is that he assumes that it is meaningful to associate a well defined frequency, i.e. he assumes that $\Delta\omega = 0$, to each material 'corpuscle' (this is Cormier-Delanoue's word) according to formula (4.15). That means that we would have determined the energy of the corpuscle with infinite accuracy, ($\Delta E=0$) which requires an infinite long measurement according the time-energy indeterminacy. But $\Delta\omega=0$ entails $\Delta p=0$, which gives $\Delta x = \infty$, i.e. the position of the corpuscle is completely indeterminate. Clearly, this cannot be a wave packet.

However, if we interpret the frequencies in equation (4.32) as mean frequencies in a wave packet the situation is different. In order meaningfully to use equation (4.32) as a criterion for the property to be real we must think of objects occupying small regions in space-time, which amounts to using, not single matter waves with well-defined wave numbers, but wave packets with a non-negligible dispersion of frequencies. Now we have no problem with equation (4.32), for in the wave packet the mean frequency is proportional to the energy of the object (again, taking indeterminacy into consideration) and since energy transforms correctly between different frames of reference, the mean frequency does that also. Wave packets therefore fulfil the reality criterion.

But, then what about the Cormier-Delanoue's derived Doppler formula conflicting with equation (4.32)? What is wrong with it?

I think there is nothing mathematically wrong with the formula, but the question is what it applies to. As we have seen, it cannot fit the behaviour of wave packets. Let us then assume that it applies to the partial waves making up the wave packet. According to Cormier-Delanoue's criterion, these partial waves cannot be real. I agree: partial waves are not identifiable *objects* that can be found using identity criteria, in contrast to classical waves. The notion of partial waves should not be interpreted as entailing that wave packets are made up of other objects. Instead, a description of a wave packet in terms of partial waves must be construed as a structural description of the features of a real and indivisible object, the wave packet. It has no parts being individual objects, but it has an internal structure. Its structural description can be given in different inertial frames and Cormier-Delanoue's formula relates descriptions in different frames to each other. The fact that this formula conflicts with the usual Doppler formula (4.32) is irrelevant.

Summarising, Cormier-Delaoune's argument is valid as an argument against de Broglie's conception of matter waves but not against my interpretation of the quantum realm as consisting of wave packets.

We must clearly differentiate between two uses of the notion of a matter wave: either as referring to an individual object, or as referring to *a part* of this individual, a partial wave, distinguished either spatially or with reference to wave number. When we use the term 'matter wave' as an alternative to a 'quantum system' – which means

that we view the object as a wave packet – no objections can be raised against calling them real. However, partial waves are different. No identity criterion can be found for the partial waves of a wave packet, hence partial waves cannot be viewed as individual objects. Therefore, talk about partial waves should be construed as talk about structural properties of those waves that are identified as real and referents of expressions such as ‘quantum systems’.

One might be tempted to talk about identifiable partial waves when these are spatially separated, for example, in an interference experiment. It appears perfectly admissible to say things like ‘the partial wave taking the upper path will hit the upper detector’. However, this is not appropriate. On closer examination, the conditions for full reification (discussed in Section 3.2) are not fulfilled; one cannot attribute a further definite property to the purported referent of ‘the partial wave taking the upper path’. The reason is that the partial wave cannot be attributed any quantitative properties of its own because any quantitative property manifests itself in an interaction and such an interaction involves the entire system. The phrase above, ‘the partial wave taking the upper path will hit the upper detector’, is seriously misrepresenting the real situation because the expression ‘a detector being triggered’ does not only refer to an interaction between the partial wave being close to that detector but to an interaction between the detector and the entire system; hence all conclusions drawn from the detection concerns the entire system. (This will be more thoroughly discussed in Chapters 6 and 7.) The condition for full reification, namely to be able first to identify an object and then to attribute another property to *it* (observe the essential pronoun!), namely a property not utilised for the identification, cannot as a matter of principle be fulfilled by parts of quantum systems. The reason is that interactions are always quantised; a quantum system interacts as a matter of principle as *an undivided whole* with other objects.

However, dismissing partial waves as real independently existing entities does not entail that sentences containing expressions such as ‘the partial wave having wave number k ’ must be false. Such a phrase cannot serve as a noun phrase, but it can be part of a predication. (Example: ‘The system Ψ has a partial wave with wave number k ’.) Although a quantum system by its very nature interacts as an undivided whole, it is not without inner structure. The wave function is a structural description of a quantum system that conforms to that wave function and the resolution of the wave function into a set of orthogonal wave functions is a way of providing that very structure.

It follows that any talk about the ‘velocity of partial waves’ is a second-order predication. It is an indirect way of saying something about the entire system in question, namely talking about the velocity with which structural changes in the system are transmitted to all its parts.

In conclusion, if the term ‘matter wave’ is interpreted as a single mode vibration within a wave packet, Cormier-Delanoue is right in claiming that such a matter wave is not real, i.e. it is no object in the full sense of the word. However, if ‘matter wave’ is interpreted as referring to an individual quantum system with a non-negligible dispersion, the matter wave does exist. ‘Matter wave’ is another word for the object in question, a word that gives a better description of its nature.

4.9 Interaction between matter waves and electromagnetic waves

Yet another aspect of Schrödinger's real wave interpretation of the wave function is of interest in this context, namely the interaction between matter waves and electromagnetic waves. Schrödinger, thinking of this interaction as analogous to any other kind of interaction between continuous fields, tried to analyse the situation by expressing the electron as a self-field (*Eigenfeld*) and adding this field to the electromagnetic field. The total field, i.e. the total potential, is then used in the Schrödinger equation, but the resulting energy levels, for example for the hydrogen atom, are completely wrong. Schrödinger commented:

Fragt man sich nun, ob diese in sich geschlossene Feldtheorie ... der Wirklichkeit entspricht in der Art, wie man das früher von dergleichen Theorien erhofft hatte, so ist die Frage zu verneinen.

[If one asks oneself if this self-consistent field theory ... represents reality in the way one had hoped for, the answer must be negative.] (Schrödinger, 1926b, p. 270)

He proposed two reasons why this was so:

Gerade die Geschlossenheit der Feldgleichungen erscheint somit in eigenartiger Weise durchbrochen. Man kann das heute wohl noch nicht ganz verstehen, hat es aber mit folgenden zwei dingen in Zusammenhang zu bringen.

1. Der Austausch von Energie und Impuls zwischen dem elektromagnetischen Feld und der 'Materie' vollzieht sich eben in der Wirklichkeit nicht in der kontinuierlichen Weise, wie die Feldmässige Aussage (24) glauben macht.

2. Auch in der Lorentzschen Theorie hat man in die Bewegungsgleichungen des einzelnen Elektrons zunächst nur die Felder der übrigen Elektronen, nicht das Eigenfeld einzusetzen. Die Rückwirkung des letzteren ist zum überwiegenden Teil *schon bei Aufstellung der Bewegungsgleichungen* als elektromagnetische Masse in Rechnung gestellt. (Schrödinger, 1926b, p. 271)

[Thus, in a peculiar way precisely the consistency of the field equations are broken. One can presently not fully understand this, but it has to do with the following two things taken together:

1. The exchange of energy and momentum between the electromagnetic field and 'matter' do not in reality occur continuously, which one would have guessed from the field like form of equation (24).

2. One has even in Lorentz's theory only introduced the field from the other electrons, but not the self-field, into the equations of motion for the individual electrons. The reaction from the latter is to a great extent taken into account as the electromagnetic mass in the equations of motion. (My translation)]

The equation (24) referred to is

$$\frac{\partial}{\partial x_\sigma}(T_{\rho\sigma} + S_{\rho\sigma}) = 0 \quad (4.33)$$

where $T_{\rho\sigma}$ is the electromagnetic field tensor and $S_{\rho\sigma}$ is the charge distribution tensor, the latter being the ‘self potential’, the potential expressing the distribution of electronic matter. This equation does not yield correct predictions.

Barut (1988) concurs with Schrödinger in thinking that an electron should be viewed as a continuous wave phenomenon, the self-potential of which should be added to the external potential when solving Schrödinger’s equation. But Schrödinger’s observation that this yields infinite energy levels, Barut claims to have avoided by adding a non-linear term $R(\Psi)$ to the Schrödinger equation:

$$i \frac{\partial \Psi}{\partial t} = H\Psi + R(\Psi) \quad (4.34)$$

This non-linear term represents the radiation reaction and the problem is to find its form. Barut uses a classical analogy to solve this problem. In Lorentz theory for the classical radiating electron there is such a non-linear term ($c = 1$):

$$m \ddot{x}_\mu = eF_{\mu\nu} \dot{x}^\nu + \frac{2}{3} e^2 (\ddot{x}_\mu + (\ddot{x})^2 \dot{x}_\mu) \quad (4.35)$$

Barut starts by expanding the wave function $\Psi(x)$

$$\psi(x) = \sum_n \psi_n(r) \exp(-iE_n t) \quad (4.36)$$

This equation is then inserted into Schrödinger’s equation with the self-potential included. The result integrated over time is an infinite set of coupled equations:

$$(\gamma^0 E_n - \boldsymbol{\gamma} \cdot \mathbf{p} - eA^{ext} - m)\Psi_n(\mathbf{r}) = -\frac{a}{2\pi^2} \sum_{rsm} \int d\mathbf{k} d\mathbf{r}' \frac{\exp(i\mathbf{k}(\mathbf{r}-\mathbf{r}'))}{(E_r - E_s)^2 - \mathbf{k}^2} \gamma^\mu \psi_m(\mathbf{r}) \psi_r^*(\mathbf{r}') \gamma_\mu \psi_s(\mathbf{r}') \quad (4.37)$$

Barut claims that the terms on the right-hand side can be respectively identified as representing spontaneous emission, Lamb shift and vacuum polarisation in the external field. Furthermore, he claims that the self-terms can be calculated by an iterative procedure by which the infinite terms can be separated and absorbed into the mass and energy terms.

Barut points out that equation (4.36) cannot be interpreted as usual: ‘The ψ in ... $\psi = \sum \psi_n(x) \exp(-iE_n t)$ is not normalized and should not be confused with a general superposition state of the linear theory in the absence of self-energy. The field equation tells us that we have only to interpret the Fourier coefficients of ψ as possible states of the system which are now, except for the ground state, not stable but metastable due to spontaneous emission’ (Barut, 1988, p. 34).

Several conclusions are possible according to Barut. First, it has been consistently established that matter in the form of matter waves and the electromagnetic field interact with each other. Secondly, the electromagnetic field is not introduced separately but is expressed in terms of its source, which makes second quantisation

unnecessary: if the matter field is quantised, then so is the emerging electromagnetic field.

I have two comments to Barut's analysis.

(1) Inserting the sum of the potentials emanating from external objects and from the electron itself into the Schrödinger equation for the electron in an electromagnetic field is, as far as I can see, inconsistent. The potentials used when writing down the Schrödinger equation are representations of the influence from other objects, the sources of the potentials, on the system described by the solution to the Schrödinger equation. We can, alternatively, neglect the sources and regard the corresponding potentials as entities in their own right. In both cases they are employed to describe the action of a potential, something else, on an object. But no object can act on itself or, within the same object, one aspect cannot act on another aspect. The only possibility to give a meaningful interpretation of the concept of selfpotential is if the charged matter and the selfpotential are conceived as different parts of the object and that these parts interact. This possibility can, however, be ruled out since charge and electromagnetic field emanating from this charge are two different descriptions of one and the same thing. The notion of a selfpotential must then be seen as an alternative description of the very charge whose evolution we try to describe by solving the Schrödinger equation. It thus makes no sense to add the selfpotential and the external potential and use it as a totally external potential in the Schrödinger equation. The complicated (and to me far from perspicuous) equations Barut uses in his theory are therefore unnecessary.

(2) Quantisation is in Barut's and Schrödinger's approach an effect of a resonance between the interacting waves. It is thus legitimate to describe the interaction as built up by infinitesimal interactions between parts of the interacting objects. But suppose we start by viewing quantisation as an ultimate trait of nature; then it is simply wrong to describe the interaction between the charged matter and the electromagnetic field as an integral over infinitesimal interactions. (This was precisely Planck's assumption in his 1900 paper.) That is to say, the first option mentioned above would be right and the non-linear correction to the Schrödinger equation is unjustified.

In my view, both the remarks made by Schrödinger are correct: exchange of energy is quantised and the self potential cannot be added to the external potential when solving the wave equation. Barut is on the wrong track.

4.10 Space-time metric and complex phases of wave equations

As promised in Section 4.2.2, I will now elaborate on the discussion of why quantum waves are complex. We saw that the evolution of real physical entities, such as potential and kinetic energy densities, can be described by real valued functions, but that these functions are coupled in such a way that the complete description requires a complex function. This is, by contrast, not the case for electromagnetic waves in classical electromagnetism. The energy in the electromagnetic pulse is proportional to the square of the electric field vector only. As the magnetic field is a relativistic manifestation of the electric field, or vice versa, there is in reality only one field, the

electromagnetic field, although it has different appearances in different frames of reference. Therefore we need only one wave function.

In quantum mechanics, however, the matter is different. Alfred Landé (1965, p. 68) has discussed the reasons why complex functions are needed. Landé reasons as follows.

Consider the matrix elements $A(p, p')$ of any operator $\mathbf{A}$:

$$A(p, p') = \int \Psi(q, p) \mathbf{A} \Psi(p', q) dq \quad (4.38)$$

These transition values are probabilities that fulfil the conditions on a probability measure, if the function $\Psi(q, p)$ is normalised. Now it seems natural to impose the restriction that these transition values shall depend only on the differences $p-p'$; why should these measures be dependent on our choice of a zero point on the scale? If this restriction is assumed to be valid for all operators, this must be so also for the delta function $D(q) = \delta(q-q_0)$, where q_0 is any point outside of which $D(q)$ vanishes. Then we have

$$T(p, p') = \int \Psi(q, p) \partial(q - q_0) \Psi(p', q) dq = \Psi(q_0, p) \Psi(p', q_0) \quad (4.39)$$

Now, which functions $\Psi(p)$ will satisfy the condition that the product of the function with itself depends only on the difference $p-p'$? The answer is that there is no such real function. We must take a complex function of the type

$$\Psi(p, q) = C \exp(iff(q)p) \quad (4.40)$$

where $f(q)$ is any real-valued function of q . Exchanging q and p we get

$$\Psi(p, q) = C \exp(ig(p)q) \quad (4.41)$$

and it is readily confirmed that $f(q) = kq$ and $g(p) = kp$, where k is a real constant. Hence

$$\Psi = C \exp(ikpq) \text{ (The constant } k \text{ is later to be identified with } \hbar) \quad (4.42)$$

In conclusion then, if the transition elements $T(q, q')$ are to be dependent solely on the differences $q-q'$, the state functions must be complex.

This argument does not really answer the philosophical question of why complex functions are needed, because it starts by taking the probabilities as given by $A(p, p') = \int \Psi(q, p) \mathbf{A} \Psi(p', q) dq$. The natural question is: why should the transition elements have exactly this form? What we need is a physical argument and not one to the effect that the theory fits the facts.

Perhaps more insight could be gained if we reformulate the question thus: measurable quantities such as density distributions are real-valued; why must we use a function, which multiplied by its own complex conjugate gives the density distribution, as the fundamental dynamical mathematical object? The answer is that the density function is bounded from below, which means that destructive interference between densities is impossible. But destructive interference occurs and

hence we need the quantum wave with its complex amplitude in order to account for destructive interference.

Another way of expressing the same thing is to say that a quantum object can divide itself into spatially distinct portions during motion through space. When these parts are united again we observe interference effects and at certain points the density, i.e. the wave intensity, becomes zero, due to destructive interference; this might be the result of different parts travelling different paths.

4.11 The path integral picture and the wave interpretation

An alternative exposition of quantum mechanics which for some purposes gives understanding of the nature of quantum objects is that of using path integrals. Although Feynman, who invented the path integral formalism for quantum mechanics, never claimed that it was anything more than an alternative description of quantum mechanics, which for some practical purposes is very useful, I think we could in fact give a realistic interpretation of its core concepts. In his book *Quantum Field Theory*, Michio Kaku lists a number of advantages of the path integral method. Most of these advantages are calculational but the following one is of deeper interest because it relates to questions of interpretation and motivation:

The path integral formalism is based intuitively on the fundamental principles of quantum mechanics. Quantisation prescriptions, which may seem rather arbitrary in the operator formalism, have a simple physical interpretation in the path integral formalism. (Kaku, 1993, p. 261)

Another benefit of this approach is that Schrödinger's equation can be derived from postulates about transition probabilities.

Kaku starts his exposition by stating two fundamental postulates.

Postulate 1. The probability $P(a,b)$ of a particle moving from point a to point b is the square of the absolute value of a complex number, the transition function $K(b,a)$:

$$P(b,a) = |K(b,a)|^2 \quad (4.43)$$

Postulate 2. The transition function $K(b,a)$ is given by the sum of phase factors $e^{iS/\hbar}$, where S is the action, taken over all possible paths from a to b :

$$K(b,a) = \sum_{\text{paths}} k e^{2\pi i S / \hbar} \quad (4.44)$$

The constant k is determined by

$$K(c,a) = \sum_{\text{paths}} K(c,b)K(b,a) \quad (4.45)$$

where the summation is done over all intermediate points b between a and c .

Kaku supports these postulates by reference to interference experiments such as the double slit experiment. Here we can supplement this support with the claim that interference can only be understood as a consequence of the objects being waves.

Furthermore, the motion of each tiny part of an object is determined by the action on this part. This is a classical principle, but in contrast to classical mechanics in which the path of the object can be determined as the path along which the action integral is minimal, the quantum object, being a wave, spreads out and goes all admissible paths. This is precisely the behaviour of, for example, classical waves on a water surface or in air, albeit these waves do not transfer ‘stuff’, only energy. However, as we now know, ‘stuff’, i.e. matter, is just a form of energy; hence the step taken by saying that matter can be transferred in space–time as interfering waves is a natural generalisation from classical physics.

What could be the support for the first postulate, then? Why must the transition probability be the square of the transition function?

There is a classical analogy which could be useful. The intensity distribution of electromagnetic radiation is proportional to the square of the field vector E . There is an analogy between the wave function and the electric field vector in so far as the evolution of both is wavelike and both are the fundamental quantities we use to describe the dynamics of the systems of interest. Could there be found an analogy between the intensity of the electric field and the probability for transition? Yes, a very close one. The square of the wave function is interpreted as the energy distribution of the quantum object, just as the square of the electric field vector is the energy distribution of the electric field. The square of the wave function in turn determines the probability for exchange of energy with an external object. Why? Well, I take the probabilities calculated in quantum mechanics to reflect an irreducible randomness in nature, not mere lack of information, as in most other statistical theories. This view was argued for in Chapter 3. There I also argued for interpreting probabilities as propensities for exchange of energy. Given these views, what, then, is more natural than saying that the probability of exchange of energy is proportional to the energy density? Thus, given the wave picture of things and the propensity interpretation of quantum probabilities, we can give a reason for these postulates. They do not appear as mere brute facts which happen to account for empirical regularities, but are motivated by the very nature of the things they describe. In short, the wave interpretation of quantum objects motivates these postulates.

Starting with these two postulates, Kaku discretises the path by taking the time differential dt to be a small interval ε and rewriting the classical kinetic energy:

$$dt \rightarrow \varepsilon, \quad \frac{1}{2} \dot{x}_i^2 dt \rightarrow \frac{1}{2} m(x_n - x_{n+1})_i^2 \varepsilon^{-1}$$

This gives an explicit expression for the transition function:

$$K(b, a) = \lim_{\varepsilon \rightarrow 0} \iint \dots \int dx_2 dx_3 \dots dx_{N-1} k \exp \left(\frac{im}{2\varepsilon} \sum_{n=1}^{N-1} (x_n - x_{n+1})^2 \right) \quad (4.46)$$

Unfortunately, this integral is hard to calculate. For Gaussian integrals however, it is possible and that is useful because many real situations can be described by Gaussian functions. In those cases the final result becomes

$$K(b, a) = \left| \frac{m}{2\pi(t_b - t_a)} \right|^{1/2} \exp \left(\frac{(1/2)im(x_b - x_a)^2}{t_b - t_a} \right) \quad (4.47)$$

which is the Green function propagating the Schrödinger waves.

This transition function is an analogy to Huygen's principle, according to which each point on a wavefront can be viewed as the starting point of a new wave. The integration over all these infinitesimal wavefronts gives an overall evolution that can be expressed as

$$\Psi(x_j, t_j) = \int_{-\infty}^{\infty} K(x_j, t_j; x_i, t_i) \Psi(x_i, t_i) dx_i \quad (4.48)$$

The time evolution from t to $t+\varepsilon$ is then

$$\Psi(x, t + \varepsilon) = \int_{-\infty}^{\infty} A^{-1} \exp\left(\frac{im(x-y)^2}{2\varepsilon}\right) \Psi(y, t) dy \quad (4.49)$$

where $A = \left(\frac{2\pi i\varepsilon}{m}\right)^{1/2}$.

A variable substitution $\eta = y - x$ gives

$$\Psi(x, t + \varepsilon) = \Psi \int_{-\infty}^{\infty} A^{-1} \exp\left(\frac{im\eta^2}{2\varepsilon}\right) \Psi(x + \eta, t) d\eta \quad (4.50)$$

Both sides of the equation can now be Taylor expanded, the left-hand side in terms of t and the right hand-side in terms of η :

$$\Psi(x, t) + \varepsilon \frac{\partial \Psi}{\partial t} = \int_{-\infty}^{\infty} A^{-1} \exp\left(\frac{im\eta^2}{2\varepsilon}\right) \times \left(\Psi(x, t) + \eta \frac{\partial \Psi}{\partial x} + \frac{1}{2} \eta^2 \frac{\partial^2 \Psi}{\partial x^2} + \dots \right) d\eta \quad (4.51)$$

The integration of the linear term in η vanishes because it is linear and the integration of higher terms vanishes in the limit $\varepsilon \rightarrow 0$. The only remaining term is that of second order in η which gives us:

$$i \frac{\partial \Psi}{\partial t} = \frac{1}{2m} \frac{\partial^2 \Psi}{\partial x^2} \quad (4.52)$$

which is the Schrödinger equation for a free field in one-dimension. It is easy to insert a potential and go through the same reasoning.

There is a further bonus of the path integral approach, namely that the operators for position and momentum can be introduced in a natural way. This is satisfying because in orthodox quantum mechanics the introduction of operators appears rather arbitrary. Without going into details one can see that if we define the momentum operator $\mathbf{p}$ by the equation

$$\mathbf{p}|p\rangle = p|p\rangle$$

It is easy to see that this operator can be identified with $i\hbar\partial/\partial x$ if $|p\rangle$ has the form of an exponential function $\exp(i\hbar x)$.

Chapter 5

Particle Behaviour of Waves

5.1 Introduction

How is it possible that entities with distinct wave properties, i.e. having considerable spread in space and governed by the superposition principle, in some situations behave just like particles? One possible answer, suggested by Schrödinger, is that the real entities must be wave packets of such a minute dimension that they sometimes act as point-like particles. For example, if the effective extension of the wave packet is less than the size of a film grain in a film emulsion, then each wave packet will most often hit only one film grain and we would observe singular dots on the film, i.e. we observe such a particle-like behaviour.

However, as discussed in the preceding chapter, this assumption leads to problems when the phenomenon of dispersion is considered; any wave packet will disperse in the direction of its propagation as well as in the transverse direction. If the film is sufficiently far from the preparation site, the wave will, on detection, hit many silver bromide grains however small it was from the start. Schrödinger never managed to solve this problem, as mentioned in Chapter 4. If, however, quantisation of action is assumed, a large wave packet will *interact* with only one object at a time, even if the wave packet is in contact with several other objects. And since only one object (for example a single silver bromide grain) is triggered while its neighbours are unaffected, it looks like a particle has hit the screen. This holds quite generally, since quantisation of action is the reason why waves show particle-like behaviour in all interactions. I will now give a more detailed argument, proceeding as follows: first the concept of quantisation is analysed, then it is shown that this quantisation of interaction between wave packets implies their collapse, irrespective of whether any measurements are performed. The collapse is then shown to be responsible for particle behaviour of quantum systems. Finally, several experiments purported to prove that electrons are particle-like objects are discussed and it is shown that the particle character can be explained as said.

5.2 What is meant by quantisation?

The idea that quantisation is a fundamental feature of the physical world was first formulated in Max Planck's famous lecture '*Zur Theorie des Gesetzes der Energieverteilung im Normalspektrum*', held 14 December 1900 and printed shortly after (Planck, 1900). In this lecture we meet for the first time the assumption that the energy levels of an oscillator coupled to a radiation field are not continuous but discretised.

However, Planck's text has been subject of various interpretations by the historians of physics in recent years. Thomas Kuhn, to mention a well-known author, claims that Planck in fact did not assume quantisation as a new fundamental principle. Kuhn claims that Planck did not view the assumption that the energy of the radiation field is divided into cells of a size proportional to the frequency of the oscillator coupled to the field as a new physical principle but as a mere calculational device. Only much later was it understood that this assumption was in outright contradiction to classical physics and therefore must be conceived as a new physical principle. Furthermore, as Kuhn points out, Planck overlooked an important condition in the derivation of the radiation law named after him, namely that the derivation is valid only for frequencies such that $h\nu \ll kT$.¹ Hence, the introduction of quantisation as a new physical principle must be credited as much to Einstein as to Planck, because Einstein (1905, 1965) was the first who clearly saw the physical implications of assuming energy cells of magnitude $\varepsilon=h\nu$, namely that *exchange of radiation energy between the field and an atom does occur only in discrete quanta*.

This quantum postulate means that any exchange process is indivisible, i.e. no further analysis in constituent steps is possible, and that this is true for the states of the absorber as well as the emitter. Despite other disagreements, both Einstein and Bohr had this conception, as can be seen from the following two quotations. First Bohr (1928, p. 581):

... its [quantum theory] essence may be expressed in the so-called quantum postulate, which attributes to any atomic process an essential discontinuity, or rather individuality, completely foreign to classical theories and symbolised by Planck's quantum of action.

Einstein (1965, p. 373) writes:

It should be strongly emphasised that according to our conception the quantity of light emitted under conditions of low illumination (other conditions remaining constant) must be proportional to the strength of the incident light, since each incident energy quantum will cause an elementary process of the postulated kind, independently of the action of other incident energy quanta. In particular, there will be no lower limit for the intensity of incident light necessary to excite the fluorescent effect.

Thus we see that quantisation actually consists of two closely related, but distinct ideas.

- (1) Every individual exchange process involves one object giving away energy and one object taking up energy and this exchange is independent of every other process. It is this aspect Bohr refers to when he writes that 'the quantum postulate attributes to any atomic process an essential individuality'. The individual exchanges between atoms and field are all 'event atoms'. There may occur partitioning of emitted energy into smaller parts, but that requires that the emitted photon first is absorbed by some piece of matter. This is the

1 See Kuhn (1987), the afterword gives a very good summary of the issue.

function of so-called ‘down converters’.

- (2) The exchange process cannot be analysed further as a sequence of continuous changes. It is a discontinuous quantum jump either from the state (in quantum field notation) $|\text{ground state atom} + \text{excited field}\rangle$ to $|\text{excited atom} + \text{de-excited field}\rangle$ or vice versa. No intermediate states are possible. This does not necessarily mean that the change between the states is instantaneous; rather it means that we cannot *observe* any intermediate states, no matter how short the interval is between two successive measurements. This in turn implies that if we observe it repeatedly at shorter and shorter time intervals, the probability for change of state of the observed system will approach zero; hence repeated observation within short time intervals will lock the system into the initial state, a phenomenon called the ‘quantum zeno effect’.

During the first decade of the 20th century it became clear that quantisation can account for a number of facts about absorption and emission of electromagnetic energy, such as the following.

- (1) If light is shed on a surface of a metal, electrons will be ejected from the metal if and only if the frequency of light is above a certain limit, typical for that metal.
- (2) The maximum possible kinetic energy for photoelectrons is determined by the energy of each photon, but independent of the intensity of the light.
- (3) If a piece of fluorescent material is illuminated, the frequency of the outgoing light is never higher than the frequency of the incoming light; $\nu_{\text{in}} > \nu_{\text{out}}$.
- (4) A correct frequency spectrum in black-body radiation cannot be derived if continuous exchange is assumed.

5.3 Quantisation in other processes

The considerations above concern only quantisation of interactions between matter and electromagnetic fields and the reader might ask about the explanation of quantisation of other interactions, for example momentum and energy exchange when two particles collide. Such an interaction is usually described in non-relativistic quantum mechanics in terms of a direct coupling of two quantum systems referring to these particles. The interaction between these systems is described by the time evolution of the wave function and this time evolution is continuous, reversible and deterministic. Quantisation is assumed to occur only when a measurement is performed and the measurement operation is represented as the application of a Hermitian operator representing an observable to the wave function. The purpose of this mathematical representation is to calculate measurable quantities in the simplest possible way. However, this simplicity obscures the real nature of the process. I believe that quantisation is a feature of the interaction, irrespective of any measurement. One reason for this belief is that, on closer reflection, this interaction, just as every exchange of energy or momentum between two objects, must be an exchange of an exchange particle such as a photon. This is so because every interaction between two objects is transmitted by one of the four fundamental

forces in nature, namely gravitation, electromagnetism, weak or strong force, and all four forms of interactions are transmitted by quanta (or so it is generally believed), which in the case of electromagnetism are the photons. Hence, the analysis of quantisation of interaction in the previous section applies to all forms of interaction. The conclusion is that not only processes which we ordinarily describe as involving radiation are subject to quantisation, but all processes involving exchange of energy, without exception.

Quantisation of interaction was characterised by Bohr as an individual event, meaning that it can be described without taking into account emission- or absorption-processes in other atoms. But what about the direct interaction between two atoms, one emitting radiation, the other absorbing it? Are not these two events strongly dependent on each other?

It is clear that if the two interacting objects are close enough, the emission of a photon from one of the objects and the absorption of this photon by the other object could be seen, from a practical point of view, as one single process, because the time needed for the photon to travel between the two objects safely can be neglected. But this statement does not conflict with the statement that every exchange process is independent of every other process; an emitted photon must not necessarily be absorbed by a neighbouring object, although the probability that a neighbouring atom absorbs the photon is much higher than the probability for absorption by more distant objects. That two elementary processes really are two individual events should not be interpreted as that they are statistically independent. Each interaction event is logically independent of every other, but not statistically independent. By logically independent I mean that each interaction event can be fully characterised without describing other interactions.

In this perspective it is also clear that the interaction process is in one sense reversible; a photon released from a particle can just as easily return to that particle as it can be absorbed by another neighbouring particle. However, it is not reversible in the absolute sense that it is possible to arrange the situation such that the probability for re-absorption is equal to one. More of this in Chapter 6.

5.4 Quantisation of interaction between waves explains the particle aspect of waves

Assume, for the sake of argument, that exchange of energy and other conserved quantities, such as momentum, were *not quantised*. Assume further that, for example, electrons are wave packets with finite extension, which entails that their momenta to a certain extent are indefinite. It should be pointed out that each such wave packet is, by assumption, one and only one electron, a definite portion of charge. This wave packet could now, assuming continuous exchange of momentum between it and, for example, a measuring device be attributed a definite momentum, namely the weighted average of the momenta of the component waves in the wave packet. A successful measurement of the momentum ought to give us precisely this value. (The fact is that we will not get the expectation value but one of the eigenvalues

– if they differ – in a momentum measurement, but we are here considering the consequences of assuming that interaction is not quantised.)

Continuing this counterfactual continuity assumption, let us consider the details of a measurement interaction with such an electron wave packet. Assuming that an electron has travelled a considerable distance from source to target, the wave packet has undergone dispersion and is spread out over a considerable volume. Hence, the measurement of the momenta of the component waves requires many detectors and the electron must be in contact with several detectors at slightly different places. This is so because, as will be argued in Chapter 6, any measurement is in the final analysis a position measurement and the purpose of the preparation of the measurement set up is to correlate different detector positions with different possible values of the measured observable. We have thus the following situation: the electron is in contact with a number of detectors and, by assumption, interaction between objects is not quantised. Such a situation allows the exchange of energy and momentum in indefinitely small portions and two or more exchange processes with the electron can take place simultaneously. Hence, if the detectors are 100% effective, quite a number of them will be triggered. Then a measurement of the momentum of an electron, i.e. a wave packet, would give as a result not one single value but several ones, a distribution of momenta reflecting the distribution of the wave vectors in the wave packet. However, such a situation is never observed; we know that in directing one unit of charge, i.e. one electron, towards a detector system, it cannot in any circumstance give more than one response at each moment of time. Hence, at least one of the assumptions above must be wrong. Almost everybody has concluded that the electron must be a point particle, not a wave. However, we know that another assumption is false, namely the assumption that interaction is continuous, and that will suffice as an argument for the point-like character of the electron in this interaction. Even if interacting objects are vastly distributed in space they cannot interact with more than one detector, due to quantisation of interaction. In short, if interaction between wave-like phenomena is quantised, then observable results are just as if the interacting objects were point particles. We need not assume that electrons are point-like objects all the time.

So far the arguments have been given solely in physical terms; there has been no mention of mathematical representations such as wave functions or operators. However, if the wave ontology and the ensuing explanation of particle appearance of these waves is accepted, we have an obvious way of attacking the measurement problem. The thought that the collapse of the wave function reflects a real collapse of a real wave is imminent. This idea will be thoroughly discussed in Chapter 6.

What does it mean to claim that a conserved quantity such as charge is distributed in space? Such properties are primarily attributed to objects, not to parts of space. But objects with no properties whatsoever are inconceivable. I accept the view that an object is a bundle of properties, and if stripped of all its properties there is nothing left. (The alternative theory that an object is made up of a substance and a number of properties is in my view incoherent.) Hence, saying that conserved properties are distributed in space means that the object made up of these portions of quantitative properties is distributed in space.

Now it is time to consider some well-known physical experiments purporting to show that the constituents of matter are particles. I will show that a wave ontology can account for these experiments just as well.

5.5 Some examples of particle behaviour

5.5.1 The Compton effect

The prime example purported to show the particle character of radiation is perhaps the Compton effect. The particle analysis goes as follows. A photon with energy $E = h\nu$ and momentum $p = hk$ is incident upon an electron at rest in its own rest frame. In the collision the electron receives some energy and momentum from the photon. After the collision the photon has energy $E' = h\nu'$ and momentum $p' = hk'$. Because of conservation of energy and momentum the electron has energy $E_e = E - E'$ and momentum $p_e = p - p'$.

