



A proposal to get some common-sense intuition for the paradox of the double-slit experiment

Gerrit Coddens

► To cite this version:

Gerrit Coddens. A proposal to get some common-sense intuition for the paradox of the double-slit experiment. 2016. cea-01383609v5

HAL Id: cea-01383609

<https://cea.hal.science/cea-01383609v5>

Preprint submitted on 5 May 2017 (v5), last revised 26 Oct 2021 (v10)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

A proposal to get some common-sense intuition for the paradox of the double-slit experiment

Gerrit Coddens

Laboratoire des Solides Irradiés, Ecole Polytechnique, CNRS, CEA, Université Paris-Saclay, 91128-Palaiseau CEDEX, FRANCE

Abstract. We argue that the double-slit experiment can be understood much better by considering it as an experiment whereby one uses electrons to study the set-up rather than an experiment whereby we use a set-up to study the behaviour of electrons. We also show that the concept of undecidability (like e.g. occurs in Gödel's theorem) can be used in an intuitive way to make sense of the double-slit experiment and the quantum rules for calculating coherent and incoherent probabilities. We meet here a situation where the electrons always behave in a fully deterministic way, while the detailed design of the set-up may render the question about the way they move through the set-up experimentally undecidable. It is very important to make a distinction in quantum mechanics between the determinism of nature (Einstein) and the decidability of a question within an experimental set-up (Bohr). The former is about the absolute truth of an answer to a yes or no question, and follows binary logic (true or false), the latter about what an experimental set-up can decide and tell about the truth of that answer, and follows ternary logic (true, false or undecidable). Binary and ternary logic are incompatible. The viewpoints of Bohr and Einstein are thus operating on different levels and it is only by confusing these two levels that these two viewpoints seem to be irreconcilable. We show that the expression $\psi_1 + \psi_2$ for the wave function of the double-slit experiment is numerically correct, but logically flawed. It has to be replaced in the interference region by the logically correct expression $\psi'_1 + \psi'_2$, which has the same numerical value as $\psi_1 + \psi_2$, such that

$$\psi'_1 + \psi'_2 = \psi_1 + \psi_2, \text{ but with } \psi'_1 = \frac{\psi_1 + \psi_2}{\sqrt{2}} e^{i\frac{\pi}{4}} \neq \psi_1 \text{ and } \psi'_2 = \frac{\psi_1 + \psi_2}{\sqrt{2}} e^{-i\frac{\pi}{4}} \neq \psi_2. \text{ Here } \psi'_1 \text{ and } \psi'_2 \text{ are the correct contributions}$$

from the slits to the total wave function $\psi'_1 + \psi'_2$. We have then $p = |\psi'_1 + \psi'_2|^2 = |\psi'_1|^2 + |\psi'_2|^2 = p_1 + p_2$ such that the paradox that quantum mechanics (QM) would not follow the traditional rules of probability calculus disappears. The paradox is rooted in the wrong intuition that ψ_1 and ψ_2 would be the true physical contributions to $\psi'_1 + \psi'_2 = \psi_1 + \psi_2$ like in the case of waves in a water tank. The solution proposed here is not *ad hoc* but based on an extensive analysis of the geometrical meaning of spinors within group representation theory and its application to QM. Working further on the argument one can even show that an interference pattern is the only way to satisfy simultaneously two conditions: The condition obeying binary logic (in the spirit of Einstein) that the electron has only two mutually exclusive options to get to the detector (viz. going through slit S_1 or going through slit S_2) and the condition obeying ternary logic (in the spirit of Bohr) that the question which one of these two options the electron has taken is experimentally undecidable. A very important element in the analysis is the problem of the non-existence of *common* probability distributions for different experimental set-ups. And this is a recurrent theme in other situations that are felt as paradoxical. The CHSH inequality used in the experiments of Aspect *et al.* clashes with quantum mechanics because it is based on the assumption that there would exist a *common* probability distribution for the hidden variables in the different experiments that have to be made to determine the various quantities that intervene in the inequality. That common probability distribution for measured quantities does not exist is well known from quantum mechanics itself (because the operators that come into play do not commute and therefore do not have common eigenstates, i.e. common probability amplitudes). But this non-existence of common probability distributions applies also to certain variables in the set-ups of the different experiments and the latter fact can also be explained by purely classical reasoning.

PACS. 03.65.Ta Foundations of Quantum Mechanics; measurement theory – 42.25.Hz Interference – 02.20.-a Group Theory – 03.65.Pm Relativistic wave equations

1 Introduction

In reference [1] we derived the Dirac equation by just expressing that the electron is a spinning particle.¹ The derivation is entirely classical in the sense that we start by considering a particle in free space that travels in uniform motion on a straight line according to $\mathbf{r}(t) = \mathbf{r}_0 + \mathbf{v}t$. A Dirac-like equation (which is still different from the Dirac equation) is then derived on the set

¹ The current dogma is that spin does not correspond to spinning motion because such an assumption cannot explain the magnetic moment of the electron. But this conclusion is based on a poor understanding of the mathematics and not compelling as discussed in [2].

$\mathcal{P} = \{(\mathbf{r}, t) \in \mathbb{R}^4 \mid \mathbf{r}(t) = \mathbf{r}_0 + \mathbf{v}t\}$. A last step in the derivation (which we will explain below in Section 2) permits to obtain on $\mathcal{P}$ the Dirac equation from the Dirac-like equation. On reading these lines one may think that it cannot possibly be right that the quantum mechanical Dirac equation can be derived from a classical reasoning: It just sounds crank. However, we have argued in reference [1] (see p. 125) and we will illustrate in this paper how *it is not the way we derive it but the way we solve it, that turns the Dirac equation into quantum mechanics*. In fact, when we solve the wave equations in quantum mechanics we transgress the definition domain of the classically derived equations in two instances,² and it is through these transgressions that we introduce quantum mechanics.

The first transgression occurs when we derive the Dirac equation from the Dirac-like equation. In this step, we replace a spinor by a sum of spinors. While it is algebraically perfectly feasible to make a linear combination of two spinors, the meaning of such a procedure is not defined by the group theory. The reason for this is that spinors represent group elements and that the rotation group and the homogeneous Lorentz group we use in physics are not vector spaces but curved Lie group manifolds.³ Making a linear combination of spinors is thus *a priori* a meaningless operation within the pristine conceptual frame work of the group theory.⁴ One could qualify using such linear combinations as a mindless use of the algebra whereby one does not realize that one transgresses the limits of the special definition domain and conceptual framework of the group theory in carrying out such calculations. In doing so, one complies exactly with the motto: “shut up and calculate”. Fortunately, one can find an *a posteriori* justification for this procedure, by finding a meaning for it in terms of sets. These sets can be seen to represent statistical ensembles and their use turns therefore the Dirac equation into statistical physics. We will explain this in Section 2.

The second transgression occurs when we solve the Dirac equation. In fact, following the classical ideas that underly the derivation, we should solve the Dirac equation on the set $\mathcal{P}$. Instead of that, we solve it generously and thoughtlessly over $\mathbb{R}^4$, without realizing that we only needed to solve it over $\mathcal{P}$. While solving the equation over $\mathbb{R}^4$ is algebraically possible, it is once more a shut-up-and-calculate approach of which the meaning has *a priori* not been defined, such that we must again try to find an *a posteriori* justification for it. As we claimed that it is the way we solve the Dirac equation that turns it into quantum mechanics, we can anticipate already that this will prove a very difficult task. And as a matter of fact, it is through this kind of over-zealous extrapolation of the definition domain of the equation from $\mathcal{P}$ to $\mathbb{R}^4$ that we allow all the trickery of quantum mechanics to enter the scene.⁵ It is thus very important to figure out what this extrapolation exactly means. We will see that there is no unique answer to this question and that the solution may very much depend on the specific physical context considered, such that it requires a discussion on a case-by-case basis. The fact that such extrapolations always seem to work confers a deep beauty to QM.

² Throughout the paper, we rely on the fact that the Schrödinger can be derived from the Dirac equation, such that we cover both equations by discussing everything in terms of the Dirac equation.

³ E.g. in SU(2), spinors represent rotations as explained in references [1] (see pp. 48-52) and [2]. We can then get a feeling for the fact that a linear combination of two spinors will not be a new spinor by an analogy: A linear combination of two rotation matrices in SO(3) is not a new rotation matrix.

⁴ The example of 3×3 rotations matrices suggests that we could use it as a source of inspiration to give a meaning to a linear combination $\mathbf{M} = c_1 \mathbf{R}_1 + c_2 \mathbf{R}_2$ of rotation matrices $\mathbf{R}_1$ and $\mathbf{R}_2$. It is the matrix that transforms a vector $\mathbf{v}$ expressed as the 3×1 column matrix $[\mathbf{v}]$ to $c_1 [\mathbf{v}_1] + c_2 [\mathbf{v}_2]$ whereby $[\mathbf{v}_1] = \mathbf{R}_1 [\mathbf{v}]$ and $[\mathbf{v}_2] = \mathbf{R}_2 [\mathbf{v}]$. This is a valid procedure for the 3×3 representation, because the 3×1 matrices represent elements of a vector space. But in SU(2), the spinors are just a stenographic notation for the SU(2) matrices obtained by taking their first columns (see reference [1]). The procedure to give meaning to linear combinations of SU(2) matrices in analogy with the procedure in SO(3) can then not be used because the second-stage question what a linear combination of spinors would be is exactly the same question as the original question what a linear combination of SU(2) matrices would be. This turns the attempt to solve the problem in analogy with SO(3) into a vicious circle. The analogous procedure fails because spinors (which are the 2×1 matrices) do not belong to a vector space like the vectors of $\mathbb{R}^3$ (which are the 3×1 matrices). We could of course argue that we must search for the meaning of the sum of spinors by translating the problem to $\mathbb{R}^3$, solve it there, and then translate the solution backwards to $\mathbb{C}^2$, but this has several problems in its own right:

(1) First of all it leads exactly to the conceptual difficulty that in Atiyah's words: “A spinor is the square root of a vector” [3]. An illustration of this fact is that probabilities (which from the relativistic point of view are the time-component of a probability charge-current four-vector) are squares of probability amplitudes ψ (which relativistically are also a component of a four-component spinor-like quantity $\mathcal{P}$). Understanding the difference between summing probabilities and summing probability amplitudes is the very hard nut we have to crack if we want to make sense of the double-slit experiment, which in Feynman's words is the only mystery of quantum mechanics [4]. As this quadratic relationship is very difficult to understand, the clear meaning of the sum of vectors in SO(3) gets lost in the attempt to translate it backwards to SU(2) and to solve the riddle of the meaning of summing probability amplitudes. The relation between the SO(3) matrices and the SU(2) matrices is also quadratic.

(2) The second problem is that we are straying out of the conceptual framework we set up to study the group. This conceptual framework of SU(2) contains all we need to know about the group in a self-contained way. We should thus not look for solutions elsewhere. It is completely artificial to draw in *external* considerations from SO(3) into SU(2). By external considerations we mean here that we are using arguments from another representation SO(3) to settle problems in SU(2). The logic of SU(2) should be self-contained (i.e. internal) and not rely on arguments drawn in from SO(3). Moreover, the *internal* approach we will develop and which is based on sets and statistical ensembles reproduces exactly (within a completely self-consistent internal logic) the quantum rules that have already been proposed in physics to solve the riddle how we must interpret a superposition of two states. These rules can be formulated within SU(2) without invoking SO(3) at any time.

⁵ This can already be appreciated from the fact that we solve the Dirac or the Schrödinger equation also over the classically forbidden regions under and at the other side of the barrier when we solve these equations for a tunneling experiment. It is certainly legitimate to ask why we should consider such a startling extrapolation of the definition domain.