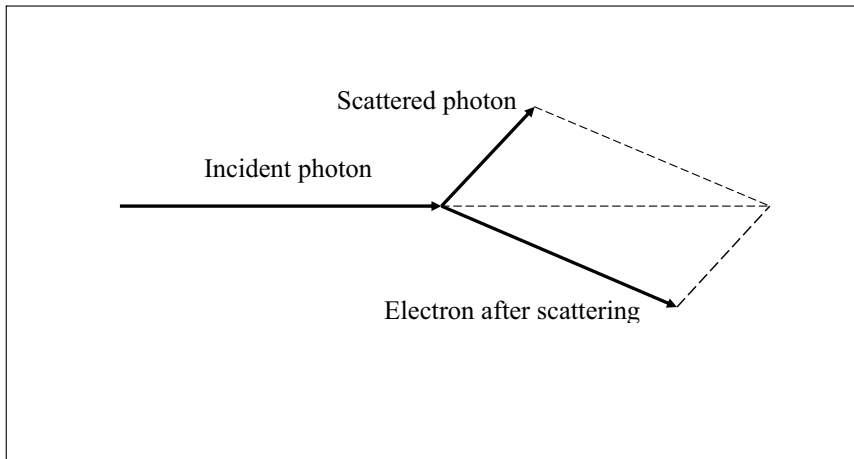


Figure 5.1 Compton scattering

Figure 5.1 exhibits a paradigm example of an inelastic collision between two particles. It is, however, possible to describe the process as a plane wave (rather than a wave packet) being reflected by a scatterer. In the paper ‘*Über den Comptoneffekt*’ Schrödinger showed that this collision can be just as well accounted for by assuming it is light waves scattered against a lattice of standing electron waves, as by assuming it is electron waves scattered against a lattice made up of standing light waves. The momentum vectors can be interpreted as representing the direction of propagation of plane waves and energy and momentum conservation is still valid. He concludes:

Die Richtungs- und Frequenzgesetze des Comptoneffektes sind vollkommen gleichbedeutend mit der Aussage, dass das beteiligte Lichtwellenpaar und das beteiligte

Ψ -wellenpaar zu einer und derselben “Netzebenenschar” in der (auf bewegten Kristall verallgemeinerten) Bragg’schen Beziehung für Reflexion erster Ordnung stehen; wobei jene gedachte Netzebenenschar von vornherein beliebige Stellung, beliebigen Ebenenabstand und beliebige translatorische Unterlichtgeschwindigkeit haben kann. (Schrödinger, 1926a, p.263) [‘The laws for direction and frequency in the Compton effect has the same content as the claim that the participating light waves and Ψ -waves would be first order Bragg reflected in one and the same grating; and there are no restrictions on the choice of place, slit separation and subluminal velocity.’ My translation.]

Thus, the Compton effect is neutral as regards the nature of matter and light; both views are compatible with the outcome of this experiment.

5.5.2 Particle tracks in cloud chambers

We have all seen numerous pictures of particle tracks in scintillation detectors, cloud chambers etc. They give us a very vivid picture of a particle travelling through the chamber along a path; for example, an electron in a magnetic field makes a conspicuous spiral track as its velocity is slowing down. This fits perfectly well into the classical electromagnetic theory of point charges. But the perception of a point charge track is misleading, for it only shows that the electron *interacts* with the surroundings, i.e. an over-saturated vapour, at well-defined places at each moment of time. The track is in fact a series of remnants from position measurements; the electron has repeatedly given away energy quanta to the vapour, resulting in a contiguous sequence of bubbles or droplets showing the path of the electron. Hence, this experiment does not contradict the notion that the electron is a wave between moments of interaction.

5.6 Calculations of the size of the electron in collisions

An argument in support of the particle nature of electrons is given by Malcolm MacGregor in his book *The Enigmatic Electron*. He considers the experimental outcomes of collisions between charged particles such as electrons. An electron colliding with another electron will be deflected and the distribution of the angular deviations can be measured. Knowing the kinematics of the collision, one can calculate the force required to produce a certain scattering angle. The results can then be compared with the predictions made using Coulomb’s law and assuming that electrons are point charges. If the force is weaker than expected, the proportion of large angle deviations will be smaller than expected. And, as weakening of the Coulomb force would occur if the particles penetrated into each other’s interior, one can ascertain to what extent the inverse square law is valid by observing the angular distribution. Hence, one can estimate the size of the electron by calculating the limit for the correctness of Coulomb’s law. Thus, MacGregor (1992, p. 48) argues:

If the force is weaker than expected, then the number of large angle scatterings will be smaller than expected. Hence the way to discover if the electron-electron scattering process has caused the two particles to interpenetrate into the interior of each other’s

charge distribution, and hence weaken the Coulomb force below the value indicated in Eq. (7.14) [i.e. $F = eq/r^2$], is to look for a fall off of the wide angle scattering.

The result is that the particle theory and experiment accord down to the limit of measurability, presently 10^{-18} m. MacGregor is very impressed:

We know that the $1/r^2$ force which operates in electron–electron and electron–positron scattering is strictly electromagnetic and is maintained down to distances of less than 10^{-16} cm. *This is one of the most decisive experimental results in particle physics!* (MacGregor, 1992, p. 50, italicised by MacGregor).

There is, however, a serious flaw in MacGregor's argument for the conclusion that this shows the size of the electron charge. (He thinks that the electron matter is bigger than its charged volume.) In collision experiments the colliding objects exchange momentum and energy. The exchange of momentum and energy is quantised: it is a discontinuous one-step process. The two electrons do *not* exchange momentum piecewise and continuously, as would be the case if they were classical objects whose parts continuously interact when they penetrate into each other. Both the colliding objects behave as indivisible units when exchanging momentum. Hence, even if the colliding objects are not rigid bodies but penetrable clouds of charge and matter, they will exchange momentum and energy in a one-step process. It follows that, in this exchange process, the interacting objects, whatever their constitution, behave as billiard balls, as rigid spheres. It further follows that if there is no limit for the penetration of them into each other there is no lower limit for the impact parameter, the centre-of-mass distance at which the interaction occurs; it can even be zero, in which case the two objects are totally overlapping. Hence, the argument from the angular distribution shows *not* that the colliding objects always are particles with negligible size, since that conclusion would follow only if interaction was not quantised. Hence, the exchange process does not tell us anything about the size of the object just before or after the interaction.

The concept of the size of an electron is not applicable to electrons during interactions involving exchange of conserved quantities. This is so because in order to be able to talk about the size of an object we must presuppose that it occupies a well-defined portion of space and that is not the case. Moreover, the size of an electron (interpreted as the bulk of the wave packet) in other circumstances is not an invariant magnitude; it depends on the circumstances, for example the potential in which the electron is moving. Hence, different methods for calculating the size of the electron will result in diverging results, and that is precisely what MacGregor found in his book. Finally, the boundaries of the charged cloud that makes up the electron are vague and so is the notion of its size. In short, the concept of the size of the electron makes no sense.

Chapter 6

The Measurement Problem

6.1 Introduction

The measurement problem is the most central and persistent issue in the interpretation of quantum mechanics. It has been described in many ways, the most eloquent being that given by Wigner in the following way: quantum theory says that quantum systems evolve *linearly, deterministically, continuously and reversibly* according to the Schrödinger equation, *except* when we perform a measurement on the system. When measuring the value of an observable, whose wave function is not an eigenfunction to the operator representing the measured observable, the measurement is *non-linear, discontinuous, non-deterministic and irreversible*. The superposition state changes, as a result of the measurement, into one of its components. In other words, the measurement induces a collapse of the wave function (Wigner, 1983). Why so? What is so special about measurements? After all, measurements are, disregarding our cognitive use of them, ordinary physical interactions and the Schrödinger equation offers a precise theoretical description of all physical interactions. But then, measurements ought to be continuous, deterministic and reversible, which they are not.

Quite a number of proposals have been made to solve the measurement problem, all of which fall into one of three categories. Common to proposals in the first category is the view that the collapse of the wave function during the interaction is a mere change of our state of knowledge. The argument for this view is twofold: (i) the result of the measurement will finally become knowledge in the mind of the observer, and (ii) the measurement process has no physical properties differing from other interactions, hence the only relevant property of a measurement is its resulting in cognition. Assuming that the knowledge process in the mind of the observer cannot affect a physical process in the laboratory, it follows that the change of the state (the described state in quantum theory) is merely a change of knowledge about the real system. It follows that the wave function describes what is known about a system and not its physical state *per se*.

In the second category of proposals it is assumed that the collapse is a mere approximation of the real evolution. Accordingly no collapse occurs; the superposition characterising the state before measurement is a correct description even after the measurement, but we cannot observe the interference effects in a real measurement because the interference terms become suppressed during the measurement interaction. It merely *seems* as if all but one component have disappeared. Bell (1987a) has labelled this view FAPP (For All Practical Purposes).

A third, more radical stance, is to assume that collapses occur all the time, not just in measurement situations, and for some reason we do not observe any effects

of collapses not being measurements. In this view the problem is not so much to find the special features of measurements but to explain why superpositions collapse at all.

It is worth noting that the original Copenhagen interpretation may fall outside this categorisation. According to Bohr, we must sharply distinguish between the macro and the micro domain and that makes stating this very problem impossible; a measurement must be described in classical terms, whereas the evolution of quantum states can only be made in terms of concepts applicable to the micro domain. Describing the measurement interaction as a collapse of the superposition is an illegitimate mixing of micro and macro concepts. However, contemporary adherents to the Copenhagen view claim that the wave function is a mere expression of our knowledge about the system, a view clearly belonging to the first category above.

The de Broglie–Bohm pilot wave interpretation can also, albeit in a quite different way, be said to belong to this first category, for it says that the wave function describes a pilot wave which guides the particle, whereas the measurement reveals the real position of it. Thus, the collapse is not a real event but a change of the observer's knowledge. Similarly, the many-worlds interpretation must be sorted into this category because it says that no collapse occurs, it is just that we happen to live in one branch of the world and thus we observe only one of the possible values of the measurement, that which happens to be actual in our world. The key notion for the many-worlds thinkers is thus 'branching', which replaces the notion of collapse. It is an open question whether branching should be interpreted as a real change in the world or a mere change in the state of our minds.

Adherents to the decoherence programme say that the collapse is a rapid diminishing of interference terms in the superposition, which for all practical purposes looks like a sudden and complete collapse. Thus it is a FAPP solution, so long as we are not given any account of why interference terms diminish. There are of course more ideas on the market, but those mentioned above are the most well known at present.

Personally I subscribe to the third category of views: collapses occur all the time. My argument is this: as a realist I hold that scientific theories purport to describe the real world; saying that quantum mechanics describes our knowledge about an aspect of the physical world amounts to a form of idealism, hence the first option is ruled out. My reason for rejecting the views in the second category, i.e. the FAPP solutions, is that they do not address the real problem. For what we need is not primarily a theory which fits the facts, that we already have, but an explanation of a perplexing feature. We need conceptual analysis and clear distinctions, which FAPP solutions fail to give. Thus, by default I have landed on the third position, the one that says that collapses should be taken as real events in the real world. This view can be backed by the observation that measurement interactions do not in any fundamental way differ from other physical interactions as regards their physical characteristics. Few philosophers would disagree with this statement, but my conclusion is somewhat different from that of the mainstream. Most people think that all state evolutions ought to be governed by the linear Schrödinger equation, whereas I believe that Schrödinger evolution is only appropriate when the system is isolated. (This stance needs to be backed by a discussion of the two notions 'system' and 'isolated' which

will appear later in this chapter.) I believe that the linear, deterministic, unitary and reversible evolution is broken whenever exchange of quanta occurs, irrespective of whether a measurement is performed. There seems to be no good argument for saying that collapses occur *only* when we perform measurements, contrary to what is commonly assumed. Thus, my guess is that collapses are normal physical events occurring repeatedly, whether or not we perform any measurements. The measurement problem thus transforms to the question: what are the necessary and sufficient conditions for a collapse to occur? Being measurements of the *first kind* (to be defined in the next section) is a sufficient criterion for collapse, but not necessary one.

Our first task is to analyse in some detail such measurement processes, which will enable us to give a criterion for collapses in purely physical terms. After that is done the concept of collapse will be discussed and it will be shown that a collapse is a result of the random behaviour of systems in certain interactions.

6.2 Measurements

To begin with there are a number of distinctions that will be useful for our purpose in this chapter. First we have von Neumann's distinction between *measurements of the first* and *of the second kind*. The difference between the two kinds of measurement depends on whether the wave function is an eigenfunction to the applied operator or not. If the wave function is not an eigenfunction to the operator, we have a measurement of the first kind. During such a process the wave function changes into one of its components, i.e. it collapses. A measurement of the second kind is a measurement where the wave function is an eigenfunction to the applied operator, hence no collapse occurs and the result (the *eigenvalue*) can be predicted with complete certainty. Only measurements of the first kind are of interest in this context since measurements of the second kind do not cause any conceptual difficulties.-

Secondly, we could distinguish between *measurements in the broad sense* and *collection of data*. What this distinction amounts to can be explained by an example.

Suppose that we want to measure the energy of particles emitted during the radioactive decay of a certain type of nuclei. If the emitted particles are electrons, the measurement procedure is called β -spectroscopy. Roughly, the idea is to utilise the Lorentz force on charged particles moving through a magnetic field and measure the deflection caused by the magnetic force. Since the force is proportional to the velocity, it is possible to calculate their velocity and hence their kinetic energy. It is obvious that the measurement of the kinetic energy of these particles is indirect; primarily, we measure the *position* of the particles after they have passed the magnetic field. The primary data collected in one such experiment are either the counts in a movable detector as the function of position, or spots on a film covering the whole range of possible positions for the deflected electrons. Hence, conceptually it is a large gap between such reports about position registrations and the 'observed' and reported fact, i.e. the value of the measured observable, in this case energy.

This point about scientific practice is quite general, as has been discussed at length by several authors. Observations of the values of many quantities are most often indirect, i.e. conclusions drawn using elaborate mathematics from position registrations. The reason for this practice of calling the conclusion of a rather elaborate calculation an observation is that the theory behind this calculation (and the measurement procedure behind it) is uncontroversial, at least in comparison with the tested theory. The conclusion to be drawn is that the distinction theory/observation is a relative one; parts of science which once were treated as a highly theoretical matter and a matter of controversy become, in the course of time, universally accepted. Hence, what once was thought of as a theoretical statement (in this case about the energy of electrons) is nowadays seen as an observation report; it is commonly trusted as safe knowledge compared with the current theories. Thus, the border between what are called observation statements and what are called theoretical statements is constantly moving, normally towards the theoretical endpoint, thus increasing the category of observation reports. In short, the distinction between theory and experiment is mainly a pragmatic distinction. Most measurements proceed through a series of steps, where one signal is converted into another, the last nowadays being a digital signal for computer analysis.

A *measurement in the broad sense* is my name for what in scientific practice is called measurements, whereas *collection of data* is my name for recording the result of the final interaction between the measured object (or a secondary object) and the detector.

Clearly, in all measurements in the broad sense the final step is a data collection procedure. These data are of a very simple nature. Usually they can be described as ‘a detector at position p was triggered at time t' ’. Of course this final data collection is also a measurement, a measurement of the *position* of the object triggering the detector. Hence ultimately all detections are position measurements, a point that has been made by several authors, e.g. Margenau (1958, 23 ff), Cartwright (1980) and Bell (1987b). Bell (1987b, p. 166) expressed the point in this way:

In physics the only observations we must consider are position observations, if only the positions of instrument pointers. It is a great merit of the de Broglie-Bohm picture to force us to consider this fact. If you make axioms, rather than definitions and theorems, about the “measurement” of anything else, then you commit redundancy and risk inconsistency.

Thirdly, a relevant distinction is that between *measurements on ensembles* and *measurements on single objects*. When comparing quantum mechanical predictions with the result of measurements we are most often interested only in ensemble properties, such as probability distributions or expectation values. But sometimes the focus of interest is the value of an observable attributable to one single object. The question is whether it is necessary to assume that the wave function collapses if we confine ourselves to statistical properties of ensembles. According to some, the measurement problem does occur in both situations, but in measurements of ensemble parameters the collapse is more or less hidden in the assumptions that underlie the rules of calculation. Others claim that the collapse does not need be assumed at all on an ensemble interpretation of quantum mechanics. The measurement problem

comes to the fore only when we consider the value of an observable attributable to one single object. We need not decide on this issue for two reasons: first, there is agreement about the necessity to take the collapse of the wave function into account when we consider a single object; second, the option of confining the interpretation to ensembles is not open to the realist who wants to know what properties *individual* objects have.

6.2.1 Measurements as attributions of numerical values to objects

Luce & Suppes (1975) characterise measurement theory in the following way:

Measurement theory studies the practice of associating numbers with objects and empirical phenomena. It attempts to understand which qualitative relationships lead to numerical assignments that reflect the structure of these relationships, to account for the ways in which different measures relate to one another, and to study the problems of error in the measurement process. (Encyclopedia Britannica, entry 'Measurement, theory of')

To *associate* a number with an object or an empirical phenomenon sounds like a cognitive process, something that happens in our minds. However, for our present purpose this is besides the point. The crucial point is that a measurement results in a stable state of a measurement device and that the set of possible states *can* be mapped one-to-one onto a set of numbers. This mapping is a cognitive, associative process, but it is of any interest only if different people arrive at the same observation statement due to the stability of the measurement device. If different people produce different observational statements, the measurement is simply a failure. Thus, in order to be fully objective, the outcome of the measurement, i.e. the final state of the measurement device, must be accessible for observation by more than one person. This implies that the final pointer state of the measurement device must be insensitive to human observations. We must be able to look at the measurement device without changing its pointer state. It sounds quite trivial, but the point is crucial in order to block von Neumann's argument purporting to show that the collapse is a cognitive process.

A measurement is thus a physical interaction, i.e. an exchange of a physical quantity between the measured object and a measurement device. And, as all interactions are mediated by one of the four fundamental forces of nature, we can conclude that all knowledge about an object is obtained through either gravitational, electromagnetic, weak or strong interaction, which, as far as we know, are the only types of physical interactions. For example, the interaction between a measurement device and our eyes is an electromagnetic interaction. Hence, the measurement process must belong to one of these four (or three, or two, if we take into account the unification of forces) types of interactions and as all these interactions are essentially exchanges of a quantum of the appropriate type (graviton, photon, vector boson or gluon), all measurements are exchange of quanta. This line of reasoning sounds plausible but there is some room left for doubt. The weakest step is the assumption that all interactions must be exchanges of quanta. No doubt, in most measurements quanta of the appropriate type are exchanged, mostly photons, but is that absolutely

necessary? Recently this assumption has been discussed and rejected by some authors, e.g. by Elizur & Vaidman (1993). A similar idea – measurements do not affect the object in any way – was discussed by Renninger (1960).

Elizur & Vaidman claim that it is possible to decide the value of an observable without interacting in any way with the object in question, i.e. without even exchanging one single photon with the object. They call this an *interaction-free measurement* and the question is: is this interaction-free measurement associated with a collapse of a superposition or not? If the answer is yes, the assumption that all measurements of the first kind are exchanges of quanta must be wrong. It will turn out that this is not so, which will be shown in Section 6.11; in Section 6.12 I will discuss and reject Renninger's conclusion as well.

Sometimes the word 'measurement' is used in spite of the fact that nobody was around to observe the outcome of the event called 'measurement'. This usage implies that cognition has nothing to do with measurements. But why, then, talk about measurements at all? The answer is that in such cases it is assumed that the evolution of the system is not unitary but must be described by a projection operator, which is to say that a collapse has occurred. But according to the received wisdom, i.e. the projection postulate, collapses can only occur during measurements, hence a measurement must have occurred. That in turn implies that the environment must have 'measured' the state of the system. An example of this line of reasoning is given through the quotation below. The paper, from which the quotation is taken, analyses the interactions between an electron passing through a metal in which the electron's wave function is present in two macroscopically distinct channels. Essentially, the passing electron interacts with free electrons in the metal, conceived as an electron bath. The question is then under what conditions does the wave function collapse during such a passage:

As emphasised above, the interaction of the environment with the interfering partial wave changes the state of the environment, so that the environment acquires information on the path taken by the interfering particle. ...But due to the bath's continuous spectrum, this polarisation has to involve an excitation of the bath, and this excitation does not disappear when the interaction of the bath with the interfering electron is over. (Stern *et al.*, 1991, p. 103)

Using the word 'polarisation', the author refers to the fact that the electron bath must change into one of two possible states, representing the electron having passed in one or the other channel. If it is taken for granted, as in the quotation above, that information in principle can be extracted and used to decide whether the particle took this or that path, it follows that the wave function for the particle must have collapsed into one such path, independently of whether any knowledge is obtained or not.

Similarly, Zurek uses the word 'information' instead of 'measurement' when discussing the physics of the collapse:

The natural sciences were built on a tacit assumption: Information about a physical system can be acquired without influencing the system's state. Until recently, information was regarded as unphysical, a mere record of the tangible material universe, existing beyond

and essentially decoupled from the domain governed by the laws of physics. This view is no longer tenable. ...Conscious observers have lost their monopoly on acquiring and storing information. The environment can also monitor a system, and the only difference from a man-made apparatus is that the records maintained by the environment are nearly impossible to decipher. Nevertheless, such monitoring causes decoherence... (Zurek, 1991, p. 44)

As decoherence is supposed to account for the collapse of a superposition, Zurek believes that the collapse is caused by the environment. Thus, both these authors say that a superposition collapses into one its components as a result of an interaction that increases information in the environment independently of whether someone actually observes the state of the environment or not. The crucial point in this context is whether the word 'information' should be interpreted intensionally or not. If it is interpreted intensionally it belongs to the same category of words as 'knowledge' and 'belief', i.e. having to do with mental states and human cognition and it is pretty clear that this is not Zurek's usage. Rather, he appears to adhere to the extensional reading of the word 'information' which means that cognitive allusions are beside the point and information has to do with physical states only. This is precisely the sense in which Shannon & Weaver (1949) use the word 'information' in their mathematical information theory. Roughly, information is a relation between a certain state description and the number of real states that are possible given that description. If only one state is possible, given a certain description, the amount of information in the description is maximal and complete. Information is thus essentially the same as negative entropy. Given this reading, information can be said to be stored in the environment if the state of the environment is correlated in such a way that a description of the state of the environment makes it possible to infer the state of the object of interest.

This could be seen as a realistic interpretation of the projection postulate: any process in which information is increased and stored in a detector system is a measurement, independently of whether anyone is aware of the result. Some participants in the discussion even explicitly write that the environment *measures* the state of the system. It is pretty clear that they intend to say that when the quantum system undergoes a collapse to a certain state among a number of alternatives, there must be something external, the environment, which undergoes a correlated change. But it is rather misleading to say that the environment *measures* the state of the system, because it suggests that the result of the measurement, some sort of pointer state, *always* can be used as input for cognitive process resulting in knowledge. However, this is not so; there is no guarantee that the environment is stable enough to allow extraction of any result. Information may dissipate rapidly, thus making it practically impossible to gain any knowledge. The problem could be described as that in order to obtain information in the sense of knowledge from information in the sense of negative entropy, the pointer state must be relatively stable.

This way of talking about measurements shows that the projection postulate often is regarded as an implicit definition of the concept of (quantum) measurement. Often, the use of the term 'measurement' is such that a measurement of the first kind is performed if and only if the wave function has collapsed, irrespective of any other

requirements on measurements. For example, Zurek in the quotation above, says that information is stored in the environment whenever decoherence (i.e. collapse) obtains, even if it is practically impossible to obtain any knowledge about the state.

It can be concluded that the term ‘measurement’ is used in a number of different ways, making it difficult to state a set of general conditions covering all cases. It seems to me that scientists usually understand the word ‘measurement’ as a physical interaction, the outcome of which it is *possible* to observe. This leads me to propose the following definition of measurement:

Def. A measurement is a physical interaction between a measured object and a measurement device such that the final state of the measurement device is a possible object for objective human cognition.

The second clause in this definition needs some elaboration. A measurement result, i.e. the final state of the measurement device, must be publicly accessible, or at least it must be physically possible to make it publicly accessible, i.e. the outcome of the measurement must be observable by several different observers. That implies that the final state of the measurement apparatus must be stable during several observations made by several different observers. Furthermore, since observations are exchanges of photons, the pointer state of the measurement device must be insensitive to repeated exchange of photons with observers. Whether this is a mere practical point showing that we need to amplify and store the outcome in some type of record, or if this has severe implications for when a definite state of the measured object obtains, is a matter of dispute. I will return to this question at the end of this chapter.

6.2.2 Measurements are energy exchange processes

The measurement process is finished as soon as the measurement device has acquired a definite pointer state, whether observed or not. Some theorists describe the measurement process as a two-step process: first, the measurement device and the measured object are physically connected so that the superposition of the state of the measured object is correlated to a superposition of pointer states of the measurement device. Then, in a separate process, the superposition of the total system of object+measurement device collapses into one of its components. This description obscures essential facts. A better way of describing the measurement interaction is as follows.

The part of a measurement process properly called a physical interaction proceeds as follows. First, by a suitable preparation, different possible values of the measured observable are correlated to different possible *positions* of the measured object (or a secondary object). The second step is the determination of the position of the measured object (or a secondary one) which is done by producing a spot on a film, a position of a needle on a meter, the triggering of a detector at a certain place, etc.

The final step in this process of measurement is an electromagnetic interaction. The detector absorbs energy from the measured object by an electromagnetic interaction. This interaction is an exchange of at least one photon, the electromagnetic quantum. It is possible that the exchange goes the other way, i.e. that the detector

emits a photon which is absorbed by the detected object, but for present purposes this doesn't matter. Since we believe that all forces are quantised, the argument does not really depend on the interaction being electromagnetic.

Now the question is: does the collapse of the wave function take place in the first or the second stage of the measurement? (Remember that we concentrate on measurements of the first kind.) At first sight it might seem plausible that the collapse occurs in the first stage when the different values of the measured observable are correlated to different positions of the measured object – the second stage being a mere registration of the presence at a particular place of the object. That is, however, a mistaken view and the following example provide a convincing counter argument. Although it is a thought experiment, it is commonly regarded as possible to perform. Wigner, for example, is of the opinion that it can be trusted as an argument in reasoning about the interpretation of quantum mechanics.¹

Assume that a collimated beam of electrons (or some other spin half particles) is

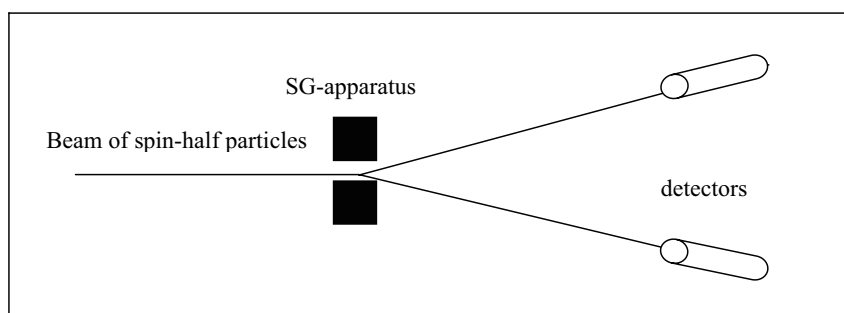


Figure 6.1 A measurement on spin-half particles

directed through a Stern–Gerlach apparatus, i.e. essentially an inhomogeneous magnetic field. (In the following I will neglect the fact that spin and the internal magnetic moment of the particles are not identical quantities.) If the electrons in the beam have been prepared to have identical spins perpendicular to the magnetic field in the Stern–Gerlach apparatus, the beam will split into two as shown in Figure 6.1. The common interpretation of this experiment is as follows: before the passage, the spin state of the electron is perpendicular to the magnetic field. Thus, its state can be described as a superposition of *spin up* and *spin down* in the direction defined by the magnetic field. When passing the field, each electron is thrown into a new state; those detected in the upper counter acquired a definite state *spin up* when passing

¹ Wigner, who has discussed this experiment writes: ‘Even though the experiment indicated would be difficult to perform, there is little doubt that the behaviour of particles and of their spins conforms to the equations of motion of quantum mechanics under conditions considered’ (Wigner, 1963; reprinted in Wheeler & Zurek, 1983, p. 331).

through the magnetic field and if it is detected in the lower counter it acquired the spin state *spin down*.

Now, insert a second Stern–Gerlach apparatus perpendicular to the first one. It will produce a new splitting. Each particle is now assumed to have *spin up* or *spin down* in the *x*-direction, defined by the direction of the magnetic field of the second apparatus.

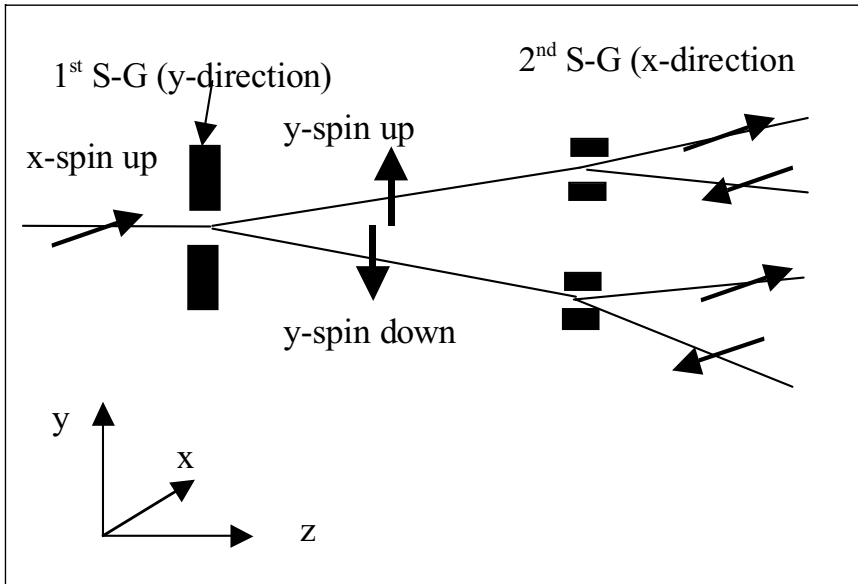


Figure 6.2 Particles with spin up in the *x*-directions pass two S-G magnets, the first oriented in the *y*-directions and the second in the *x*-direction

The situation is shown in Figure 6.2. Thus, the spin of each particle appears to change every time it passes through a magnetic field not parallel or anti-parallel to the spin. This analysis is however incompatible with the outcome of the following experiment. Assume as before a beam of particles being prepared to have spin up in, say, the *y*-direction entering a Stern–Gerlach apparatus oriented in the *z*-direction. The beam will split into two distinct parts and, according to the previous interpretation, we would say that the particles in the upper beam have spin up and those in the lower beam have spin down in the *z*-direction. If we now reunite these two beams (which can be done without affecting the spin states by an adiabatic process) and direct them into a second S-G magnet, oriented in the *y*-direction, we would expect a 50:50 distribution of spin up and spin down in this direction. But that is not the case; all particles will go the upper way, i.e. having spin up in the *y*-direction. Astonishingly, the original spin orientation is preserved during the experiment – see Figure 6.3.

Obviously, the interaction between particles and magnetic field in the first S-G apparatus is completely reversible. Hence, the wave function cannot have collapsed when passing the first S-G apparatus, since if the state of the particles is definitely spin up or spin down in the z -direction, it is impossible to return the particles back to their original y -spin-up state.

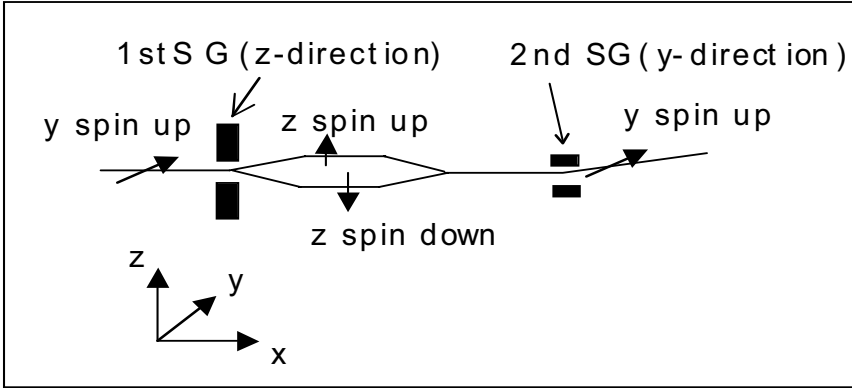


Figure 6.3 Reuniting two separated beams will result in the revival of the original spin state

Now let us look at the formalism. The prepared state is (taking spin up/down in the z -direction as base states):

$$\Psi_1 = \frac{1}{\sqrt{2}} \begin{pmatrix} i \\ -1 \end{pmatrix} \quad (6.1)$$

When passing into the first Stern–Gerlach apparatus this state is not changed. We may write the *same* state as a superposition of two states, *spin up* and *spin down*, in the z -direction:

$$\Psi_1 = \frac{i}{\sqrt{2}} \begin{pmatrix} 1 \\ 0 \end{pmatrix} + \frac{-1}{\sqrt{2}} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \quad (6.2)$$

However, this state description should not be interpreted as if the first component refers to the upper beam and the second to the lower beam. The complete wave function is a mathematical representation of the state of the *entire* wave made up of two spatially separated parts, the upper and the lower part, but the parts of the mathematical expression do not refer separately to spatial parts of the wave. Rejoining the two beams (which can be done adiabatically and without affecting spin because spin and linear momentum operators commute) and measuring the spin states of the particles in this beam confirms that the state is the same as the prepared state.

This analysis shows that the collapse cannot occur when the particles pass through the magnetic field. The evolution of the wave function up to the detectors can be described by a unitary operator. Since unitarity guarantees the existence of an inverse operator, the reversed evolution is possible. Hence the collapse must occur in the detection stage. A similar conclusion can be drawn regarding the previously discussed example of β -spectroscopy. The wave function in that case a spherical wave, does not collapse into one definite position before the detection of the electron is made.

Take a look once more at Figure 6.1. Since the wave function does not collapse when passing the S-G magnet, it must collapse when entering the counters. The physical object represented by the wave function must get into contact with both counters since both parts of the wave function still exists. Then the wave triggers only one of the detectors. This is the collapse; of two possible interactions, one is actualised.

By now it should be obvious that the collapse of the wave function is not only a collapse of an abstract mathematical superposition into one component, but also a physical contraction of an entity more or less spread out in space into one narrowly defined region in space. This way of seeing things amounts to accepting that the wave function describes a kind of real wave in physical space. At the first stage of the measurement the wave is split into two spatially separate parts by the inhomogeneous magnetic field, but still, these two partial waves make up one electron! Hence, we must accept that the electron, i.e. the electron wave, *arrives at both the detectors*. But a reaction cannot be induced in both detectors, since such a process would need two distinct objects.

Thus, the measurement (of the first kind) is a collapse of the wave function to one of its components and this mathematical collapse represents a spatial contraction of an entity spread out in space. This is in fact only a rehearsal of the argument, given in Chapter 4, showing that wave-like entities manifest themselves as particles, i.e. as localized entities, in interactions. As was proved in that chapter, the collapse is a logical consequence of the fact that an entity which travels in space in the form of a wave-like entity, i.e. spread out in space, interacts, if it interacts at all, as one individual object at one single place. Thus, the sufficient criterion for particle behaviour of a wave, namely irreversible exchange of energy with a macroscopic object, is also a sufficient criterion for collapse.

This view has an astonishing consequence. Suppose that we insert a screen into one of the paths before the detectors in a Stern–Gerlach electron experiment.

Obviously, then, a detector behind the screen cannot be triggered. It follows that the wave collapsed to one component when passing the screen; either the entire object was absorbed in the screen or the entire object continued in the upper path. Suppose the latter case obtains. That means that the object's momentum changed when passing the screen, since no exchange of momentum occurred in the S-G apparatus. The screen must be seen as 'pushing away' the part of the wave hitting the screen towards the other path. Since this is an exchange of momentum and energy, it must be possible in principle to detect a change in the state of the screen (but I doubt it is possible in practice).