It is the aim of this paper to try to clarify the issues which arise in these two transgressions. We will discuss them at the hand of a result of quantum mechanics (QM) that is particularly counter-intuitive and difficult to understand, viz. the double-slit experiment. The discussion about the two transgressions of the definition domains in the mathematics is of a type that may look like splitting hairs to many a physicist. But they really lead straight into the heart of the matter. The analysis we present here of the double-slit experiment is meant to show, we hope in a convincing way, the truth of the daring thesis that the Dirac equation is classical and that it is only the way we solve it that renders it QM. *The readers who want to make the effort to study both the derivation of the Dirac equation in reference [1] and the present paper, will be able to verify that the whole presents a rigorous (but rather long) step-by-step mathematical argument that starts from scratch and whereby every step is justified and intuitive.*

We hope the reader will recognize the importance of this approach. Its power resides in the fact that one can exactly spot where and how we cross the frontier between classical mechanics and QM. We can use the knowledge about the exact location of this frontier as a bistoury to analyze in detail what the transgressions may mean and this way gain a better understanding of QM. This is something that we may never have guessed or anticipated, because at face value it just does not sound right to claim that the Dirac equation can be derived classically. We must also point out that in the traditional approach to QM, one cannot possibly become aware of the two transgressions we are discussing here. They become only apparent in our approach, which may illustrate why it could be important. Any possible perceived irony about these transgressions is thus just for didactical reasons. It may already have transpired from the preceding lines, but we must warn the reader that it requires a special mindset to meet with our approach. Its merits cannot be evaluated by promoting rigidly (a plethora of) strong beliefs anchored in the traditional analysis to peremptory touchstones. That would be as clueless as criticizing hyperbolic geometry from the stronghold of the premisses of Euclidean geometry. He should be prepared to question the current dogma as possibly *ad hoc*.

2 The first transgression: From a Dirac-like to the Dirac equation

With the Cartan-Weyl choice for the gamma matrices, the free-space Dirac-like equation reads:

$$\left[\begin{array}{c} -\frac{\hbar}{i} \frac{\partial}{\partial ct} \mathbb{1} + \frac{\hbar}{i} \nabla_{\parallel} \cdot \sigma \\ -\frac{\hbar}{i} \frac{\partial}{\partial ct} \mathbb{1} - \frac{\hbar}{i} \nabla_{\parallel} \cdot \sigma \end{array} \right] \left[\begin{array}{c} \Psi \\ \Psi^{-1\dagger} \end{array} \right] = m_0 c \left[\begin{array}{c} s_{ct} \mathbb{1} + \mathbf{s} \cdot \sigma \\ -(s_{ct} \mathbb{1} - \mathbf{s} \cdot \sigma) \end{array} \right] \left[\begin{array}{c} \Psi \\ \Psi^{-1\dagger} \end{array} \right]. \quad (1)$$

Here $\mathbf{s}$ is the spin axis. Very often one supposes that $\mathbf{s} = \mathbf{e}_z$ as we identify $\mathbf{s}$ with the basis vector $\mathbf{e}'_z$ of some triad of basis vectors that have been originally attached to the electron to permit following its spinning motion. The directional partial derivative $\nabla_{\parallel}$ is here taken in the direction of a boost velocity vector $\mathbf{v}$. The fact that $-(s_{ct} \mathbb{1} - \mathbf{s} \cdot \sigma)$ occurs here with a minus sign is due to the fact that $\mathbf{s}$ is a pseudo-vector. This equation is derived from the Rodrigues formula in SU(2) for a rotation around an axis $\mathbf{n}$ over an angle φ within a frame at rest:

$$\mathbf{R}(\mathbf{n}, \varphi) = \cos \frac{\varphi}{2} - i \sin \frac{\varphi}{2} [\mathbf{n} \cdot \sigma]. \quad (2)$$

The substitution $\varphi \mapsto \omega_0 \tau$ (where τ is the proper time) transforms this into the description of a particle that spins around $\mathbf{n}$. After using the 2×1 spinor ψ as a stenographic notation for the 2×2 matrix $\mathbf{R}(\mathbf{n}, \varphi)$ by taking its first column (see references [1]-[2]), and taking the time derivative, one obtains:

$$\frac{d}{d\tau} \psi(\tau) = -i \frac{\omega_0}{2} [\mathbf{n} \cdot \sigma] \psi(\tau). \quad (3)$$

The fact that we use $\mathbf{s}$ instead of $\mathbf{n}$ in Eq. 1 hides a tedious technicality we discuss in reference [1]. The idea is to render the equation independent from the choice of a reference frame by rendering it covariant. This does not work with $\mathbf{n}$ because it is not a covariant quantity and it does not transform like a vector, while it works for $\mathbf{s}$ which is a covariant quantity and does transform like a vector. Eq. 1 is just a covariant reformulation of Eq. 3 in a moving frame with boost velocity $\mathbf{v}$. We have then $(\frac{d}{d\tau}, 0) \rightarrow (\frac{\partial}{\partial ct}, \nabla_{\parallel})$. To obtain Eq. 1 from Eq. 3, one must introduce $m_0 c^2 = \hbar \omega_0 / 2$ (with the effect that $e^{-i\omega_0 \tau / 2}$ becomes $e^{-\frac{i}{\hbar}(Et - \mathbf{p} \cdot \mathbf{r})}$), and lift the equation from SU(2) to the representation of the Lorentz group in terms of gamma matrices. All these steps are discussed in full detail in reference [1]. We obtain this way the description of a spinning particle that moves at a constant velocity $\mathbf{v}$ along a straight line $\mathbf{r} = \mathbf{r}_0 + \mathbf{v}t$, $t \in \mathbb{R}$. If we can assume that $\mathbf{s} \perp \mathbf{v}$ then $s_t = 0$ and:

$$\left[\begin{array}{c} -\frac{\hbar}{i} \frac{\partial}{\partial ct} \mathbb{1} + \frac{\hbar}{i} \nabla_{\parallel} \cdot \sigma \\ -\frac{\hbar}{i} \frac{\partial}{\partial ct} \mathbb{1} - \frac{\hbar}{i} \nabla_{\parallel} \cdot \sigma \end{array} \right] \left[\begin{array}{c} \Psi \\ \Psi^{-1\dagger} \end{array} \right] = m_0 c \left[\begin{array}{c} \mathbf{s} \cdot \sigma \\ \mathbf{s} \cdot \sigma \end{array} \right] \left[\begin{array}{c} \Psi \\ \Psi^{-1\dagger} \end{array} \right]. \quad (4)$$

The fact that $\mathbf{s}$ is an axial vector is even more obvious in this equation. We have now arrived at the point where we must transgress the definition domain of the algebra. To obtain the Dirac equation from Eq. 4, one would need that:

$$\left[\begin{array}{c} \mathbf{s} \cdot \sigma \\ \mathbf{s} \cdot \sigma \end{array} \right] \left[\begin{array}{c} \Psi \\ \Psi^{-1\dagger} \end{array} \right] = \left[\begin{array}{c} \Psi \\ \Psi^{-1\dagger} \end{array} \right]. \quad (5)$$

But this is just not possible. An operator:

$$\mathbf{S} = \begin{bmatrix} & \mathbf{s} \cdot \boldsymbol{\sigma} \\ \mathbf{s} \cdot \boldsymbol{\sigma} & \end{bmatrix} \quad (6)$$

can after operating on a group element:

$$\boldsymbol{\Psi}_1 = \begin{bmatrix} \Psi & \\ & \Psi^{-1\dagger} \end{bmatrix} \quad (7)$$

never yield the group element $\boldsymbol{\Psi}_1$ again because $\mathbf{S}\boldsymbol{\Psi}_1$ must have an off-diagonal block structure while $\boldsymbol{\Psi}_1$ has a diagonal block structure. (Matrices representing group elements obtained from an even number of reflections, i.e. rotations and boosts, always have a block-diagonal structure in the Cartan-Weyl choice for the gamma matrices). To obtain the Dirac equation, we must thus introduce a second quantity:

$$\boldsymbol{\Psi}_2 = \begin{bmatrix} & \mathbf{s} \cdot \boldsymbol{\sigma} \\ \mathbf{s} \cdot \boldsymbol{\sigma} & \end{bmatrix} \begin{bmatrix} \Psi & \\ & \Psi^{-1\dagger} \end{bmatrix} = \mathbf{S}\boldsymbol{\Psi}_1, \quad (8)$$

and introduce a superposition of states: $\boldsymbol{\Phi} = \boldsymbol{\Psi}_1 + \boldsymbol{\Psi}_2$:

$$\begin{bmatrix} \Psi & \\ & \Psi^{-1\dagger} \end{bmatrix} + \begin{bmatrix} & \mathbf{s} \cdot \boldsymbol{\sigma} \\ \mathbf{s} \cdot \boldsymbol{\sigma} & \end{bmatrix} \begin{bmatrix} \Psi & \\ & \Psi^{-1\dagger} \end{bmatrix} = \begin{bmatrix} \Psi & [\mathbf{s} \cdot \boldsymbol{\sigma}] \Psi^{-1\dagger} \\ [\mathbf{s} \cdot \boldsymbol{\sigma}] \Psi & \Psi^{-1\dagger} \end{bmatrix} = \boldsymbol{\Phi}, \quad (9)$$

such that $\mathbf{S}\boldsymbol{\Phi} = \mathbf{S}(\boldsymbol{\Psi}_1 + \boldsymbol{\Psi}_2) = \mathbf{S}(\boldsymbol{\Psi}_1 + \mathbf{S}\boldsymbol{\Psi}_1) = \mathbf{S}\boldsymbol{\Psi}_1 + \mathbf{S}^2\boldsymbol{\Psi}_1 = \boldsymbol{\Psi}_2 + \boldsymbol{\Psi}_1 = \boldsymbol{\Phi}$ and:

$$\left[-\frac{\hbar}{i} \frac{\partial}{\partial ct} \mathbb{1} + \frac{\hbar}{i} \nabla_{\parallel} \cdot \boldsymbol{\sigma} \quad -\frac{\hbar}{i} \frac{\partial}{\partial ct} \mathbb{1} - \frac{\hbar}{i} \nabla_{\parallel} \cdot \boldsymbol{\sigma} \right] \boldsymbol{\Phi} = m_0 c \boldsymbol{\Phi}. \quad (10)$$

This is algebraically feasible, but geometrically meaningless. First of all, it is even not clear that $\mathbf{S}$ is a group element, but even if it is, such that $\boldsymbol{\Psi}_2$ represents a group element, a quantity $\boldsymbol{\Psi}_1 + \boldsymbol{\Psi}_2$ is *a priori* not defined as we already mentioned in the Introduction.

In fact, in group representation theory the meaning of $\mathbf{D}(g_2)\mathbf{D}(g_1) = \mathbf{D}(g_2 \circ g_1)$, where the representation matrices $\mathbf{D}(g_j)$ represent the group elements $g_j \in G$ and $\circ$ represents the group operation of the group $(G, \circ)$, is well defined. But the meaning of a *purely formal* [5] linear combination $c_2\mathbf{D}(g_2) + c_1\mathbf{D}(g_1)$, where $(c_1, c_2) \in \mathbb{C}^2$ or $(c_1, c_2) \in \mathbb{R}^2$ is not defined. The tentative homomorphism between $c_2\mathbf{D}(g_2) + c_1\mathbf{D}(g_1)$ and $c_1g_1 + c_2g_2$ has not been defined, and even the operation $c_1g_1 + c_2g_2$ itself on the group has *a priori* not been defined. To illustrate this, we could ask the question what the meaning of the sum of two permutations p and q :

$$\begin{pmatrix} 1 & 2 & \cdots & j & \cdots & n \\ p_1 & p_2 & \cdots & p_j & \cdots & p_n \end{pmatrix} + \begin{pmatrix} 1 & 2 & \cdots & j & \cdots & n \\ q_1 & q_2 & \cdots & q_j & \cdots & q_n \end{pmatrix}, \quad (11)$$

is supposed to be.⁶ This leads us straight into the thorny problems evoked in Footnote 4.