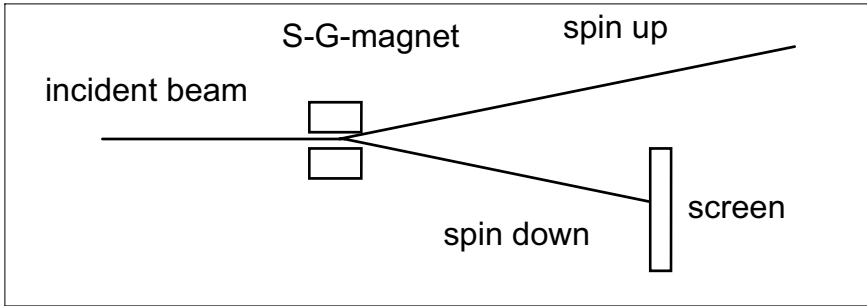


Figure 6.4 Preparation of a definite spin state

An opponent might say that this interpretation gives us more problems than solutions. This pushing away by the screen is a kind of non-local effect since the two partial waves are spatially separated and no force acts between them.

This kind of non-locality is an essential part of the notion of a physical collapse, i.e. a discontinuous contraction of an entity into a more narrow volume. It might further be argued that such an idea is so counterintuitive that we have not gained any explanatory value by assuming such processes. However, I think we have, for, as will be argued in the next section, the physical collapse is a consequence of quantisation of interaction and quantisation is in any case accepted as a real trait of the world. And furthermore, the non-local effect on the wave from the screen is but another aspect the general non-local trait of quantum theory which we must confront whatever interpretation we adopt. We do not get any additional problems.

Another problem with saying that all measurements are energy exchange processes arises in contexts where the energy of the system is not well defined. It could be claimed that in those cases it is not meaningful to talk about energy exchange. To take a well-known example, the spin- x -component of a spin half particle being in an eigenstate to the spin- z -operator does not commute with energy, hence a measurement of spin x cannot reveal its energy state.

This problem can be resolved if we note that the amount of energy exchange or the new energy state of the measured object need not be known, albeit it is a definite energy exchange. The point is that the detector must be in a stable state for some time, both before and after the measurement interaction, stable enough to be observed by several humans so that the objectivity of the result, i.e. the final state of the detector, can be ascertained. We may, however, be left ignorant of the energy state of the measured object, before and after, as well as the amount of energy exchanged in the measurement; still, energy must have been exchanged. The point is simply that the detector cannot react on the mere presence of the object; something must be exchanged, and this something must be a quantum - in the electromagnetic case, a photon.

6.3 Collapses

So far we have found that the collapse of the wave function during the measurement is a mathematical representation of a real collapse, a contraction of a wave-like entity spread out in space. But why does the wave collapse? Or better yet, why does the wave collapse in some situations and not in some others? In order to answer that question, let us begin by characterising the collapse in the following way. A collapse of a superposition to one of its components is a *random, discontinuous and irreversible* change of state, whereas ordinary state evolution of a quantum system is deterministic, continuous and reversible. Thus, what we need is an account of why *some* processes are random, discontinuous and irreversible and others not.

Collapses once seemed implausible and unphysical; for a long time the general outlook of many physicists were that all real physical events are deterministic (randomness is due to lack of information), continuous (all real events are describable by continuous and derivable functions) and reversible (which was interpreted to mean symmetric under the time reversal operation $t \rightarrow -t$). Hence, the collapse seemed for a long time to be an anomaly, an artefact of the theory, not reflecting the real nature of the world. The task seemed to be to explain away the collapse as an approximation or a mere change of information state on the part of the observer. (The ultimate reason for taking continuity, reversibility and determinism as fundamental traits of the natural world was, I think, the success of classical dynamics and classical electromagnetism which all had these ‘classical’ properties. In this respect, the theory of relativity marked no change.) Hence, it seemed reasonable to believe that the collapse of the wave function did not describe a real process. *A propos* randomness, Einstein once famously exclaimed ‘Der liebe Gott würfelt nicht’ (‘The good Lord doesn’t dice.’)

But our outlook has changed. Nowadays it appears to be a prejudice to assume that randomness, irreversibility and discontinuity are mere artefacts of the theory. The distinction between natural processes and mere artefacts of theory is a metaphysical one; it is the question of what kind of processes we view as unnatural, requiring explanation, and which are considered natural, not requiring explanation. This issue is never a purely scientific question, but often the result of extrapolations from common sense (another word for old prejudices) and since long well-established scientific achievements.

We can record several changes of metaphysical outlook in the history of science. One example is what kind of motion is natural: once philosophers wondered about the cause of uniform motion, but nowadays we take uniform motion as a natural state not requiring scientific explanation. A second example is the fate of the ether: once it seemed obvious that waves must be in some medium and hence there must be an ether for electromagnetic waves to be waving in. But we have been accustomed to the fact that electromagnetic waves make sense without any ether. In a similar move, I suggest that we should call into doubt the assumption that continuity, reversibility and determinism are natural and fundamental features of the physical world, whereas discontinuity, irreversibility and randomness needs explanation. For the unprejudiced mind, neither continuity nor discontinuity is more in need of explanation than the other feature; we cannot claim to have any a priori insight into the fundamental traits of nature.

Once we take a realistic attitude towards collapses, i.e. viewing the collapse of the wave function as a representation of a real physical collapse, there seems to be little reason to assume that collapses can happen only during measurements. If we do not perform any measurement on a quantum system, it can just as well collapse and we will never notice it. Nevertheless, we must confine the discussion of collapses to measurements in the rest of this chapter for the obvious reason that a collapse that is never observed is not very useful as an object for analysis. (But the resulting view should be an analysis of collapses irrespective of whether they are part of measurement processes or not.) So, when discussing the causes of collapses we need to investigate measurements of the first kind. Hence our question is what *physical* traits in the measurement process (but not uniquely restricted to processes we call measurements) are responsible for the collapse, i.e. for the discontinuous, non-deterministic and irreversible change? As usual, such a question leads to further questions like, ‘what exactly is meant by the terms “discontinuity”, “randomness” and “irreversibility”’.

Another aspect of the collapse is that it breaks the linear state evolution and one might wonder whether it would be possible to give a unified account of both linear and non-linear state evolution in one unified theoretical framework. Although this train of thought has some initial attractiveness it is a misconception of the situation, as I will argue in the next section. After that I will discuss discontinuity and randomness, and finally irreversibility, the latter being a consequence of the former.

6.3.1. The non-linearity of the collapse

The collapse of a superposition is, as pointed out in Chapter 3, an actualisation of potentialities: indefinite properties become definite in the process. But at the same time there is a de-actualisation of actual properties, turning definite properties into indefinite ones. Thus, it is misleading to say that a quantum measurement is a process in which we increase our knowledge about a constant state of the system. In a maximal measurement (i.e. a measurement after which the wave function for the system is completely known) on a pure state we know all there is to know about the state (the real physical state) before the measurement and similarly after the measurement. Consequently, entropy should be constant during a maximal measurement of the first kind. That is also the case, as was shown by Misra *et al.* (1979, 1983).

It is not difficult to show that no linear evolution operator can account for the collapse and the natural question is then: it is possible to give a physical description of the collapse in terms of a non-linear but continuous evolution?

It seems physically reasonable to regard the description:

$$\sum c_n \phi_n \rightarrow \phi_k \quad (6.3)$$

of the collapse as a mere phenomenological description. If so, it ought to be possible to set up a non-linear Schrödinger equation of the type

$$\frac{d\Psi}{dt} = H\Psi + R(\Psi) \quad (6.4)$$

where $R(\Psi)$ is a non-linear term that accounts for a very rapid change. Several authors have pondered upon this possibility, for example Barut and Shimony. Shimony (1986, p. 199) writes:

It is premature to pronounce judgement upon the program of nonlinear variants of the Schrödinger equation. Nevertheless, my personal view is pessimistic. In particular, there is great tension between desiderata (iii) [that the dynamical equation should be general and not restricted to measurement situations] and (v) [approximate agreement with standard linear dynamics in the case of micro systems].

The conflict lies in the fact that on the one hand the non-linearity must be large in order to account for such a rapid process as the collapse, but on the other hand it must be small in order to retain the approximate correctness of the linear Schrödinger equation. In fact, the correctness of the linear Schrödinger equation is established to a very high degree.

In my view we should give up the generality requirement. Unitary and linear evolution are, as has been argued in Chapter 3, no change of any *internal state*, just ways of re-describing one and the same internal state at different points in time. When we apply the Hamiltonian to a system and use Schrödinger's equation to calculate its evolution we do not get different internal states but one and the same state at different points of time. However, exchange of energy with other systems results in a change of internal state. Since exchange of energy (and other conserved quantities) is discretised, no linear and continuous function can describe that process correctly. In this perspective the demand for a unified account of dynamical change is wrong-headed.

6.4 Discontinuity

The discontinuity–continuity dichotomy belongs primarily to mathematical discourse. A function is continuous if, roughly, for any two values of the function there exists a third value between them. (In several dimensions this needs elaboration of the concept ‘between’.) The evolution of a physical system is continuous if we can represent its sequence of states by a continuous function. A discontinuous evolution means that we cannot *represent* its state evolution by continuous functions. It further implies that if we observe the state of the object within shorter and shorter time intervals we will never observe smaller and smaller changes; either we observe no change at all or one and the same amount of change. Assuming that the probability for change is a positive function of time, the probability for a change of state will diminish if we shorten the time interval between successive observations. This entails that the evolution of an observed system will slow down when observed during shorter and shorter time intervals and this is exactly what has been observed – a feature that is called the *quantum Zeno effect*. The argument can be illustrated by assuming that the probability for change is proportional to elapsed time. This means that there is a definite probability per unit time for a given change. Let us further assume, contrary to fact, that changes of state are continuous. If such a process is divided into consecutive steps, the probability for each step would then be higher

than the probability for the total change. For example, dividing a given change of state with probability p in n consecutive steps, the probability for each step must be $p^{1/n}$. This number is closer to 1 than p and if we take the limit of an infinite number of steps, i.e. continuity, the probability for a minute change of state would approach unity asymptotically. Hence, if we observe the state within shorter and shorter time intervals we would observe smaller and smaller changes and no total slowdown would result. But we do observe such a slowing down, which shows that the probability for a change of state approaches zero if the time intervals are made shorter. The reason is that quantisation of interaction coupled with probability for change is a positive function of time: the entire change is either not performed at all or the entire change is taken in one single step in the time allowed. It follows that changes of internal state are discontinuous, they can occur only in discrete steps.

It might be tempting to ask for the underlying reason for quantisation, but that is outside the scope of this book. We are, in this context, aiming at an interpretation, an explanation within the settings of present quantum mechanics. In this theory, quantisation of interaction is a basic postulate. To ask for the reason for quantisation is to ask a physical question, whose proper answer requires new physics. Philosophical analysis cannot provide that.

6.5 Randomness

A measurement of the first kind is a chancy event; there is more than one alternative outcome and the theory is silent as regards what precise outcome will occur in a particular case. Is this silence due to the theory being incomplete or are natural events ultimately random? This is a metaphysical question that transcends all possible knowledge, for even if we, without knowing it, happen to invent a complete and correct theory about nature – a theory which, like quantum mechanics, is non-deterministic – we can never prove that it is in fact a complete theory; hence, we can never be sure that this indeterminism is a genuine trait of the world. But there are inductive reasons to believe that quantum mechanics is complete (we have for decades looked in vain for ways of adding deterministic hidden variables) and hence that genuinely random events do occur. We have a considerable amount of evidence for the empirical correctness of quantum mechanics and none whatsoever against.

The probability for a specific result of a measurement is given by the so-called overlap integral. Suppose the final possible states after the measurement are given by a set of functions $\{\varphi_i\}$ and the initial state is given by the function Ψ . If these functions are normalised, the probability p_i for the result a_i corresponding to the wave function φ_i is given by $p_i = \int \Psi^* \varphi_i dx$. This integral expresses the amount of overlap of these two functions, hence the name ‘overlap integral’. The overlap can take any value between 0 and 1. Hence it is impossible that the overlap represents the actual amount of exchange of any quantity since exchange processes are quantised.

In fact, there is a methodological argument for this identification of probabilities with overlap integrals; the general rule for the construction of quantum theory is the *principle of correspondence*. This principle says that in the limit, where quantisation of interaction can be neglected (i.e. in situations where Planck’s constant can be

approximated to be zero), quantum theory and classical theory should give the same predictions. If overlap integrals are identified with probabilities for individual exchange processes, the probability distribution calculated according to quantum rules will then equal the classical one.

The same argument can be given a little more formally as follows. Consider two non-commuting operators $\mathbf{O}_1$ and $\mathbf{O}_2$, such that the prepared state is an eigenstate to $\mathbf{O}_1$, and $\mathbf{O}_2$ represents the measured observable. This means that more than one outcome of the measurement is possible and we can calculate the probabilities for the possible outcomes, none of which equals one. Now we have a correspondence between quantum operators and classical variables given by²

$$[\mathbf{O}_1, \mathbf{O}_2] = i\hbar \{O_1, O_2\}$$

where ‘ $\{O_1, O_2\}$ ’ is the classical Poisson bracket, for the variables O_1 and O_2 , whose quantum counterparts are the operators $\mathbf{O}_1$ and $\mathbf{O}_2$. If we make the thought experiment of letting Planck’s constant go to zero, we will get the result that the commutator $[\mathbf{O}_1, \mathbf{O}_2]$ also becomes zero, which means that we can simultaneously decide with certainty both the values of O_1 and O_2 . Hence, randomness is removed. But assuming that Planck’s constant is zero in a thought experiment is just another way of expressing the thought that interaction is not quantised. This is so because the values of all quantised observables are expressed as multiples of Planck’s constant; Planck’s constant is the spacing unit in the value spectrum of quantised observables. When Planck’s constant is zero, the value spectrum is continuous. Hence, quantisation of interaction is a sufficient condition for randomness.

6.6 Irreversibility

What do we usually mean by saying that the measurement process is irreversible? Certainly, we do not mean that it is completely impossible to return the object+measuring apparatus to their original states, since there are easily recognized examples in which it is possible and still the interaction is said to be irreversible. Consider, as an example, a measurement of the x -component of the spin on an ensemble of spin-half particles. Assume that the spin state of all the particles are $|z+\rangle$ (spin up in the z -direction) before the measurement. This means that we will get a 50/50 distribution of spin up/spin down in the x -direction. Let us now close the spin-up channel so that the remaining ensemble consists solely of particles in the state $|x-\rangle$. Assume then that we detect the presence of the particles by a measurement not disturbing their spin states, which is possible. Since the position is correlated with spin, the measurement results are thus spin down in the x -direction because the measurement reveals the state immediately after the measurement has been concluded.

Surely, spin-half particles will not return to their former state $|z+\rangle$ by themselves. However, it is possible to bring about conditions so as to increase the probability for

² This is a basic postulate of quantum mechanics. See for example, Dicke & Wittke (1960, p. 102).

returning to their former state. We can repeat the process with a new Stern–Gerlach apparatus oriented in the z -direction. Just as before, we close one channel, namely that corresponding to spin down, leaving only the particles in a definite spin up state. These remaining particles, 25% of the original number, have now returned to the initial state. Hence, if we say that the collapse of the wave packet during the first measurement is an irreversible process, it cannot be taken to mean that it is completely impossible to return the state of an individual system to its original state. In fact, the probability of such a return in this case is 25%. Irreversibility must mean something weaker.

The conclusion is that when the term ‘irreversibility’ is said to be a property of quantum measurements, or more generally of collapses, this does not exclude the possibility of bringing the measured object back to its original state. Neither does it mean that the measuring device has undergone a totally irreversible process. What does it mean then?

A plausible interpretation might be that a measurement process is said to be irreversible if it does not have an inverse operator. The argument would run something like this: unitary operators describe ordinary dynamical evolution, and unitary operators all have inverses. A dynamical evolution is hence reversible, at least in principle, because an operator describing the reversed process exists. In contrast, the collapse of the wave function during a measurement is described by a projection operator that lacks an inverse because it is non-unitarian. Therefore, measurement interactions are irreversible.

This train of thought confuses *physical states* with *state descriptions*. It is important to keep in mind that one and the same state can be given different descriptions. More precisely, one and the same physical state at one and the same point in time can be described using different coordinatisations of time. The evolution operator $\exp(-i\mathbf{H}t/\hbar)$ is used for transforming state descriptions differing with a time interval t . The minus sign signifies a change of the zero point to earlier times with the amount t . If we know the state of a system at a certain time $t_0 = 0$, i.e. $\Psi(0)$, we can calculate the state description at any time t . But this is only another description of *the same internal state*, for – as was shown in Chapter 3 – the evolution operator does not describe real internal changes of states but parameter transformations. The evolution operator can also be used for calculating states at different times, given a fixed coordinatisation of time. Since these two operations, changing the time coordinatisation and calculating states at different times using a fixed coordinatisation, are not distinguished in the formalism, these two operations must give the same result. This means that time evolution of an isolated system is no real change of state.

It follows from this, as we have seen in Chapter 3, that non-existence of an inverse does not exclude the possibility of reversing the state of a system to a former state. Thus, applying the inverse of the evolution operator is our method for pushing the zero point on the time scale forward in time; thus we get the *state description* at a previous time $-t$. All depends of course on the assumption that the system has not interacted by exchanging energy with its surroundings, but that assumption is represented mathematically by the time evolution of the system governed by

$$\exp\left(\frac{-i\mathbf{H}t}{\hbar}\right)$$

where $\mathbf{H}$ contains no interaction terms. But what does that tell us about reversibility, i.e. the possibility of forcing a system to return to a state that differs from its present one?

Reversibility is a modal concept. It means the *possibility* of reversing a system to its former states. Possibility comes in degrees, there is no strict reversibility–irreversibility dichotomy, but a continuum from 100% reversible changes to 0% reversible changes, i.e. complete irreversibility. Claiming that the collapse is an irreversible process could be interpreted as the impossibility of designing an experiment which enables us to bring 100% of the objects in an ensemble of identically prepared objects in the final state back into the initial state. This interpretation, I think, catches exactly the crucial feature of measurements in which the wave functions for the initial and final states are eigenfunctions to non-commuting operators.

It is, of course, right to say that in those measurements where the measured object is destroyed, e.g. when a photon is absorbed, the measurement process is irreversible by necessity. But this is not always the case; in many measurement processes the measured object is simply lost out of control, which is a purely practical affair. (The Copenhagen interpretation which says that the measurement necessarily introduces an uncontrollable disturbance will be discussed in Chapter 9.)

6.7 Irreversibility as a result of randomness

Irreversibility, as defined above, is in quantum mechanics the inevitable effect of two properties of quantum interaction processes, namely discontinuity and randomness. This can be seen as follows. The state of a quantum system is represented as a ray in a Hilbert space. A measurement of the first kind on such a state results in a collapse, a reduction of the superposition of eigenfunctions to the applied operator to one of its components, which means that the measurement is represented as a projection on one of a set of base functions spanning the Hilbert space. These projections cannot be the result of a continuous evolution from an earlier state. Furthermore it is impossible to predict which of the base functions the projection will fall on. Hence, the change is both discontinuous and random.

Now, the number of possible results equals the dimension of the Hilbert space. As a measurement of the first kind will have more than one possible outcome, the dimension of the Hilbert space is two or higher. That in turn implies that when we try to return the system to the original state by a suitable preparation, this change is, just as the first change, discontinuous and random, which means that we cannot with any certainty bring the system back to its original state. Hence, the evolution from the initial state to the later state was irreversible in the sense given above. We can state this as a theorem of quantum theory:

Randomness entails irreversibility: If an object being in a state ϕ changes state into ψ_j belonging to a set of states $\{\psi_j\}$, and if it is impossible to devise experiments in which we can predict with certainty which state in the set the object will evolve

into, then the reversed process is also impossible to monitor so that the object with certainty will return to its original state.

Summarising so far, we have the following picture. Measurements of the first kind can be characterised as discontinuous, random and irreversible changes of state. In the settings of quantum theory, discontinuity is a consequence of quantisation of action, and so is randomness. Irreversibility is a consequence of randomness. Hence, *the collapse of the wave function during measurements of the first kind is an effect of quantisation of interaction.*

6.8 Collapse as loss of coherence

When the wave function collapses into a small volume in the vicinity of a detector absorbing energy from the detected object, the part of the wave overlapping with other detectors disappears. This is necessary because otherwise there would still be a certain probability for the other detectors to be triggered and such an event does not occur; if it did, the exchange of energy would not be quantised. Given the assumption that one single object triggers the detector, it is impossible that another detector is simultaneously triggered by the same object, because of quantisation of interaction. Then these latter unaffected detectors are no longer phase correlated, neither with the wave function of the measured object, nor with the wave function for the triggered detector. This is so because the phase is given by the function $\exp(-i\omega t)$ where $\omega = E/\hbar$ and thus, if the wave does not overlap with the detectors not being triggered, the evolution of the different detectors is determined by different energies. Thus, even if the triggered detector a moment later gives back a quantum of energy to the incoming object, thereby reversing the energy exchange, the correlation with the other detectors will not be recovered. Hence, a return of the energy quantum to the incoming object will not be sufficient for a complete revival of its former state; the wave function has irreversibly collapsed, and started a new evolution.

How does all this accommodate for the now well-known phenomena called *quantum eraser*, in which a complete erasure of a measurement result is said to recover the original state before the collapse? Quantum erasure appears to provide a counter-argument to my view, but, actually, it does not. I will discuss the matter in Section 6.14.

6.9 Superposition of macroscopic states

Bohr repeatedly maintained that, in order to talk meaningfully about nature, and more specifically, to talk meaningfully about experiments and their results, we must use classical physical concepts. Hence, in his opinion it is simply meaningless even to ask the question whether a macroscopic object such as a measuring instrument could be in a state of superposition or not, for the concept of superposition can be used meaningfully only in microphysics and there is a sharp distinction between micro and macro physics. I am not convinced by his arguments, although I think his conclusion is correct. I cannot really see what it would mean to claim that a detector is in a superposed state of $|\text{fired}\rangle \otimes |\text{not fired}\rangle$. Or, still worse, what would it mean to

say that, for example, the height of a person is a superposition of, say 1.70 m and 1.80 m? If the concept of a superposed state is generally applicable to macroscopic objects, it would at least be possible to describe a non-actual state of affairs which is described by saying that it is a superposition. That is not the case and thus the concept is not generally applicable.

On the other hand, we know perfectly well what is meant by a superposition of two macroscopic waves, for example two waves on a water surface. It means that, at each particular place and time where the two waves are present, the actual wave amplitude is the sum of the amplitudes of the two waves. It does not mean, however, that the state, in this case the level of the sea surface, is indefinite.

The difference between a superposition of macroscopic classical states and a superposition of quantum states is, once more, quantisation. Quantisation implies that it is impossible to observe the weighted sum of the values of two observables, contrary to the case of classical wave mechanics. Superposition of quantum states manifests itself as interference effects. But what exactly are the conditions for observing a superposition in a macroscopic device? As the macroscopic–microscopic distinction is only a pragmatic distinction without any fundamental importance, it should in principle be possible to observe interference effects in macroscopic objects after all. This deserves a closer look and I will utilise Anthony Leggett’s discussion of the matter.

6.10 Leggett’s argument

My argument in Sections 6.3–6.8 depends on the assumption that two macroscopic objects such as two detectors cannot really be in a common state longer than a short moment. In other words, I have taken for granted that the superposition $|\text{detector 1 triggered and detector 2 not triggered}\rangle \otimes |\text{detector 1 not triggered and detector 2 triggered}\rangle$ does not exist. Is this correct? Anthony Leggett (1986) has discussed this issue and concluded that macroscopic objects as a matter of principle can be in a superposed state. I will now review his argument.

Leggett starts by writing down a superposition of two macroscopically different states Ψ_1 and Ψ_2 of a system, say a Geiger counter:

$$\Psi = a\Psi_1 + b\Psi_2 \tag{6.5}$$

He then rehearses the majority opinion of physicists, namely that (a) the superposition is, under given conditions, a physically correct description, but (b) it is impossible to detect any observable consequences of the superposition because the complexity of the system makes it indistinguishable from an incoherent mixture of Ψ_1 and Ψ_2 . Leggett’s point is that (b) is not necessarily correct; hence it should be, at least in principle, possible to perform experiments to test the validity of the superposition principle on a macroscopic level.

There are three arguments in favour of (b) and here is how Leggett rebuts them.

The first argument is that complexity makes it impossible to observe a superposition. Leggett thinks this is invalid. He asks us to imagine a macroscopic object such as a Geiger counter with two pointer states: triggered and not triggered.

Besides its two pointer states the Geiger counter has a great number of macroscopic degrees of freedom, labelled ξ_i ($i = 1, \dots, M$). The total number of degrees of freedom is N and $N-M$ of these behave in the same way in the two pointer states. Thus Ψ_1 and Ψ_2 will be

$$\Psi_1 = X_{N-M} \prod_{i=1}^M \phi_i(\xi_i) \quad (6.6)$$

$$\Psi_2 = X_{N-M} \prod_{i=1}^M \psi_i(\xi_i) \quad (6.7)$$

where each ϕ_i is orthogonal to the corresponding ψ_i and X_{N-M} is the wave function for the last $N-M$ degrees of freedom. It is obvious that the linear superposition $\Psi = a\Psi_1 + b\Psi_2$ cannot be distinguished from a mixture of Ψ_1 and Ψ_2 . But, as Leggett observes, this is too simple an argument:

The point we have missed is that we do not have to detect the simultaneous existence of the two states directly; rather we let nature do the work for us by applying an operator, namely the time-dependent operator $U(t) = \exp(-iHt/\hbar)$ which does have a 50-particle (and higher) correlations built into it. (Although H contains only one and two particle operators, arbitrarily high powers H occur in $U(t)$). The result of applying this operator is that the 50-odd microscopic (electronic and nuclear) co-ordinates are all locked adiabatically to a single degree of freedom, namely the centre-of-mass co-ordinate. Since this is a sum of one-particle operators, not a product, there is no problem about measuring it. Of course, we still cannot detect the existence of the superposition at the diffraction directly, any more than we could for a single electron or photon; but just as in that case, we can detect its effects later, when the quantum mechanical time development operator has recombined the two waves at the detecting screen. (Leggett, 1986, p. 31)

Thus, many degrees of freedom are not sufficient to destroy interference effects.

The second argument for the impossibility of observing a macroscopic superposition state, of, say, a billiard ball, is that the de Broglie wavelength of a macroscopic body under all reasonable conditions is so small that the interference effects would be totally unobservable. The third argument is that the energy level spacing of macroscopic objects is smaller than the thermal energy kT at all attainable temperatures; hence thermal noise would destroy all interference effects. Leggett claims that both these arguments are valid as applied to such things as billiard balls, but both fail when we consider more general types of collective (macroscopic) coordinates. Leggett gives as example an induction coil in a LC -circuit which oscillates harmonically with the classical frequency $\omega_0 = (LC)^{-1/2}$. Hence, the quantum description presumably implies a level spacing of $\hbar\omega_0$. This need not be small as compared with kT even for coils of dimensions of 1 cm. Moreover, the flux uncertainty in the ground state is $\hbar/(\omega_0 C)^{1/2}$, a magnitude that easily can be of the order of 10^{-17} Wb. This is large enough to be detected by modern magnetometers. Hence, a superposition of states of this system seems to be observable and, Leggett claims, the possible states must be considered as macroscopically different.

Unfortunately, this example is not suited for observing interference effects, because there is no difference between the classical and quantum dynamics of a harmonic oscillator.

More generally, a minimum condition for observing characteristically quantum effects is that we are sufficiently far away from the correspondence limit, i.e. the limit where quantum and classical mechanics coincide. Roughly speaking, the characteristic energy scale V_0 of the potential in which the system moves must be comparable to the level spacing, which in this case, the LC -circuit, is of the order of $\hbar\omega_c$, where ω_c is the classical frequency of the motion. Now the actual potential associated with a macroscopic variable is also macroscopic whereas the quantity $\hbar\omega_0$ of course is small. Thus, we are normally well beyond the correspondence limit. This argument shows that it is the macroscopic energy scale rather than the macroscopic geometrical scale that prevents us from seeing interference effects for macroscopic variables.

There is, however, at least one significant escape from this argument. In the Josephson effect, the potential implicitly contains $\hbar$, whereas the macroscopic quantity is the current carried by Cooper pairs (of the order of 10^{-9} A), or the related trapped flux of the order of 10^{-15} Wb. A general condition for observing interference effects can now be calculated. The potential V_0 must not be too large compared with $\hbar\omega_c$. The characteristic frequency of classical motion of a system of mass m in a potential V_0 which varies over a region of dimension l is of the order of $(V_0/m^2)^{1/2}$. Applied to a Josephson junction in a SQUID (Superconducting Quantum Interference Device) this quantity becomes $I_c/(C\Phi_0)^{1/2}$, where C is the capacitance of the ring and Φ_0 is the flux quantum $h/2e$. Thus the condition for the potential V_0 not be too large compared with $\hbar\omega_c$ can be expressed as

$$8CI_c\Phi_0^3/\pi^3\hbar^2)^{1/2} \text{ not much larger than } 1,$$

which can be reasonably well fulfilled using attainable parameter values. As before, the thermal noise must not be allowed to blur the interference effect, hence $kT \ll \hbar\omega_c$.

However, one more condition must be satisfied in order to observe quantum coherence of two or more states. The two conditions given above are conditions for it being a difference between quantum and classical dynamics. But these are only necessary, not sufficient for really observing macroscopic coherence. As is claimed by many authors, macroscopic systems show highly irreversible behaviour, which means that the phase information stored in the macroscopic degrees of freedom rapidly dissipates into the environment. A number of authors have attributed this to the fact that the environment continually observes the system, and if this watched-pot effect is effective, we will never be able to observe characteristically quantum behaviour. Leggett concludes:

Since the interactions which destroy the phase information are precisely those which fail to satisfy the adiabatic principle, and these in turn are the ones which dissipate the energy of the macroscopic system, the problem is to determine *the effect of dissipation on macroscopic quantum tunnelling and coherence*. (Leggett, 1986, p. 37)

Leggett is then able to give a condition that guarantees that the dissipation of phase information does not destroy quantum coherence effects in a SQUID. This condition is:

$$(\Delta\varphi/\varphi_0)^2(R_0/R_n)\langle 1/2$$

$\Delta\varphi$ is the spacing between consecutive minima of flux, φ_0 is the flux quantum, R_n is the shunting resistance of the Josephson junction and R_0 is the characteristic quantum unit of resistance, $h/4e^2 \approx 7 \text{ k}\Omega$ (Leggett, 1986, p. 39). Is it possible to satisfy this condition within the constraints of present technology? According to Leggett this is an open question.

In the present context, the question whether it is really possible to make measurements is of minor importance. The crucial thing is that Leggett has given numerical conditions, which for a macroscopic system, a SQUID, tells us:

- (1) when the thermal noise destroys interference effects,
- (2) when quantum dynamics deviates from classical dynamics,
- (3) when the coupling constant between the SQUID and the environment is negligible so that we can ignore dissipation of phase information.

It is easy to see that these conditions are not contradicting each other. Hence, as a matter of principle, macroscopic and microscopic objects are not different. The reasoning also suggests a general methodology for answering questions of the type ‘why does not this particular macroscopic object show quantum coherence effects?’. For each particular object one must make calculations analogous to those given by Leggett for the SQUID, and show that one or two or all three of the conditions are violated.

When using a macroscopic object as measuring apparatus, we want its pointer states to be effectively orthogonal. We can now see that this requirement can be met if certain microscopic conditions are fulfilled and we need not depend on the macroscopic/microscopic distinction.

Is Leggett’s analysis by itself a solution to the measurement problem? No, it is not, because it takes for granted that the environment show irreversible behaviour. Hence, the ordinary quantum dynamics given by the Schrödinger equation cannot account for the evolution of the environment. Thus, Leggett presupposes that there are collapses in the environment, without giving any closer analysis.

6.11 Interaction-free measurements

I have argued that all measurements necessarily involve exchange of conserved quantities. This stance appears to be undermined by the possibility of interaction-free measurements, which have been discussed in a number of papers in recent years. However, these purported interaction-free measurements involve exchange of quanta after all. In what follows I will first review the core of the important paper by Elitzur & Vaidman (1993) and then give my arguments for not regarding these measurements as interaction-free.

Elitzur & Vaidman (1993) start by considering a Mach–Zehnder interferometer through which single photons are directed, as shown in Figure 6.5. The point of their argument is to show that one can measure the presence of an object (indicated as an elliptical ball in the figure) in one of the paths without interacting in any way with this object.

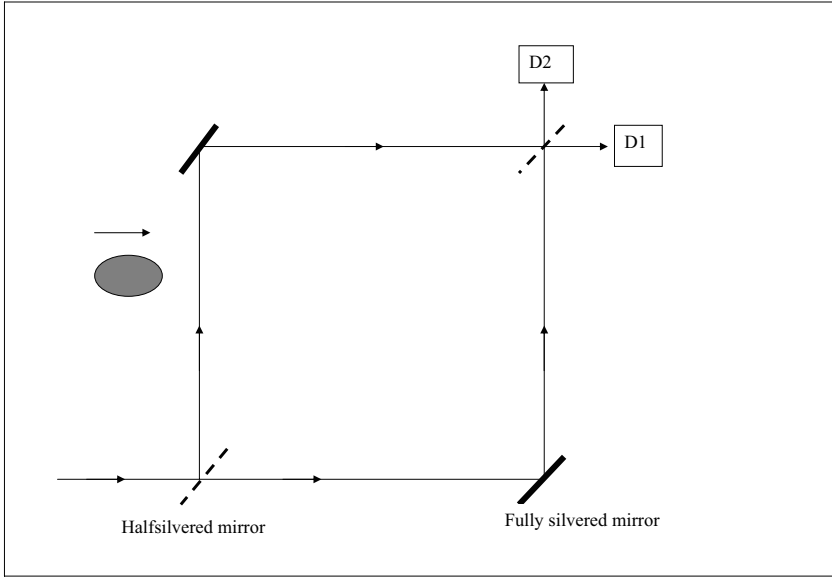


Figure 6.5 Interaction-free measurement of the position of a ‘bomb’

The photon state after passage of the first half-silvered mirror is a superposition of two states $|1\rangle$ and $|2\rangle$, where $|1\rangle$ denotes horizontal and $|2\rangle$ vertical motion. When a photon is reflected in the half-silvered mirror its phase is changed by $\pi/2$, hence the operation of the half-silvered mirrors can be symbolised as

$$|1\rangle \rightarrow \frac{1}{\sqrt{2}} [|1\rangle + i|2\rangle] \quad (6.8)$$

and

$$|2\rangle \rightarrow \frac{1}{\sqrt{2}} [|2\rangle + i|1\rangle] \quad (6.9)$$

and the effect of the two fully-silvered mirrors is described as

$$|1\rangle \rightarrow i|2\rangle \quad (6.10)$$

$$|2\rangle \rightarrow i|1\rangle \quad (6.11)$$

Due to interference, the detector D2 will never fire and D1 will fire with 100% certainty, if both paths of the interferometer are open, i.e. the object is outside the path. Now, the object is inserted into the interferometer as shown in the figure. There are three possible outcomes:

- (1) no detector clicks,
- (2) detector D1 clicks,
- (3) detector D2 clicks.

If no detector clicks we can infer (assuming 100% efficiency of the detectors) that the photon has been absorbed by the inserted object. The probability for this outcome is $p=0.5$. If D1 clicks ($p = 0.25$) nothing can be inferred about the presence of the object. If D2 clicks ($p = 0.25$) we can infer with certainty that there is an object present in the path and we can also infer with certainty that this object has *not* interacted with the photon. Hence, Elitzur & Vaidman (1993) claim, we have an interaction free measurement of the position of the inserted object.

The quantum mechanical description of the state of the photon is as follows. If the object is absent from the path, the evolution of the photon state is described by

$$\begin{aligned} |1\rangle &\rightarrow \frac{1}{\sqrt{2}} [|1\rangle + i|2\rangle] \rightarrow \frac{1}{\sqrt{2}} [i|2\rangle - |1\rangle] \rightarrow \\ &\frac{1}{2} [i|2\rangle - |1\rangle] - \frac{1}{2} [|1\rangle + i|2\rangle] = -|1\rangle \end{aligned} \quad (6.12)$$

Therefore, the photon leaves the interferometer moving to the right towards D1. If however, the object is present the evolution is described by

$$\begin{aligned} |1\rangle &\rightarrow \frac{1}{\sqrt{2}} [|1\rangle + i|2\rangle] \rightarrow \frac{1}{\sqrt{2}} [i|2\rangle + i|\text{scattered}\rangle] \rightarrow \\ &\frac{1}{2} [i|2\rangle - |1\rangle] + \frac{i}{\sqrt{2}} |\text{scattered}\rangle \end{aligned} \quad (6.13)$$

This superposition then collapses into one of the possibilities at the detection stage and the probabilities become 0.25 (D2 clicks), 0.25 (D1 clicks) and 0.5 (no click) respectively.

In this formalism the inserted object is not described as a quantum object but rather as an external condition. If we instead think of it as a quantum object with two possible states, $|A\rangle$, which signifies that it is present in the path of the photon, and $|B\rangle$, which signifies the state of not-present in the path, the total state of these two possibilities make up a superposition:

$$|\Psi\rangle = \alpha|A\rangle + \beta|B\rangle \quad (6.14)$$

The combined system of object + photon then makes up an entangled state:

$$|1\rangle|\Psi\rangle \rightarrow \alpha \left[\frac{1}{2} [i|2\rangle - |1\rangle] + \frac{1}{\sqrt{2}} |\text{scattered}\rangle \right] |A\rangle + \beta |1\rangle |B\rangle \quad (6.15)$$

As before, according to the standard interpretation of quantum mechanics, the collapse occurs when the detector D2 clicks, which means, because of entanglement, that the state of the inserted quantum object collapses into a definite position when the detector D2 clicks. According to Elitzur & Vaidman (1993), this is an interaction-free collapse of the state of the particle both when the particle is present and when it is not present in the photon path.

Comments

Consider first the case when the inserted object is not treated quantum mechanically but rather as a macroscopic object. By assumption, such an object is not in a superposition of two position states. Let us assume that the object is inserted into the path as shown in the figure. The photon amplitude then becomes zero behind the object, because the possibility for detecting the photon behind the object is zero. (Remember that we have treated the object as an opaque macroscopic object.) The photon is either absorbed or it is not; in the latter case it would be detected with certainty in the other path (100% efficiency of the detectors is assumed). Thus, there exists in neither case a partial wave behind the object, which is to say that a collapse has occurred. Hence, Elitzur & Vaidman's claim that the usual interpretation is that the photon state collapses in the detectors, is not correct.

I do not mind calling this an interaction-free measurement of the position of the object, but it should be observed that it is not a measurement *of the state of quantum object being in a superposition*. I cannot see that it has any relevance for the interpretation of the measurement problem. Similar situations occur also in the macroscopic world. As an example, imagine that you participate in a TV programme and you are invited to a lottery game: in front of you there are two doors, behind one of them there is a car which you win if you open that door. Now, suppose you open the wrong door. Then you know with certainty that the car is behind the other door. You have made an interaction-free measurement of the position of the car, just because you knew in advance that there were only two possible positions.

Secondly, let us consider the case when the 'inserted' object is indeed a quantum object, for instance an electron or an atom, being in a superposition of several position states. In this case we must be careful with the notion of the object *being at a particular place*. As has been discussed earlier, we cannot equate the probability for an object being *detected* at a particular place with the probability for the object *being* at that particular place independently of any measurements. The notion of position is simply not well defined for quantum objects, except at moments of interactions with macroscopic objects.

In order to have a detection in D2 the partial wave going the left-upper path must be reduced in intensity. If that is going to happen the inserted quantum object must 'push the wave function for the photon away' from the left-vertical path, i.e. the spatial distribution of the passing photon wave must change without any exchange of any conserved quantity with the inserted object. Let us suppose that this is what happens. And let us further suppose that we have a click in D2 which confirms that the photon wave in the upper-left path was reduced in intensity behind the inserted object. This detection is, according to E&V a collapse of the superposition

$$|1\rangle|\Psi\rangle \rightarrow \alpha \left[\frac{1}{2} [|i2\rangle - |1\rangle] + \frac{1}{\sqrt{2}} |\text{scattered}\rangle \right] |A\rangle + \beta |1\rangle |B\rangle \quad (6.16)$$

to the state

$$|\text{scattered}\rangle |A\rangle \quad (6.17)$$

Thus, the state of the quantum object has collapsed to the state $|A\rangle$, i.e. to the definite state of being in the path. Using my analysis of measurements of the first kind we can say the following.

- (1) The collapse of the state occurs when an exchange of energy occurs, namely when the photon triggers a detector.
- (2) The measurement is ultimately a position measurement, in this case a position measurement of a photon, i.e. a secondary object.
- (3) The position of the secondary object is correlated with the observable of interest, in this case the position of the inserted object.
- (4) As the properties of the inserted object and the positions of the photon are correlated, neither of them can be determined independently of each other, which means that photon and inserted object are not two independent objects but one single object.
- (5) It follows that it is wrong to claim that the detection of the position of a photon is not an interaction with the object inserted in the path.

The apparent plausibility of Elitzur & Vaidman's analysis depends on the fact that they discuss the situation semi-classically using quantum mechanics to describe the states mathematically, but their informal reasoning is classical, *taking for granted* that objects and events are individuated classically. They assume that we can treat as distinct the two events (i) the photon passing and interacting with the inserted object and (ii) the interaction with the detector. But the first event is not a measurement; it is not even an identifiable event because it lacks identity criteria.

It might be thought non-physical to consider the detection of a photon at a particular place as an interaction with another object situated some distance away. But that is not so; we have considerable evidence for the view that photons are not localised, except at moments of creation and destruction, which entails that just before the collapse the photon under consideration is present all over the interferometer, perhaps in the entire space. We have in fact here an example of a non-local phenomenon, the general analysis of which is to be undertaken in Chapter 8.

6.12 Renninger's negative result experiment

The crucial point in my interpretation of the collapse is that every measurement is an interaction whereby quanta are exchanged. Despite appearance, Elitzur & Vaidman's interaction-free measurements do not make up not any counter argument. Another purported counter argument would be the so-called negative result experiment described by Renninger (1960). He has described a type of experiment which he claims would enable us to draw conclusions about the state of a system without interacting with this system in any way at all.

Renninger imagined an α -radiating source in the centre of a sphere, E2, whose interior surface is covered with fluorescent material. An α -particle absorbed by the material will produce a scintillation that can be recorded by a photo-multiplier. Inside this sphere a concentric semi-sphere E1 is placed. This too is covered with a

fluorescent material. The emitted particles will hit either E1 or E2. The wave function before the collapse can be written

$$\Psi(t) = a G_1(t) + b G_2(t) \quad (6.18)$$

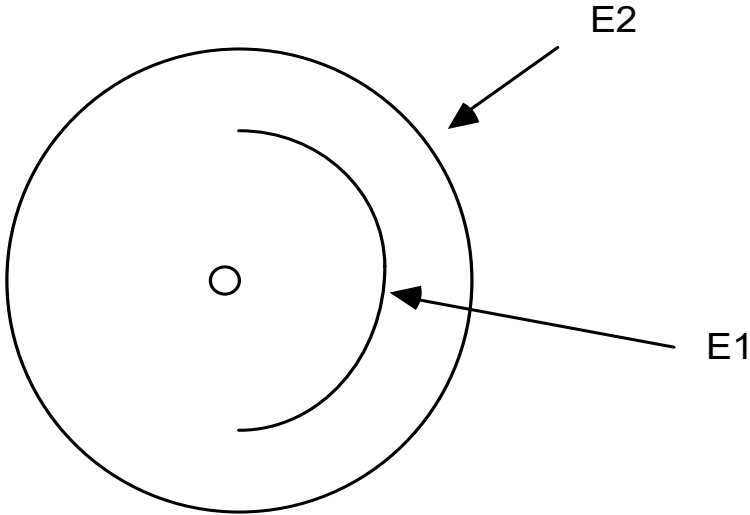


Figure 6.6 Renninger's negative result measurement

where G_1 and G_2 refer to the right and the left parts of the total wave function, respectively. The total wave function represents in this case a spherical wave expanding from the source. Hence, the part G_1 will hit E1 and the part G_2 will hit E2. The probability that a flash from E1 will be recorded is $|a|^2$ and the probability for a flash from E2 is $|b|^2$. Suppose that in a certain case we do not record a flash from E1. That must be interpreted as a collapse of the wave function at the moment when a part of it hit E1. Hence, Renninger claims, we are forced to the conclusion that when the wavefront is between the spheres E1 and E2, the wave function must be

$$\Psi(t) = G_2 \quad (6.19)$$

which immediately is confirmed when a flash from E2 is seen. In conclusion, in those cases when there are no flashes from E1, the wave function collapses when passing this semi-sphere and no interaction between the semi-sphere and the object takes place. Renninger concluded that the observation of E1 was a measurement even in those cases when there were no flashes. He considered this thought experiment as an argument against the view that any measurement necessarily disturbs the measured object. If his conclusion is correct, the claim that all measurements are position measurements must be rejected.

But the validity of Renninger's conclusion depends on how the concept of measurement is defined. Which interaction is supposed to be the measurement: the

wave collapse at E1 or the object hitting E2 and triggering a flash? Renninger argues that the experiment shows that the measurement does not disturb the measured object, which implies that in his view the event at E2 cannot be the measurement under discussion because this event is a disturbance of the object, a change of its energy state. He thus takes for granted that the first event, the wave collapse at the passage of E1, is a measurement. Hence he must admit that even if there were no sphere E2 he still would have to maintain that the object, the α -particle, was measured upon. But if there was no sphere E2, it would be impossible for us to know if there has been any collapse. I have no problem with this line of thought, but it contradicts the very starting point of Renninger's argument, which is that wave functions collapse *only* during measurements of the first kind. His argument is hence not coherent.

In my view, Renninger's conclusion that there could be measurements without there being any interaction at all is a consequence of conflating (a) the collapse, which occurs at the inner semi-sphere (and it occurs both in the case when the α -particle is absorbed by E1 and when not) and (b) the registration of the result, which occurs later, at E2. There will either occur a flash on E2 or not; these two alternatives do not make up a superposition, because the collapse has already occurred. Hence, the interaction between the α -nuclei and E2 is not a measurement of the first kind, it is a mere registration of a state. Such events may be interaction free, there are plenty of such examples in everyday life as already discussed.

6.13 Delayed choice

John Archibald Wheeler has argued that the dual character of quantum objects is an effect of our way of observing things and that it is impossible to ascribe to quantum objects either wave or particle properties independently of our way of obtaining knowledge. His argument is built on the delayed choice experiment described in Chapter 1. As was explained in that chapter, it is a double-slit experiment with *two* detecting set-ups: one photographic film, which can be moved in and out of the electron beam and a pair electron counters, i.e. Geiger counters. These are placed so that one can infer through which slit in the screen a particular electron has passed, if the electron is conceived as a particle. The principle of the experimental set-up is shown in Figure 1.1 (Cf. Wheeler, 1983, pp. 182–184). The experiment is conducted in the following way: the electron source emits electrons at a rather low rate, low enough to admit the conclusion that at each moment only one electron is on its way from the source to the target. This restriction is imposed in order to exclude the possibility that the electrons interfere with each other. The photographic film is moved up and down with a frequency such that an electron detected on the film must have passed the double slit before the film was in the detecting position. The 'decision' to move the film into the detecting position is hence made after the electron has passed the screen with the double-slit. If the electron has travelled far enough when the film is moved into place it will miss it and hit one of the Geiger counters. As before, the film will show a typical interference pattern, suggesting that electrons are waves having passed through both slits. But some electrons are counted in the upper detector and some in the lower one and, so it is said, these electrons must have

passed one slit only. Wheeler accepts this conclusion. He maintains that if the upper counter has fired, the electron has passed the lower slit, if the lower counter has fired, the electron has passed the upper slit and if the film is in place the electron must have passed both slits. This is quite awkward and completely unnecessary if one accepts that in all cases there is a real wave which collapses either at the screen or in one of the counters. The wave always passes both slits and there is nothing mysterious with this experiment.

6.14 Quantum Eraser

Several authors, for example Kwiat *et al.* (1992) have claimed that it is sometimes possible to ‘erase’ the information stored in a detection equipment and that this erasure retrieves the superposed state of the object. If this is correct it implies that the irreversibility of the collapse is not due to the storage of information in some medium, but the final human observation of the outcome. Thus interpreted, quantum erasure conflicts with my views. Either I have to give up the stance that irreversibility is caused by the physical interaction between object and measurement device, not by human cognition, or reinterpret the thought experiment ‘quantum eraser’. Needless to say, the second alternative is my choice, and I will discuss interference in a Hong–Ou–Mandel interferometer, as described by Kwiat *et al.* (1992).

A Hong–Ou–Mandel interferometer consists of:

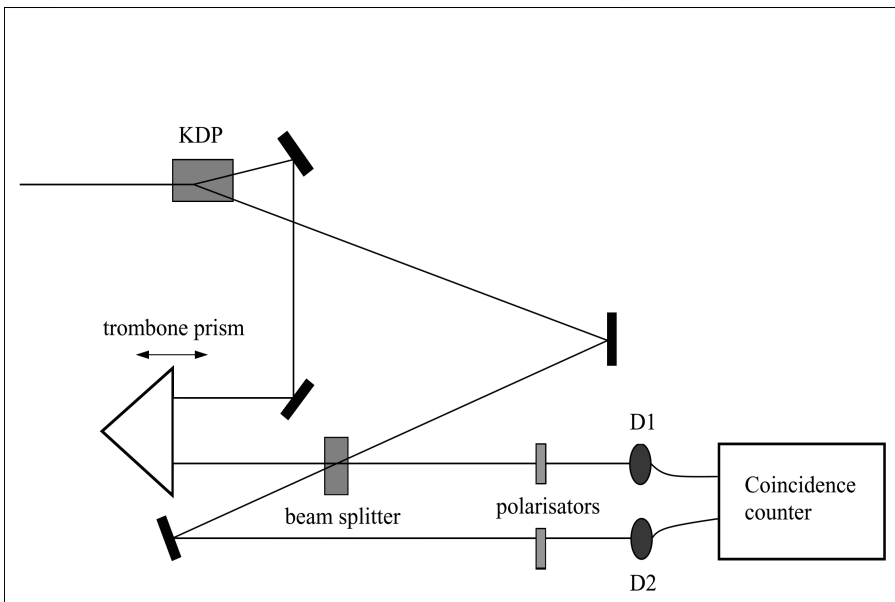


Figure 6.7 A Hong-Ou-Mandel interferometer

Picture adapted from Phys. Rev. A, vol 45, no 11, June, 1, 1992, pp. 7729–39.

- (1) a non-linear crystal (KDP in the Figure 6.7) which convert incident photons into two correlated photons, each with half the energy of the incident photon,
- (2) a beam splitter,
- (3) two photo-detectors and
- (4) a coincidence counter.

The principal set-up is shown in the Figure 6.7.

When a photon passes the beam splitter it is divided in two parts, a transmitted part and a reflected part. The state of the photon can be represented as

$$|\Psi\rangle = \frac{1}{\sqrt{2}} |t\rangle + \frac{1}{\sqrt{2}} |r\rangle \quad (6.20)$$

The reflected part undergoes a phase shift, which is represented by the i -factor. Now, if two photons arrive simultaneously to the beam-splitter we should multiply their wave functions:

$$|\Psi\rangle = \left(\frac{1}{\sqrt{2}} |t_1\rangle + \frac{i}{\sqrt{2}} |r_1\rangle \right) \left(\frac{1}{\sqrt{2}} |t_2\rangle + \frac{i}{\sqrt{2}} |r_2\rangle \right) = \frac{1}{2} |t_1\rangle |t_2\rangle + \frac{i}{2} |t_1\rangle |r_2\rangle + \frac{i}{2} |r_1\rangle |t_2\rangle - \frac{1}{2} |r_1\rangle |r_2\rangle \quad (6.21)$$

Coincidence in the two detectors occurs either when both photons are reflected or when both are transmitted, represented by the states $|t_1\rangle |t_2\rangle$ and $|r_1\rangle |r_2\rangle$ respectively. As photons are mere quanta of energy, the state description $|t_1\rangle |t_2\rangle$ and $|r_1\rangle |r_2\rangle$ denote one and the same state. Hence the probability for correlation is

$$P_{\text{corr}} = \left| \frac{1}{2} |t_1\rangle |t_2\rangle - \frac{1}{2} |r_1\rangle |r_2\rangle \right|^2 = 0 \quad (6.22)$$

Thus, a correlation equal to zero indicates interference, contrary what is common to the majority of experiments showing interference phenomena.

All this requires that the photons arrive more or less simultaneously at the beam-splitter so that their wave functions overlap. By inserting a so-called ‘trombone prism’ in one of the paths by which one can increase or decrease the path length, Kwiat *et al.* could show that the correlation function increases in proportion to the increased difference in path length. When the difference in path length exceeds the coherence length of the photons there is full correlation, which in this case means that no interference effect occurs in the experiment, exactly as expected.

Now, a fundamental fact of quantum mechanics is that if we somehow could obtain information about which way a photon took, the interference would be destroyed. In order to test this, a half-wave plate is inserted in one of the beams. This plate rotates the polarisation plane twice as much as the angle between the polarisation plane of the incident photons and the axis of the plate. Hence, if the axis of half wave plate is set to an angle of 45° relative to the polarisation plane of the incident photons, their polarisation plane will rotate 90° . Now it is possible to find out which way a photon took because we can place polarisation filters just before the detectors and choose to detect only horizontally or only vertically polarised photons. That means that the interference pattern should be destroyed, i.e. the correlation

increases well above zero. This is also precisely what happens. Observe that it is not necessary actually to check the polarisation of the incoming photons: the mere possibility to do so destroys the interference.

Finally, we have come to the discussion of erasure. Obviously, a detection of the polarisation in one of the paths will destroy the interference, which is trivial. Kwiat *et al.* instead tested the erasure of the *possibility* of knowing which way a photon took. They inserted therefore a polariser oriented in an angle of 45° towards the horizontal plane before each detector. These polarisers will erase the difference between vertical and horizontal polarisation, thus erasing the information about which way a single photon took. Hence this measure should restore the interference effect, i.e. the zero correlation figure again. This is also corroborated without any doubt in the experiment.

Viewed from the perspective of a wave ontology the entire experiment is almost trivial; nobody would have thought about performing it, had it not been assumed that photons are distinct and discrete objects, not mere portions of energy that are discretised only in interactions. It is the conflict between a wave interpretation and a particle interpretation that makes this experiment interesting. Had people fully understood that photons are wave phenomena without individual existence when passing from source to detector, there would have been no talk of erasure of which way information, because no information is stored anywhere. What is called erasure by Kwiat *et al.* could better have been described as the destruction of the possibility for a future correlation. It is obvious that the outcome of this experiment by no means causes any trouble for the interpretation suggested in this book.

There is one further point to be observed, namely how the concept of information is used by Kwiat *et al.* The authors write that the wave function contains information about which way a particular photon took when the half-wave plate was inserted in one of the beams. Information is in this context not used as a cognitive concept but as referring to physical structures. This usage conforms to Shannon & Weaver's analysis according to which information is the logarithm of a probability function, taking the number of possible micro states as arguments. The idea is simple enough: assume that we have a description of a system in terms of macroscopic variables. Usually this state description is compatible with a number of different microscopic states. The logarithm of the number of microscopic states that are possible in the light of this macroscopic description is a measure of the lack of information about the real state of the system, *given the macroscopic description*. It is natural to say that the theoretical concept of information is a measure of our amount of knowledge of a particular state, given the input values. No reference to subjective beliefs is implicated by using this concept.

Returning to the quantum eraser: what is erased is not actual information, i.e. a certain structure in a human being or in the memory of a computer, but the *possibility* to extract information about the path from the state of the quantum system. It seemed reasonable to Kwiat *et al.* to say that there was information in the system and that this information was erased, but that is wrong! What has been erased is the correlation between a quantum state and a macroscopic state of the detecting device. One could alternatively describe the situation by saying that the conditional 'if there is a click in the upper counter, the measured photon took the upper path' is no longer true.

This does not mean that any information state is erased. The polarizer decreases the intensity of the wave, which is to say that there is a possibility of absorption of a quantum of energy in the polarizer.



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Chapter 7

Quantum Mechanical Spin

7.1 Introduction

The fact that quantum objects exhibit the exclusively non-classical property of spin has often been used to rule out the possibility of a realistic pictorial representation of quantum phenomena. This intuition got strong support by the famous Kochen–Specker theorem, which by most people has been taken to prove that a realistic interpretation of spin is impossible. I think this is completely wrong; all spin properties of quantum systems are possible to give a pictorial, fully realistic interpretation and Kochen–Specker’s theorem does not exclude a realistic understanding of spin properties.

7.2 A visualisable model of spin-half objects

From a naive realistic outlook, the most astonishing feature of spin-half particles, no matter what direction we choose as a reference for a measurement, the spin is either parallel or anti-parallel to the chosen direction. (I will call these objects ‘particles’ for convenience only; they are waves showing particle properties only in interactions.) As this choice can be made at will at any moment before the detection it seems as if the spin particle has no intrinsic spin. A particle can be said to have spin up (or spin down) when interacting with a detector, otherwise not. On the other hand, however, under certain circumstances we can anticipate the result of a spin measurement, which seems to suggest that spin is attributable to the object independent of any measurements.

A spin measurement is always a two-step process; first, the object to be measured upon is directed through an inhomogeneous magnetic field, and then it is registered in a detector. The first step is necessary because any measurement is ultimately a position measurement, hence in order to measure spin we must arrange a one-to-one correlation between spin and position. (This should not be interpreted as if the objects get definite positions when passing the magnetic field, because the collapse of the superposition does not occur in the magnetic field, as was argued in Chapter 6.)

In outline, my view is as follows; the spin $\mathbf{s}$ is connected to magnetic momentum $\boldsymbol{\mu}$ by the relation $\mathbf{s} = -g\boldsymbol{\mu}$, which is to say that these two observables are always anti-parallel. But magnetic momentum, conceived as an actual property, can be attributed to an object only if the object is or recently has been in an external magnetic field. The same goes for spin; it can only be attributed to an object which is or recently has been in an external magnetic field. This magnetic field defines a direction in space, usually labelled the z -direction. When it is said that the magnetic momentum and spin point in the z -direction,

the reference is by necessity made to the direction of magnetic field. Given this condition we can construct a vector model of spin-half particles as follows.

7.3 A vector model of spin-half¹

The spin operator

$$\mathbf{S} = \mathbf{S}_x + \mathbf{S}_y + \mathbf{S}_z$$

has the property that only one component of the spin has a definite value, $\pm\hbar/4\pi$, i.e. in units of $\hbar$ the value is either $+1/2$ or $-1/2$. Usually the z -component is chosen as the definite one and the proper value as well as the expectation value of the z -component is accordingly $+1/2$ or $-1/2$. The expectation values of the x - and y -components are both zero. However, the expectation values of the squared x - and y -operators are not zero but $1/4$, just as the expectation value for $\mathbf{S}_z^2$. Thus we have:

$$\langle \mathbf{S}_x^2 \rangle = \langle \mathbf{S}_y^2 \rangle = \langle \mathbf{S}_z^2 \rangle = \frac{\hbar^2}{4} \quad (7.1)$$

Now the following relation holds:

$$\langle \mathbf{S}^2 \rangle = \langle \mathbf{S}_x^2 \rangle = \langle \mathbf{S}_y^2 \rangle = \langle \mathbf{S}_z^2 \rangle \quad (7.2)$$

Equations (7.1) and (7.2) imply that

$$\langle \mathbf{S}^2 \rangle = \frac{3\hbar^2}{4} \quad (7.3)$$

All these relations can be visualised by a vector model, *provided we assume that expectation values in an ensemble of identically prepared systems can be put equal to time averages within each single system*. This is an important assumption because it enables us to infer the state of individual systems starting from observed quantities, i.e. expectation values referring to ensembles. Since we want a model of the spin of a single object, we need this assumption to connect the individual properties with ensemble properties.

The length of the spin vector is $\sqrt{3}/2$ and the projection of the vector on the z -axis (defined as the direction of the external magnetic field) is $1/2$. The fact that the expectation values of the x - and y -components in the z -spin up-state both are zero can be understood as the vectors in the ensemble of systems being randomly distributed around the z -axis. Using the assumption above, this means that in an individual object the spin vector is moving around the z -axis in such a way that the time average of both the x - and y -components equal zero. This interpretation accords with the fact that the expectation values $\langle \mathbf{S}_x^2 \rangle$ and $\langle \mathbf{S}_y^2 \rangle$ are both $1/4$, implying that the length of the projection on the xy -plane is $\sqrt{1/4 + 1/4} = 1/\sqrt{2}$. The situation is shown in Figure 7.1.

1 A discussion of the vector model of spin can be found in most textbooks, see for example Dicke and Wittke (1960, 194–201).

I have, in effect, suggested that the spin vector is rotating around the z -axis. Let us now check whether this interpretation is consistent. If we make a classical calculation of the mean value of the projection on the positive side of x - and y -

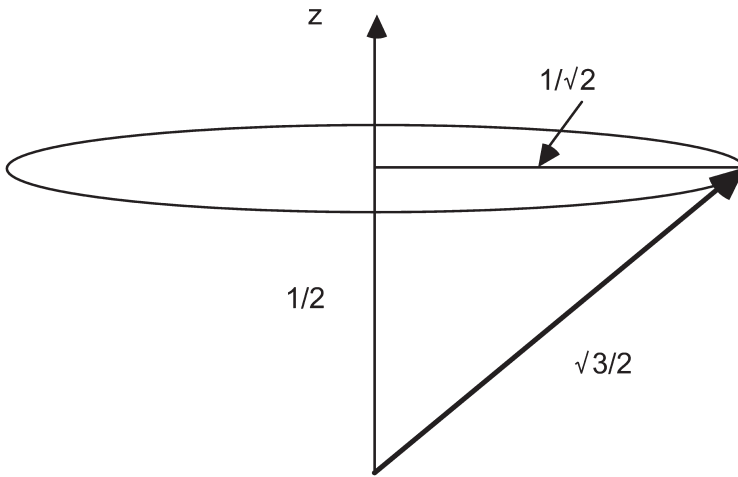


Figure 7.1 A vector model of spin half

axes of a rotating vector in the xy -plane of length $1/\sqrt{2}$ the result is $1/(\sqrt{2}\pi) \approx 0.45$. However, according to the model, all measurements made when the spin-vector happens to be on the positive side of the x - or y -axis will result in ‘spin up’ i.e. the value is exactly $1/2$ in the measured direction. So one might conclude that the rotating model is not consistent with quantum theory. But this is too hasty a conclusion however, for the rotating model by itself does not take into account of the effect of quantisation of interactions. A spin measurement in the x -direction upon an object which is in the z -spin up state induces a collapse. This can be understood as a sudden change of rotation axis when passing into a magnetic field. Before the interaction it rotates around the z -axis, a result of the preparation. The first step in the measurement of the spin- x component (or spin- y component) is to bring the object into a magnetic field where the magnetic field vector points in the x -direction. This will tilt the rotation axis and that is a perfectly classical behaviour of a magnetic vector; it will immediately align its axis of rotation either parallel or anti-parallel with an external magnetic field when passing into it. If the projected x -component is positive it will be parallel, otherwise anti-parallel. Hence, a proper understanding of the spin state, together with the recognition that we cannot observe the state without interacting with it, explains the statistical relations between the spin components of a quantum object.

When measuring the spin in a direction other than the direction of the prepared spin state, we are changing the stationary component of the spin vector. Using the distinction introduced in Chapter 3, one component of spin is an *actual property*, the other two are *potential properties*, and the measurement can be described as a swap between an actual and a potential property; what was an actual property

before measurement becomes a potential property after measurement and one of the potential properties becomes an actual property after measurement. We could now say that the spin vector is a function of time,

$$\mathbf{S} = \mathbf{S}_x \cos \omega t + \mathbf{S}_y \sin \omega t + \mathbf{S}_z \quad (7.4)$$

This is precisely the same as in classical mechanics, but there is little point in writing down this formula. We cannot use this equation for calculating the values of the x - and y -components at any point of time, because we have no initial values to insert. The only empirical conclusions which can be drawn using the formula (beside the constant $\mathbf{S}_z$) is that when measuring the spin in the x - or y -direction we will get a 50/50 distribution of positive and negative values, a result that is precisely what quantum mechanics in the Born interpretation tells us. And this is precisely as it should be; the goal of interpretation is to increase our understanding, not to add new predictions.

A related but different question is whether the spin vector, just as in the classical case, will precess around the magnetic field vector. The answer is that it depends on the situation: if the prepared spin state is perpendicular to the magnetic field, the notion of classical precession is applicable, but if the prepared spin state is parallel or anti-parallel to the magnetic field, it is not. I will discuss these cases in turn, beginning with definite z -component spin in a $\mathbf{B}_z$ -field.

In order to talk about a precession of the spin vector we must be able to specify its angular velocity. This can be done through a measurement of the interaction between the magnetic moment and an external magnetic field. The Hamiltonian for this interaction is given by

$$H = -\mu B = \frac{e}{m} BS \quad (7.5)$$

which in the case of $\mathbf{B} = B_0 \mathbf{B}_z$ and $\mathbf{S} = S_z$ becomes

$$H = \frac{eB_0\hbar}{2m} \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \quad (7.6)$$

The time-dependent Schrödinger equation

$$\mathbf{H}\Psi(t) = i\hbar \frac{\partial}{\partial t} \Psi(t) \quad (7.7)$$

has the solution

$$\Psi(t) = \exp\left(\frac{-iHt}{\hbar}\right) \Psi(0)$$

By identifying coefficients we arrive at the conclusion that the phase velocity of the wave function is (Dicke & Wittke, 1960, pp. 197–198)

$$\omega = \frac{H}{\hbar} = \frac{eB_0}{2m} \quad (7.8)$$

which could be identified with the angular momentum in equation (7.4). However, there is, as already said, a difference between the classical precession and this case in so far as this model cannot be used to calculate the precise direction of the vector at a certain time. Therefore, one cannot attribute an observable precession to this particle in this situation and hence no precession operator exists.

However, in a different circumstance it is possible to attribute precession to the particle. When the eigenstate of the spin particle is perpendicular to the external field one can talk about the precession of the determinate component of the spin, that component which is an actual property. Assume that particles in an eigenstate to S_x at $t = 0$ are directed into a B_z -field. If they have spin up in the x -direction at $t=0$, the wave function is

$$\Psi(t) = \frac{1}{\sqrt{2}} \begin{pmatrix} \exp\left(\frac{-i\omega t}{2}\right) \\ \exp\left(\frac{i\omega t}{2}\right) \end{pmatrix} \quad (7.9)$$

Consider now the operator

$$S^+ = S_x \cos \omega t + S_y \sin \omega t \quad (7.10)$$

which represents the component of the spin in the xy -plane along a line rotating with the frequency ω (Dicke & Wittke, 1960, pp. 197–199). The wave function (equation (7.9)) is an eigenfunction to the operator S^+ :

$$S^+\Psi(t) = \frac{\hbar}{2} \Psi(t) \quad (7.11)$$

Thus, in this case it is legitimate to say that the spin component precesses around the magnetic field vector with a frequency ω , i.e. twice the angular frequency of the wave function.

A measurement of the observable corresponding to the operator S^+ must give the result $\hbar/2$. However, the Hamiltonian does not commute with this operator, implying that the energy is not a constant of the motion. The energy corresponds to the projection of the spin vector on the z -axis, and it follows that the spin vector has no definite projection on the z -axis. The situation can be visualised as in Figure 7.2.

In conclusion, the vector model must be interpreted with some caution. Only one component of the spin vector is an actual property, the others are potential properties.

Now it is time to proceed to the most intriguing case, which was depicted in Figure 6.3 in the previous chapter. Is it possible to explain the occurrence of the definite y -spin after the joining of the two partial beams having z -spin up and z -spin down respectively?

The first thing to notice is that the prediction that rejoining the two beams will yield the y -spin up state is true only under very specific conditions. The first SG-apparatus produced a B_z -field, and we have just seen that such a field causes a precession of the spin component in the xy -plane. The expectation value of the y component then equals $\hbar/2 \cos \omega t$, hence the y component is not a constant of the motion. However,

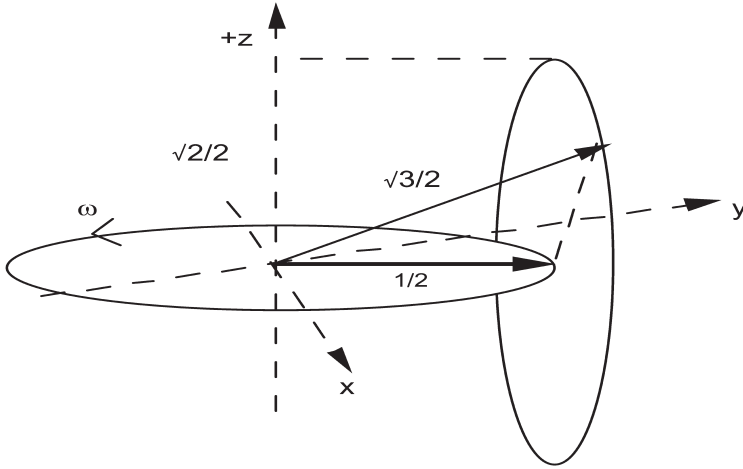


Figure 7.2 Vector model of a S^+ eigenstate rotating around a B_z -field

as we saw, the observable $S^+ = S_x \sin \omega t + S_y \cos \omega t$ has the constant value $1/2$. Hence, we must say that generally the y spin up state will not be measured. All we can say is there is a direction in the xy -plane and a point of time such that all objects have spin up.²

Next we must consider the complication that as the SG-apparatus necessarily has an inhomogeneous magnetic field, the two parts of the wave might be subject to different magnetic field strengths. How does this affect the possibility of finding a direction into which the spin component is predetermined?

Recalling that the angular velocity ω depends on the $\mathbf{B}$ -field we can write the time dependent wave function as

$$\Psi_y(t) = \frac{i}{\sqrt{2}} \begin{pmatrix} 1 \\ 0 \end{pmatrix} \exp\left(\frac{-i\omega_+ t}{2}\right) + \frac{-i}{\sqrt{2}} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \exp\left(\frac{i\omega_- t}{2}\right) \quad (7.12)$$

where ω_+ and ω_- represent the angular velocity in the two beams respectively. The expectation value for the spin- y component is

$$\langle S_y \rangle = \frac{\hbar}{2} \cos \frac{1}{2}(\omega_+ + \omega_-)t \quad (7.13)$$

Just as anticipated, the expectation value of the y -component precesses with an angular velocity which is the sum of the angular velocities in the two wave components. The situation is essentially similar to that of the homogeneous case. The angular frequency of the expectation value is the sum of the angular frequencies of the components of the wave function.