We must thus try to give a meaning to the undefined procedure of calculating $\boldsymbol{\Phi} = \boldsymbol{\Psi}_1 + \boldsymbol{\Psi}_2$ and more generally $c_2\mathbf{D}(g_2) + c_1\mathbf{D}(g_1)$. We will define this way, what has been called the group ring [7]. We can illustrate the idea for Pauli's formalism for the spin in $SU(2)$. For a reflection operator $[\mathbf{s} \cdot \boldsymbol{\sigma}]$ we can consider the eigenvalue equation:

$$[\mathbf{s} \cdot \boldsymbol{\sigma}] \psi = \lambda \psi. \quad (12)$$

Here $\mathbf{s}$ is a unit vector, while ψ is again a 2×1 spinor. The unit vector $\mathbf{s}$ defines the reflection operator, whose reflection plane is orthogonal to $\mathbf{s}$. This reflection operator is proportional to the spin operator $\frac{\hbar}{2} [\mathbf{s} \cdot \boldsymbol{\sigma}]$. The reflection operator represented by $[\mathbf{s} \cdot \boldsymbol{\sigma}]$ is a group element S , and in principle, a spinor ψ also represents a group element g . But $Sg = g$ is not possible on the rotation group as S transforms a rotation into a reversal and vice versa. Although the eigenvalue equation Eq. 6 has an algebraic solution ψ , it is therefore not possible to interpret ψ as a spinor representing a group element. The problem is thus what the meaning of ψ could be. Let us therefore consider a set $\mathcal{S} = \{g_1, g_2\}$, whereby g_1 is a group element and $g_2 = Sg_1$. We have then $S(\{g_1, g_2\}) = \{g_1, g_2\}$, because S transforms g_1 into $g_2 = Sg_1$ and $g_2 = Sg_1$ into g_1 . The set $\mathcal{S} = \{g_1, g_2\}$ is then a kind of generalized eigenvector of the reflection operator S . This leads to the idea that $\psi = \psi_1 + \psi_2$, whereby ψ_1 represents g_1 and $\psi_2 = [\mathbf{s} \cdot \boldsymbol{\sigma}] \psi_1$ represents g_2 could represent the set $\mathcal{S} = \{g_1, g_2\}$. This is further confirmed by the explicit value for ψ when $\lambda = 1$:

⁶ In reality, one could give a meaning to this operation in the same spirit in the way one can give a meaning to the sum of two 3×3 matrices in $\mathbb{R}^3$ as discussed in Footnote 4. In fact, the permutations can be modeled geometrically as the symmetry group of a regular simplex in $\mathbb{R}^{n-1}$. A regular simplex $A_1A_2 \cdots A_n$ with edge length 1 in $\mathbb{R}^{n-1}$ is defined [6] as the set of n points A_k , $k \in [1, n] \cap \mathbb{N}$, with position vectors $OA_k = \mathbf{a}_k = (x_1, x_2, \dots, x_n) \in \mathbb{R}^{n-1}$, whereby: $x_{2r-1} = \frac{1}{\sqrt{n}} \cos \frac{2(k-1)r\pi}{n+1}$, $x_{2r} = \frac{1}{\sqrt{n}} \sin \frac{2(k-1)r\pi}{n}$, $\forall r \in [1, \lfloor \frac{n-1}{2} \rfloor] \cap \mathbb{N}$, if $n-1$ is odd, and $x_n = \frac{1}{\sqrt{n}} \frac{(-1)^{k-1}}{\sqrt{2}}$, if $n-1$ is even.

$$\psi = \begin{pmatrix} 1 + s_z \\ s_x + i s_y \end{pmatrix}, \quad (13)$$

which corresponds to the choice $g_1 = \mathbb{1}$.⁷ For a different choice of g_1 , we will obtain a different quantity $\zeta = e^{i\chi/2}\psi$. It must be noted that this idea to associate $\psi = \psi_1 + \psi_2$ with a set can only be introduced under the form of a new definition, because a linear combination of spinors is not a new spinor, as we already explained. We may note that in reference [5] a formal sum of group elements also corresponds to the notion of a set $\mathcal{S}$. E.g. when one proves $g \circ (g_1 + g_2 + \cdots g_n) = (g_1 + g_2 + \cdots g_n) \circ g$, whereby $g_j, j = 1, 2, \cdots n$ are all the members of $\mathcal{S}$, one only expresses that $g\mathcal{S} = \mathcal{S}g$. More generally, we could consider an ensemble of $2N$ objects O_j whereby N objects have the same orientation and parity as an object $g_1(O)$ and N objects the same orientation and parity as an object $g_2(O)$, whereby O is a reference object. When we call the numbers of objects $N_1 = N_2 = N$, the set can then be associated with $\psi = N_1\psi_1 + N_2\psi_2$. But if we normalize this quantity like a true spinor, we obtain then: $\psi = c_1\psi_1 + c_2\psi_2$, whereby $c_1 = c_2 = \frac{1}{\sqrt{2}}$ and $|c_1|^2 = |c_2|^2 = \frac{1}{2}$ represent the probabilities to find the objects in the orientations of $g_1(O)$ and $g_2(O)$. We see that this way we obtain a mixed state with exactly the interpretation given to it in QM, viz. that a mixed state:

$$\psi = \sum_{j \in \mathcal{B}} c_j \psi_j, \quad \text{where:} \quad \sum_{j \in \mathcal{B}} |c_j|^2 = 1, \quad (14)$$

where $\mathcal{B}$ is a set, and where the sums must be replaced by integrals if the set is not countable, represents a statistical ensemble whereby the particles can be in one of the states ψ_j . The probability for this to happen is $|c_j|^2$. We can consider this as an *a posteriori* justification of the superposition principle for the wave function solutions of a Dirac or Schrödinger equation.

We must note in anticipation that this rule, which is completely intuitive, will force us to consider the wave function ψ in situations where the probabilities are not summed according to the:

$$\text{incoherent rule:} \quad p = \sum_{j \in \mathcal{B}} |c_j|^2 |\psi_j|^2, \quad (15)$$

but according to the much less intuitive:

$$\text{coherent rule:} \quad p = \left| \sum_{j \in \mathcal{B}} c_j \psi_j \right|^2, \quad (16)$$

as not being obtained by the superposition principle. We will argue that the wave function is then not constructed according to a superposition principle but according to a Huyghens' principle.

The motivation for defining such a distinction is a concern about mathematical rigor and clarity. First of all, there can be only one probability rule for one construction principle. By introducing the distinction we will have two different words for two different concepts, where using a same single word for two different concepts can lead to confusion, especially if the difference between the concepts is subtle enough to escape attention. But there is a second, much more fundamental reason to introduce this distinction, viz. that Huyghens' principle stealthily introduces a third transgression of a definition domain, viz. the definition domain of the method we had to invent in order to make sense of the superposition principle in terms of sets. In fact, in certain cases one can obtain $\psi_1(\mathbf{r}) + \psi_2(\mathbf{r}) = 0$, and the union $S_1 \cup S_2$ of two non-empty sets S_1 and S_2 cannot be an empty set! The reason why this paradox occurs is that $\psi_1(\mathbf{r})$ and $\psi_2(\mathbf{r})$ do not have the right physically meaningful phases. We will see that the physically meaningful result is $\psi'_1(\mathbf{r}) + \psi'_2(\mathbf{r})$, while $\psi_1(\mathbf{r}) + \psi_2(\mathbf{r}) = e^{-i\alpha_1}\psi'_1(\mathbf{r}) + e^{-i\alpha_2}\psi'_2(\mathbf{r})$. The reason for this is that not all linear combinations of spinors are meaningful. Linear combinations $c_1\psi'_1(\mathbf{r}) + c_2\psi'_2(\mathbf{r})$, with $c_j \in \mathbb{C} \setminus \mathbb{R}^+$ have not been defined in the preceding lines. Only linear combinations $c_1\psi'_1(\mathbf{r}) + c_2\psi'_2(\mathbf{r})$ whereby $c_j = \frac{n_j}{\sqrt{n_1^2 + n_2^2}} \in \mathbb{R}^+$, and $n_j \in \mathbb{N}$ have been defined by the construction in terms of sets. Here $\mathbb{R}^+$ is the set of positive real numbers. We will refer to this third transgression in terms of a difference between a Huyghens' principle and a superposition principle. As we will see, we will be forced to consider the Huyghens' principle as a very useful numerical technique without physical meaning. This will play a major rôle in the double-slit experiment.

We have discovered simultaneously in this section how the necessary introduction of the mixed states replaces the description of a single electron in the Dirac-like equation by the description of a statistical ensemble of electrons in the Dirac equation. This should not be misunderstood. It does not describe a many-particle problem with wave function $\psi(\mathbf{r}_1, \mathbf{r}_2, \cdots \mathbf{r}_n, t)$ whereby the particles could interact. It describes statistical ensembles of single-electron histories in terms of probabilities. To define these probabilities we have to count frequencies within the statistical ensemble. It is in this sense that the wave function describes a statistical ensemble of electrons rather than a single electron, as is the case in the Dirac-like equation. This point will be worked out further in Subsection 7.1. This statistical issue remains hidden in the traditional presentation of the Dirac equation, which renders it very hard to imagine what could be going on behind the scenes of this equation. One almost has the impression that Dirac must have received the idea for the stroke of genius that lead to his equation in a phone call from God. This impenetrability in turn leaves the door open for speculations leading to the conviction that the Dirac equation could describe a single electron.

⁷ The case of the eigenvalue $\lambda = -1$ is analogous. One just must make a different choice $\psi_2 = -[\mathbf{s} \cdot \boldsymbol{\sigma}] \psi_1$ for ψ_2 .

We may note that the *a posteriori* justification of the superposition principle solves also the paradox of Schrödinger's cat. The cat is not in a superposition state where it would be half dead and half alive, because the wave equation does not describe a single cat. The wave function describes a statistical ensemble of cats, whereby half of the cats are dead and half of them alive.

3 The second transgression: From orbits to orbitals

We will argue below that the second transgression discussed above, boils down to a further extension of the sets considered. Instead of the history of particles on one set $\mathcal{P}$ we will consider histories of particles on many alternative sets $\mathcal{P}'$, such that we obtain something that resembles Feynman's all-histories formulation of QM. Instead of an all-histories approach, it will be rather a many-consistent-histories approach [8]. In fact, what we do for the case of the free-space Dirac equation in order to extrapolate the equation to $\mathbb{R}^4$ is to translate the straight line $\mathcal{P}$ over all vectors $\mathbf{w} \in \mathbb{R}^3$. This corresponds to substituting $\mathbf{r}_0$ by $\mathbf{r}_0 + \mathbf{w}$. The Dirac equation is then stipulated to be true on the sets $\mathcal{P}_{\mathbf{w}} = \mathbf{w} + \mathcal{P}$. Of course the equation becomes then valid over $\cup_{\mathbf{w} \in \mathbb{R}^3} \mathcal{P}_{\mathbf{w}} = \mathbb{R}^4$. In Eq. 4 we can then replace $\nabla_{\parallel}$ by ∇ because by construction $\nabla_{\perp} \psi_* = 0$. Here ψ_* is the extrapolated wave function. For the original wave function ψ on $\mathcal{P}$, the expression $\nabla_{\perp} \psi$ was not defined. We could have tried to define it, after introducing the *Ansatz* $\psi(\mathbf{r}, t) = 0, \forall (r, t) \notin \mathcal{P}$ but this would then have led to singularities, because the result would not have been a smooth function. For this reason, the extrapolation ψ_* cannot be considered as the superposition of the wave functions $\psi_{\mathbf{w}}$. We have $\psi_* = \sum_{\mathbf{w} \in \mathbb{R}^3} \psi_{\mathbf{w}}$ but the wave functions $\psi_{\mathbf{w}}$ are solutions of mutually different wave equations defined over different sets $\mathcal{P}_{\mathbf{w}}$, and these wave functions $\psi_{\mathbf{w}}$ are also not solutions of the wave equation for ψ_* defined over $\mathbb{R}^4$.