² However, Dr John Rundgren at Royal Institute of Technology in Stockholm has pointed out to me that due to fuzzy boundaries of the magnetic field in the Stern–Gerlach apparatus this might be practically impossible.

But what is real in the two beams? Is it the z spin up/down or the spin along a direction in the xy -plane? I propose the following interpretation: just as in the two-slit experiment, the splitting of the beam into two parts is a splitting of each single object into two spatially separated parts. Only together do these parts make up one object with definite properties. No property, actual or potential, can be ascribed separately to any of its parts. Further, as the wave function did not change irreversibly when passing through the SG-apparatus, the object (both parts together!) still has the actual property of spin half in the direction defined by $x\cos\omega t + y\sin\omega t$. This explains the possibility of joining the two parts together and measuring a definite spin in the xy -plane.

It is illuminating to compare spin measurements with double-slit experiments. In the double-slit experiment as well as in the spin detection experiment the particles are detected at places behind some equipment. Further, in both cases it is commonly assumed, first that the objects undergo a definite change when passing the equipment (the double slit or the SG apparatus) and, secondly, that they follow a definite path to one of the detectors. These assumptions are equally wrong in both cases. Objects are waves that show particle properties in irreversible interactions with other objects. But the object's passage through an inhomogeneous magnetic field is not an irreversible interaction: energy is not exchanged and the influence of one SG-apparatus can be erased by another one as we have seen. Neither is momentum exchanged: the two parts of the wave gain opposite momenta during the passage of the SG-apparatus, i.e. the total wave does not change momentum. Hence irreversible interaction occurs not until the object reaches the detecting equipment (or, in the case with the stopping screen in one of the paths, when one part of the object hits this stopping screen). It follows that after the passage of a SG-apparatus the object is still the same wave, albeit separated into two spatially distinct parts. It should be remembered that the principle of individuation requires that the resulting individuals can be predicated definite properties of their own, as was argued in Chapter 3. As already remarked, individuation of objects in terms of position presupposes that they are impenetrable, hence such a principle of individuation is ruled out in micro physics.

This comprehensive analysis of spin measurements has, as far as I know, not been proposed earlier, although the vector model is often used to describe spin properties. But no author has, as far as I know, assumed that the vector model can be used as a complete explanation. For example A. Rae (1986, p. 149) writes:

It is tempting to conclude from this vector model that the angular momentum vector of the particle precesses about the field direction with this angular velocity, which is what would happen in the similar classical situation. But it is important not to pursue this analogy too far. In classical precession, the direction of the angular-momentum vector, and hence the magnitudes of all three of its components, always have known values; but in quantum mechanics only one angular-momentum component can be measured at any given time. For example the precession model would imply that the y component of the angular momentum would be zero at the times when φ is zero, but we know from quantum mechanics and experiment that a measurement of this quantity always yields a result equal to either plus or minus $\hbar/2$ and never zero. The application of this precession model to a quantum mechanical system therefore attributes properties to the system that are additional to those predicted by quantum mechanics.

Rae is of course right in claiming that a vector model by itself does not fit quantum phenomena. What is needed in addition is postulating that spin is quantised, i.e. that a spin measurement (and more generally that all irreversible interactions, whether measurements or not) always yields discrete values. When combining three ideas, namely spin particles being waves, the vector model and quantisation of spin, we get an explanation of quantum mechanical spin.

7.4 The Kochen–Specker Theorem

The Kochen–Specker theorem is an impossibility theorem; it says that if the dimension of the Hilbert space is greater than two, it is impossible to consistently assign definite spin values to one and the same state vector described in different choices of coordinate systems. This is often used to justify the claim that spin values are not real properties of the objects. This conclusion is completely wrong. In order to show that, let us begin by stating the theorem more precisely. In order to do that we need the concept of *value function*. This is a function which, to each physical quantity A and each state Ψ , assigns a unique value $V_\Psi(A)$ of this quantity. In a measurement of the second kind, i.e. where the wave function for the measured object is an eigenfunction to the applied operator, the value of the quantity represented by the operator is simply its eigenvalue. However, in a measurement of the first kind it is not obvious how to assign a value to the measured property. As Isham (1995, p. 162) puts it: ‘Nothing useful can be said about the existence of quantum value-functions without postulating further properties for them.’

The most natural condition to impose is to require that for any real valued function F

$$V_\Psi(F(A)) = F(V_\Psi(A)) \quad (7.14)$$

This means that the value of a function of a physical quantity equals the function evaluated on the value of the quantity, in any given state Ψ . However, this assumption leads to contradictions, if the dimension of the Hilbert space is greater than two. This is Kochen–Specker’s theorem (Kochen & Specker, 1967; discussed in Redhead, 1985, ch. 5):

KS-theorem: There does not exist any value function V_Ψ if $\dim(H) > 2$.

An informal proof can be given as follows. The general equation for the total angular momentum $\mathbf{J} = \mathbf{L} + \mathbf{S}$ is:

$$\mathbf{J}^2 = \mathbf{J}_x^2 + \mathbf{J}_y^2 + \mathbf{J}_z^2 \quad (7.15)$$

For a spin 1-system in the $l = 0$ state this yields $j = 1$ and the operator $\mathbf{J}^2$ then has the eigenvalue two (in units of $\hbar$). Then the eigenvalue for two of the operators $\mathbf{J}_x^2$, $\mathbf{J}_y^2$ and $\mathbf{J}_z^2$ must be equal to one, and for the third equal to zero. Hence it is certain that a measurement of the total angular momentum in three orthogonal directions (which is possible to do simultaneously because the operators commute) will yield one value

which is zero and two which are equal to unity. This is true for any choice of directions for the orthogonal unit vectors x, y and z . Now, suppose we tilt our coordinate system slightly, getting new directions x', y' and z' . The angular momentum equation is still valid; it does not depend on our choice of coordinate system. Hence one of the operators $\mathbf{J}_{x'}^2, \mathbf{J}_{y'}^2$ and $\mathbf{J}_{z'}^2$ has the eigenvalue equal to zero, i.e. a measurement of these will yield one value equal to zero. Kochen and Specker were now able to show that if two vectors, (one from each coordinate system) differ by an angle less than $\arcsin(1/3)$ (approximately 19°) and the angular momentum is zero in one of these directions, it follows that the angular momentum in the other direction also is zero. But five consecutive rotations of 18° ($\approx 19^\circ$) around, for example, the z -axis will make the x' -axis take the position of the y -axis. Hence, two orthogonal directions must both have their angular momentum equal to zero. That however contradicts equation (7.15) and some assumption must be wrong. Which one?

7.5 The KS theorem does not contradict realism

The step not justified by quantum mechanics is the simultaneous assignment of definite values to angular momentum observables in different coordinate systems. If one coordinate system (x, y, z) is chosen and the angular momentum observables J_x^2, J_y^2 and J_z^2 relative to this system are assigned definite values, this means that we have performed a spin measurement. That in turn means that there has been an interaction between an external magnetic field and the spin object. The KS theorem tells us that we cannot generally use functions of these values to calculate values of other observables defined with respect to other coordinate systems, realised as external fields. Why is that impossible? The reason is that these other observables do not have definite values simultaneously with the observables J_x^2, J_y^2 and J_z^2 and this fact in turn is nicely explained by the rotating vector model. Assigning definite values to a set of angular momentum observables can only be done when the system is forced to pass through an external magnetic field in which the field vector defines the z -axis of the coordinate system. Having done this, the system is forced to rotate around an axis parallel to this field. The total angular momentum vector $\mathbf{J}$ for a spin 1-system must then be orthogonal to the axis of rotation because J_x^2, J_y^2 and J_z^2 represents the square of the lengths of the projections of the rotating vector and as one is zero and two have equal value, ($=1$), this is only possible if the projection on the axis of rotation is zero.

The system cannot possibly rotate around two different axes and if it passes a magnetic field it is forced to rotate around an axis parallel to that field. It is immediately clear from the rotating vector model that the squared components of angular momentum must have the predicted values. Only the angular momentum values determined in a particular experiment, i.e. relative to the directions determined by the magnetic field used for the measurement, have definite values. Hence, quantum systems having passed the defining magnetic field do not have definite values of observables defined with respect to another external field, i.e. to another coordinate system.

Another way of analysing the fault in the argument for the conclusion that two orthogonal vector components both have spin value zero, is by observing the use of the term 'coordinate system'. Usually it means just a mathematical tool enabling us to attribute coordinates to observed entities and the choice of directions in a coordinate system is usually held to be insignificant. However, in the argument above, the directions are determined empirically, by reference to a magnetic field. The apparent innocence in the argument comes from the pure mathematical sense of 'coordinate system', whereas the fatal conclusion depends on the implicit physical premises.

The crucial feature of the rotating vector model is that a measurement of the first kind is not just a passive registration of a pre-existing value but rather a process where one among a number of potentialities is actualised. Then, the seemingly natural assumption (7.14) is far from natural and its rejection can be motivated by the rotating vector model.

Chapter 8

Non-locality

8.1 Introduction

A main topic in the discussion about the foundations of quantum mechanics is the conundrums associated with EPR phenomena, Bell's inequalities and non-locality. In this discussion one can identify four fundamental questions:

- (1) Under what conditions can we derive a Bell-like inequality which conflicts with quantum mechanics?
- (2) Which empirical evidence do we have for the correctness of quantum mechanics and against Bell's inequalities?
- (3) Is violation of the locality condition an inevitable consequence of violation of Bell's inequalities?
- (4) Can we explain non-local correlations?

Of these four questions the last one has evoked surprisingly little interest; perhaps because most philosophers interested in non-locality believe that there must be a loophole in the argument leading to non-locality, which has to be detected, whereas physicists are more concerned with checking whether quantum mechanics or Bell's inequality fits their empirical data. But it now seems to me that these questions are answered: there are no loopholes in the derivation of non-locality from the fundamental postulates of quantum mechanics, and Bell's inequalities do not fit the facts. It is clear that as far as non-local phenomena are concerned, quantum mechanics is a correct theory and we must accept that nature shows non-local features. Therefore, the most pressing philosophical task is to try to explain these non-local correlations. This task could be stated a bit more precisely in the form of two related questions:

- (i) How is it possible that two objects can affect each other without anything at all transmitting the influence between them?
- (ii) How is it possible that the states of two distant objects can be instantly correlated?

These two questions are asked against two presuppositions which appear almost trivial. These presuppositions are by Redhead (1987, p. 75) labelled *Bell locality* and *Einstein locality* respectively. Bell locality is the assumption that if the states of two distinct objects are correlated and the correlation is not caused by a common cause, then some real physical thing must transmit the influence between them. Mere empty space cannot have this capacity.

Einstein locality is the assumption that if the states of two distinct objects are thus correlated without a common cause, the influence from one to the other cannot go faster than the speed of light. Quantum mechanics appears to violate both these conditions.

Bell himself used to illustrate the paradoxical nature of non-local correlation using a little story that goes something like this.

Assume that you have two colleagues who you meet every morning. You observe that the colours of their socks are peculiarly correlated; whatever colour A's socks have, B wears socks with the complementary colour. This occurs without exception. Surely, you think, they must somehow communicate each morning before leaving home. However, this possibility can be rejected, a thorough investigation reveals that they have no physical means for such communication. Then you think that they must have met once and agreed on a complicated rule governing the choice of socks. As long as they both follow this rule, the colour of their socks is correlated. However, this possibility can also be rejected, for example by studying the sequence of colours on one of the colleagues, and this sequence appears to be completely random (This conclusion can of course never be conclusively proved from observing a limited number of cases, but let us allow for inductive inferences.) Now, almost everyone would say that such a situation simply is not possible; either they communicate or they follow an agreed rule, however complicated. This conclusion is an application of Reichenbach's principle that every correlation has a causal explanation in terms either of direct causal links between the correlated events or in terms of common causes.

This little story is easy to translate into quantum language if we replace 'socks' with 'electrons' and 'colours' with 'spins'. If we prepare an experiment in order to determine the spins of pairs of electrons in singlet states, we will find precisely the astonishing result that for each electron the sequence of spin detections is random (although the probability distribution is determinate), but the spins of each pair of electrons are perfectly anti-correlated, none of the known forces connect the electrons, and no common cause is possible.

Many would be inclined to say it is simply not possible that two physical systems at whatever distance are perfectly anti-correlated without any common cause or any physical connection between them. But such a stance is a piece of a priori metaphysics and other people with more empiricist inclination disapprove of the use of a priori statements in the natural sciences, especially if these are incompatible with an established theory. Instead they just accept non-local connections as a brute fact. Although I certainly feel the strength of this empiricist argument, I cannot give up wanting some sort of explanation of non-local connections. Good or bad, that is the way I feel and, evidently, I am not alone. If people generally took the empiricist stance and accepted non-locality as a brute fact, there would hardly be any particular interest in EPR phenomena and Bell's inequalities. But as there is great deal of interest in this topic, people are obviously not prepared to accept non-locality as a brute fact.

In fact, most have thought that non-locality is an artefact of an incomplete theory and consequently have tried to find loopholes in the argument leading up to non-locality. It seems to me, however, that it is time to give up these efforts and accept

that non-local correlations are real traits of nature. Therefore we should focus our efforts on the conflict between non-local correlations and our metaphysics. I think that this conflict can be dissolved, i.e. that there is an explanation of non-locality to be found.

There are at least two other phenomena that exhibit a kind non-locality. One is the Aharonov–Bohm effect, where the motion of charged particles is influenced by fields, through which they could not have possibly passed. The other non-local type of phenomenon is the interaction between atoms via the electromagnetic field, the non-local character of which has been discussed in several papers by Gerhard Hegerfeldt (1974, 1994) in Göttingen. The meaning of ‘non-locality’ is however a bit different in the different contexts (EPR phenomena, the AB-effect and the case discussed by Hegerfeldt) but there is still reason to use the common label ‘non-locality’. Hence, I will use ‘non-locality’ in a slightly more general sense; the common denominator for the events so described is the fact that although they appear to be strictly localised to well defined places they have consequences at distant places without there being any transmitting force. First I will discuss non-local correlations which occur in EPR-type experiments. Before giving my explanation I will rehearse the argument that EPR non-locality is a consequence of quantum mechanics.

8.2 Derivation of non-locality

Let us adopt the following shorthand expressions:

QM: The set of all statements derivable from the axioms of quantum mechanics.

100% corr.: In a singlet state made up of two spin-half particles their observed spin values in the same direction are perfectly anticorrelated – if one is found to have the value ‘up’ in a chosen direction, the other has the value ‘down’ in the same direction, and vice versa.

Determinism: The state of a physical system at any given time is completely determined by laws and initial conditions, which is to say that upon conditionalising on all factors, known or unknown, all probabilities for measurement outcomes collapse into zero or one.

Locality: Elements of reality (values of observables) pertaining to one system cannot be affected by measurements performed ‘at a distance’ on another system. This definition is adapted from Redhead (1987, p. 75). For some refinements see Section 8.3 below.

Using these definitions we have the following argument showing that quantum mechanics is a non-local theory:

- | | |
|---|------------------|
| 1. $QM \rightarrow \{100\% \text{ corr.}\}$ | Premise 1 |
| 2. $\{100\% \text{ corr.}\} \wedge \text{Locality} \rightarrow \text{Determinism}$ | Premise 2 |
| 3. $QM \wedge \text{Locality} \rightarrow \text{Determinism}$ | (by 1 and 2) |
| 4. $QM \rightarrow \neg \text{Locality} \vee \text{Determinism}$ | (by 3) |
| 5. $\text{Locality} \wedge \text{Determinism} \rightarrow \text{Bell's inequality}$ | (Bell's theorem) |
| 6. $QM \rightarrow \neg \text{Bell's inequality}$ | Premise 3 |

7. $QM \rightarrow (\neg \text{Locality}) \vee (\neg \text{Determinism})$ (by 5 and 6)
 8. $QM \rightarrow \neg \text{Locality}$ (by 4 and 7)

There are four assumptions in the argument, namely premises 1, 2, 3 and Bell's theorem. The arguments for premises 1 and 2 will be discussed in this section and in the next one I will give a proof of Bell's theorem and show that premise 3 is true.

Premise 1 is a straightforward consequence of the wave function for a singlet state made up of two spin particles:

$$\Psi_{\text{tot}} = \frac{1}{\sqrt{2}}(|up\rangle \otimes |down\rangle - |down\rangle \otimes |up\rangle) \quad (8.1)$$

where the first ket in each term refers to the spin state of the first electron and the second ket to the other electron. This function is an application of the superposition rule together with the rule for joining two wave functions when systems have been in interaction. Upon measurement the wave function will collapse to one of its components, which is to say that for every pair of particles making up a singlet state, their spin will be found to point in opposite directions without exception. This is a straightforward consequence of quantum mechanics.

Premise 2 is quite well-known and is given in a slightly different form in the original EPR-paper (Einstein, Podolsky & Rosen, 1935) and later rehearsed many times, for example by Redhead (1987). The original EPR argument goes like this: If the states of two systems are perfectly correlated and we measure the state of one of the systems, we can predict with complete certainty the state of the other system. (The word 'system' should here be understood as 'single particle'.) If this measurement does not in any way influence the state of the other system (locality), it must have been in the predicted state before the measurement on the first system was performed. But before the measurement the quantum formalism cannot provide this information, hence the theory is incomplete.

For our purposes it is, however, determinism and not completeness that is the issue. Thus, when we perform a measurement on the first particle and predict the state of the other particle, the predicted state must have been obtained some time before the prediction (because nothing happened to that system according the locality assumption), which is to say that it was determined. Another way of deriving this conclusion is by observing that a stochastic local hidden variable theory conflicts with quantum mechanics, which can be seen from the following consideration. It is straightforward to infer from the state function (8.1) that the conditional probability $\text{prob}(\text{particle 2 has spin up} / \text{particle 1 has spin down}) = 1$. Hence, the addition of a stochastic hidden variable $\lambda = \lambda_i$ to the formalism cannot change the conditional probability that the particle 2 has spin up. Thus, $\text{prob}(\text{particle 2 has spin up} / \text{particle 1 has spin down} \cdot \lambda_i) = \text{prob}(\text{particle 2 has spin up} / \text{particle 1 has spin down}) = 1$. Similarly, $\text{prob}(\text{particle 1 has spin up} / \text{particle 2 has spin up} \cdot \lambda_i) = \text{prob}(\text{particle 1 has spin up} / \text{particle 2 has spin up}) = 0$. This shows that upon conditionalising on the result of a measurement of the other particle we get probabilities of zero or one independently of the value of any hidden variables, if such things exist, which is to say that we have a deterministic theory, if locality is assumed. Let us not forget

that the states ‘spin up’ and ‘spin down’ refer to one and the same randomly chosen direction (Redhead, 1987, p. 102).

8.3 Bell’s theorem

Suppose we perform a series of measurements on pairs of singlet state particles with half integer spin, such as electrons. Each measurement will return the value +1 or –1 in units of $\hbar/2$. Let us assume determinism and locality.

The measurement on the left wing electron can be performed in two alternative directions denoted a and a' , and the two possible spin directions on the right wing electron is denoted b and b' . Let a_n be the spin value in the a direction of the n th left particle, and correspondingly for b_n . The product $a_n b_n$ will have either the value +1 or -1. We can now define a correlation function

$$C(a, b) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N a_n b_n \quad (8.2)$$

$C(a, b) = 1$ means that the correlation is perfect and the spin values are the same, whereas $C(a, b) = -1$ means that the spin values always are opposite.

Let us further define the following function

$$g_n = a_n b_n + a_n b'_n + a'_n b_n - a'_n b'_n \quad (8.3)$$

As $(a_n)^2 = (a'_n)^2 = (b_n)^2 = (b'_n)^2 = 1$, it follows that the last term equals the product of the first three. Therefore, we need only to investigate all possible combinations of the first three terms in order to decide the possible values of g_n . The result is displayed in Table 8.1.

Table 8.1 The possible values of the function g_n

$a_n b_n$	$a_n b'_n$	$a'_n b_n$	$a'_n b'_n$	g_n
1	1	1	1	2
1	1	-1	-1	2
1	-1	-1	1	-2
-1	1	1	-1	2
-1	-1	1	1	-2
-1	1	-1	1	-2
-1	-1	-1	-1	-2

Obviously, the mean value of the function g_n as we let $N \rightarrow \infty$ must be a number between –2 and 2. It then follows that

$$\left| C(a,b) + C(a,b') + C(a',b) - C(a',b') \right| \leq 2 \quad (8.4)$$

which is one version of Bell's inequality (adapted from Isham, 1995, pp. 181–184).

Premise 3. I will now perform a quantum calculation of the value of the left-hand side of equation (8.4), which will show that Bell's inequality conflicts with quantum mechanics.

Suppose we measure S_1 , the spin of the left particle, along the z -axis and the spin of the other particle along an axis with an inclination of ϕ degrees from the z axis. The spin matrix for this latter observable is

$$\mathbf{S}_2 = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \cos \phi + \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \sin \phi = \begin{pmatrix} \cos \phi & \sin \phi \\ \sin \phi & -\cos \phi \end{pmatrix} \quad (8.5)$$

There is no loss of generality by assuming one of the axes to be the z -axis because we can always rotate the system by a unitary operation. The quantum mechanical correlation function is the expectation value of the two spin operators $\mathbf{S}_1 = \hbar/2\boldsymbol{\sigma}_z$ operating on the first particle and $\mathbf{S}_2$ operating on the second particle

$$C(a,b) = \left(\frac{2}{\hbar} \right)^2 \langle \Psi | \mathbf{a} \mathbf{S}_1 \otimes \mathbf{b} \mathbf{S}_2 | \Psi \rangle \quad (8.6)$$

which in matrix notation is

$$\begin{aligned} C(a,b) &= \left\langle \begin{pmatrix} 1 \\ 0 \end{pmatrix} \otimes \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right| \left(\begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \otimes \left(\begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \cos \phi + \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \sin \phi \right) \right| \begin{pmatrix} 1 \\ 0 \end{pmatrix} \otimes \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\rangle = \\ &= \left\langle \begin{pmatrix} 1 \\ 0 \end{pmatrix} \otimes \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right| \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} \otimes \begin{pmatrix} \cos \phi & \sin \phi \\ \sin \phi & -\cos \phi \end{pmatrix} \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right\rangle = \\ &= \left\langle \begin{pmatrix} 1 \\ 0 \end{pmatrix} \otimes \begin{pmatrix} 0 \\ 1 \end{pmatrix} \right| \begin{pmatrix} 1 \\ 0 \end{pmatrix} \otimes \begin{pmatrix} \sin \phi \\ -\cos \phi \end{pmatrix} \right\rangle = -\cos \phi \end{aligned} \quad (8.7)$$

Let us now calculate the quantum mechanical version of the function g_n (equation (8.3)) with the following assumptions: a and b are parallel, a , a' , b and b' are all in the same plane and the angles between a and b' and between a' and b are equal, namely ϕ . Then we have

$$\left| C(a,b) + C(a,b') + C(a',b) - C(a',b') \right| = \left| 1 + \cos \phi + \cos \phi - \cos 2\phi \right| \quad (8.8)$$

and this quantity exceeds 2 if $0 < \phi < 90^\circ$, which is in plain contradiction with Bell's inequality. The quantum mechanical correlation function is given in Figure 8.1.

One possible way to explain non-local correlations is that settings of the measuring devices might influence (by an hitherto unknown mechanism) the state of the singlet pair when produced in the preparation site so as to adjust the states of two particles to

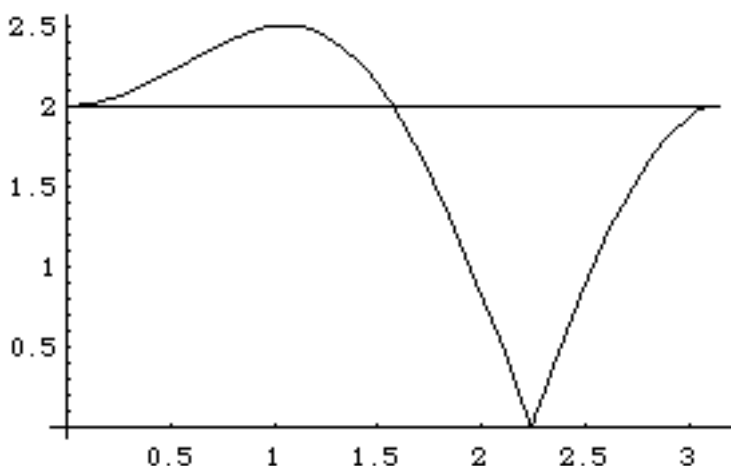


Figure 8.1. The function $|1 + \cos \phi + \cos \phi - \cos 2\phi|$. There is a contradiction between Bell's inequality and quantum mechanics when the correlation function is above the line $y = 2$

the settings. This possibility is discussed and tested by Aspect and his collaborators (Aspect *et al.*, 1981, 1982a, 1982b). They performed their experiment in such a way that the settings of the measuring instrument are made *after* the two particles (in their case photons) have left the preparation site and thus are separated. Hence, no causal influence from the settings of the detectors to the particles being in the preparation site is possible. Non-local correlations still are observed, which means that there is little hope of explaining the correlations as a result of some influence from the measuring devices.

8.4 Bell locality and Einstein locality

The notion of locality has been explicated in a number of different ways. In this section it is suitable to use Redheads formulation of the locality principle:

PL: *The value of an observable obtained by a measurement on one particle is not affected by the measurement which was made on another spatially distant particle.*¹

Of course, two objects can affect each other at great distance by exchanging field quanta of any of the fundamental forces. That goes without saying, so the locality principle should be read as saying that if there are no forces present, then two objects cannot affect each other. Moreover, interaction mediated by forces is propagated, at most, at the speed of light. Hence the locality condition has two components: (1)

1 Bell cited Einstein who formulated the locality condition thus: 'But one supposition we should, in my opinion, absolutely hold fast: the real factual situation of the system S2 is independent of what is done with the system S1 which is spatially separated from the former' (Einstein in Schilpp, 1951, p. 85).

one object can affect another one only if there is a fundamental force connecting them, called the *Bell locality* by Redhead, and (2) the connection is transmitted, at most, at the speed of light, called the *Einstein locality*. It is logically possible that a connection between two objects violates Bell locality without violating Einstein locality and vice versa. But quantum mechanics violates both Bell and Einstein localities. In what follows I will first concentrate on violations of the Bell locality and propose an explanation of this phenomenon. In Section 8.6.2, I will discuss the violations of Einstein locality, the conclusion being that the violation of this aspect of locality does not contradict relativity theory.

8.5 Interference of wave packets

Let us start with the assumption, already argued for, that quantum objects are wave packets. The component waves in the wave packet are much more extended than the composite wave packet, which implies that two such wave packets will have component waves that overlap even if their centres are very far apart. But mere overlap is not sufficient. In order to account for non-local correlation we must also explain why these overlapping wave packets interact with each other.

The crucial thing is that a singlet state wave function, which entails non-local correlations of the EPR type, is not separable into a sum of two one-particle wave functions. It follows that we cannot predicate definite values of spin to one of the particles independently of the other. This means, as was shown in Chapter 3, that criteria of identity are not fulfilled and we do not have two individual objects, but one. We can only attribute definite spin (helicity in the case of photons) to the total system.

The sceptic might argue that it is rather a matter of words whether we say that the system consists of one or two objects. They are mistaken: this distinction is imperative. If we have two objects we can exchange energy (or spin, or momentum or any conserved quantity) with one of them without affecting the other, whereas that is impossible if the purportedly different objects really are two parts of one and the same object. Saying that a system is made up of, for example, two electrons does not necessarily mean that we have two individual objects, only that the system contains more than one quantum of the relevant kind (cf. the discussion of individuation of objects in Chapter 3).

Consider a pair of particles, i.e. two wave packets, in a singlet state. Being in the singlet state implies that the partial waves of the two objects are phase correlated. It follows that if two such constituent waves, one from each wave packet, have the same wave numbers, or nearly so, they will interfere. It means that we cannot interact with only one of the wave packets without interacting with the other.

In Chapter 4, I referred to the experiments by Rauch *et al.*, which convincingly showed that a particle, in their case a neutron, does have a partial wave structure, and also that its partial waves have extensions far beyond the ‘bulk’ of the wave packet. There are reasons to suppose that this is true for all kinds of quantum objects, not just neutrons.

A singlet state is thus better described as constituted of two phase-correlated wave packets, not as two point particles. No matter how far apart the centres of these wave packets are situated, their partial waves are still overlapping and the tails of one wave packet is present at the same place as the centre of the other wave packet. This view of the matter removes the weirdness of non-local correlations, since it entails that when we perform an experiment on one of the constituent wave packets at one place, we are in fact also interacting with the other. No matter the distance between the centres of mass, the two parts in the singlet pair are never effectively spatially isolated.

All material objects are, according to this picture, a kind of wave and as such more or less spread out in space. But why, then, should the presence of *some* partial waves have significance and not some others, for example those belonging to macroscopic devices in the laboratory? The answer is of course phase correlation.

I have claimed that a quantum system in a singlet state is one object because its parts cannot be attributed properties individually and independently of each other. This means that the probabilities attributed to the two-particle system cannot be separated into two independent one-particle probabilities. This means in turn that there are interference terms in the wave function description that do not cancel out when integrated over a time interval of some length. In order to achieve just that, the partial waves in the two wave packets – and thus the entire wave packet – must be phase correlated. This is the physical effect of producing a singlet state pair of two ‘particles’. The presence of a multitude of other waves (constituting other objects) is of no importance because their contribution cancels out effectively when integrated over even a very short time.

A mathematical analysis of the situation strengthens this informal argument. Assume two wave packets travelling in opposite directions with velocities

$$\pm v = \frac{\hbar k_0}{m}$$

and having a Gaussian envelope. (This last assumption is chosen only for convenience of calculation.) The velocity v is the group velocity of the wave packet v_g . According to equation (4.23), the phase velocities have opposite signs in the two wave packets when they move in opposite directions. Thus, the total wave function can be written

$$\begin{aligned} \Psi_{\text{tot}} = & \int_{-\infty}^{+\infty} \exp\left(-\frac{(k-k_0)^2}{a}\right) \exp[ik(x-v_p t)] dk \\ & + \int_{-\infty}^{+\infty} \exp\left(-\frac{(k+k_0)^2}{a}\right) \exp[ik(x+v_p t)] dk \end{aligned} \quad (8.9)$$

At the time t_0 the two wave packets are situated at positions $x = \frac{\omega t_0}{k}$ and $x = -\frac{\omega t_0}{k}$

with the wave vectors peaked at k_0 and $-k_0$ respectively. In order to calculate this sum we first reverse the k -axis in the second integral. This gives us the same Gaussian in both integrals and the result is

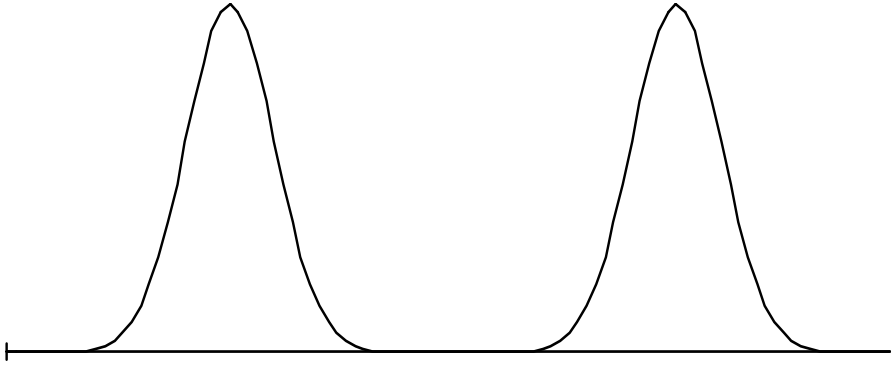


Figure 8.2 Two wave packets travelling in opposite directions

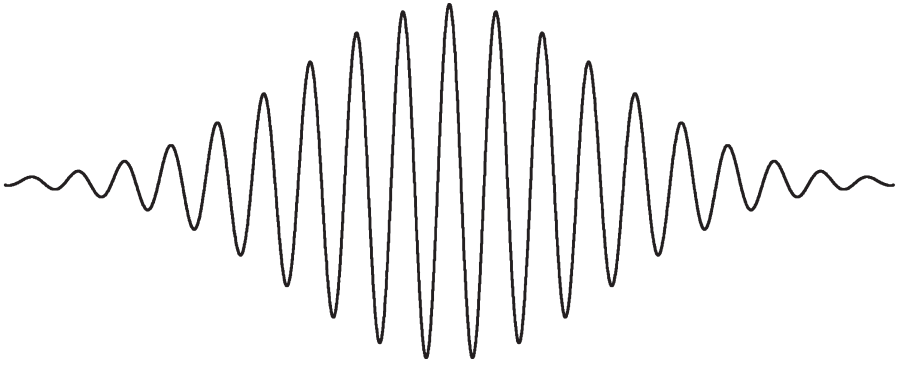


Figure 8.3 The function $\phi(k) = \exp\left(-\frac{(k-k_0)^2}{a}\right) 2 \cos kx$ for given $x \neq 0$.

This figure represents two wave packets as seen in k -space. The scale is arbitrary.

$$\begin{aligned} \Psi(x,t) &= \int_{-\infty}^{\infty} \exp\left(-\frac{(k-k_0)^2}{a}\right) \left[\exp[i(k(x-v_p t))] + \exp[i(-k(x+v_p t))] \right] dk = \\ &= \int_{-\infty}^{\infty} \exp\left(-\frac{(k-k_0)^2}{a}\right) \left[\exp[i(kx - kv_p t)] + \exp[-i(kx + kv_p t)] \right] dk = \quad (8.10) \\ &= \int_{-\infty}^{\infty} \exp\left(-\frac{(k-k_0)^2}{a}\right) \exp(ikv_p t) 2 \cos(kx) dk \end{aligned}$$

This is a Fourier expansion in the function series $\exp(ikv_p t)$ and the Gaussian coefficients are modulated by $\cos kx$:

$$\phi(k) = \exp\left(-\frac{(k - k_0)^2}{a}\right) 2 \cos(kx) \quad (8.11)$$

We have thus in k -space a typical interference pattern as is shown in Figure 8.3.

The conclusion is evident: two separated wave packets whose bulks do not overlap at all in ordinary space, nevertheless overlap and can interfere in momentum space. It means that the moment of two spatially separated objects influence each other and this could only be possible if the individual waves in each of the wave packets are infinitely extended in ordinary space and phase correlated.

This shows that our perplexity when confronted with non-local correlations depends on an unnoticed assumption, which is in fact false. When we ask, ‘how can two objects at two different places be correlated without anything at all bringing about the contact between them?’, we assume that the correlated objects really are well separated from each other. This is not true; these objects are made of partial waves that can be indefinitely extended in space and are thus in direct contact with each other, no matter the distance between the centres of the wave packets. And if these two waves are phase correlated they interact with the environment in certain situations as one single object. If we insist on talking about two objects (but as I have stressed several times, quantum ‘particles’ are not individuals), we could describe the situation in the following way: when we measure the spin state of one object we cannot avoid interacting with the other object, and this is so because neither object is confined to distinct and separate regions in space. As has already been argued, the question ‘is this system made up of one or two objects’ has no general answer because the individuation of objects is relative to the applied predicate. If we measure, for example, the spin angular momentum there is one object, because only the particle pair as a whole, not each particle, has a well-defined spin angular momentum, namely zero. But if we measure other quantities, such as charge, we have two objects; each particle can be attributed a definite charge independently of the other one.