For the free-space Dirac equation, the meaning of the extrapolation is nevertheless extremely simple. The histories become possible alternative histories over $\mathbf{w} + \mathcal{P}$ and the equation describes potentialities. The spinor $\psi(\mathbf{r}, t)$ describes the motion the spinning electron would display, if it were in $(\mathbf{r}, t)$. It is thus normal that also this procedure corresponds to introducing probabilities. However, we show in reference [1] (on p. 204) that we introduce this way also unexpected solutions.⁸ The reason why the algebra permits to carry out the extrapolation of the definition domain unwittingly is that we only have to perform the time part of a Lorentz transformation, because the Rodrigues equation Eqs. 2-3 is purely temporal and does not depend on the position in space of the electron. This time part of the Lorentz transformation does not reflect the full information about the Lorentz transformation.

To obtain an equation that provides the same kind of description for an electron moving within a potential, we can introduce the minimal substitution, as explained in reference [1]. We can do this before or after the extrapolation procedure. When we do it before the extrapolation, the classical Dirac equation will be defined on a classical orbit, and the meaning of the extrapolation will not be as clear. The non-triviality of the procedure transpires here in several issues. The extrapolation replaces the classical orbit by a quantum mechanical orbital. The extrapolation can also entail self-consistence conditions that lead to quantization as explained in Chapter 8 of reference [1]. In the case of tunneling, the extrapolation to regions of space that are classically forbidden also gives rise to conceptual difficulties. The answer to the question what the extrapolation means will vary from case to case as we already pointed out in the Introduction. The reason for this is that different cases will involve different mechanisms.⁹ This is not surprising because the solution of a wave equation is largely dictated by the symmetry of the equation and because symmetry arguments do in general not permit to nail down the underlying mechanism. This is due to the circumstance that there can exist several wildly different mechanisms that all share the same symmetry, e.g. the Bloch wave eigenfunctions $e^{i\mathbf{k}\cdot\mathbf{r}}$ used for phonon propagation and for diffusion in a crystal are both defined by the translational symmetry of the crystal lattice, but the mechanisms involved in these two processes are entirely different (see reference [1], pp. 32-40, p. 46, pp. 337-340). The problems become less obvious when we introduce the minimal substitution after the extrapolation, but this can only mean that we are sweeping then the difficulties surreptitiously under the carpet, because the end result of both approaches is the same. It must also be pointed out that the minimal substitution is not rigorously exact, because it fails to account for Thomas precession as discussed in [2]. This point has escaped the attention of generations of physicists, due to a lack of understanding of the geometrical meaning of the algebra of group representation theory.

4 Conventions and description of the double-slit experiment

4.1 Classical and quantum mechanical, one slit or two slits

The paradox of the double-slit experiment has been described in a very detailed way by Feynman in his famous lectures [4]. Let us point out the main features, assuming we are performing a double-slit experiment with electrons. In Fig. 1 we show the

⁸ It can be shown that the free-space Dirac equation allows also for circular orbits, because it is only based on the temporal part of the Lorentz transformation.

⁹ In reference [1] we have tried to give a tentative explanation for the meaning of the extrapolation of the definition domain of the Dirac equation in the case of the calculation of the energy spectrum of the hydrogen atom. The argument runs over pages and pages (pp. 206-247). But understanding the double-slit experiment and tunneling proved to be an even much harder problem, which is why we managed to address these problems only recently.

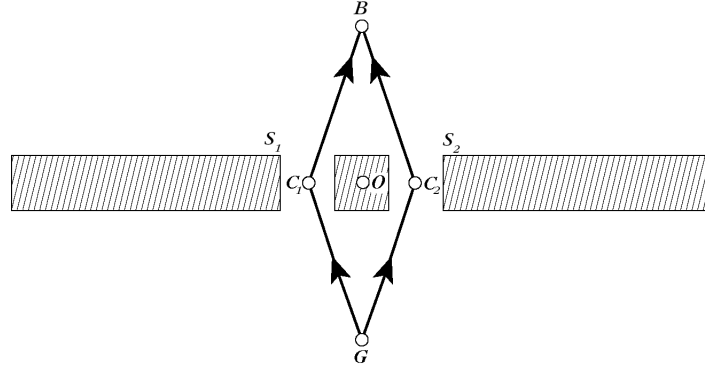


Fig. 1. Drawing of the double slit experiment, showing the notations used in the text.

set-up of this experiment and summarize the meaning of the various symbols we will use. We call the slits S_1 and S_2 , and their equal widths $w_1 = w_2 = w$. The distance between the centres C_1 and C_2 is called $D > w$. The centre of symmetry O is defined by $C_1O = OC_2$. We can define a reference frame $Oxyz$ whereby the x -axis coincides with the line C_1C_2 , while the y -axis is perpendicular to the plane that contains the slits. The z -axis is perpendicular to the figure and parallel to the extension of the slits outside the plane. The coordinates of C_1 are thus $(-D/2, 0, 0)$, those of C_2 are $(D/2, 0, 0)$. We imagine the source G of the electrons positioned at some point $(0, d_G, 0)$ on the y -axis, with $d_G \ll 0$. We can actually put symbolically $d_G = -\infty$. In Fig. 1 the scales used for the x -direction and for the y -direction are thus completely different. With $d_G = -\infty$ the wave we would use to describe the wave function before the slits could thus be considered as a plane wave propagating along the y -axis.¹⁰ There is a planar detector screen for the electrons far beyond the slits. The detector plane is parallel to the Oxz plane. The point B in the figure can be considered to be in the plane of this detector. When only one slit S_j , $j \in \{1, 2\}$ in the experiment is open, we will assume the source is positioned at $(-D/2, d_G, 0)$ or $(D/2, d_G, 0)$ rather than $(0, d_G, 0)$, where again $d_G = -\infty$. We will call the wavelength of the electrons λ . Let us first consider the case that only one slit S_j , $j \in \{1, 2\}$ is open. We can then consider two cases:

- (C1) When the energy of the electrons is high enough, such that $w \gg \lambda$, we will be in the classical regime and the form of the distribution $p^{(incoh)}(\mathbf{r})$ of the impacts of the electrons on the detector screen we will measure will be similar (on a smaller length scale) to what we would measure in a macroscopic experiment whereby we try to send tennis balls through a hole in a brick wall. This is the classical case where the probability distribution for the impacts of the electrons displays “particle behaviour” and just corresponds to the distribution for tennis balls without spin. The corresponding wave function will then be $\psi_j^{(incoh)}(\mathbf{r})$ and the probability distribution will be $p_j^{(incoh)}(\mathbf{r}) = |\psi_j^{(incoh)}(\mathbf{r})|^2$.
- (Q1) But if the energy of the electrons is low enough, such that $w \lesssim \lambda$, we will be in the quantum regime and the distribution will also show diffraction fringes. In this case the probability distribution for the impacts of the electrons displays “wave behaviour”. The corresponding wave function will then be $\psi_j^{(coh)}(\mathbf{r})$ and the probability distribution will be $p_j^{(coh)}(\mathbf{r}) = |\psi_j^{(coh)}(\mathbf{r})|^2$.

Let us now also consider the case when both slits are open. Also here we will consider two cases, because the double-slit experiment must be set up with $w < D \lesssim \lambda$:

- (C2) When the energy of the electrons is high enough, $D \gg \lambda$, we will be in a situation that is the analog of a classical-physics experiment with tennis balls, when there are two identical holes at the same height in the brick wall. We will obtain

¹⁰ In traditional QM this would be interpreted literally as a plane *matter wave* ψ . Here we consider the wave function as the description of a spinning motion, which has been defined on a single path and then extrapolated to the whole Oxy -plane. We could think that the original path takes the electron through slit S_1 and then has been extrapolated to $\mathbb{R}^2$, yielding a wave function ψ_1 , which is a pure state. One could argue that there is a second pure state ψ_2 obtained by symmetry $x \rightarrow -x$ from ψ_1 . The true wave function $\psi = \psi_1 + \psi_2$ is then a mixed state. However, within certain limits the phase $e^{-\frac{i}{\hbar}(Et - \mathbf{p}_1 \cdot \mathbf{r}_1)}$ of the plane wave allows for values $\mathbf{p}_2 \cdot \mathbf{r}_2 = \mathbf{p}_1 \cdot \mathbf{r}_1$, such that the path that takes the electron through S_2 is a history that is compatible with ψ_1 . We could thus consider a description of the physics that uses only ψ_1 . However, such a description of the two compatible histories is not symmetrical while the experimental set-up is. In fact, the path through S_1 has a compatibility with ψ_1 that is superior to the compatibility of the path through S_2 . Such a superiority is not present in the experimental set-up. We therefore risk to introduce some bias by using just ψ_1 . It is therefore better to take a wave ψ that propagates along the y -axis with a phase $e^{-\frac{i}{\hbar}(Et - \mathbf{p} \cdot \mathbf{r})}$ to describe the history before the slits and to consider the two actual electron paths (through S_1 or S_2) as compatible histories. We do then not consider ψ as a superposition state. Alternatively, one can introduce the mixed state $\psi = \psi_1 + \psi_2$ in order to eliminate the bias, but we will see that this is *not exact*.

a distribution of impacts that is the superposition $\frac{1}{2}p_1^{(incoh)} + \frac{1}{2}p_2^{(incoh)}$ of the individual probability distributions, where the normalization factor $\frac{1}{2}$ is introduced to take into account that there are now twice as many impacts, because there are two slits.¹¹ We could define a wave function that is a *mixed state* $\psi^{(incoh)} = c_1\psi_1^{(incoh)} + c_2\psi_2^{(incoh)}$ for this experiment, whereby $c_1 = c_2 = \frac{1}{\sqrt{2}}$. As explained in Section 2, QM tells us that the probability that the electron is in state $\psi_j^{(incoh)}$ is given by $|c_j|^2 = \frac{1}{2}$ and that the distribution of impacts is obtained by summing the probability amplitudes *incoherently* according to the rule $p = |c_1\psi_1^{(incoh)}|^2 + |c_2\psi_2^{(incoh)}|^2$. This is the classical case where the probability distributions for the electrons (and the tennis balls) display particle behaviour.

- (Q2) When the energy of the electrons is low enough, $D \lesssim \lambda$, the probability distribution of the electron impacts on the detector screen will show interference fringes, analogous to what we observe in a macroscopic experiment with waves, e.g. macroscopic waves in a water tank wherein the waves can propagate through two openings in a wall in the middle of the tank.¹² Most importantly, the probability distribution observed is different from the sum of the probability distributions obtained in the two single-slit experiments. Text books generally claim that the wave function is now $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$, whereby $c_1 = c_2 = \frac{1}{2}$ (or $c_1 = c_2 = \frac{1}{\sqrt{2}}$ after normalization). However, this is not rigorous and in general *not exact*. It can under certain circumstances be a good approximation, but the scope of validity of the *Ansatz* is never discussed. What is always correct however, is that the probabilities are now given by $|\psi^{(coh)}|^2$. This implies that when the approximation $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ is valid, then the probability amplitudes have to be summed *coherently* according to the rule $p^{(coh)} = |c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}|^2$, which implies that we no longer sum probabilities but probability amplitudes. Here one must also take $c_1 = c_2 = \frac{1}{\sqrt{2}}$ to obtain normalization.

4.2 Caveat

The fact that text books present $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ as an *exact rule* is disturbing, because it serves as a starting basis for introducing false concepts. We have anticipated this in Section 2 by introducing the distinction between the superposition principle and Huyghens' principle. In fact, $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ seems to embody the superposition principle, but if the result is not exact, then it is not obtained from the true superposition principle. It is also difficult to reconcile that we could explain the rules $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ and $p = |\psi^{(coh)}|^2$ by using the superposition principle and simultaneously explain the rules $\psi^{(incoh)} = c_1\psi_1^{(incoh)} + c_2\psi_2^{(incoh)}$ and the corresponding probabilities $p = |c_1\psi_1^{(incoh)}|^2 + |c_2\psi_2^{(incoh)}|^2$ by using the superposition principle. We cannot credibly justify a rule that stipulates that we should sum coherently on Mondays, Wednesdays and Fridays and incoherently on the other days of the week.

As pointed out in the Introduction, the superposition principle is not justified by the group theory, because spinors belong to a curved manifold rather than to a vector space. We have been able to rationalize the superposition principle and give a precise meaning to it in terms of sets (i.e. statistical ensembles) in the case of Pauli's definition of spin (see Section 2). This rationalization leads to the rule $p = |c_1\psi_1^{(incoh)}|^2 + |c_2\psi_2^{(incoh)}|^2$. The rules $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ and $p = |\psi^{(coh)}|^2$ must therefore be wrong in principle (in the sense that they are not truly based on the superposition principle). The probabilities $p = |\psi^{(coh)}|^2$ exhibit then perhaps some oscillatory behaviour, but different from the one that could be deduced in a simple way by an exact calculation from the superposition principle. A rigorous basis for explaining the oscillatory behaviour can thus not rely on the superposition principle.

In [1] we ran into this problem. We found a very different rationale to explain the oscillatory behaviour in the double-slit experiment but were unable to prove the rule $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ given by the "superposition principle". The point is that we could prove ([1], pp. 330-333) that $\psi = \sum_{j \in \mathbb{Z}} c_j \chi_j$, whereby χ_j is a wave function for which the phases built up over the paths GC_1B and GC_2B are different by an angle $2\pi j$ and the different pure states χ_j do not overlap, whereby we mean that $(\forall \mathbf{r} \in \mathbb{R}^3)(\chi_j(\mathbf{r})\chi_k(\mathbf{r}) = \chi_j^2(\mathbf{r})\delta_{jk})$.¹³ The calculus for the probabilities yields then $|\psi|^2 = \sum_j |c_j|^2 |\chi_j|^2$, whereby there is no difference between the coherent and the incoherent calculation of the probabilities. But the textbook rule $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$

¹¹ When incidentally $w \approx \lambda$ (such that we are in the case Q1C2) the superposition $\frac{1}{2}p_1^{(incoh)} + \frac{1}{2}p_2^{(incoh)}$ will exhibit some diffraction fringes.

¹² This is of course only true when $w \lesssim \lambda$. If incidentally $w \gg \lambda$ such that we are in the case C1Q2, the pattern will look more like a diffraction pattern.

¹³ There is an alternative logical option, which is to consider that there is only a single wave function whose definition domain can be subdivided in patches $\mathcal{S}_j$ over which the wave function is different from zero and the phases built up over the paths GC_1B and GC_2B are different by an angle $2\pi j$. In both options the phase difference cannot make a jump of 2π when we change one of the paths over an infinitesimal amount. Therefore, if there is only one wave function, then these patches must be separated by regions of space where the wave function vanishes. This single wave function can then be broken up in different wave functions χ_j that are non-zero over the interior of their definition domains $\mathcal{S}_j$. When there are different wave functions χ_j , the absence of overlap between the patches seems to be no longer a compelling requisite, as the probabilities must then anyway necessarily be calculated according to $|\psi|^2 = \sum_j |c_j|^2 |\chi_j|^2$. However, overlap can also be excluded here based on the fact that on the patches the electron is traveling in free space. From a classical viewpoint (which we adopt in the

cannot be justified by such a classical argument, because $\psi_1^{(coh)}$ and $\psi_2^{(coh)}$ do overlap. The coherent calculation would yield then interference fringes while the incoherent calculation would not. The two calculations are thus not equivalent when we apply them to the sum $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$, such that the incoherent calculation could then no longer be used to justify the coherent one.

We must and will justify the coherent rule in a completely different way. We will claim that it is just the rule $p = |\psi|^2$ applied to a single wave function. The rule $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ is not exact as $\psi^{(coh)}$ is not obtained by a superposition principle in the sense we have used in Section 2. In Section 2 we have used the superposition principle for solutions of a same wave equation. The rule $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ is all together different as it combines solutions from two different wave equations (corresponding to two different single-slit potentials) to propose a solution for a third wave equation (corresponding to a double-slit potential).¹⁴ The falsification of presenting the superposition principle as an absolute truth can badly mislead somebody who is interested in the foundations of physics underlying the double-slit experiment. It could side-track him in trying to derive $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ as an exact result, while it is not. He will then run in circles for ever because the proof he tries to find simply does not exist. It would be like trying finding a rigorous proof for a wrong theorem. As long as you stay rigorous and do not make an error, you will never find such a proof, because it is a chimera. It simply does not exist.

4.3 Summary

In summary, in certain double-slit experimental set-ups we must add probability amplitudes quadratically or “incoherently”. We first calculate the squares and then take the sum. The observations are then classical and correspond to “particle behaviour”. In other experimental set-ups, textbooks tell us that we must add the probability amplitudes linearly or “coherently”. We first take the sum and then square the result. The observations are then quantum mechanical and correspond to “wave behaviour”. However, the latter result is only qualitatively correct and not exact.

Feynman points out how classical and quantum mechanical behaviour can be observed e.g. for neutrons, which can undergo both coherent and incoherent scattering. That an electron can display both particle and wave behaviour, depending on the details of the experimental set-up is the famous “particle-wave duality”, which Feynman called the only mystery of QM. What is really mysterious is that when we are in case (Q1), then the distribution observed in case (Q2) is different from what one would obtain by summing the two-single slit distributions observed in case (Q1). Feynman highlighted this mystery by asking laconically how an electron that travels through slit S_1 to the detection screen can know that slit S_2 is open or otherwise.

5 Feynman’s analysis: Knowing or not knowing which way the electron has traveled

Feynman further explained that the crucial point that determines whether the probabilities display particle or wave behaviour in a configuration where both slits are open is if the experimental set-up permits us to find out through which one of the two slits the electron has travelled. Of course we can know this, when one of the slits is closed, but as we are discussing here the case where both slits are open, the criterion at stake here must be a different one than just knowing which slits are open. It is a criterion for quantum behaviour, that is more refined than the inequality $w < D \lesssim \lambda$ given above.

E.g. if we shine a light beam on the region of space just behind slit S_1 , the interaction of the electron with a photon of the light beam may tell us that the electron has traveled through slit S_1 . The probability amplitudes will then exhibit particle behaviour. In the case of neutron scattering we may also discover which trajectory a neutron has traveled when it has flipped its spin and simultaneously flipped the spin of a nucleon in one of the atomic nuclei of the experimental set-up in a spin-spin exchange interaction. The change of total spin of the nucleus is a mark left by the passage of the neutron revealing that the neutron has interacted with the given nucleus. The scattering is then incoherent. If there is no way to find out through which one of the two slits the electron has travelled, then probabilities will exhibit wave behaviour. The analogous situation in the neutron scattering experiment will be that the neutron has interacted with the nuclei in the set-up without inducing any spin flip. There is then not a single trace the neutron has left behind of its passage and the scattering is coherent.

An experiment whereby we put the detector screen close to the slits will also permit us to find out to a certain extent through which slit the electron will have travelled. In the region close to the slits we will thus not so much observe wave behaviour. This could be due to two reasons. The first one is that ψ_2 becomes very small close to slit S_1 and ψ_1 very small close to slit S_2 . A second reason could be that the rule $\psi^{(coh)} = c_1\psi_1^{(coh)} + c_2\psi_2^{(coh)}$ does not hold in this region, and becomes more accurate at a long distance behind the slits. The idea that the rule is not exact is in agreement with the results of the work of Sinha *et al.* [9].

When an electron does not leave any trace of its passage through the device, in analogy with the case of slow-neutron scattering without any spin-flip, then its interaction with the set-up is coherent. It is then impossible in principle to know which

whole paper) the result within a putative overlap between two patches should therefore be equal because identical causes must produce identical results. But in the overlap of different patches, the results would by definition be different by a phase difference that is a non-zero multiple of 2π . This way considering that there is only one or several wave functions boils down to the same thing. Note that the idea used in reference [1] remains correct in a formalism whereby the wave functions are real, $\cos(Et - \mathbf{p} \cdot \mathbf{r}/\hbar)$, instead of complex, $e^{-i(Et - \mathbf{p} \cdot \mathbf{r}/\hbar)}$, as needed in Section 12.

¹⁴ Note that we encountered this situation already in the description of the extrapolation procedure in Section 3.

way the electron has traveled. Even Hercule Poirot will not be able figuring it out. In the double-slit experiment, this means that not a single atom or electron in the set-up carries a mark of the passage of the electron. This can only be true if not a single atom or electron of the device has been scattered by its interaction with the electron, because the recoil produced by the scattering would create a phonon or an electronic transition within the measuring device and at least in principle this could be detected. The first excited electronic level might be too high to allow a transition or the interaction must be recoilless like in the Mössbauer effect, where it is the set-up as a whole which recoils instead of a single atom (Such a recoil of a macroscopic device is undetectable, because the device is too massive, and even if it could be detected it would not be able to tell us through which slit the electron has traveled). Or it should not induce a spin flip in a neutron scattering experiment. When the interaction of the electron with the measuring device does leave a mark in the device, e.g. in the form of a recoil, then the interaction process is incoherent. In a neutron scattering experiment, the incoherent interaction that left its mark on the device could be a spin flip rather than a recoil. However, when the interaction in the double-slit experiment is coherent, then still the probability distribution is not the superposition of the coherent single-slit probability distributions.

6 Thesis: The quantum weirdness does not reside in the properties of the particles

6.1 Two principles

This looks very mysterious indeed and seems to defy all daily-life intuition or common sense. We are used to describe this as weird quantum behaviour. The weirdness is often considered as a quantum paradox. We must point out that this paradox does not reside in the physics. The basic cause for it is not a weird property of the electron. To explain this we must develop two points, that we introduce here shortly.

(1) The first one will be amply discussed In Subsection 7.1 and is of a general nature. We will see that we can summarize it by stating that:

*We are not measuring electrons with the experimental set-up,
We are measuring the experimental set-up with electrons !*

Moreover, the complete information about the set-up we can obtain from the electrons (if we use enough of them) is non-local, because the geometry of the set-up is non-local!

(2) The second point is very specific for the double-slit experiment and must be combined with the general ideas evoked under point (1). For the double-slit experiment, we will argue in addition that the paradox is a probability paradox that comes about when we can no longer use binary logic when it comes down to answering simple yes-or-no questions, and we are forced to admit in all honesty for the possibility that the answer to a question can also be *undecidable*.¹⁵ We will here shortly introduce the concept of undecidability, before applying it to the double-slit experiment in Subsection 7.3. Subsection 7.2 will explain in detail the difference we want to define between a superposition principle and a Huyghens' principle.

6.2 Undecidability

Questions that are undecidable are well known in mathematics. Examples occur e.g. in Gödel's theorem or in the question if there exists a cardinal number between the cardinal $\aleph_0$ of the set of the integers and the cardinal $2^{\aleph_0}$ of the set of real numbers [10]. Paul Cohen [11] has shown in 1963 that this question is undecidable within the Zermelo-Fraenkel set theory with the axiom of choice included.¹⁶ The existence of such undecidable questions may look hilarious to common sense but this does not need to be. In fact, the reason for the existence of such undecidable questions is that the set of axioms of the theory is incomplete. We can complete then the theory by adding an axiom telling the answer to the question is "yes", or by adding an axiom telling the answer to the question is "no". The two alternatives permit to stay within a system based on binary logic ("*tertium non datur*")

¹⁵ Let us state right ahead for the reader who may feel that this is farfetched, that we will show that under certain conditions this expresses something he is much more familiar with, viz. Heisenberg's uncertainty principle. The undecidability does not apply to the electrons but to the set-up (see Section 12).