This way of answering the ontological question ‘how many things are there?’ is an application of Quine’s doctrine that individuation of a domain of discourse into different objects depends on the individuating principles connected to our predicates (cf. ‘Any coherent general term has its own principle of individuation, its own criterion of identity among its denotata’ (Quine, 1981, p. 12)).

Most philosophers in the area use the expression ‘entangled state’ as the label for the situation where the state of a single quantum ‘particle’ depend on other ‘particles’. It appears to me better to change individuation instead, especially since there are independent arguments against viewing quantum ‘particles’ as individual objects. So in my view there are no entangled states, simply because individuation and state attribution are geared to each other.

Waves do not interfere unless they are coherent. Preparation of a singlet state by necessity produces coherent waves. This is a consequence of conservation of angular momentum. If the total angular momentum in a given direction is to be

conserved in a system made up of two parts, and if this property cannot be attributed to these two parts separately, then the total system must react as one single object when exchanging momentum with the surroundings. Hence, interference in spin measurements is the consequence of conservation of angular momentum and the fact that a conserved quantity cannot be divided between two constituent parts of an object.

8.6 Einstein locality and relativity theory

It is often assumed that Einstein's Special Theory of Relativity (STR) implies Einstein locality. However, that is correct only if one adopts a certain interpretation of relativity theory, an interpretation that by no means is mandatory. In fact there are good reasons for adopting an alternative interpretation that does not imply Einstein locality. Here is my view on the matter.

8.6.1 The notion of frame of reference and STR

The special theory of relativity is based on two postulates:

1. The laws of physics are the same in all inertial frames of reference.
2. The speed of light in free space, c , is the same in all systems.

These postulates are firmly believed by any competent physicist. It follows immediately that any interpretation of quantum mechanics must comply with these postulates in order to be worthy of serious consideration.

It is important to note that any frame of reference has to be anchored in the physical world. Objects have to be chosen as reference points and as standards for measurements of time and spatial intervals. At the very least we need one object in order to be able to apply the concepts of position and velocity. If we want to use the concept of angular momentum we need either four objects lacking internal structure, or an object with structure enough, in order to be able to define directions. Furthermore, it is not enough that the reference object is identifiable once only; as we are talking about motions it must be possible to identify and re-identify the reference object at all times in the time interval of interest. This implies that only material objects, not too small, can be used, which rules out photons in this context. There is a similar argument for time measurements: we need a physical object with a periodic time evolution, where the period can function as a unit of time; all objects in relativity theory can be attributed a proper time.

From the two postulates above it is easy to derive the following three relations:

$$t' = \frac{t}{\sqrt{1 - v^2/c^2}} \quad l' = l \cdot \sqrt{1 - v^2/c^2} \quad m' = \frac{m}{\sqrt{1 - v^2/c^2}} \quad (8.12)$$

where the unprimed variables are measured in a frame following the body in motion with the velocity v , and the primed variables in the frame of reference of the observer.

From these formulae it follows that no body with *rest mass* differing from zero can be accelerated to the speed of light, because it would then have infinite mass and infinite energy. However, as already remarked, we must use a body – not an empty point in space, or a point on the light wave front – in defining a frame of reference. For example, if we direct a laser beam on to the moon and pass a screen in front of the laser, the shadow on the moon should move faster than the speed of light, if you move the screen with a speed of approximately 1 m/s and if it is less than 1 metre in front of the laser. The edge of this shadow cannot be used as a reference point, even if it were sharp enough. Nor can we use a point on the light wave front, even if it could be well defined. Suppose you did, however; you could then ask, what are the rulers of length and time in that system? Clearly there are none, so our concepts of space and time intervals would not be well defined in such a frame and transformations from and to such a system are not well defined.

Our conclusion is thus: in order to use the concept of a frame of reference it must be possible to give an operational definition of a unit length and unit time interval in every such frame. (Presently, physicists have reversed the natural order between fundamental and derived units in kinematics: the velocity of light is nowadays used as a definition and the unit length is defined in terms of the distance travelled by light during a specified time interval. But this reversal is irrelevant in this context.) That means that it must be possible to have a material object resting in its own frame of reference. A part of a wave packet will not do because such a part cannot be identified and re-identified as an individual object.

8.6.2 Violation of Einstein locality and relativity theory

Now the ground is laid for discussing measurements on singlet state pairs of electrons. Suppose, as before, that quantum objects are wave-like entities that have spatial extension. Let us take the simplest possible case, one single electron moving freely. As before, a reasonable model of the object is a wave packet. Suppose further that the collapse of the wave function during measurement is a real collapse, a sudden contraction of the spatial extension of the wave packet. The wave has no individual parts with dynamics of their own, which implies that we cannot talk about different events. Hence we have no violation of relativity theory because we discuss only one event.

When we proceed to the collapse of a singlet state made up of two electrons we must distinguish between two situations: when we measure the spin of both the correlated particles, and when only one particle is measured.

As before, inducing the collapse of the entire wave by one single measurement is one event only. Hence, such an event cannot possibly be a violation of STR. But measuring the spin of the second particle is of course a separate, second event; however, it is no collapse of a wave (it has already collapsed!) and STR cannot possibly be violated.

As the state of the second particle changes at the same moment as the collapse of the entire wave takes place, this is a violation of Einstein locality. But *not* a violation of relativity theory, since identifying an individual object passing from the first measured object to its correlated partner is not possible. Both particles are in fact waves, phase correlated waves; and in spin interactions with the environment they behave as one single object, irrespective of their extension in space. The propagation of the state change from the measurement site to the rest of this system cannot be described as a chain of interactions, because the concept of interaction presupposes that we have several distinct objects that will affect each other, which we have not. The collapse is not transmitted by any object which can be used for the definition of a frame of reference. Therefore, no matter how fast the collapse takes place, no frame of reference can be defined such that it moves with the collapse and thus no violation of STR will take place.

8.6.3 Collapse of the wave function in different frames of reference

It is well-known that two events simultaneous in one frame of reference are not simultaneous for an observer in another frame which moves with a velocity $v \neq 0$ relative to the first one. Hence, if it is assumed that the measurement on a particle in an EPR experiment instantly makes the entire wave function collapse, thus simultaneously changing the state of the correlated distant particle in the same moment, the measurement on the first particle and the change of state of the correlated particle would not be simultaneous for an observer in another system. In some systems the change of state of the correlated particle would in fact take place before a measurement on the first particle is done! These consequences, however remarkable, do not lead to any difficulties.

First, let us observe that if we measure the spin of one particle only and from this result merely infer that the spin of the other particle has changed accordingly, we have only one event and no violation of relativity can occur. Secondly, let us assume that we measure the correlated particles' spin within a time interval shorter than what is needed for the transfer of a photon between the two measuring devices. In other words, these two detection measurement events are space-like separated, they are situated outside each other's light cones. As space-like separation is a Lorentz invariant property, all observers, independently of their chosen frame of reference, will judge the events space-like separated. Thus, all observers could justifiably say that the collapse of the wave function is instantaneous, if that is taken to mean space-like separated.

8.7 Collapse and phase velocity

In Chapter 4, it was argued that quantum objects are made up of wave packets in which the mean phase velocity is c^2/v , where v is the velocity of the wave packet in a given frame of reference. Then it seems reasonable to assume that the collapse, i.e. the change of the spin orientation in the case of spin measurements, is propagated from the measured particle to its correlated partner with the phase velocity. This

velocity is higher than the velocity of light in free space, but only infinite in the frame of reference in which the measured particle is at rest (cf. equation (4.23)). This is a modification of the assumption made in the preceding section. Can it be consistent with what we know about non-local correlations?

Assume that particle 1 is measured before particle 2, as seen from the laboratory frame. In the frame of reference where the wave packet is at rest the phase velocity is infinite, and seen from this frame of reference no exception from complete anti-correlation could arise. However, let us look at the situation in the laboratory rest frame, where the left-wave packet for particle 1 moves with the velocity v . Then the collapse would occur somewhat later at particle 2. Suppose that we start a clock in the laboratory rest frame when the spin state of particle 1 is measured. Then particle 2 changes its spin state at a later time $\Delta t = \Delta x v / c^2$, Δx being the distance between the two particles in the laboratory frame. Thus, if we measure the spin state of particle 2 after such a time interval, the result is certain to be the opposite spin, but if we measure the spin of particle 2 before that time, the result is not likely to be certain because the wave has not yet collapsed in that position. This would contradict the 100% correlation and conservation of angular momentum. I conclude that such a situation could not occur and some assumption must be wrong.

A possible way out is to assume that the determination of the spin state in the measuring device is not instantaneous, and is not finished until the state of distant correlated object has been changed. Thus, the collapse time, i.e. the time it takes for the detector to arrive at a definite state, depends on the distance between the correlated objects. Imagine two electrons in a spin singlet pair being very far from each other, say 10 light-years away, and the velocity of the measured electron being, say $0.1c$. The collapse of the wave function at the distant particle will occur 1 year later, and, consequently, if we want to avoid violation of conservation of angular momentum, a definite result in a detector cannot be obtained before such a time has elapsed. However, the electron had a velocity of $0.1c$ (i.e. the group velocity of the wave packet) and therefore it disappears from the vicinity of the measuring device very quickly, which means that it will not trigger the detector; the time needed to collapse the entire wave by far exceeds the time during which the particle is in detector. If the detector were a perfect one registering the presence of all objects of the relevant sort, it would be possible to observe a decrease in the detection rate in such an experiment. However, all measuring devices are far from perfect, so such a consequence would presumably pass unnoticed. But as a matter of principle it appears possible to observe a decrease in detector efficiency in proportion to the distance between the singlet pair particles. Is such an experiment feasible?

All experiments so far made are confined to distances within laboratories, which is to say that the distances are of the order of magnitude of 1 metre, in which case the time elapsed during collapse, assuming as before that the group velocity is $0.1c$, is of the order of 10^{-9} s. An effect on the detection rate could occur if the measured object has left the detector within such an interval. It means that we might be able to observe the effect if the time evolution of the overlap integral for the wave packet and the detector goes from nearly zero to maximum and to nearly zero again within a time interval of this length, for then the completion of the collapse cannot occur in shorter times than that. It means that the bulk of the wave packet needs to be

confined within a volume of less than 3 cm in diameter. I have neglected the spatial extension of the detector, which is defensible because it is only a rough estimate. It does not seem impossible to detect a decrease in detection rate depending on the distance to the correlated particle.

All this is of course very speculative. However, we have at least seen that the assumption that the collapse, as seen from the laboratory frame, takes some time, namely the time it takes for the propagation of the partial waves, does not lead to any contradiction, neither with the fundamental assumptions of the orthodox theory, nor with experiments performed at present.

Now, assume that experiments of this sort some day will be done. We would either get the result that the predicted decrease in detector efficiency depending on distance to the correlated object is confirmed or we will get a violation of predictions. If experiments accord with the predictions it would, I guess, be counted as strong support for the analysis given above. But if such a dependence were not observed, several different conclusions are possible. I would suggest that the collapse is instantaneous in the entire wave field. The collapse need not have anything to do with the phase velocity.

It seems to be an open question whether the collapse takes some time or not. It further seems possible to decide the matter by experiments. In either case, no problems for my interpretation will arise.

8.8 Violation of locality in interactions between two two-level atoms

Gerhard Hegerfeldt of University of Göttingen has, in two articles (Hegerfeldt, 1974, 1994) discussed the atomic emission-absorption process and the consequences of using quantum mechanics for the analysis. For the present discussion, the most interesting result is his discussion of the interaction between two two-level atoms at some distance. Now, we do not yet have a renormalisable quantum field theory, but, he claims, we only need to assume that the states (he means state descriptions) of the theory form a Hilbert space and that the Hamiltonians are bounded from below (i.e. that the eigenvalues to the Hamilton operator have a minimum). These are very weak conditions and physically well motivated. He is then able to show that if an atom A in an excited state de-excites, the probability for excitation of a second atom B becomes immediately non-zero, no matter the distance between the two particles. This is a new aspect of violation of Einstein locality, which will be discussed after giving some more details of the argument.

Assume that two systems A and B are situated at a distance R , that system A is in an excited state, B in its ground state and no photons are present. Call this state Ψ_0 . Let $P_B^e(t)$ be the probability of finding B excited at time t :

$$P_B^e(t) = \langle \Psi_t | \mathbf{O}_{eB} | \Psi_t \rangle \quad (8.13)$$

where $\mathbf{O}_{eB}$ is a projection operator that is bounded from below.

Then either:

- (1) the excitation probability of B is non-zero for almost all t and the set of such

- t s is dense and open, or
 (2) the excitation probability of B is identically zero for all t .

Alternative (1) means that B starts to move out of its ground state immediately, i.e. before any photon coming from A could arrive at B. This is a contradiction of Einstein locality – Hegerfeldt calls it a contradiction of ‘Einstein causality’. It is clear that we have a contradiction between alternative (1) and the conjunction of the following assumptions:

- (i) Every part of system A is separated at least by a distance R from any part of system B.
- (ii) System B cannot be excited except by taking up a photon from system A.
- (iii) Photons travel, at most, at the speed of light.

These assumptions appears indeed well motivated and so the alternative (1) above should be rejected. But alternative (2), i.e. that the probability for excitation is identically zero, is even worse, because it implies that the state of B is totally unaffected by the presence of A. This could not possibly be the case, Hegerfeldt claims, so the first alternative must after all be accepted. But how should we understand that? Hegerfeldt lists three alternatives:

- (a) assumption (i) is wrong, which means that systems localized in different and distinct regions of space might not exist as a matter of principle. Hence two systems do always overlap and immediate excitations may occur.
- (b) Renormalisation will introduce a photon cloud around each system. This implies an overlap of the two systems, which leads back to case (a).
- (c) The notion of a ground state of B in the presence of A may not make sense. Without A present one will expect a lowest energy state to exist for the system B plus a radiation field, with no real photons. However, with A present, the lowest state of the complete system may change. Thus, the ground state of B may not be preparable independently of A. This again leads back to case (a).

Hegerfeldt concludes that the only viable interpretation of this result is that the two systems are not confined to two distinct regions in space, no matter how distant they are. This is exactly the same conclusion as I have argued for as explanation of non-local correlations in EPR experiments.

8.9 The Aharonov–Bohm effect

The Aharonov–Bohm effect, hereafter to be called ‘AB-effect’, is a different kind of non-locality. It does not result in correlations and has nothing to do with conservation of momentum or energy. Thus, as will be seen, this kind of non-locality must be defined differently.

The AB-effect can be described thus (adapted from Aharonov, 1986): a beam of electrons is incident on a detecting screen. A long cylinder of metal is inserted perpendicular to the beam. In this cylinder a thin solenoid is placed. The magnetic

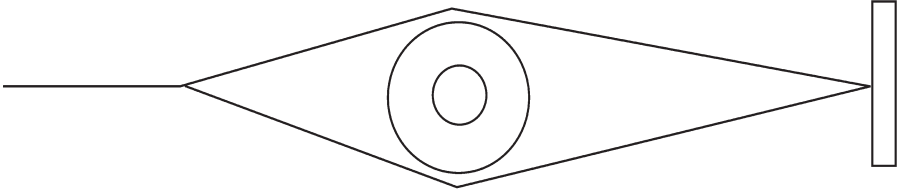


Figure 8.4 A principal set-up for testing the AB-effect

field produced by the current in the solenoid is effectively screened by the cylinder being a good conductor of electricity, so the magnetic field is zero outside the cylinder. The metallic cylinder also prohibits the electrons from entering into the interior. Hence, the electrons cannot possibly come into contact with the magnetic field coming from the current. And yet they are so affected!

Assume that the wave function for an electron splits into two spatially separated parts Ψ_1 and Ψ_2 travelling on each side of the cylinder:

$$\Psi = \Psi_1 + \Psi_2 \quad (8.14)$$

The wave function at a point on the target is

$$\Psi = \Psi_1 \exp(-iS_1/\hbar) + \Psi_2 \exp(-iS_2/\hbar) \quad (8.15)$$

where

$$S_i = -\frac{e}{c} \int A_i(x) dx, \quad i=1, 2 \quad (8.16)$$

A_i is the electromagnetic vector potential along the path i . The probability distribution for the electron is given by $|\Psi|^2$, which contains an interference term depending on the phase difference

$$(S_1 - S_2)/\hbar = \frac{e}{\hbar c} \oint A dx = \frac{e}{\hbar c} \Phi \quad (8.17)$$

Two of the well known Maxwell equations are

$$\mathbf{E} = \nabla \Phi - \frac{1}{c} \frac{\partial \mathbf{A}}{\partial T} \quad (8.18)$$

$$\mathbf{B} = \nabla \times \mathbf{A} \quad (8.19)$$

from which we can infer that both the electric and magnetic field can be zero outside the cylinder, even if $\mathbf{A}$ is non-zero there. Hence the presence of non-zero vector potential affects the probability distribution, in spite of no fields present. The effect is measurable and the peculiar thing is that the metallic cylinder shields the electric

and the magnetic field produced by the current in the wire, but it does not shield the vector potential. It is indeed remarkable.

For a long time the vector potential was considered a mathematical tool and completely lacking physical significance. The AB-effect shows that such is not the case. But surely it invites questions about the nature of the vector potential?

It is clear that we cannot exchange any conserved quantity, such as momentum or energy, between the coil inside the cylinder and the passing electrons. This is another way of saying that the electromagnetic field from the current is zero outside the cylinder, or, what amounts to the same thing, the electromagnetic field arising from the passing electrons is zero inside the cylinder. But the two objects, electron passing outside, and moving charges inside the cylinder, affect each other's motion via the vector potential. Contrary to the case of EPR correlations, this does not demand any previous phase correlated fusion of two systems into one. We can, without ambiguity, say that the current inside the cylinder and the charge passing outside the cylinder are two distinct systems. This is a profound difference between non-local effects in EPR-like experiments and in AB-experiments, respectably. The magnetic field from the current inside the cylinder affects the motion of charges outside the cylinder. It is a general effect acting on all charges and this suggests an analogy between electromagnetism and gravitation: both forces can be viewed as geometrical properties of space. More specifically the analogy goes as follows. The vector potential affects the phase of wave functions, which means that the vector potential is a description of the metric of electromagnetic space. It plays the same role as gravitation in general relativity, in which the mass density distribution defines the space-time metric. In the AB-case, the current in the wire defines the metric of the space in which charges move.

There is a further similarity between the gravitational effects of matter and the phase effects of charges, namely that it is impossible to shield a part of space from the effects of metric-shaping sources, matter and charge, respectively.

The AB-effect could be described as an interaction between the current and the passing electrons, but one must observe that interaction then means something different from the notion of interaction used in this book, namely, exchange of quanta.

8.10 Conclusion

If we assume that real objects are wave-like entities which could be described as a superposition of partial waves with infinite extension, we can give an explanation of non-local correlation. Two objects making up a singlet state are phase correlated and have overlapping partial waves independently of how far apart their centres of mass are. We cannot interact by exchanging momentum with only one of them without interacting with the other, because as regards momentum and spin their waves are coupled.

The collapse of the wave function during a measurement is a real physical process. The internal collapse mechanism cannot be analysed as a continuous series of interactions because of quantisation. No part of the wave can be isolated and

regarded as a particular object whose motion defines a frame of reference. This makes possible distant correlations within indefinite short time intervals without contradicting Lorentz invariance.

If we perform two spin measurements, one on the left side and one on the right, we have two well-defined events and the space–time interval separating these two events is Lorentz invariant. However, if we measure the spin state of only one side we have only one well-defined event. This event, i.e. this measurement interaction, determines the spins of both particles, due to their being in a singlet state. But this event cannot possibly be a violation of relativity theory because no space–time interval between two events can be defined; we have but one event. It follows that Einstein locality does not contradict relativity theory.

Many would, I guess, request a complete physical explanation of the collapse mechanism, thereby implying that I have not given one. The implicit assumption behind such a request is that only a mathematical account in terms of a differential equation describing the time evolution of the state function is fully satisfactory. Such a request cannot, as a matter of principle, be met, because the state change is discontinuous; hence the state as a function of time is not differentiable. This is an argument for rejecting the conception of explanation as requiring continuous evolution.

In all three examples of non-locality – EPR-type experiments, AB-effect and interaction between two two-level atoms – the non-local feature is connected to the phase information in the wave. We can summarise by saying that non-local phenomena are phenomena in which the phase information of a wave extends to places far outside the ‘bulk’ of the wave packet.

Chapter 9

Gentle Criticism

9.1 Introduction

In this chapter I will discuss some other interpretations of quantum mechanics. I have selected those having some affinity with my own view in order to bring it out in relief. Therefore, I will discuss three interpretations expressing clearly realistic attitudes, namely de Broglie–Bohm’s theory, the GRW theory and the many-worlds interpretation. In addition, I take up the decoherence interpretation because, according to this view, we need not assume quantisation. It seems to me that the decoherence interpretation is close to Schrödinger’s own idea that interactions are not discrete, it only so appears, and my view is the opposite one, namely that quantisation cannot be explained away.

For a long time the Copenhagen interpretation was the received view regarding quantum mechanics. That is not so any longer and the debate is intense. But still, most philosophers feel the need to motivate why they reject the Copenhagen view, and so do I. In addition, I think it wrong to identify Bohr’s views with what is commonly called the ‘Copenhagen interpretation’; Bohr had a much more realistic outlook than others claiming adherence to the Copenhagen view. Therefore, I will start by discussing the Copenhagen interpretation.

There are several other interpretations which certainly deserve attention, in particular the modal interpretation. But discussing and comparing all interpretations means discussing criteria for interpretations, for example realism versus empiricism or whether we really should require visualisability of a physical theory. But where to stop arguing? I’m rather convinced that most philosophers and physicists deep in their heart agree with me in requiring (minimal) realism and visualisability, but have thought, in my view mistakenly, that these demands on interpretation of physical theories cannot be met. Therefore, I believe the best way is to offer an interpretation that fulfils these goals.

9.2 The Copenhagen interpretation

The Copenhagen interpretation – CI for short – is a widely used but rather vague term. It is generally used to cover some of Bohr’s central ideas: first and foremost the principle of complementarity. Another idea of Bohr, usually included in CI, is the sharp distinction between the macro and the micro world. Bohr’s own argument for this distinction was that the macro domain is the only place where unambiguous communication is possible. According to Bohr, experimental results expressed in ordinary language are the domain where intersubjectivity is guaranteed. This does

not imply that the micro domain is less real, but some adherents to CI seem to draw this conclusion. Thirdly, most (all?) believers in CI adopt the Born interpretation of the wave function. (Whether Bohr himself accepted Born's interpretation is an open question.) Hence, I assume that CI consists essentially of three ideas: (1) a sharp division between the micro and the macro world; (2) the principle of complementarity and (3) the Born interpretation of the wave function in one of its versions.

The distinction between the macro and micro world suffers from the same diffuseness as the theoretical/observational dichotomy: where is the sharp dividing line? Without a clear-cut criterion a distinction is obviously useless for an ontological or semantical analysis. As we all know, this was one argument used against the old logical positivists, who by saying that theoretical statements lacked truth value, whereas observational statements have truth value, tacitly assumed that there was a sharp observational/theoretical border. A similar argument can also be used against the denial of the reality of the micro world, a stance sometimes made part and parcel of the CI.

The second component of the CI, the complementarity principle is commonly thought to be the main idea – and the main problem: what does it really mean and does it not imply idealism? To my knowledge, the best discussion of the principle of complementarity has been given by Hooker (1972, p. 136) and Murdoch (1987, ch. 4). The following definition is essentially Murdoch's:

CP (complementarity principle): Two concepts are said to be complementary if

- (i) they are both necessary for a complete description of an entity
- (ii) the conditions for the application of one of the concepts contradict the conditions for the application of the other concept.

An example might illustrate the meaning of statement (ii). As is well known, position and momentum are complementary concepts. This is so because applying the concept of position, a fixed frame of reference must be laid down, preferably the laboratory frame. Probing the momentum of an object invariably means an exchange of momentum, which in turn requires that the detector can be accelerated. In both cases the detector constitutes a frame of reference (cf. for example Bohr, 1949, p. 218). But since the detector cannot at the same time be movable and fixed, position and momentum cannot simultaneously be measured. The condition for using the concept of position is a fixed physical frame of reference, and the condition for exchanging momentum is a movable frame of reference.

On the face of it, complementarity leaves two different interpretations open: either the objects considered cannot simultaneously have definite values of complementary observables (the ontic interpretation), or they cannot simultaneously be measured with arbitrary accuracy, although the observables have definite values (the epistemic interpretation).

Bohr would not have accepted any of these options; his stance was that it is meaningless to ascribe a definite value to an observable if the constraints on the situation effectively rule out the possibility of a precise measurement. This view seems to indicate that Bohr would hold that objects cannot simultaneously *have*

precise values for two complementary observables, but I doubt he would engage in such an ontic mode of speech. He consistently talked about what can be known, about the limits of observation and application of concepts. Perhaps we could call this a semantic interpretation of complementarity.

In Quine's, view, which is also mine, we display our ontology when we accept what objects to take as values of variables in a regimentation of the language. Ontology and semantics are thus closely connected. In this perspective, saying that objects in some situations do not have precise values of certain variables is the same as saying that the corresponding concepts cannot be applied. Thus, Bohr could be interpreted as adopting the ontic version of the complementarity principle.

So far, the complementarity principle by itself does not violate any requirements on an acceptable interpretation laid down in Chapter 2. Furthermore, neither the epistemic, nor the ontic or the semantic interpretation of complementarity violates the stated demands on interpretation. It is not against realism to claim that an attribute of an object sometimes is indefinite and sometimes definite. The conflict between realism and CI arises only if the objects change from definiteness to indefiniteness or vice versa as a result of a change in human perspective only, i.e. without any physical interaction. This interpretation of the complementarity principle has been suggested (for example by Wigner), but according to Hooker (1972, pp. 134–135) and Murdoch (1987, ch. 4) this was not Bohr's view. Hence, although many writers have thought that complementarity and realism are incompatible,¹ they are not, according to the present analysis. Complementarity does not even contradict strongly realistic hidden variable interpretations, according to which every object has definite but unknown values for all its observables.

The third component of the Copenhagen interpretation is the interpretation of Born's rule, i.e. formulas of the type $P(x)dv = |\Psi|^2 dv$. Born interpreted this formula as 'the absolute square of the amplitude of the normalized wave function in a point is the probability to *detect* the object in that point when a position determination is performed'.² Similar statements for the measurements of other quantities are easily constructed. The crucial point is that the Born interpretation does not say anything about where the object *is* or *was*, but only tells us about the probability of a *detection* at a certain point. The difference is that the later formulation is an epistemic statement whereas the former one is an ontic. The difference seems to be subtle indeed: if an object is detected at a point it is indeed reasonable to claim that the object was at that point when the detection was performed. Two successive detections of the position of an object will give us two points and thus the object must have travelled

1 One example is Lakatos, who claimed that 'the complementarity principle enthroned (weak) inconsistency as a basic ultimate feature of nature and merged subjectivist positivism and antilogical dialectic and even ordinary language philosophy into one unholy alliance' (Lakaos, 1970, p. 145).

2 The statement is nowhere to be found in Born's writings, it is Jammer's reconstruction of Born's view. Cf. the following quotation: 'Summarizing Born's original probabilistic interpretation of the Ψ -function, we may say that $|\Psi|^2 dt$ measures the probability density of finding a particle within the elementary volume dt , the particle being conceived in the classical sense as a point mass possessing at each instant both a definite position and a definite momentum' (Jammer, 1974, pp. 42–43).

along a path connecting these two positions. Then it seems innocuous to infer that the object must have occupied a number of places in the interval between the two determinations. This conclusion is, as we have seen, impossible to reconcile with interference phenomena. Hence, adherents to CI prefer the epistemic statement that the square of the amplitude tells us the probability of *finding* the object at the position in question when a suitable measurement is made, and refrain from any statement about the object's whereabouts between observations. Cf. the following statement of Heisenberg:

the conception of the objective reality of the elementary particles has thus evaporated in a curious way, not into the fog of some new, obscure, or not yet understood reality concept, but into the transparent clarity of a mathematics that represents no longer the behaviour of the elementary particles but rather our knowledge of this behaviour. (Heisenberg, 1958, p. 100)

Heisenberg thus clearly says that the quantum mechanical theory (what he calls 'mathematics') do not refer to physical objects but to our knowledge about these objects. This stance conflicts with minimal realism and if Heisenberg's position is part of the Copenhagen interpretation, this view also conflicts with minimal realism.

On the one hand this is a reasonable reconstruction of CI because it is often understood as an anti-realistic position, but on the other hand I doubt that most physicists claiming adherence to CI really would fully subscribe to Heisenberg's view. It appears to me that in most circumstances they have a rather realistic stance, more like Bohr's. Thus, the well-known label 'the Copenhagen Interpretation' is a bit unclear on a crucial issue. In one reading it contains an anti-realistic element, namely the interpretation of Born's rule.

Summarising my criticism, it is not complementarity but the interpretation of probability statements which is the problematic aspect of the Copenhagen interpretation. In this respect my position differs from other realists who usually hold that complementarity is the main defect of the Copenhagen view.

9.3 The measurement problem in the Copenhagen interpretation

When it comes to the analysis of the measurement problem, the Copenhagen interpretation is not a suitable label. The reason is that those physicists who adhere to CI usually accept von Neumann's account of quantum measurement. But the expression 'Copenhagen interpretation' also leads one to think of Bohr and his position is definitely not the same as that of von Neumann. Hence, as regards the measurement problem there are two positions that could be said to belong to the Copenhagen interpretation, but these positions are very different indeed. Von Neumann's analysis of the measurement problem conflicts with realism: as was shown in Chapter 1, it is difficult to avoid the conclusion that the famous Schrödinger's cat will be definitely dead or alive as a result of the cognition process within the observer. Bohr's view, however, is not that paradoxical.

Bohr's theory of measurement is centred upon his sharp division between the macro- and the micro-world. Every measuring instrument must be a macroscopic

object describable in classical terms, because otherwise one scientist would not be able to convey to others his measured results. Bohr had a strong conviction that only classical concepts have a clear empirical meaning and therefore they are absolutely necessary for an objective science. Hence, the quantum mechanical formalism cannot meaningfully be applied to the measurement apparatus. It is a mistake to use the quantum mechanical interaction formalism to describe the measurement interaction.

But we might still want to have some physical arguments, not just a semantical argument, showing why we cannot apply a quantum mechanical formalism to macroscopic objects. Bohr answered something like this: the measurement process induces an uncontrollable interaction between object and measurement device. This disturbance is not a process that is governed by the usual interaction formalism. Hence, the measurement problem will never arise because a quantum mechanical description of the measurement process is not possible.³

As can be seen from the preceding chapters, I believe that Bohr was on the right track. The measurement interaction is a discontinuous change of the state of the object, which could be called a disturbance,⁴ and the macroscopic variables of the measuring device cannot be described as observables governed by quantum rules. On this point I think Bohr's stance was correct but incomplete. But his semantical argument – that we must use classical concepts for description of objective results – does not convince me.

The version of CI including von Neumann's analysis emphasizes that a cut between object and measurement device must be made, but where exactly to draw the line is insignificant. If, for example, the measurement apparatus is described quantum mechanically, it is *ipso facto* regarded as an object. Hence, the measurement interaction must now be seen as occurring when an observer looks at the result, i.e. the person who observes the result is the measurement device. But again, one can describe the human body, or any part of it, quantum mechanically, thereby regarding it as an object. Von Neumann's conclusion was that the mind was the ultimate limit; every material object can consistently be treated quantum mechanically. Hence, the collapse of the state vector occurs as a result of the cognition process. This last step is not explicitly taken by von Neumann but by a number of his followers, such as London & Bauer (1939, 1982) and Wigner (1961). Obviously, this view conflicts with realism.

In conclusion, the most objectionable component of the Copenhagen interpretation is the Born interpretation and not the complementarity principle. The

3 This is essentially Murdoch's reconstruction of Bohr's stance, Cf. Murdoch (1987, p. 113).

4 On this point I differ from Murdoch's interpretation of Bohr. Murdoch takes the disturbance as implying that after the measurement interaction the measured observable has another value than that obtained. In my interpretation the measurement interaction changes the measured observable (and some other) from a potential to an actual property. This transition from potential to actual is not what we usually call a disturbance, but lacking a term for this transition, why not extend the scope of the word 'disturbance'?

third component, i.e. the analysis of the measurement process, is either incomplete (Bohr's view) or contradicting realism (the von Neumann version).

9.4 The many-worlds interpretation

The many-worlds interpretation is not a complete interpretation of quantum mechanics. The adherents of this position have never aimed at anything more than an explanation of the measurement problem. The fundamental idea in this interpretation is to assume that quantum systems do not change discontinuously during measurements, instead they are said to evolve continuously all the time. When a measurement is performed, the system is said to split into a number of 'copies' or branches, one for each component in the superposition of the wave function, i.e. one for each possible outcome of a measurement. This is the central idea of the so called *relative state formulation*, another name for the many-worlds interpretation, of quantum mechanics elaborated by Everett (1957).

The term 'relative state' is chosen for the following reason: when two systems S^1 and S^2 interact to form one conjoined system S described by $\Psi^S = \sum a_{ij} \phi_i \eta_j$, where the sets $\{\phi_i\}$ and $\{\eta_j\}$ are complete orthonormal sets of eigenfunctions for the systems S^1 and S^2 respectively, neither of the subsystems can be said to be in a definite state independently of the other subsystem. But for any choice of a state of one subsystem a corresponding state of the other subsystem can be assigned. For example, if ϕ_k is the state of system S^1 after the measurement, the corresponding relative state of S^2 will be

$$\Psi(S^2 / \phi_k) = N_k \sum a_{jk} \eta_j \quad (9.1)$$

where N_k is a normalisation constant. This is, so far, a mere reformulation of the quantum formalism, but it points to the many-worlds interpretation of the collapse of the wave function as a branching of the world. The state of the total system is definite as soon as the subsystem S^1 is definite. Hence, the conclusion that the state of the object system collapses during the measurement is premature, according to Everett. The only thing that happens is that a splitting of the universe puts the observer into a situation where he observes only one of the possible states of the measuring device, namely, the one belonging to his branch of the universe.

9.4.1 Criticism of the many-worlds interpretation

The main problem with this interpretation is how to understand the expression 'branching of the universe'. I can imagine two main alternatives: either it should be understood as some sort of *real* physical event, implying that the different branches of the universe are equally real. The other alternative is to understand 'branches of the world' in analogy with possible worlds as this term is used in possible world semantics, i.e. as essentially a modal notion. In that case the splitting of the universe is a transition from possible to actual. This latter alternative seems plausible, but Everett rejects it explicitly. In a footnote in his seminal paper where he introduced

the many-worlds interpretation (Everett, 1957, p. 459) he discusses this alternative, claiming that from the viewpoint of the theory all elements of a superposition, i.e. all branches are actual and none is any more real than the rest. It is not necessary to suppose that all but one element in the superposition are destroyed since they all individually obey the wave equation. As there is no effect whatsoever of one branch on any other, no observer will ever be aware of any splitting process (Everett, 1957, p. 459 fn). The modal reading is thus rejected by Everett, and in fact also the physical interpretation of the branching. He claims that the branching process is not a branching of the physical systems but a branching of the *states* of a physical system subject to measurement: 'Throughout all of a sequence of observation processes there is only one physical system representing the observer, yet there is no single unique *state* of the observer...' (Everett, 1957, p. 459).

As the observer does not branch, only his state, there is no reason to suppose that the observed physical system branches, so Everett's stance seems to be that also the *state* of the observed system branches. But this is not consistent with the formerly stated viewpoint. That the observer is represented by one physical system not being in any unique state must be interpreted as saying that to each possible state there corresponds a possible world. However, in the former quotation, Everett rejects this, arguing that 'all separate elements of the superposition individually obey the wave equation with complete indifference to the presence or absence of any other elements.'⁵ I am unable to integrate these statements into a coherent interpretation. As an 'explanation' of the (perhaps only apparent) collapse of the wave packet this is a case of *obscurum per obscurus*. Using somewhat different arguments, Healey (1984) and Stein (1984) arrived at the same conclusion.

The main motive for Everett and his followers to adopt the many-worlds interpretation was the desire to apply quantum mechanics to cosmology. Assuming that quantum mechanics is a universally valid theory it must be possible – they claim – to ascribe a wave function to the entire world, and then it could not be observed from 'outside'. And as there is nothing outside the system which could produce irreversible state transitions, i.e. collapses of wave functions (Everett, 1957, p. 455), there is a need for reinterpreting quantum mechanics.

This argument is weak because it rests on a drastic extrapolation. Quantum mechanics has proven correct and useful in micro-physics and from this Everett and others conclude that it is applicable to the entire universe. There is little, if any, evidence for this inductive inference.

Another argument for the many-worlds interpretation, closely related to the former one, was the need for a quantum theory without external observers. As the collapse of the wave function is caused by an external observer in the orthodox interpretation, a realist might wish to accomplish a similar transformation of the wave function without external and human interactions. This demand is indeed reasonable, but it does not pick out the many-worlds interpretation as the only possibility. The same goal is reached in the interpretation presented in this book.

5 The statement is to be found in a note added in proof on p. 459 as a reply to criticism of a draft of the paper.

9.5 The de Broglie–Bohm interpretation

All attempts at realistic interpretation of quantum mechanics have, so far, taken the form of adding ‘hidden’ variables to conventional quantum mechanics. The general idea seems promising from a realistic viewpoint, because the problems of understanding quantum mechanics could be due to a superficial description of the real phenomena. A hidden variable interpretation can have various goals: to explain the EPR correlation phenomena, to remove indeterminism or to make the theory fulfil some realistic constraints. The opponents of hidden variables have construed a series of no hidden variable proofs purporting to show that these (hidden variables) are impossible. As far as I know the only hidden variable interpretation that is unaffected by these no-hidden-variable proofs is the pilot-wave interpretation of de Broglie and Bohm.

De Broglie’s and Bohm’s guiding idea was to give a realistic interpretation of wave-particle dualism by assuming that there exists two sorts of entities, namely both particles and waves. The latter have dynamical significance by acting on the particles and are described by the wave functions of quantum mechanics. The particles, being objects that are guided by waves and are thus distinct from the waves, have no representation in conventional quantum mechanics. Hence, the hidden variables introduced by de Broglie and Bohm are the positions of the particles, as distinct from the position parameters being arguments of the wave functions. Bohm’s (1952) argument for the plausibility of this idea goes as follows.

The time-dependent Schrödinger equation for a one-particle system is

$$\frac{i\hbar\partial\Psi}{2\pi\partial t} = -\frac{\hbar^2}{2m}\nabla^2\Psi + V(\mathbf{x})\Psi \quad (9.2)$$

where $V(\mathbf{x})$ is an external potential function in the position coordinates $\mathbf{x} = (x_1, x_2, x_3)$. It has the solution

$$\Psi(x) = R \cdot \exp\left(\frac{2\pi i S}{h}\right) \quad (9.3)$$

where R and S are real functions. The square of R is the probability density, so we put $P(x) = R^2(x)$. Then we obtain

$$\frac{\partial P}{\partial t} + \nabla \left(P \frac{\nabla S}{m} \right) = 0 \quad (9.4)$$

$$\frac{\partial S}{\partial t} + \frac{(\nabla S)^2}{2m} + V(x) - \frac{\hbar^2}{4m} \left(\frac{\nabla^2 P}{P} - \frac{(\nabla P)^2}{2P^2} \right) = 0 \quad (9.5)$$

Bohm observes that in the classical limit ($\hbar \approx 0$) the above equations have a simple interpretation. The second equation can be identified as the Hamilton–Jacobi equation, which determines the dynamical conditions on an ensemble of non-interacting particles subject to the external potential $V(x)$. P is the classical probability

density and S is the so-called generator function, the gradient of which is equal to the momentum of the individual particles.

However, in non-classical situations where Planck's constant cannot be neglected, the second equation is given a different meaning by Bohm, who interprets the fourth term in the equation as a *quantum potential* $U(x)$:

$$U(x) = \frac{-\hbar^2}{4m} \left(\frac{\nabla^2 P}{P} - \frac{(\nabla P)^2}{2P^2} \right) = \frac{-\hbar^2 \nabla^2 R}{2mR} \quad (9.6)$$

which is possible because P is a function of the coordinates only. Then the two potentials $V(x)$ and $U(x)$ could be joined together to make up one single potential, which puts the equation back to the Hamilton–Jacobi form. Hence, it is possible to interpret quantum mechanics in analogy with classical mechanics even for an ensemble of particles, provided that a potential can be found that also includes the specific non-classical interactions such as spin–orbit interactions.

This interpretation immediately solves the measurement problem because the collapse of the wave function is now an ordinary statistical collapse, through which one gains knowledge about the particle's state prior to the measurement. In modern versions of this interpretation (cf. for example Vigier, 1985) it is possible to account for non-local correlations in many-body states by writing the quantum potential as a function of all the coordinates of the participating particles. Thus, Bohm claims to be able to explain EPR-correlation phenomena. However, the pilot wave theory fails to explain away non-locality because the theory itself is a non-local theory.

There is, in my opinion, a taint of ad hocness in this interpretation. This has to do not so much with the assumption of hidden variables, i.e. variables for the positions of the particles, but with the ontological status given to the quantum waves in the de Broglie–Bohm interpretation. These waves are causally dependent on the four well-known forces of nature. This causal dependence is represented as a logical dependence of the wave functions on the total potential representing these forces. By solving Schrödinger's equation, one can derive the wave functions. According to de Broglie and Bohm the wave functions represent the evolution of real quantum fields and these fields are distinct from the particles they act upon. Hence, the quantum wave fields are not logically independent entities but derived from the classical fields. The particles in turn are sensitive both to classical and quantum fields alike. The structure could be described as follows.

Particles are associated with classical force fields, which produce quantum wave fields. Both the classical force fields and the quantum wave fields in turn act on other particles.

This detour to quantum waves viewed as separated from the particles can be criticised on two points. First, it seems vulnerable to Occam's razor: why introduce an entity when the dynamics of particles can be accounted for without it? An adherent to de Broglie–Bohm might reply that the additional entity is needed, not in order to account for physical phenomena, but to make the theory compatible with realism. My counter argument is of course that that goal can be accomplished without adding new additional entities, witness the present book.

The second point is that in divorcing the particles from the waves one must assume a new kind of interaction, i.e. an interaction between the waves and the particles. This interaction must be something quite different from ordinary physical interactions that commonly are understood as exchanges of particles: photons carry electro-magnetic interaction, W-bosons carry the weak interaction, gluons the strong one and the assumed gravitons carry gravitation. (Perhaps one should use the verb 'constitute' rather than 'carry', since 'carry' indicates that the exchange particle and the interaction are two distinct things; but the common verb for stating these facts is 'carry'.) The interaction between the particles and the quantum waves seems not to allow of such an analysis, because the quantum waves are conceived of as truly continuous entities. Hence, there seems to be no analogy between classical force fields and the quantum field. Of course, new types of phenomena cannot be excluded as a matter of principle. However, the question is whether this interpretation is sufficiently explanatory to be acceptable. My verdict is that it is not. Assuming a completely new type of force, totally different from those hitherto known, is a big step and there is no independent support for such a bold assumption. Summarising, the interpretation of de Broglie and Bohm fulfils standards of realism but has an explanatory power less than what could be wished. The assumption of pilot waves and their accompanying forces looks like an *ad hoc* assumption.

Bohm was not unaware of the possibility of this type of criticism, it seems. At the end of the first of his two articles he wrote:

Now, there are an infinite number of ways of modifying the mathematical form of the theory that are consistent with our interpretation and not with the usual interpretation. We shall confine ourselves here, however, to suggesting two such modifications namely to give up the assumption that v is necessarily equal to $\nabla S(x)/m$ – and to give up the assumption that Ψ must necessarily satisfy a homogeneous linear equation of the general type suggested by Schrödinger... (Bohm, 1952, p. 179)⁶

Without going into the details of these suggestions, it is clear that by adding any of these suggestions we have no longer a mere interpretation of quantum mechanics but a modified theory, which is inconsistent with the orthodox theory. This modified theory makes empirical predictions deviant from those of orthodox quantum mechanics in certain circumstances and is thus not *ad hoc*. However, by taking this step, Bohm has rejected the assumption that orthodox quantum mechanics is an empirically correct theory – in contrast to my own point of departure.

In summarising: as long as Bohm accepts quantum mechanics without modification his interpretation can be criticised as being *ad hoc*, and when his suggested modifications are included, his theory is no longer an interpretation but an alternative theory.

A further criticism of de Broglie–Bohm's pilot wave theory is that matter and energy appears much more dissimilar than in my conception. De Broglie and Bohm clearly do not assume that electromagnetic radiation consists of two kinds of objects, electromagnetic field and photons, but appear to accept the common view that an electromagnetic field and field quanta, photons, are two descriptions of one and the

⁶ In the reprint in Wheeler & Zurek (1983) it is on page 382.

same entity. It seems natural to me to take the same stance regarding matter, namely that the wave properties and the particle properties are but two aspects of one and the same fundamental entity. This analogy argument is by itself not very strong, but if we further take seriously the equivalence between matter and energy, which can be observed as gravitational effects on electromagnetic waves, the argument is significantly stronger.

9.6 The GRW Theory

In an article published in 1986, Ghirardi, Rimini & Weber (GRW) assumed that the collapse associated with measurement is not merely a change of information obtained by the experimenter, but a real physical change. They also took for granted that all measurement devices are macroscopic objects made up of a great number of elementary particles, 10^{22} or more. They also took for granted that every measurement is a localisation of the object in question; whatever property we measure, the measurement interaction is performed at some well-defined place containing the measuring device and the result of the measurement is that the object measured upon is found at that particular place. Hence, the collapse must be a *spatial* contraction of the wave function into this particular place in real three-dimensional space. Hence, the measurement problem can be formulated: what is the cause of collapse, i.e., this spatial contraction?

The GRW paper did not really answer that question. It merely postulated that for all objects there is per unit time a very small probability for a collapse. In mathematical language, it means that for every wave function there is per unit time a small probability for being multiplied by a narrow wave packet. GRW postulates this as a new fundamental law. The probability for this multiplication is very low, and proportional to the number of particles accounted for by the wave function, so for single particles this spontaneous collapse of the wave function is very rare. However, if we consider a macroscopic device, the probability increases by a factor equal to the number of particles involved. GRW calculated the magnitude of the probability that would yield an observable effect on single particles. They found that for a localization probability to be of the order of 10^{-16} s^{-1} per particle the localization width has to be of the order of 10^{-7} m . From that follows that for a single particle the collapse will occur once in 10^8 – 10^9 years. On the other hand, the wave function of a macroscopic object being made up of 10^{22} particles will collapse within 10^{-7} s , and we will never observe the superposition. So far this theory accounts for the facts.

It should be observed that GRW does not give any reason for supposing that the localization probability is 10^{-16} s^{-1} per particle. They frankly admit that they cannot give any deeper reason for proposing this value and that this collapse probability is an entirely new fundamental law of nature. On the other hand, there is some slack here; if the GRW theory is correct, the localization probability could deviate several orders of magnitude from this value without being in conflict with observations.

9.6.1 Albert's criticism

There are some technical difficulties with the GRW theory, but David Z. Albert (1992) leaves these aside while concentrating on what he sees as the main problem with the GRW theory. The problem is this: GRW assumes that measuring instruments generally have some kind of pointer being able to have macroscopically different positions. The pointer positions are in the measurement interaction correlated to the measured observable of the object, i.e. the pointer positions are the measurement outcomes. However, this is not always the case. Albert considers a thought experiment, a device for measuring the hardness of a particle. Only two outcomes are possible: a particle is 'hard' or it is 'soft'. Particles are fed into a hardness box which separates hard particles from soft ones by means of magnetic fields. The particles are directed towards a TV screen and 'hard' and 'soft' particles hit the screen at different places. When a particle strikes the screen at a point A certain electrons in certain atoms on the screen in the vicinity of point A are excited into higher energy states. Soon thereafter they de-excite and emit excess energy in the form of photons. The spot A becomes luminous and we can observe that a particle has hit the spot.

The question is now: will the GRW theory entail that such a hardness measurement has a definite outcome? That depends on whether it necessarily comes a time when the hardness of the particle gets correlated with the position of a macroscopic object, i.e. the pointer.

It is evident that the correlation cannot obtain in the hardness box, because nothing is yet correlated with the hardness, save the position of the particle and the particle is not a macroscopic object. Next, the particle will hit the screen and the electrons in the fluorescent atoms get involved. The entangled wave function can be written:

$$\Psi_{\text{tot}} = 1/\sqrt{2}|\text{hard}, X=A\rangle|ex\rangle_{e_1}\dots|ex\rangle_{e_n}|unex\rangle_{e_{n+1}}\dots|unex\rangle_{e_{2n}} \\ + 1/\sqrt{2}|\text{soft}, X=B\rangle|unex\rangle_{e_1}\dots|unex\rangle_{e_n}|ex\rangle_{e_{n+1}}\dots|ex\rangle_{e_{2n}}$$

where $|ex\rangle_{e_1}$ refers to electron e_1 being in an excited orbit and $|unex\rangle$ to an electron being in a stable orbit. In other words, the two properties *hard* and *soft* are correlated to two different distributions of energy and not two different distributions of objects. Thus, the problem for the GRW theory is that the two possible outcomes of the experiment concern exactly the same electrons. Hence, a GRW collapse will not choose between the two alternatives.

GRW predicts that the wave function of the collection of electrons collapses onto (nearly) position eigenstates. However, it is the *energies* of the electrons that get correlated here, not their positions. The GRW collapse is not the right *sort of collapse* in this situation, Albert claims. The difference between excited and unexcited electrons is not correlated with the position of anything, because the electrons do not change place when being excited (Albert, 1992, pp. 101–103).

Although this is a thought experiment, Albert has a point here I think. Not all measurements change the position of anything, although the starting point for GRW theory, namely that all measurements are in the final analysis position measurements,

is correct. It does not follow, however, from this assumption, that any macroscopic object must *change* position.

(1) Albert notes in passing that in order to observe the luminous spot on the TV screen it is sufficient if some 10 photons are emitted, because that is the sensitivity of our eyes, and a collapse involving so few particles will, according to GRW, hardly ever occur as the collapse frequency according to GRW is approximately 10^{-15} s^{-1} . But if we take energy exchange as the crucial feature there is nothing problematic about allowing involvement of only 10 particles. One can strengthen this argument still more. Perfect photodetectors cannot be ruled out as a matter of principle, and if such a device were in use, only one emitted photon would be sufficient for measuring whether position A or B is luminous. But even one single photon carries an amount of energy (or rather *is* a certain amount of energy).

(2) The GRW theory implies that energy is not conserved. The collapse is a contraction of a wave well spread out in space and this contraction will also result in an increase of momentum indeterminacy. Hence, there will be a statistical increase in the spread of momentum as time passes, and that in turn will also result in an increase of energy. This energy increase is, according to GRW,

$$\partial E = \frac{\hbar^2 \lambda \alpha}{4m} t$$

The parameter λ is the frequency with which the collapses occur. With the proposed values for the parameters this energy increase is small, approximately 10^{-15} K per year in an ideal atomic gas. However, no matter how small it is, this is really a very much unwanted effect; energy conservation is a consequence of the requirement of invariance under time translations, and this invariance principle together with invariance under rotations and translations in space are the most fundamental principles in physics. The GRW theory thus conflicts with a very fundamental symmetry principle which, in my opinion, strongly counts against it. Even if the other problems could be solved I still would regard the GRW theory as not satisfactory due to this conflict with time translation symmetry. Of course, people could argue that you could never observe the violation of time translational symmetry, so why bother. But remember that the motive for the interest in the measurement problem is matters of principle; we are not satisfied with a solution which works well *for all practical purposes*.

(3) The third argument in favour of my proposal and against the GRW theory is that postulating a universal collapse frequency and a universal size of the width of the collapsed wave function such that these two parameters fits the facts has an taint of ad hocness. I do not claim that the proposal falls short of Popper's criterion of independent testability, but rather that there is absolutely no reason to believe in a universal collapse frequency, save for the measurement problem. It is simply too closely tailored to the measurement problem to be a telling explanation. Empiricists like van Fraassen would certainly dismiss this argument as only an expression of personal taste, but as already remarked, the measurement problem is a pressing problem only for realists.

9.7 The decoherence theory

The decoherence theory is also only aimed at solving the measurement problem. It could be given as follows. Consider a quantum object being in a superposition of two states, for example a spin particle being in a superposition of spin up and spin down. When we measure the spin of that particle the detector system may enter into one of two possible states, corresponding to the two spin states. The total wave function for the combined system is then

$$\Psi_{\text{tot}} = \alpha|\text{up}\rangle|d_1\rangle + \beta|\text{down}\rangle|d_2\rangle$$

where d_1 and d_2 respectively refer to the two possible states of the detector. We will never observe such a superposed state of the detector. According to the decoherence theory, that depends on the influence from the environment. The environment is also a quantum system and correlated to the apparatus. Hence a more complete wave function would be

$$|\Psi_{\text{tot}}\rangle|E\rangle = \alpha|\text{up}\rangle|d_1\rangle|E_1\rangle + \beta|\text{down}\rangle|d_2\rangle|E_2\rangle$$

where E refers to the total state of the environment and E_1 and E_2 refer to the two states of the environment correlated to the two states of the measurement device. The argument, as given by Zurek goes:

The final state of such a combined ‘von Neumann chain’ of correlated systems SDE (system-detector-environment) extends the correlation beyond the SD pair. When the states of the environment E_i corresponding to states d_1 and d_2 of the detector are orthogonal, $\langle E_i|E_j\rangle = \delta_{ij}$, the density matrix that describes the detector-system combination obtained by ignoring (tracing over) the uncontrolled (and unmeasured) degrees of freedom is

$$\rho_{\text{sd}} = \text{Tr}|\Psi\rangle\langle\Psi| = \sum \langle E_i|\Psi\rangle\langle\Psi|E_i\rangle = \rho'$$

The result is precisely the reduced density matrix... (Zurek, 1991, p. 39)

To ‘ignore’ or trace over a number of degrees of freedom is a mathematical process: what physical assumption does it correspond to? This question is not explicitly answered, although the reasonable interpretation amounts to the same stance as that taken by E. B. Davies (1976): any measurement device is made up of a great number of atoms, usually in a condensed state. Therefore, the energy spectrum of this device is practically continuous, and that in turn implies that even the slightest external fields can interact with the measurement device. Hence, there will be constant interaction between the measurement device and its environment and that effectively removes the phase correlation between the different possible pointer states of the measurement device. This is so because the environment has an enormous number of degrees of freedom and the phase information will, after a very short time, dissipate into all these degrees of freedom. Hence, when we observe the detector and treat it as an individual system, we do not observe any effects at all of the superposition. In short, the phase information has diffused from the pointer states to unobserved states of the total system.

This theory deserves two comments.

(1) Do we really need to introduce the environment in the discussion? All measurement devices are rather big objects with an enormous number of degrees of freedom all by themselves. When we use an object as a measurement device we utilize only one out of all these degrees of freedom as the pointer variable. Hence, if the phase information can dissipate out into the environment, it certainly will dissipate into all of the degrees of freedom of the measurement device. And just as it is impossible for us to gather complete information about the environment, it is usually impossible to collect information about the total state of the measurement device.

(2) The argument is given at a very abstract level and when one tries to apply it to a concrete measurement it becomes less convincing. Take, as an example, once more, the measurement of spin. In such measurements we use two detectors placed in the two partial beams behind a SG apparatus. As already discussed, the collapse must occur when one of the detectors fires. Hence, in the decoherence view, if the detectors were completely isolated from the environment, neither of them could be attributed a definite state of fired or not fired when the particle arrives, both would still be in a superposed state of $|\text{fired}\rangle \otimes |\text{not fired}\rangle$. Hence, if the decoherence argument were valid, a necessary condition for a definite outcome of the measurement is that the detectors are not completely isolated from the environment. This seems implausible in my opinion. As far as I can see, even a completely isolated photographic film would change its state when hit by an incoming photon with sufficient energy.



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Chapter 10

Summary and Conclusions

10.1 Realism and quantum mechanics

Quantum mechanics has for decades been used to rule out the possibility of a realistic interpretation of scientific theories. This is due to two facts, one historical and one conceptual. The historical fact is this: when quantum mechanics entered the stage in the late 1920s and in the 1930s, there was a strong preponderance for instrumentalism in philosophy; hence, quantum mechanics in its matrix mechanics formulation accorded with the received view, logical positivism. Minority physicists dissatisfied with the Copenhagen orthodoxy were accused of adhering to old-fashioned requirements of visualisability and understanding of physical theories. Predictive accuracy was the only requirement.

The conceptual fact is the mismatch between the structure of quantum mechanics and the common notion that physical objects are made up of particles. Although the official stance was that interpretation of theoretical terms was considered impossible, the people in the trade constantly interpreted their mathematical formulas in terms of the dynamics of particles, and they still do. This forced them to attribute some rather weird properties to these particles, such as not having a position or modifying their properties depending on the experiments they were going to be the subject of. No wonder that realism is deemed impossible, but it should be carefully observed that it is realism about particles that is held to be impossible. (And to that I agree!) Thus, the conceptual mismatch consists of the fact that quantum mechanics at the core is a theory describing the dynamical evolution of systems in terms of mathematical waves and such a theory cannot consistently be interpreted in terms of particles because – among other things – waves obey the superposition principle and particles do not. This is a conceptual necessity. We therefore have a structural discord that cannot be resolved. My solution in this book is that we reject the particle ontology. It is perfectly possible to be realist and at the same time deny the existence of particles as perdurable objects. Realism does not force us to subscribe to any particular ontology.

10.2 The equivalence between matter and energy

Classically, i.e. before the advent of relativity theory and quantum mechanics, we distinguished between on one hand a ‘billiard ball’ picture of matter as consisting of portions of some kind of rigid stuff and, on the other a wave picture of electromagnetic radiation. The world was thus believed to consist of two completely different kinds of stuff: matter and radiation. Living with this dichotomy was no problem. The

negative outcome of Michelson–Morley’s experiment, aimed at measuring the ether wind and consequently the absolute motion of the earth, was interpreted as a matter of technological shortcoming and no real cause for anxiety. But the tides turned during the first decade of the 20th century when Planck introduced quantisation of radiation and Einstein presented his theory of relativity.

With the advent of relativity we got a completely new picture of the world in several respects. In this context we may concentrate on the implications of the formula $E = mc^2$. This formula equates matter and energy, two types of phenomena believed to be completely different. But how could one substance become a completely different one? It appears as a kind of transubstantiation in the physical world. However, if we conceive of matter as a wave phenomenon, a change from energy to matter and vice versa appears less radical. It is in this view a change from one species of wave to another.

If we think of matter as built up out of matter waves we might ask: ‘what is that thing that performs oscillatory motion?’ Well, one answer could be that energy is oscillating between potential and kinetic energy. As we saw in Chapter 4, the real and the imaginary part of the wave function could be interpreted as the amplitudes of the potential and the kinetic energy, respectively. But energy of what? Here we strain the conceptual scheme and our grammar, for energy is traditionally conceived of as an attribute to objects. When energy becomes rather the constitutive feature of individual objects we can no longer ask the question ‘to what is energy an attribute’. Energy has taken the same conceptual role as that formerly occupied by matter.

The (relativistic) mass of the object is equal to – or rather the same as – its total energy. The mass of a body is on the other hand a measure of its inertia, its resistance against changes of state of motion. When we measure the inertia of a body we measure its mass.

Energy in turn is related to, or rather identified with, frequency through the relation $E = h\nu$. Thus, we have two forms of energy, mass and frequency. We can use one of them as part of the identification of a particular object and then the other form provides the possibility for predication. But we could just as well swap the semantic roles.

The matter wave angular frequency ω is given by $\omega = E/\hbar$, hence, the energy-frequency relation $E = hf$ which originally was applied only to radiation, is valid also for matter waves. From this viewpoint it is a natural question to ask what fundamental difference there is between things with and without rest mass (except this one). If we formulate this question as ‘why have some objects rest mass?’ we have come to the same question as some elementary particle theorists. As far as I know, nobody yet knows the answer.

Heisenberg took as the departure in his version of quantum theory the observation that the only properties associated with an atom we actually can observe are the frequency and intensity of emitted radiation.¹ Asim Barut (1992) has pointed out the same thing and shown that we can formulate quantum mechanics as a pure-wave mechanics without using any constants for mass and charge in the differential

1 I have not found this statement in the writings of Heisenberg himself, but in Frank (1947, p. 212).

equations; these constants can be introduced later, as a mere matter of convention. The possibility of doing so shows that it is coherent to view all objects as fundamentally wave objects. In this view it is not at all astonishing that we get radiation with a certain frequency when objects lose energy, because an energy decrease is a decrease in frequency of vibration and the loss is emitted as vibration energy.

10.3 Objects in quantum mechanics

The phenomenon of quantisation has important implications for the notions of systems and particles. The individual objects in microphysics are not the so-called ‘particles’. Rather, the term ‘particle’ (and ‘photon’, ‘electron’ etc) denotes mere portions of ‘stuff’: of energy, mass or charge. These portions cannot be identified as objects in the sense of being re-identifiable from one moment to another. The individual objects are the systems that are made up of one or several ‘particles’. When we use particle terms as count terms, they must be interpreted as expressing quantities. What we count by these count terms is not the number of individuals but the amount of charge, energy, etc in the system of interest. The reason it is possible to determine a quantity without counting individual objects is quantisation. If exchange of energy and other conserved quantities were not quantised we could not talk about the amount of a quantity as the number of some sort of ‘particles’.

Individuation of objects is a logico-semantic affair. The principle of individuation and the criteria for identity tells us under what conditions we coherently can talk about a particular object being identifiable as the *same* object referred to by two or more different descriptions, names or pronouns. These criteria are implicit in the way we construct definite descriptions of individuals. When we do that we must use general terms. Such a description is of the form ‘that thing which is ...’ where the blank is filled in with at least one general term and, when physical objects are concerned, a space–time reference.

Criteria for identity are associated with our way of construing descriptive singular terms, i.e. definite descriptions. The question whether these singular terms really refer to physical objects is partly an empirical question. But we need not perform any new experiment to decide this question. It suffices to investigate the theory, namely quantum mechanics, since this theory is generally believed to be empirically adequate. And quantum theory tells us that fermions and bosons are not individuals. We cannot in general apply identity criteria for electrons, for example. Hence they are not individual objects, but mere portions of charge, and analogously for photons and other ‘particles’.

10.4 Wave–particle duality and complementarity

The wave–particle duality astonishes most people, at least when they first hear about quantum theory. In the physics community the accepted way of ‘explaining’ wave–particle duality has for many decades been to refer to Bohr’s principle of complementarity, which says that quantum objects show wave or particle properties depending on what experiment we prepare. Once we have described an experiment

in such a detail as to allow other physicists to repeat it, our verbal communication is unambiguous. When that is achieved we have effectively chosen perspective; we will without ambiguity observe either the particle or the wave properties of the object in question. Bohr continued by saying that a complete description of all phenomena comprises both aspects but a unified picture is impossible.

I think Bohr is partly right on this matter. A careful description of an experiment will reveal whether we are interested in describing the motion of systems or their interactions. When our interest is in the motion of an object, we use the wave picture, when interaction is of interest we use the particle picture, since interactions are quantised. But in my view this *is* a unified picture, namely a coherent interpretation of quantum mechanics, provided we have a purely *physical* criterion telling us whether the wave or particle aspect will be visible. What more could be wished for by way of a unified picture? Hence, I think a unified picture is possible. Plausibly, the difference between Bohr and me is that we differ on the meaning of the expression ‘unified picture’.

My position differs from that of Bohr also in not thinking that the complementarity principle is sufficient as an explanation, even if it marks the limit for our use of classical physical concepts. Complementarity is a feature of our world, and that feature needs an explanation. I look upon my own views not as replacing Bohr’s but as complementing them.

However, I differ more radically from Heisenberg because he and many others express themselves as if a realistic view of the world is impossible.

10.5 Historical remarks

From the discussion so far it is obvious that among the founding fathers of quantum mechanics, Schrödinger is, in my view, nearest a correct interpretation. But it is well known that he met severe criticism and he gave up, at least publicly, although he never accepted the Copenhagen Interpretation. (For a thorough account of Schrödinger’s evolving views, see Bitbol, 1996.) Why was he not able to pursue his ideas to a coherent interpretation?

It seems to me that the answer comprises of two components. One is that he did not fully endorse Einstein’s interpretation of Planck’s radiation law; the other is Schrödinger’s strong classical conception of wave dynamics.

As regards the first point it is relevant to remember that Planck did not really introduce the concept of quantised radiation in his famous paper of 1900. What he did was to introduce a calculational device, namely to divide the energy spectrum in a number of intervals (cf. for example Kuhn, 1978). There were absolutely no theoretical reasons for this move: its sole purpose was to get a radiation law in accordance with observations. The question ‘what is the physical reality behind this partition of the spectrum’ was later discussed by Einstein (1905). In this paper, Einstein introduces the concept of quanta of energy, i.e. that radiation is made up of distinct portions that interact and move independently of each other. One of his arguments is built on the assumption that the statistical behaviour of a collection of

these quanta enclosed in a container is analogous to the statistical behaviour of a collection of molecules. Einstein's conclusion in this paper is:

Monochromatic radiation of low density (within the range of validity of Wien's radiation formula) behaves thermodynamically as though it consisted of a number of independent energy quanta of magnitude $R\beta/N$.

Einstein did not introduce the word 'photon' in this paper – it came later – but clearly it is the photon he is talking about.

Thus, Einstein observes that radiation *behaves thermodynamically*, as if it consisted of photons. He did not claim that it behaves as photons in *all situations*. The difference is subtle but important, for radiation does not behave as a flow of individual objects in all circumstances. My view is that radiation shows particle behaviour only in interactions, but wave behaviour when not interacting. It seems to me that it accords with Einstein's claim because the thermodynamic behaviour is the collective behaviour of a number of quanta in interaction with external things, such as the walls of containers.

Does quantisation of interaction necessarily have to be caused by the discrete nature of radiation, or could it result from the fact that the energy states of matter are discretised? This question seems to lie behind the following discussion between Bohr and Schrödinger, which was held in Copenhagen in 1927 when Schrödinger paid a visit to Bohr's institute:

Schrödinger:

Surely you realise that the whole idea of quantum jumps is bound to end in nonsense. You claim first of all that if an atom is in a stationary state, the electron revolves periodically but does not emit light, when, according to Maxwell's theory, it must. Next, the electron is said to jump from one orbit to the next and to emit radiation. Is this jump supposed to be gradual or sudden? If it is gradual, the orbital frequency and energy of the electron must change gradually as well. But in that case, how do you explain the persistence of fine spectral lines? On the other hand, if the jump is sudden, Einstein's idea of light-quanta will admittedly lead us to the right wave number, but then we must ask ourselves how precisely the electron behaved during the jump. Why does it not emit a continuous spectrum, as electromagnetic theory demands? And what laws govern its motion during the jump? In other words, the whole idea of quantum jumps is sheer fantasy.

Bohr:

What you say is absolutely correct. But it does not prove that there are no quantum jumps. It only proves that we cannot imagine them, that the representational concepts with which we describe events in daily life and experiments in classical physics are inadequate when it comes to describing quantum jumps. Nor should we be surprised to find it so, seeing that the processes involved are not the objects of direct experience.

Schrödinger:

I don't wish to enter into long arguments about the formation of concepts; I prefer to leave that to the philosophers. I wish only to know what happens inside an atom. I don't really mind what language you choose to discuss it. If there are electrons in the atom and if these are particles – as all of us believe – then they must surely move in some way. Right now I

am not concerned with a precise description of this motion, but it ought to be possible to determine in principle how they behave in stationary state or during transition from one state to the next. But from the mathematical form of wave or quantum mechanics alone it is clear that we cannot expect reasonable answers to these questions. The moment however, that we change the picture and say that there are no discrete electrons, only electron waves or waves of matter, then everything looks quite different. We no longer wonder about the fine lines. The emission of light is as easily explained as the transmission of radio waves through the aerial of the transmitter, and what seemed to be insoluble contradictions have suddenly disappeared.

Bohr:

I beg to disagree. The contradictions do not disappear: they are simply pushed to one side. You speak of the emission of light by the atom or more generally of the interaction between the atom and the surrounding radiation field, and you think that all the problems are solved once we assume that there are material waves but no quantum jumps. But just take the case of thermodynamic equilibrium between the atom and the radiation field – remember, for instance, the Einsteinian derivation of Planck's radiation law. This derivation demands that the energy of the atom should assume discrete values and change discontinuously from time to time; discrete values for the frequencies cannot help us here. You can't seriously be trying to cast doubt on the whole basis of quantum theory!

Schrödinger:

I don't for a moment claim that all these relationships have been fully explained. But then you, too, have so far failed to discover a satisfactory physical interpretation of quantum mechanics. There is no reason why the application of thermodynamics to the theory of material waves should not yield a satisfactory explanation of Planck's formula as well – an explanation that will admittedly look somewhat different from all previous ones.

Bohr:

No there is no hope of that at all. We have known what Planck's formula means for the past twenty-five years. And, quite apart from that, we can see the inconsistencies, the sudden jumps in atomic phenomena quite directly, for instance when we watch sudden flashes of light on a scintillation screen or the sudden rush of an electron through a cloud chamber. You cannot simply ignore these observations and behave as if they did not exist at all.²

Rozenthal (1967, p. 103) reports from these meetings that Schrödinger desperately exclaimed 'If one has to go on with these damned quantum jumps, then I'm sorry that I ever started to work on atomic theory.'

It is not easy to see Schrödinger's line of thought here, perhaps because there was not a single one, only a number of convictions. It appears that he was disturbed by the contradiction between assuming that electrons are corpuscles moving around the nucleus, on the one hand, and on the other the electromagnetic theory, according to which an accelerated charge will emit radiation. I wholeheartedly agree with him on this point; it is not satisfying just to accept an exception from the electromagnetic theory. The wave theory of matter provides the dissolution of this conflict, since according to the wave theory nothing is accelerated; the electron is spread around the

2 Heisenberg (1971, pp. 74–75). Quoted in Mehra & Rechenberg (1987, p. 823).

nucleus as a standing wave. But Bohr was correct when responding that this move does not answer the question about the transition from one stationary state to another one; Schrödinger had asked about the dynamics of the transition of electrons and of course the question can be asked concerning the dynamics of the waves.

This transition, i.e. the discreteness of the interaction between matter and radiation is a radically new feature of physics. Before the advent of quantum theory the dynamics of all physical objects had been given in terms of differential equations, which is to say that the dynamical variables were assumed to be represented by continuous real valued functions. It seems that Schrödinger was very reluctant to accept that discontinuous changes occur in quantum physics. He believed that a wave theory of matter might overcome this. On this point he was wrong. The wave theory can explain the dynamics of both matter and energy when there is no interaction, but not the discrete interaction itself.