¹⁶ A whole school of thought in mathematics founded by Brouwer (the intuitionists) does not accept proofs based on a *reductio ab absurdo* due to the idea that a theorem could be undecidable. In a sense, the *reductio ab absurdo* can be considered as an untidy short-cut. The theorem you want to prove could be undecidable, which means that there does not exist a proof. At that stage, you need a supplementary axiom. You can use the affirmation of the theorem as the supplementary axiom or it's negation. But the negation leads to a contradiction. The only extension of the axioms that does not obviously contain a contradiction is then the one that has the axiom of the affirmation of the theorem added. The theory contains then a new truth, but this truth has the status of an axiom rather than the status of a theorem, because we could not prove it, in the sense that we proved that the theorem was undecidable. This way, we build up theories in the hope that we can avoid that contradictions will occur.

and lead to two different axiomatic systems and thus to two different theories. An example of this are Euclidean and hyperbolic geometry [12]. In Euclidean geometry one has added on the fifth parallels postulate to the first four postulates of Euclid, while in hyperbolic geometry one has added on an alternative postulate that is at variance with the parallels postulate. We are actually not forced to make a choice: We can decide to study a “pre-geometry”, wherein the question remains undecidable. The axiom one has to add can be considered as information that was lacking in the set of axioms. Without adding it one cannot address the yes-or-no question which reveals that the axiomatic system without the parallels postulate added is incomplete. As Kurt Gödel has shown, we will almost always run eventually into such a problem of incompleteness.

When the interactions are coherent in the double-slit experiment, the question through which one of the two slits the electron has traveled is very obviously also undecidable. Just like in mathematics, this is due to lack of information. There is no place for nitpicking remarks here like the ones I received from cynical referees who tried to force me to abandon using the same terminology in physics because it would be “*bad taste*” (*sic!*) and that “*I should not mix up categories*”. This is a hideous faultfinding remark. There are no different categories here. The idea of undecidability that occurs here is conceptually the very same idea as the one that underlies its formal definition in mathematics. We just do not have the information that could permit us telling which way the electron has gone.¹⁷ According to common-sense intuition this would not be too much of a problem in performing our probability calculus, as the undecidability is just experimental. We reckon that in reality, the electron must have gone through one of the slits and not through the other. We argue then that we can just assume that half of the electrons went one way, and the other half of the electrons the other way. But rigorous logic shows that the way we apply this reasoning to calculate the probabilities contains a fallacy, while by doing the logic correctly we just reproduce the quantum mechanical result!¹⁸ In order to show this we will introduce a number of assumptions that are all intuitive.

7 Explicit formulation of all assumptions underlying the philosophy of our approach

7.1 Assumption 1: Particles behave like particles, statistical ensembles of particles behave like waves

There is no particle-wave duality. Electrons are always particles, never waves. Electrons never travel like a wave through both slits simultaneously. That electrons are always particles is evidenced by the fact that a detector detects always a full electron at a time, never a fraction of an electron. It is the probability distributions of the electrons which display wave behaviour, not the electrons. In the preceding lines we have already tacitly anticipated the introduction of this postulate, by carefully phrasing our ideas to make them consistent with it. This postulate only reflects literally what QM says, viz. that the wave function is a probability amplitude. Measuring the probabilities requires measuring relative frequencies within a statistical ensemble of many electrons, such that the probability amplitude is a probability amplitude of an ensemble of electrons. Although this sharp dichotomy is very clearly present in the rules, we seem to find it difficult to perceive it. This is due to the fact that there has been a tendency to blur this very sharp image again by imposing a doctrine that wants to read more into this issue by adding additional interpretation. But there is absolutely no need for such additional interpretation: *In claro non interpretatur!* All one adds is over-interpretation

¹⁷ There exists a video entitled “*Probability & uncertainty - the quantum mechanical view of nature*”, with a lecture of Feynman on the double-slit experiment. In this lecture Feynman considers three possibilities for the history of an electron: “slit 1”, “slit 2”, and “do not know”. The third option corresponds exactly to this concept of undecidability. This is worked out with many examples in reference [4], to show that there is a one-to-one correspondence between undecidability and coherence.

¹⁸ Two remarks are due here:

(1) As evoked in Section 2, one can derive the Dirac equation in a completely classical way from the *Ansatz* that the electron spins. The equation is then defined on a classical path $\mathcal{P}_1$. We can call it the equation E_1 . However when we solve it, we search for a solution over the whole of $\mathbb{R}^4$. We extrapolate thus the definition domain of the differential equation from the path $\mathcal{P}_1$ to $\mathbb{R}^4$. Due to the difference of definition domains, we will call this extrapolated differential equation, the equation E . We might also derive an equation E_2 on a different path $\mathcal{P}_2$ and extrapolate it, and obtain the same equation E . The equation E describes then several possible histories which are mutually consistent in the sense that they can be described by the same equation E . Of course to be consistent two histories must satisfy certain conditions. We find here back the concept of consistent histories introduced by Griffiths [8]. For the equation E we can also justify that the probabilities can be obtained from the wave function ψ by using $|\psi|^2$. This is done in all textbooks by showing that one can derive a continuity equation for the probability charge-current from the Dirac equation. We can calculate this way the probabilities for the two single-slit experiments and the double-slit experiment and perfectly reproduce this way the double-slit experiment paradox within the theory. The mathematical solution of the paradox is then that slightly different boundary conditions for a differential equation can give rise to vastly different solutions (see below). The different boundary conditions (in combination with the value of the wavelength) are then the means by which we can express if a question is decidable or not decidable in the description of the experiment.

(2) It might provoke disbelief that we would have to abandon our binary logic based on the mutually exclusive concepts true and false, in order to be able to make sense of QM and the idea might meet resistance, even though we tried to iron out some of this resistance in the preceding lines. Choice of accurate terminology can be very important. Would the resistance of the reader have been the same if we had used the word *uncertainty* instead of *undecidability*? This raises the very interesting question about the exact relation between the two concepts, on which we will dwell below.

(“*Hineininterpretierung*”).¹⁹ The wave functions ψ do not describe single electrons but a probability distribution for electrons corresponding to a given set-up. Such a probability distribution is defined over the whole space accessible to the electron. It includes thus information about the whole experimental set-up. A single electron can in principle not see if the other slit is open or otherwise, because it cannot probe non-local information.²⁰ But the probability distributions and the wave functions contain this kind of global non-local information as they are defined simultaneously over whole space. This information is based on a simultaneous, global and therefore non-local description of all parts of the whole experimental set-up.²¹ The wave function is therefore defined in a non-local way while the electron cannot have non-local interactions.

That such a non-locality of the formulation is not in contradiction with relativity can be explained by pointing out how also Lorentz frames used to write the Lorentz transformations are non-local because they assume that all clocks in the frame are synchronized up to infinite distance (see e.g. p. 278 in reference [1]). It is by no means possible to achieve this, such that the very tool of a Lorentz frame conceptually violates the theory of relativity. In general, this is without consequences for the theory, because we never use the Lorentz transformations truly over the whole of $\mathbb{R}^4$ to describe some physical problem. We can always restrict the transformation to a patch of $\mathbb{R}^4$, but we never do this. We can justify this *a posteriori* by claiming that this was in order to avoid burdening the presentation. This way we catch up for the error that we did not realize that a Lorentz frame is conceptually not an exact tool. What this discussion shows is that the very concept of a wave function $\psi(\mathbf{r}, t)$ defined over whole $\mathbb{R}^3$ at an instant t is also non-local, and therefore also conceptually not an exact tool. It is not the physics here that is non-local. It is our tool we use to describe the physics which is. We find non-locality back in the description of the double-slit experiment in Bohm’s interpretation of QM [14]. Non-locality is thus *a priori* not a physical issue as one might be tempted to conclude from Bohm’s result. It would require more study to check if Bohm’s result implies some non-locality that goes beyond the kind of trivial non-locality we are describing here, but as Bohm’s approach is strictly equivalent to standard QM, and as our approach will also reach the same conclusions as standard QM, this should not be the case.

As we highlighted above, in the double-slit experiment we are not measuring electrons with the experimental set-up. We are measuring the set-up with electrons. These electrons will measure the global information about the set-up and it is therefore that the wave function must be global. The wave function must be able to account for every detail of the set-up because the

¹⁹ Perhaps this due to the historical introduction by de Broglie of the concept of matter waves, while the wave function corresponds only to the (extrapolation to $\mathbb{R}^4$ of the) description of the spinning motion of the electron over a path $\mathcal{P}$ as explained in Section 1 (see also Footnote 10).

²⁰ In reality, the electron could encounter different potentials in the two-slit and single-slit experiments. As we only describe the experiments based on symmetry arguments (whereby we do not want to refer only to symmetries of Euclidean geometry but also to physical conservation laws), we elude describing the detailed underlying mechanism. This mechanism is certainly not the same for photons and electrons. For electrons it might e.g. involve polarizing the material of the slits (see Footnote 25). The induced charge distribution may then be slightly affected by the presence of a second slit. Photons on the other hand may induce oscillations of dipoles, which then in turn could emit photons. This leads to the idea that the textbook treatment of the double-slit experiment is an idealized and simplified, abstract toy model that manages to capture the essence of some qualitative features of the physics that occur as a common denominator in both the photon and the electron experiments, but not the quantitative details of these experiments. What pleads for this is that electrons, photons and even other particles (like alpha particles, neutrons and protons) have different physics and obey different wave equations. As we will see, this essence is the undecidability built in into the set-up (see Subsection 6.2 and Subsection 7.3). A quantitative exact calculation of the interference patterns might require describing the whole case-dependent mechanism. Such situations occur in other instances in physics. E.g. Ohm’s law has a huge range of validity. In different parts of this range of validity there are certainly many different mechanisms responsible for it. The idealized description of the experimental set-up may even impede us to see the correct formalism. We discussed e.g. that for neutron scattering the reason why the question through which slit the neutron has gone is decidable or not decidable resides in the presence or absence of spin flip. The idealized description does not address this issue. It tells us if the question is undecidable or decidable for a wrong reason: the symmetry of the experimental set-up combined with some considerations ($D \lesssim \lambda$ and $w \lesssim \lambda$) about the wave length of the neutron. These considerations provide a good guess for the answer to the question if the scattering process will leave a mark on the system or otherwise, i.e. if it will be incoherent or coherent. They are not exact, but just a crude rule of thumb, and they have historically also always been presented as such (The Heisenberg uncertainty relations are also such a rule of thumb but they are covering actually only a special subset of the possible undecidable situations. They are a consequence of the Fourier transform (see Subsection 9.1)). But as the rule of thumb shortcuts the task of providing the mechanism, the result looks mysterious. The same aspects will also transpire in the description of the tunneling experiment for an electron, where the description in terms of a constant potential barrier does not allow for any description of the band structure, while this may be a fundamental ingredient to explain the process of the creation of an electron-hole pair in the solid [13]. Similarly, a potential barrier with tetrahedral symmetry might be used to describe the tunneling of a CH_4 molecule. This is absolutely clueless with respect to the possible mechanism, which for this particular case seems almost beyond imagination or guessing. Finally, without the wavelength criterion we would even not be able to distinguish between coherent and incoherent scattering when both slits are open in the double-slit experiment, because the symmetry does not provide a clue about the mechanism.

²¹ Indeed, the very description itself of the set-up we use to define the wave equation does not provide any clue as to which way an electron will have traveled. It leaves all possibilities open. This lack of information must then also show up in the wave function, because it cannot contain more information than the equation that defines it. If you do not tell in the description of the set-up which way the electron travels, why would the solution of the wave equation then tell you which way it would have gone? But the idealized description fails to describe the microscopic information that might enable to tell which way the particle has gone, e.g. a spin flip for the neutron. The macroscopic description is therefore not a real theory. It just is a very fortunate shortcut to the detailed microscopic argument that a correct theory might turn out.

electrons will explore every detail of the set-up provided you are sending out enough of them to make the full exploration. It is also therefore that you must extrapolate the wave function to whole space-time. That is the only way to be sure that it will reflect the full global information.

We can render these ideas clear by an analogy. Imagine a country that sends out spies to an enemy country. The electrons behave as this army of spies. The double-slit set-up is the enemy country. The physicist is the country that sends out the spies. Each spy is sent to a different part of the enemy's country, chosen by a random generator. They will all take photographs of the part of the enemy country they end up in. The spies may have an action radius of only a kilometer. Some of the photographs of different spies will overlap. These photographs correspond to the spots left by the electrons on your detector. If the army of spies you send out is large enough, then in the end the army will have made enough photographs to assemble a very detailed complete map of the country. That map corresponds to the interference pattern. We will argue later on in Subsection 9.1 that this interference pattern presents the information about the experimental set-up. It does not present this information directly but in an equivalent way, by a Fourier transform. The spies are not correlated, but the information about the country is correlated, it is the information you put on a map. You will see straight lines (roads). None of your spies has seen the global picture. None of them will have seen that long straight road that stretches out for thousands of miles (and is a kind of correlation). They may just have seen a kilometer of it. They have only seen the local picture. The global picture, the global information about the enemy country is non-local, and contains correlations, but it can nevertheless be obtained if you send out enough spies to explore the whole country, and it will show on the map assembled. That is what we are aiming at by invoking the non-locality of the Lorentz frame and the non-locality of the wave function. Each electron sees at the best one slit. Well, it even does not see a slit, it is just that other electrons, which you do not measure, are killed when they hit a part of the set-up. The wave function contains also information about electrons that you may not measure. And that is a point you will never be able to make sense of if you consider the wave function as information about the electrons you measure. This point could e.g. intervene in delayed-choice experiments [15] (see Footnote 31). But the global information gathered by many electrons contains the information how many slits are open. Because for each set-up that global information is the complete global information about your set-up contained in the wave function. You need many single electrons to collect that global information. You may (e.g. in the double-slit experiment) need also many electrons that may just not find their way to your detector. *The global wave function contains also information about the electrons you have not detected, because they are e.g. reflected by the set-up and end up in the reflected wave, or absorbed (which is an incoherent process), and this defines a part of the values of the transmission function C defined in Eq. 25 below.* And as far as those electrons that are being detected are concerned, a single one just gives you one impact on the detector screen. That is almost no information. Such an impact is a Dirac delta measure, which is the Fourier transform of a flat distribution. It contains hardly any information about the set-up because it does not provide any contrast.²²

The description of the experimental set-up that we use to calculate a wave function is conventionally highly idealized and simplified. Writing an equation that would make it possible to take into account all atoms of the macroscopic device in the experimental set-up is a hopeless task. Moreover, the total number of atoms in "identical" experimental set-ups is only approximately identical. What we present is always a figure like Fig. 1, and we could write down a Schrödinger or Dirac equation based on this figure, assuming e.g. that the macroscopic set-up presents an impenetrable, infinite potential barrier to the electron. But with such a simplified *Ansatz* we can never describe the atom that could bear the mark of the passage of the electron. Therefore, solving the equations with such an infinite potential exactly can only yield the coherent wave function. For large electron energies, this wave function will oscillate furiously and reach the limit of the incoherent wave function. We must then use *ad hoc* rules to obtain the result for the incoherent case or take refuge to a different description of the same set-up. The advantage of the infinite potential is however that it introduces very simple boundary conditions, which catch the essence of the probability paradox (see below in Footnote 24), and make it easy to carry out and understand the extrapolation of the wave function from $\mathcal{P}$ to $\mathbb{R}^4$. It is rather logical to assume that it is the probability distribution which "knows" if the other slit was open or otherwise. As said an electron cannot sense non-local information. Furthermore, a single electron impact on a detector plane does not tell us very much about the detailed probability distribution. Only by making the statistics of many impacts can we collect knowledge about the detailed probability distribution.²³

²² A very nice illustration of how the information gradually builds up is given in the work of Tonomura *et al.* [16]

²³ The critical reader may object that some photon correlation experiments refute the validity of our attempts to make sense of QM based on a completely classical philosophy because they violate some Bell-type inequalities. But the Bell-type inequalities are based on the assumption that there exists a unique *common hidden-variable distribution* that can be used in *all the various configurations of the experimental set-up*. These configurations are e.g. the different combinations of polarizer settings. The inequality is derived from a simpler inequality by integrating over this common distribution (see reference [18], p.385). The idea is that from $a > b$ it follows that $\int a\rho(Q)dQ > \int b\rho(Q)dQ$, when $\rho(Q) > 0, \forall Q$. However in [1] we show that under certain circumstances this assumption of a unique common distribution function ρ is not tenable. The quantities in the simpler inequality must then be integrated over different distributions $\rho_1, \rho_2, \dots$, in order to obtain the expressions for the measured quantities and the Bell-type inequality can then no longer be derived since from $a + b > 0$, it does not follow that $\int a\rho_1(Q)dQ + \int b\rho_2(Q)dQ > 0$. The fact that there does not always exist a common probability distribution is well known in mathematics as Gleason's theorem [17]. It is curious to notice that certain physicists [18] consider this as a confirmation of the claims made on the basis of Aspect's experiments [19], while in reality it refutes them because it shows that the Bell inequalities are wrong in the sense that they cannot be considered as an expression that would be universally valid for all local hidden-variables theories.

7.2 Assumption 2: There is a fundamental difference between the superposition principle and Huyghens' principle

The good question is thus not if the electron sees if the other slit is open or otherwise. That is a leading question. The electron never does. A good question is how the electron interacts with the measuring device, because even when both slits are open we can observe cases (Q2) and (C2), depending on the electron energy. Similarly, when only a single slit is open, we can observe cases (Q1) and (C1). What tells the cases (Q2) and (C2), or (Q1) and (C1) apart is the question whether the interaction was coherent or incoherent. When the interaction was coherent, it has become impossible in principle to know which way the electron has traveled. It is the full distribution that knows if both slits are open or otherwise, because it can be calculated from a wave function that is constructed taking into account whether both slits are open or otherwise, using a probability rule that will imply if the processes were coherent or incoherent. The mystery does thus not reside within the electron but in the set-up of the experiment which acts as a filter for wave functions. If we make the idealization that the material of the set-up presents an infinite potential to the electrons, then the wave function must vanish at the boundaries of the set-up. These conditions are boundary conditions for the wave function. By taking into account the boundary conditions in the Schrödinger equation, the information will be rigorously taken into account and make its way to the final solution, such that it represents exactly the solution that represents the information available.²⁴

The two single-slit experiments and the double-slit experiment result then in three different boundary conditions. As different boundary conditions can lead to wildly different solutions, the superposition principle is *a priori* not justified. Sometimes the linearity of the Schrödinger and Dirac equations is invoked to justify the superposition principle, but this can only be true for solutions of the equation that share the same boundary conditions. What we do in the double-slit experiment is adding the solution ψ_1 for the Schrödinger equation with a potential V_1 to the solution ψ_2 of the Schrödinger equation with a completely different potential V_2 and claim that the solution ψ_3 of the Schrödinger equation with yet another completely different potential V_3 is given by $\psi_3 = \psi_1 + \psi_2$. Such a procedure cannot at all be justified by the linearity of the Schrödinger equation with the potential V_3 . The procedure that tells us that we should take $\psi_3 = \psi_1 + \psi_2$ is thus in lack of proof. If we call S_j the set of points where $V_j(\mathbf{r}) = 0$, then $S_3 = S_1 \cup S_2$, and over $S_1 \cap S_2$ the wave equations are identical. However, the other points define supplementary boundary conditions and these are different for V_1 , V_2 and V_3 . Therefore the procedure $\psi_3 = \psi_1 + \psi_2$ is at least *in principle* wrong. One cannot always represent the effect of two causes as their sum.²⁵ But from now on we will assume that this rule could be a good approximation to the exact result. We will see that the justification for this is not a superposition principle but a Huyghens' principle. We think it is important to make a very clear distinction between these two principles. The superposition principle applies to two solutions of a same equation. Huyghens' principle applies within a single wave function of an equation. The idea is that the Huyghens' principle could become instrumental when we are able to meet certain boundary conditions for the wave function on a patch S_1 , without considering those on a different patch S_2 , and vice versa, while it may still not at all be obvious how to meet both boundary conditions simultaneously together on $S_1 \cup S_2$ when $S_1 \cap S_2 \neq \emptyset$.

The distinction between Huyghens' principle and the superposition principle is as follows. It is a difference between coherent and incoherent histories. When we extrapolate the wave equation to whole $\mathbb{R}^4$ we can only keep coherent (i.e. mutually consistent) histories in a single wave. If two histories cannot be combined consistently into a single wave-function we must collect them into different wave functions. We have then several wave functions to which we can apply the superposition principle. The meaning to be given to a linear combination has been explained in Section 2: The linear combination $\sum_{j=1}^n c_j \psi_j$ represents a set of spinors containing $N|c_j|^2$ spinors of the type ψ_j . Each spinor represents a possible state. These states are pure states and mutually orthogonal according to some criterion. Such a criterion could be e.g. that $\int_{\mathbb{R}^3} [\psi_j^*(\mathbf{r})\psi_k(\mathbf{r}) + \psi_k^*(\mathbf{r})\psi_j(\mathbf{r})] d\mathbf{r} = 0$, which

²⁴ In reference [1] we have derived the Dirac equation from scratch for a free electron. We can derive the Dirac equation for an electron in a potential from this by a substitution. The minimal substitution is not completely correct as it neglects e.g. the Thomas precession. But in the present context we assume that the equation with the minimal substitution is correct. The idea is then that the Dirac equation describes correctly the physical situation both in the single-slit and in the double-slit experiment. The paradox that intuitively the single-slit and the double-slit solutions seem to be incompatible is then just the paradox that different boundary conditions can lead to vastly different solutions. We encounter such a sensitivity to boundary conditions in the Dirichlet problem. The Dirichlet problem is the problem of finding a function which is the solution of a partial differential equation over the interior of a given region and satisfies the condition of taking prescribed values on the boundary of the region. The Schrödinger equations whereby we consider the potential barrier of the single-slit and double-slit experiments to be infinite correspond to such Dirichlet problems. As pointed out, these boundary conditions for the double-slit potential express in combination with the wavelength if the question through which slit the electron has traveled is decidable or otherwise.

²⁵ A good illustration of this could be the following. An electron approaching a slit S_j could attract holes and repulse electrons in the material surrounding the slit. The resulting probability distribution of the holes would have a behaviour which resembles $L_j C_j$, where L_j would be a Lorentzian centered at the position of the electron in the slit S_j and C_j a function that is zero over the slit S_j and one elsewhere. The true function could be different from a Lorentzian and what it really is does not matter. The idea is here only to describe the increasing or decreasing behaviour of the probability distribution of the holes. This probability distribution could have an effect on the wave function. Bluntly applying a sum rule for the probabilities or the probability amplitudes of the holes in the double-slit configuration would yield $L_1 C_1 + L_2 C_2$. This function would still contain a reminiscence of the two local maxima of L_1 and L_2 , which is certainly wrong. The result should only display the reminiscence of one local maximum of L_j , because an electron approaches only one slit S_j . The correct result should be $L_j C_1 C_2 \neq L_1 C_1 + L_2 C_2$.

leaves open the possibility that there could be values $\mathbf{r} \in \mathbb{R}^3$ for which $\psi_j(\mathbf{r})\psi_k(\mathbf{r}) \neq 0$.²⁶ After normalizing we can then state that this set of spinors represents a statistical ensemble where a fraction $|c_j|^2$ of the particles is in the state ψ_j . The probability to find the particle in this state is then $|c_j|^2$, which is what the quantum rule states. In [1] and Section 2 we show that we become forced to introduce this rule to make sense of Pauli's formalism of spin, and also to take the ultimate step in a rigorous, *deductive* derivation of the Dirac equation. This is evidence for the fact that the wave functions in the Pauli and Dirac equations are not pure spinors, but special linear combinations of them from the group ring. They are special because not all linear combinations are allowed, just one. Such linear combinations represent thus statistical ensembles as we just explained.