My diagnosis of Schrödinger's failure is this: Schrödinger believed that the wave theory of matter could replace the quantum postulate by providing a continuous description of the interaction between radiation and matter. He assumed that exchange of energy between matter and radiation as well as between two material objects could be analysed as a resonance phenomenon. This is incorrect and the present alternative in quantum theory is to accept the quantisation of interaction as a new physical principle. This is one of the two core ideas of quantum mechanics.

Perhaps in the future will we be able to explain quantisation by a still more fundamental principle. However, in my view, the quest for an analysis of quantisation is based on the assumption that only continuous changes are natural, and so not requiring explanation, and that discrete changes are exceptions requiring explanation. I do not agree. The preponderance for continuity is a metaphysical assumption with a long history (*natura non facit saltum*) and it gained strong support by the success of classical mechanics. Nevertheless, it appears to be wrong, and that's it. From our vantage point of present day knowledge, continuity is a mere surface feature; physical state changes appear continuous as long as we give a macroscopic description of them.

There is an analogy between, on the one hand, the transition from viewing continuity as fundamental to the quantum view that discontinuity is fundamental and on the other hand the transition from Aristotelian mechanics in which rest is the only natural state and forces are the causes of motion, to Newtonian mechanics in which uniform motion is the natural state and forces are the causes of *change* of motion. In both cases the new theory brought about a change of what requires explanation and what should be accepted as an ultimate fact of the world. After Newton we do not require explanation of uniform motion because it is a natural state. I suggest we make a similar move concerning discontinuity; continuity has been replaced by discontinuity as a natural trait, which is not more in need of explanation than uniform motion. Changes only appear continuous when we do not look into the minute details.

10.6 From orthodox quantum theory to quantum field theory

Orthodox quantum mechanics is an approximation. We know that a correct theory must be relativistic, i.e. covariant under space–time translations, and this criterion is not fulfilled in orthodox quantum theory. A fully relativistic quantum theory is the general quantum theory of fields.

Quantum Field Theory (QFT) describes the physical world as consisting of sets of interacting fields, whereas orthodox quantum theory usually has been understood as describing the dynamics of particles. It follows that if we take QFT to be the more accurate theory, we are obliged to view the world as consisting ultimately of fields, not particles. The particle ontology is thus at best viewed as an approximation, sufficient for many practical purposes. However, when discussing fundamentals we should not neglect the finer details. Thus, it could be argued that we should dismiss the whole task of interpreting quantum mechanics and go directly to quantum field theory instead.

However, there are good reasons to begin the interpretation with orthodox quantum theory. The task of giving an interpretation of a theory should start from what we understand, i.e. everyday experiences in terms of medium size objects and their observable properties, and proceed by successive generalisations when introducing new ideas and new concepts. Therefore, it is appropriate first to interpret orthodox quantum mechanics and then continue to the still more abstract theory, Quantum Field Theory. On the other hand, we should be careful and from time to time take a look at what QFT tells us so that we do not unnecessarily get into trouble because we neglect aspects that are visible only in the quantum field perspective. There are three aspects of QFT of interest in this context: the first is the individuation of local fields in QFT, the second is the QFT picture of interactions, which was discussed in the connection with the measurement problem, and the third is the principle that all interactions in QFT are local.

10.6.1 Individuation in QFT

One reason for the use of field concepts in physics is the insight that locomotion is no real change of state; it is always possible to find a frame of reference in which any given object is at rest and using the generalised relativity principle according to which there are no preferred frames of reference. We can just as well regard the studied system as being at rest. As all systems are more or less vastly distributed in space and do not occupy any well-delimited portions of space, the most natural description of such systems is in terms of functions describing the distribution of system properties over space–time points. This is precisely the character of fields and we are led to the conclusion that physical objects ultimately are fields.

A fundamental trait in all field theories is that every interaction between two fields is localised. Hence action at a distance by way of a force is replaced by interaction at the spot. The reason for this could be said to be the observed fact that interactions are always localised to particular places. However, I suspect that the ultimate reason is metaphysical; action at a distance is thought to be incomprehensible.

The general concept of a field can mathematically be characterised as a function having an infinite number of degrees of freedom. The dynamical variable Ψ describing a field takes as its parameter x , which is conceived as a variable that at the outset has no direct reference to points in physical space. The parameter can be of any kind, a scalar, a vector, a tensor or a spinor. When x is a vector in the four-dimensional space–time, $\Psi(x)$ can be conceived as a plenum filling all space. Or rather, the set of possible values of x constitutes the space; space–time is not something ‘out there waiting to be filled with fields’. Rather, the material content of space–time and the space–time continuum are two sides of the same conceptual coin.

10.6.2 Fields and events – the whole and its parts

The field concept is the modern inheritor of the old concept of a plenum. In antiquity it was discussed whether the world was made up of discrete portions of matter in a void or whether the world ultimately was a plenum. The atomists believed in discrete material objects, the atoms, distributed in the void, but the prevailing opinion during antiquity was that matter was continuous. It was thought that matter strived to fill every place, the well-known *horror vacui* idea. The modern way of thinking is on the other hand a form of atomism, a result of the impact of atomic models such as Dalton’s and Bohr’s. By atomism I do not only refer to the idea that matter is made up of indivisible least portions, but also that these portions of matter are confined to definite volumes in space and in between these atoms we have sheer emptiness. I do not claim that this is what theoretical physicists actually believe. What I mean is that the general notion of an atom, and generally the notion of a particle, as it is used in common speech in this respect resembles the old idea of portions of matter in empty space.

The notion of atoms in a void constitutes a paradigm for our understanding of the relation between parts and wholes. We tend to think that a whole of some kind, be it a material object or something else, can be analysed as constituted of a number of individual entities whose relations to each other are external. By this I mean that it is possible to identify the individual constituents in the whole without invoking the relations between these individuals in the individuating descriptions. The whole can be fully analysed as constituted of individual objects, which can be described and identified without mention of the whole of which they are part. In quantum field theory this is no longer so. A field is a whole made up of parts called *local fields* but they are not conceptually unrelated to the whole. A field is a totality, which has an internal structure, and the structural elements are the local fields. The interesting thing is that this internal structure cannot be described as a set of fully independent elements, i.e. as a set of independent local fields. The question is then if it is possible to treat the local fields as individual entities in the full sense. Sunny Auyang answers *yes* in her book *How is Quantum Field Theory Possible?* and here is how she argues.

Fields are wholes made up of related parts. The parts are not field quanta, ‘particles’, but events. Field quanta, the ‘particles’ such as electrons or photons, cannot generally be individuated and cannot therefore be the objects referred to in our

statements. However, Auyang argues that interaction events, suitably re-described, can do that. Here is her argument.

Consider a field $\Psi(x)$ where the parameter x is a vector specifying a point in the field. This is not yet identified as a point in space–time: such an identification cannot be made until we have chosen a frame of reference enabling us to get a representation of the field points.

When x is fixed we have the local field $\Psi(x)$. It is a free field if it does not interact with other local fields at the same point. However, such a case is an approximation. All fields at the same point interact more or less. Only if the coupling to other fields is negligible in the case at hand, is it reasonable to talk about a free field.

An *event* is the sum total of all fields with the same parameter value x . An event is spatially extensionless but analysable; it can be analysed into a set of free fields $\{\Psi_i(x)\}$ and interaction terms of the form $\Psi_i(x)\Psi_j(x)\Psi_k(x)\dots$ where the index runs over the free fields.

The events are, according to Auyang, the individuals in the field theory. What then is the identity criterion; what are the conditions for two event descriptions to be descriptions of the same event? The answer requires some preliminaries.

The name of an event is its representation in a specified coordinate system. Using a coordinate function f we get $f\{x''\}$ as the name of an event. It is this coordinate function that connects the theoretical concept of a local field $\Psi\{x''\}$ and hence an event, to the empirical world. When we chose another coordinate system we get another name for the very same event, and the identity of this event is given by the transformation rules when we pass from one coordinate system to another. Thus, the coordinates $\{x''\}_a$ and $\{x''\}_b$ refer to the same event, if and only if

$$f_b f_a^{-1}(\{x''\}_a) = \{x''\}_b$$

The choice of the word ‘event’ for the co-occurrence of all local fields at the same point suggests that something happens at this point. However, the expression ‘something happens’ presupposes the categories *before* and *after*, which are not applicable in a description of the world as constituted of a four-dimensional set of events in which time is one of the dimensions. Hence, we cannot say that something ‘happens’. For the same reason, we cannot really talk about changes, for a change presupposes temporality, the distinction between before and after. A change of time coordinate is a change in position in four-dimensional space–time and such a change is a change of object, not a change of the state of some object.

Thus, individuals in Auyang’s interpretation of QFT are not perdurable objects in any ordinary sense of the word. In spite of that they fulfil the general criterion for object-hood, formulated by Quine as ‘no entity without identity’. Hence these individuals are acceptable as those things we talk about in a fully interpreted Quantum Field Theory.

The dynamics of QFT cannot be a theory describing how things change in time, but must be a theory that gives the relations between local fields. What we ordinarily would call an exchange of quanta between two objects must be re-described as two local fields connected in such a way that energy conservation and other conservation rules are fulfilled. It is in this sense that the local fields are parts of a whole and

cannot conceptually be thought of as independently existing. A local field without its connections to other local fields cannot be had as a matter of principle.

The most important difference between a local field and a physical object as the latter term commonly is understood is that the former occupies only a point in space–time instead of a world-line. That means that it has no parts, no extension in space and is not perdurable.

Quanta, modes of excitation of fields, differ both from local fields and ordinary perdurable physical objects. They are not localised in space and they have no other distinguishing features. Therefore, they cannot be individuals. Thus, when we use particle terms as names for kinds of quanta these names cannot be conceived as names of individual objects and hence no names in the final analysis of the theory. Quanta are portions of substances and thus the conceptual structure best suited to them is that of mereology in terms of mass terms. Now it is well known that it is possible to analyse mass terms in the usual semantic framework of quantification theory, individuals and predicates, but this adaptation is a procrustean achievement. The very idea of a portion of a substance entails having all qualitative properties in common with the rest of the substance, making it impossible to use any of these as the individuating property. For material substances that can be attributed place, we can solve the individuation problem by using place as the individuating property for parts of substances being impenetrable. However, when we come to quanta this is ruled out by the fact that quanta are not localised and not impenetrable.

This conceptual difference between local fields and quanta lies at the root of our difficulties in understanding quantum mechanics, both the non-relativistic and the relativistic version. We tend to think of quanta as the individuals and they cannot possibly play that role; the true individuals are the local fields. In non-relativistic language: the individuals are the interaction events.

10.7 An unsolved problem

It cannot be denied that some interactions are discontinuous, random and irreversible and there is strong evidence, indeed proof, for saying that these properties are consequences of quantisation of interactions. The hope, entertained by many, that discontinuity, randomness and irreversibility in the final analysis can be shown to be the result of approximations must be given up. This conclusion was implicit already in Planck's 1900 paper, although he himself did not fully see the implications of his assumptions. (The argument is that the spectral distribution of blackbody radiation is incompatible with the assumption that interaction is continuous.) Hence there are discontinuous interactions. But what is less clear is what conditions are necessary to bring about such interactions and the ensuing collapse. I think it safe to say that being a measurement of the first kind is a sufficient but not a necessary condition. But what are the necessary conditions?

Nancy Cartwright once proposed that there is no general answer to be found; it just happens that quantum systems sometimes collapse and start a new unitary evolution, and the physicist's task is to find appropriate operators for any given situation (Cartwright, 1983, p. 205). I cannot convince myself that no general

conditions for the difference between continuous and discontinuous evolution can be found. But I have, regrettably, no suggestions. However, since this book is a contribution to philosophy of physics, not to physics, and the question about the necessary conditions for collapse is a physical question, this lack of answer can be excused.

The question about the necessary conditions for collapse is intimately connected to another fundamental issue in philosophy of physics, namely the cause of irreversibility. Why are so many real processes irreversible, despite the fact that all fundamental laws of physics which we take to be generally valid, all are time symmetric? This is a question, which, so far, has not been answered to the majority's satisfaction.

It was shown in Section 6.7 that irreversibility is a consequence of quantisation of interaction. If we further accept that quantisation of interaction is a fundamental feature of the world and occurs independently of observation, we have a partial explanation of irreversibility. This is a bonus of the present interpretation.

Bibliography

- Aharonov, Y. (1986), 'Non-local phenomena and the Aharonov-Bohm effect.' In: Penrose & Isham (Eds) *Quantum Concepts in Space and Time*, pp. 41–56 (Oxford: Clarendon Press).
- Aharonov, Y. & Vaidman, L. (1993), 'Measurement of the Schrödinger wave of a single particle'. *Phys. Lett. A* **178**, 38–42.
- Albert, D. Z. (1992), *Quantum Mechanics and Experience* (Cambridge, MA: Harvard University Press).
- Aspect, A., Grangier, P., & Roger, G. (1981), 'Experimental tests of realistic local theories via Bell's theorem'. *Phys. Rev. Lett.* **47**, 460–467.
- Aspect, A., Grangier, P., & Roger, G. (1982a), 'Experimental realization of Einstein-Podolski-Rosen-Bohm Gedankenexperiment: a new violation of Bell's inequalities'. *Phys. Rev. Lett.* **48**, 91–94.
- Aspect, A., Dalibard, J. & Roger, G. (1982b), 'Experimental tests of Bell's inequalities using time-varying analyzers'. *Phys. Rev. Lett.* **49**, 1804–1807.
- Bitbol, M. (1996), *Schrödinger's Philosophy of Quantum Mechanics* (Dordrecht: Kluwer).
- Barut, A. O. (1988), 'Schrödinger's interpretation of Ψ as a continuous charge distribution'. *Ann. Phys.* **7**. Folge, Band 45.
- Barut, A. O. (1992) 'Formulation of wave mechanics without the Planck constant h '. *Physics Letters A* **171**, 1–2.
- Barut, A. O. (1994), 'Fallacy of arguments against local realism'. In A. van der Merve & A. Garuccio (Eds) *Waves and Particles in Light and Matter* (New York and London: Plenum).
- Bell, J. (1964), 'On the Einstein-Podolsky-Rosen paradox.' *Physics* **1**, 195–200. Reprinted in Wheeler & Zurek (1983).
- Bell, J. (1987a), 'Against measurement'. In J. Bell, *Speakable and Unsayable in Quantum Mechanics* (Cambridge: Cambridge University Press).
- Bell, J. (1987b) *Speakable and Unsayable in Quantum Mechanics* (Cambridge: Cambridge University Press).
- Bohm, D. (1952), 'A suggested interpretation of the quantum theory in terms of "hidden" variables, I and II'. *Phys. Rev.* **85**, 166–93. Reprinted in Wheeler & Zurek (1983).
- Bohr, N. (1928), 'The quantum postulate and the recent development of atomic theory'. *Nature*, **121**, 580–589.
- Bohr, N. (1935), 'Can quantum mechanical description of physical reality be considered complete?' *Phys. Rev.* **48**, 696–702.

- Bohr, N. (1949) 'Discussions with Einstein on epistemological problems in atomic physics'. In P. A. Schilpp (ed) *Albert Einstein -Philosopher Scientist*, 2nd edn, pp. 199–242 (New York: Tudor, Library of living philosophers. First edition 1949).
- Cartwright, N. (1980) 'Measuring position probabilities' In P. Suppes (Ed) *Studies in the Foundations of Quantum Mechanics* (East Lansing, Ill: Philosophy of Science Association).
- Cartwright, N. (1983), *How the Laws of Physics Lie* (Oxford: Clarendon Press).
- Chen, R. (1990), 'The real wave function as an integral part of Schrödinger's basic view on quantum mechanics'. *Physica, B*, 167, 183–184.
- Chen, R. (1993), 'Schrödinger's real wave equation (continued)'. *Physica B*, **190**, 256–258.
- Cohen, R., Hilpinen, R. & Renzong, Q. (Eds) (1996), *Realism and Anti-Realism in the Philosophy of Science* (Dordrecht: Kluwer).
- Colodny, R. (Ed) (1972), *Paradigms and Paradoxes* (Pittsburgh: University of Pittsburgh Press).
- Cormier-Delanoue, C. (1996) 'Strangeness of matter waves'. *Found. Phys*, **26**, 95–103.
- Cushing, J. (1991), 'Quantum theory and explanatory discourse; endgame for understanding?' *Phil. Sci.* **58**, 337–358.
- Davies, E. B. (1976), *Quantum Theory of Open Systems* (New York: Academic Press).
- Dicke, R. H. & Wittke, J.P. (1960), *Introduction to Quantum Mechanics* (Reading, MA: Addison-Wesley).
- Dummett, M. (1991a), 'Metaphysical disputes over realism'. In M. Dummett, *The Logical Basis of Metaphysics*, pp. 1–19 (Cambridge, MA: Harvard University Press).
- Dummett, M. (1991b), *The Logical Basis of Metaphysics* (Cambridge, MA: Harvard University Press).
- Einstein, A. (1905, 1965), 'Über einen die Erzeugung und Verwandlung des Lichtes betreffenden heuristischen Gesichtspunkt.' *Ann. d. Phys ser 4*, **17**, 132–148. English translation by A. B. Arons & M. B. Peppard (1965), 'Einstein's proposal of the Photon Concept – a translation of the Annalen der Physik paper of 1905.' *Am. J. Phys.*, **33**, 367–374.
- Einstein, A., Podolsky, B. & Rosen, W. (1935) 'Can quantum mechanical description of physical reality be considered complete?' *Physical Review*, **47**, 777–780.
- Elitzur, A. C. & Vaidman, L. (1993), 'Quantum mechanical interaction free measurements'. *Found. Phys.* **23**, 987–997.
- Everett, H., III (1957), 'Relative state formulation of quantum mechanics', *Rev. Mod. Physics* **29**, 454–462.
- Feynman, R. (1967), *The Character of Physical Law* (Cambridge, MA: MIT Press).
- Frank, P. (1947), *Einstein - His Life and Times* (New York: A. Knopf).
- Fraassen, B van (1980), *The Scientific Image* (Oxford: Clarendon Press).
- French, S. & Redhead, M. (1989), 'Why the principle of the identity of indiscernibles is not contingently true either'. *Synthese* **78**, 141–166.

- Ghirardi, A. C., Rimini, A. & Weber, T. (1986), 'Unified dynamics for microscopic and macroscopic systems'. *Phys. Rev. D* **34**, 470–491.
- Good, I. J. (Ed) (1961, 1962), *The Scientist Speculates* (Heinemann: London; New York: Basic Books).
- Goodman, N. (1984) *Of Mind and Other Matters* (Cambridge, MA: Harvard University Press).
- Healy, R. (1984), 'How many worlds?' *Nous* **XVIII**, 591–616.
- Hegerfeldt, G. (1974), 'Remark on causality and particle localization'. *Phys. Rev. D* **10**, 3320–3321.
- Hegerfeldt, G. (1994), 'Causality problems for Fermi's two-atom system.' *Phys. Rev. Lett.* **72**, 596–599.
- Heisenberg, W. (1930), *The Physical Principles of Quantum Theory* (Chicago: Chicago University Press).
- Heisenberg, W. (1958), *Daedalus* **87**, 100.
- Heisenberg, W. (1962), *Physics and Philosophy* (New York: Harper & Row).
- Heisenberg, W. (1971), *Physics and Beyond: Encounters and Conversations* (New York: Harper & Row).
- Hooker, C. A. (1972), 'The nature of quantum mechanical reality: Einstein versus Bohr'. In Colodny (ed): *Paradigms and Paradoxes* (Pittsburgh: University of Pittsburgh Press).
- Isham, C. J. (1995), *Lecture Notes on Quantum Theory; Mathematical and Structural Foundations* (London: Imperial College Press).
- Jammer, M. (1974), *The Philosophy of Quantum Mechanics* (New York: Wiley).
- Johansson, L.-G. (1996), 'Van Fraassen's constructive empiricism - a critique'. In R. Cohen, R. Hilpinen & Q. Renzong (Eds) *Realism and Anti-Realism in the Philosophy of Science* (Dordrecht: Kluwer), pp. 339–341.
- Kaku, M. (1993), *Quantum Field Theory* (Oxford: Oxford University Press).
- Kochen, S. & Specker, E. (1967), 'The problem of hidden variables in quantum mechanics'. *Journal of Mathematics and Mechanics* **17**, 59–87.
- Kramer, B. (Ed) (1991), *Quantum Coherence in Mesoscopic Systems* (New York: Plenum Press).
- Kuhn, .T (1962), *The Structure of Scientific Revolutions* (Chicago: Chicago University Press).
- Kuhn, T. (1978, 1987), *Black-Body Theory and the Quantum Discontinuity 1894-1912* (Oxford: Oxford University Press; 2nd edn, Chicago: University of Chicago Press).
- Kwiat, P. G, Steinberg, M & Chiao, Y. (1992), 'Observation of a "quantum eraser": a revival of coherence in a two-photon interference experiment'. *Phys. Rev. A* **45**, 7729–7739.
- Lakatos, I. (1970), 'Methodology of scientific research programmes'. In I. Lakatos & Musgrave (Eds) *Criticism and the Growth of Knowledge* (Cambridge: Cambridge University Press).
- Lakatos, I. & Musgrave, A. (Eds) (1970), *Criticism and the Growth of Knowledge* (Cambridge: Cambridge University Press).
- Landé, A. (1965), *New Foundations of Quantum Mechanics* (London: Cambridge University Press).

- Leggett, A. (1986), 'The superposition principle in macroscopic systems'. In Penrose & Isham (Eds) *Quantum Concepts in Space and Time* (Oxford: Clarendon Press).
- London & Bauer (1939), 'La theorie de l'observation en mécanique quantique', No 775 of *Actualités scientifiques et industrielles: Exposés de physique generale, publiés sous la direction de Paul Langevin* (Paris: Hermann). English translations – including a new paragraph by London – done independently by A. Shimony, J. A. Wheeler & W. H. Zurek, and J. McGrath & S. McLean McGrath; reconciled in 1982. Printed as pp. 217–259 in Wheeler & Zurek (Eds) (1983).
- Luce, R. D. & Suppes, P. (1975) 'Measurement, theory of'. *Encyclopaedia Britannica*.
- Margenau, H. (1937), 'Critical notice in modern physical theory'. *Phil. Sci.* **4**, 352–356.
- Margenau, H. (1958), 'Philosophical problems concerning the meaning of measurement in physics'. *Phil. Sci.* **25**, 23–34.
- Maxwell, J. C (1931), *James Clerk Maxwell: a Commemoration volume 1831–1931* (Cambridge: Cambridge University Press).
- McGervey, J. D. (1971), *Introduction to Modern Physics* (New York and London: Academic Press).
- McGregor, M. H. (1992), *The Enigmatic Electron* (Dordrecht: Kluwer Academic).
- Merve, A. van der & Garuccio, A. (Eds) (1994), *Waves and Particles in Light and Matter* (New York and London: Plenum).
- Mehra, J. & Rechenberg, H. (1987), *The Historical Development of Quantum Theory*, vol. 5 (New York, Berlin: Springer Verlag).
- Meystre, P. & Sargent, P., III (1990), *Quantum Optics* (Berlin: Springer Verlag).
- Misra, B., Prigogine, I. & Courbage, M. (1979, 1983), 'Lyapounov variable: entropy and measurement in quantum mechanics'. *Proceedings of the National Academy of Sciences of Unites States of America*, **76**, 4768–4772, reprinted in Wheeler & Zurek (Eds) (1983), pp. 687–691.
- Moore, W. (1989) *Schrödinger: Life and Thought* (Cambridge: Cambridge University Press).
- Murdoch, D. (1987), *Niels Bohr's Philosophy of Physics* (Cambridge: Cambridge University Press).
- Neumann, J. von (1932), *Grundlagen der Quantenmechanik* (Berlin: Springer Verlag).
- Penrose, R. & Isham, C. J. (1986), *Quantum Concepts in Space and Time* (Oxford: Clarendon Press).
- Peres, A. (1998), Book review of J. Bub: 'Interpreting the Quantum Word.' in *Stud. History Philos. Modern Physics* **29**, 611–623. Also electronic document xxx.lanl.gov/archive/quant-ph/9711003.
- Phlegor, R. L. & Mandel, L. (1968), 'Further experiments on interference of independent photon beams at low light levels'. *J. Opt. Soc. Am.* **58**, 946–950.
- Planck, M. (1900), 'Zur Theorie des Gesetzes der Energiverteilung im Normalspektrum'. *Ver. d. D. Phys. Ges.* **2**, 237–345.

- Post, H. R. (1963), 'Individuality and physics'. Radio talk repeated in *The Listener*, 10 October 1963. Also Departmental Report, Department for History and Philosophy of Science, Chelsea College, University of London.
- Przibram, K. (Ed) (1967), *Letters on Wave Mechanics* (New York: Philosophical Library).
- Quine, W. V. O. (1981), *Theories and Things* (Cambridge, MA: Belknap Press).
- Quine, W. V. O. (1990), *Pursuit of Truth* (Cambridge, MA. and London, England: Harvard University Press).
- Quine, W. V. O. (1994), 'Promoting extensionality'. *Synthese* **98**, 143–151.
- Quine, W. V. O. (1995), *Pursuit of Truth* (Cambridge, MA. and London, England: Harvard University Press).
- Rae, A. (1986), *Quantum Mechanics*, 2nd edn (Bristol: Adam Hilger).
- Rauch, H., Wölwitsch, H., Clothier, R., Kaiser, H. & Werner, S. A. (1992), 'Time of flight neutron interferometry' *Phys. Rev. A* **46**, 45–57.
- Redhead, M. (1987), *Incompleteness, Nonlocality and Realism* (Oxford: Clarendon Press).
- de Regt, H. W. (1997), 'Erwin Schrödinger, Anschaulichkeit and quantum theory'. *Stud. Hist. Phil. Mod. Phys.* **28**, 461–481.
- Renninger, M. (1960), 'Messungen ohne Störung des "Messobjekts"'. *Z. Phys.* **158**, 417–421.
- Rosenthal, S. (Ed) (1967), *Niels Bohr – His Life and Work as seen by his Friends and Colleagues* (North-Holland, Amsterdam).
- Salmon, W. (1984), *Scientific Explanation and the Causal Structure of the World* (Princeton: Princeton University Press).
- Schilpp, P. A. (Ed) (1951), *Albert Einstein – Philosopher Scientist*, 2nd edn (New York: Tudor, Library of Living Philosophers. First edition 1949).
- Shimony, A. (1986), 'Events and processes in the quantum world'. In Penrose & Isham *Quantum Concepts in Space and Time* (Oxford: Clarendon Press).
- Schrödinger, E. (1926a), 'Über den Comptoneffekt', *Ann. d. Phys.* **82**, 257–264.
- Schrödinger, E. (1926b), 'Der Energieimpulssatz der Materiewellen', *Ann. d. Phys.* **82**, 265–272.
- Schrödinger, E. (1928), *Collected Papers on Wave Mechanics* (London: Blackie & Son).
- Schrödinger, E. (1935, 1983), 'The present situation in quantum mechanics'. In J. A. Wheeler & W. Zurek (Eds) (1983) *Quantum Theory and Measurement* (Princeton: Princeton University Press), pp. 152–167. First published as 'Die gegenwärtige Situation in der Quantenmechanik', *Naturwissenschaften* **23**, (1935).
- Schrödinger, E. (1995), *The Interpretation of Quantum Mechanics*. Dublin seminars (1949–1955) and other unpublished essays. Michel Bitbol (Ed.) (Woodbridge, CO: Ox Bow Press).
- Shannon, C. & Weaver, W. (1949), *The Mathematical Theory of Communication* (Urbana: University of Illinois Press).
- Stein, H. (1984), 'The Everett interpretation of quantum mechanics: many worlds or none?', *Nous* **XVIII**, 635–652.

- Stern, A., Aharonov, Y. & Imry, Y. (1991), 'Linear response and dephasing by Coulomb electron-electron interactions'. In B. Kramer (Ed) *Quantum Coherence in Mesoscopic Systems* (New York: Plenum Press).
- Strawson, P. F. (1959), *Individuals. An Essay in Descriptive Metaphysics* (London: Methuen).
- Suppes, P. (Ed) (1980), *Studies in the Foundations of Quantum Mechanics* (East Lansing, Ill: Philosophy of Science Association).
- Tarozzi, G. & van der Merve, A. (Eds) (1985), *Open Questions in Quantum Physics* (Reidel: Dordrecht).
- Teller, P. (1986), 'Relational holism and quantum mechanics', *British Journal for the Philosophy of Science* **37**, 71–81.
- Vigier, J. P. (1985), 'Nonlocal quantum potential interpretation of relativistic actions at a distance in many-body problems'. In Tarozzi & van der Merve (Eds) *Open Questions in Quantum Physics* (Reidel: Dordrecht), pp. 297–322.
- Weinberg, S. (1995), *The Quantum Theory of Fields*, vol. 1 (Cambridge: Cambridge University Press).
- Wheeler, J. A. (1983), 'Law without law'. In J. A. Wheeler & W. Zurek (Eds) *Quantum Theory and Measurement* (Princeton: Princeton University Press), pp. 182–213.
- Wheeler, J. A. & Zurek W. (Eds) (1983), *Quantum Theory and Measurement* (Princeton: Princeton University Press).
- Wigner, E. (1931), *Gruppentheorie und ihre Anwendung auf die Quantenmechanik der Atomspektren* (Braunschweig).
- Wigner, E. (1961), 'Remarks on the mind-body-question'. In Good (Ed) *The Scientist Speculates* (London: Heinemann). Reprinted in Wigner (1967), and in Wheeler & Zurek (1983).
- Wigner, E. (1963), 'The problem of measurement.' *Am. J. Phys.* **31**, 6–15. Reprinted in Wheeler & Zurek (1983), pp. 324–341.
- Wigner, E. (1967), *Symmetries and Reflexions* (Bloomington: Indiana University Press).
- Wigner, E. (1983), 'Interpretation of quantum mechanics'. In J. A. Wheeler & W. Zurek (Eds) *Quantum Theory and Measurement* (Princeton: Princeton University Press), pp. 260–314.
- Wolpert, L. (1992), *The Unnatural Nature of Science* (London: Faber & Faber).
- Zurek, W. (1991) 'Decoherence and the transition from quantal to classical'. *Physics Today* October, 36–44.

Index

- actual properties 47
- Aharonov & Vaidman 67
- Aharonov-Bohm effect 155
- Albert 170
- Aspect 12, 145
- Auyang 183

- Barut 16, 79, 108, 176
- Bell 96
- Bell's theorem
 - derivation of 143
- Bell locality 146
- black-body radiation 87
- Bohm 168
- Bohr 20, 24, 86, 113, 159, 178
- Bohr's theory of measurement 162
- Born 20
- Born's rule 161
- Born interpretation 4
- Bose-Einstein statistics 25, 37

- cardinal number 39
- Cartwright 96, 185
- charge density 58
- Chen 63
- cluster decomposition principle 43
- collapse
 - as spatial contraction 104
- collapse of a superposition 106
- complementarity 159
- complementarity principle
 - definition of 160
- Compton effect 90
- continuity-discontinuity 181
- Copenhagen interpretation 14, 17, 159
- Cormier-Delanoue 75
- correlation function 144
- correspondence principle 109
- Cushing 20, 24

- de Broglie 166
- decoherence theory 172
- delayed choice 123
- delayed choice experiment 6
- density current 21
- de Regt 20
- derivation of Schrödinger's equation 84
- determinism 141
- dispersion of wave packet 74
- dispositions 48
- Doppler effect for matter waves 75
- Dummett 14

- Einstein 16, 17, 20, 86, 179
- Einstein, Podolsky & Rosen 10
- Einstein locality 146
- electromagnetic interaction 101
- Elitzur & Vaidman 117
- entangled state 41
- entropy
 - constant during measurement 107
- EPR 24
- EPR argument 142
- essential pronouns 30
- Everett 164

- FAPP 93
- Fermi-Dirac statistics 25, 37
- Feynman 19
- fields and events 183
- Frege-Russell analysis of number 39
- French 36

- Gell-Mann 19
- genidentity 32
- Ghirardi, Rimini & Weber 169

- Healey 165
- Hegerfeldt 154
- Heisenberg 17, 20, 47, 162

- Hong-Ou-Mandel interferometer 125
- Hooker 160
- identity criterion 29
- identity of indiscernibles 36
- identity of properties 45
- identity of states 50
- identity postulate 35
- individuation 30
- information 99
- information theory 99
- instrumentalism 3
- interaction-free measurement 117
- irreversibility
 - degree of 112
 - due to randomness 112
- Isham 44, 136
- Josephson effect 116
- Kaku 47, 82
- Kochen-Specker's theorem 136
- Kuhn 22, 85
- Kwiat, Steinberg & Chia 124
- Lakatos 161
- Landé 47, 81
- Leggett 114, 116
- Leibniz principle 36
- local fields 183
- locality 141
- London & Bauer 163
- MacGregor 91
- Mach-Zehnder interferometer 117
- magnetic momentum 129
- many-worlds interpretation 164
- Margenau 96
- maximal observable 46
- Maxwell-Boltzmann statistics 25
- measurement by environment 98
- measurement, definition of 100
- measurement of the first kind 48, 95
- measurement theory 97
- Meystre & Sargent 34
- minimal realism 15
- Misra, Prigogine & Courbage 107
- Murdoch 17, 160
- negative result experiment 121
- neutron interference 71
- Noether's theorem 19
- non-linear equation 79, 107
- observables 46
- observation statements 96
- ordinal number 39
- path integral formalism 82
- Peres 49
- phase velocity is higher than the velocity of light 69
- Phlegor & Mandel 40
- photon 35, 86
- physical quantities 45
- pilot-wave interpretation 166
- Planck 85
 - on quantisation 178
- pointer state 97
- Poisson bracket 110
- potential properties 47
- precession 132
- projection postulate 99
- pronouns of laziness 30
- protective measurement 67
- quantisation 85, 87, 88
- quantisation of charge 36
- quantitative magnitudes 45
- quantum eraser 124
- quantum field theory 182
- quantum postulate 86
- quantum potential 167
- quantum Zeno effect 87, 108
- Quine 25, 29, 161
- Rae 135
- Rauch 71
- received view 3
- Redhead 36, 141
- relative state formulation 164
- Renninger 121
- S-matrix 43
- Schrödinger 2, 20, 22, 58, 59, 78, 90
 - letter to Lorentz 58, 62
- Schrödinger's cat 8
- Schrödinger-Bohr discussion 179
- Schrödinger equation 84

- non-linear 107
- semantic interpretation 25
- Shannon & Weaver 99
- Shimony 47, 108
- singular definite description 31
- special theory of relativity 150
- SQUID 116
- state change – real change 54
- Stein 165
- Stern et al. 98
- Stern-Gerlach apparatus 101
- Strawson 31
- superposition of macroscopic states 114
- Teller 4, 42
- theoretical statements 96
- Theseus' ship 32
- time reversal 52
- unitary/anti-unitary evolution 51
- value determinism 16
- value function 136
- van Fraassen 18
- visualisability 20
- von Neumann's cut object/subject 163
- Weinberg 43
- Weisskopf 58
- Wheeler 17, 123
- Wigner 51, 93, 101, 163
- Zurek 98, 99, 172