But in the case of the double-slit experiment the rule $\psi_3 = \psi_1 + \psi_2$ is then still not justified, because we are combining solutions from different equations. It would therefore be of interest to note such sums differently, e.g. under the form $\psi_3 = \psi_1 \boxplus \psi_2$, or $\boxplus_{j=1}^2 c_j \psi_j$, to indicate that ψ_1 and ψ_2 are not solutions of the equation with potential V_3 and that the result $\psi_3 = \psi_1 \boxplus \psi_2$ is not exact but only a good approximation. In anticipation of the forthcoming discussion, we could then say that $\psi = \psi_1 \boxplus \psi_2$ indicates that the sum is obtained from a Huyghens' principle while $\psi = c_1\psi_1 + c_2\psi_2$ is obtained from a true superposition principle. We are using Huyghens' principle to construct a single wave function that contains only coherent (i.e. consistent) histories. We may note that the mere fact that $\psi = \psi_1 \boxplus \psi_2$ is not exact, already points out that there is something wrong with our brute-force ideas about summing, not only for probabilities, but even for probability amplitudes.

7.3 Assumption 3: The double-slit experiment involves also undecidability (uncertainty)

7.3.1 Elaboration of the idea

The double-slit paradox is a probability paradox. It is a paradox about how we calculate probabilities for parameters whose values are undecidable. When we state that we cannot possibly know through which one of the two slits the electron has traveled, we must take into account this piece of information self-consistently. The answer to the question if the electron has gone through slit S_1 is then not “yes” or “no” but “undecidable”. We are not used to undecidable questions in daily-life experience and it looks tantalizing that such questions could exist. But as we pointed out, they do occur in mathematics. *Mutatis mutandis*, the same requirement of self-consistence applies when we do know through which one of the two slits the electron has traveled.

We can thus not assume at one stage of the argument that we cannot know for any history through which slit an electron has gone in that particular history because it has not left the slightest mark on the measuring device and act in another stage of the argument as though we do know through which slit the electron has gone. If we did that we would introduce a logical contradiction (we know and we do not know) such that our attitude would contain a blatant inconsistency. The major point is here perhaps that we are convinced that this would not matter or that we know how to cope with such contradictions. What we want to show is that an ensemble $\mathcal{H}^{(coh)}$ of undecidable histories $h_\mu^{(coh)}$ of passages through slits S_1 or S_2 cannot be obtained by making the conjunction $\mathcal{H}_1 \cup \mathcal{H}_2$ of two ensembles $\mathcal{H}_1$, and $\mathcal{H}_2$ of “decided histories” $h_{j,v}$ of passages through slit S_1 and passages through slit S_2 . Here the indices j refer to the slit S_j through which the electron has gone. The histories may here be artificially “decided” by a fake decision process that will be described below, whereby we act as though the situation is decidable while it is not and think that this would not matter because we can make up for it.

We have here deliberately written $\mathcal{H}_j$ without specifying if we meant $\mathcal{H}_j^{(incoh)}$ or $\mathcal{H}_j^{(coh)}$, because to discuss this issue we must consider two cases. In fact, the problem is complicated by the fact that a history whereby an electron emerges behind the single-slit or double-slit device can be “decided” for two reasons, viz. a *local* one (the interaction with the set-up is incoherent) and/or a *global, non-local* one (the other slit is closed). We want to explain that it is impossible to construct the probability distribution for the *doubly undecided* situation that both slits are open (global undecidability) and the question through which slit the electron has traveled is undecidable (local undecidability) from the probability distribution for the situation that only one slit is open (such that the path is always “decided” regardless of the local type of the interaction, coherent or incoherent).

(1) Suppose first that the interactions in the single-slit experiment of which we have taken the probability distribution are incoherent. It is then just not true that we could obtain the ensemble of undecidable histories $\mathcal{H}^{(coh)}$ in the double-slit experiment as the union of two ensembles of histories $\mathcal{H}_1^{(incoh)}$, and $\mathcal{H}_2^{(incoh)}$ in which the electron has interacted incoherently with the set-up: $\mathcal{H}^{(coh)} \neq \mathcal{H}^{(incoh)} = \mathcal{H}_1^{(incoh)} \cup \mathcal{H}_2^{(incoh)}$. In both sets $\mathcal{H}_j^{(incoh)}$ of histories where the interactions with the set-up are incoherent there is e.g. an atom recoiling (or undergoing an electronic transition) within the set-up, while in the whole set $\mathcal{H}^{(coh)}$ of histories where the interactions with the set-up are coherent, there is not a single atom recoiling (or undergoing an electronic transition). The “undecided histories” and “decided histories” are physically fundamentally different. They can be differentiated in principle by inspection of all the atoms of the set-up in order to check if the interaction has been coherent or incoherent. This is even intuitively completely clear. In terms of neutron scattering it just says that histories with a spin flip are different from histories without a spin flip. Ignoring this physical difference is an error we make when we take our macroscopic intuition as guidance to

²⁶ The circumstance that $\exists \mathbf{r} \in \mathbb{R}^3 \parallel \psi_j(\mathbf{r})\psi_k(\mathbf{r}) \neq 0$ does not lead to an interference term in the superposition principle, clearly indicates that the probability calculus according to the superposition principle is very different from the calculus we apply when we are faced with interference. As we will argue, interference is based on a Huyghens' principle, which stresses the importance of clearly distinguishing the superposition principle and Huyghens' principle.

think about the microscopic situation. The ensemble $\mathcal{H}^{(coh)}$ in the coherent case can thus not be constructed from the union of the ensembles $\mathcal{H}_j^{(incoh)}$ that occur in the incoherent case, even if we apply a statistical randomization procedure (to be described below) in order to simulate the undecidability. An “undecided” randomized ensemble of decidable histories is not the same as an ensemble of undecidable histories. We contradict ourselves when we assume that the histories are decided in one stage of the argument and that they are undecided in another stage of the argument.

(2) However, this reasoning falls apart when in the single-slit experiment of which we have taken the probability distribution the interactions are coherent. It is then only the global undecidability that makes the difference. It is in this instance that we might be very much convinced that certain differences will not matter. We must here reason in a different way. Footnote 24 already points out that what comes into play here could be just a paradox of the sensitivity to boundary conditions of the Dirichlet problem. The Dirichlet problem is difficult, because it requires satisfying the boundary conditions globally, not just locally. Therefore the three wave equations (based on the potentials V_1 , V_2 , and V_3) will have three different solutions. While this point is mathematically well taken when we consider it within the framework of a wave equation, it is much less so intuitively. We might try to invoke here that the wave function contains also information about electrons you are not detecting. In an experiment where slit S_1 is open, it contains also information about the electrons that have “died” (in terms of transmission). The decided histories occur in a context that contains information about the electrons you do not detect but have “died” on other parts of the set-up. This is exactly what the different boundary conditions imply. It may nevertheless look surprising and non-intuitive that this fact has significant consequences. The surprise resides in the fact that we can use arguments of statistical mechanics. Against all odds, we attribute an arbitrary label to each “undecided history” in $\mathcal{H}^{(coh)}$. This is a fake decision process (as mentioned above), because we have no objective criterion that entitles us to perform such an operation. Our choices can only be completely arbitrary and they contradict the fact that the histories are undecidable. The set $\mathcal{H}^{(coh)}$ cannot be split objectively into two subsets. Of course, we are then lying, because we act as though we would know through which slit the particle has gone, while we do not. But we are convinced that this does not matter. We argue that we do this only to take into account that the history must have taken the particle either through S_1 or through S_2 . There are just no other options than those two. We argue that we can make up for our inconsistency by considering all possible ways to perform this artificial labeling in order to calculate a statistical average. This procedure, we argue, will then take into account that in reality we did not know. This is the statistical randomization procedure we referred to above.

From the theoretical point of view whereby we are wary or contemptuous of intuition such that we prefer just relying on the calculations, we will not get away with such a cover-up story based on a statistical procedure because we act as though the particle that has gone through a slit S_j in the double-slit experiment would be described by an equation with the boundary values for a single-slit experiment. In other words, in the statistical procedure we ignore the warning we wanted to issue by introducing the notation $\psi_3 = \psi_1 \boxplus \psi_2$. Detailed comparisons of the calculations based on the exact solutions of the wave equations for the three potentials must and will show that the statistical averaging procedure fails. This comparison leads to the conclusion that the averaging procedure is flawed in ternary logic. One may find this astounding but it is a factual mathematical truth. Factual truth shows also that the averaging procedure does not reproduce the experimentally observed results.

From the theoretical point of view we could also point out that we are not used to logic that allows for undecidability. The idea that we can label the histories by the labels S_1 or S_2 occurs in a theory based on a system of axioms $\mathcal{A}_1$ (binary logic), while the undecided histories occur in a theory based on an all together different system of axioms $\mathcal{A}_2$ (ternary logic). The paradox results from the fact that we try to understand ternary logic from a viewpoint based on binary logic. One could compare this to trying to understand hyperbolic geometry from the viewpoint of Euclidean geometry. Just as hyperbolic and Euclidean geometry are incompatible axiomatic systems, ternary and binary logic are incompatible axiomatic systems. There is thus absolutely no ground for using binary intuition to tackle problems with ternary logic. We have not proved a theorem that would justify the algorithm based on the randomization procedure within the framework of the axiomatic system $\mathcal{A}_2$, but of course we could try our luck. We might hope that the randomization procedure would be correct in $\mathcal{A}_2$ because it is correct in $\mathcal{A}_1$. That we are hoping that the randomization procedure could work reflects that we are unwilling to believe that there could exist situations where $\mathcal{A}_2$ prevails in the real world, relying on our intuition. Or that we believe that we can ignore the fact that such situations could exist. To show that this intuition is wrong, nothing is better than giving a counterexample. The counterexample is the double-slit experiment where clearly the probability is not given by $|\psi_1|^2 + |\psi_2|^2$ but by $|\psi_3|^2 \approx |\psi_1 \boxplus \psi_2|^2$, where the index 3 really refers to the third (undecidable) option. It is then useless to insist any further. But we might be reluctant to accept this.

Whereas Kleene, Priest and Łukasiewicz have proposed truth tables for ternary logic (see reference [20]), it is not clear to us if anyone has ever derived rules for calculating probabilities according to ternary logic from first principles. We therefore do not know how the probabilities must be calculated within such logic in order to compare them with the quantum mechanical results. Quantum mechanics is all we have to calculate such probabilities in a very specific context (that relies e.g. on the existence of spin). This way, the problem becomes even more fundamental, because it evolves from a probability paradox to a paradox of ternary logic.

Despite all these correct logical arguments we have used to justify it, the final verdict that the statistical procedure is wrong in ternary logic remains certainly counter-intuitive. The idea of averaging is actually correct. It is the probabilities we average over that are changed. We must thus assume that the numerical value of the wave function is to a very good approximation equal to $\psi_1 \boxplus \psi_2$, but that ψ_1 and ψ_2 do not represent the amplitudes for the probabilities that the electron has traveled through the slits S_1 and S_2 in the double-slit experiment, as one might intuitively expect from a comparison with the single slit experiments.

Physical intuition is here a very bad guide. It is physical intuition that draws us here into the paradox. These amplitudes must be ψ'_1 and ψ'_2 , whereby $\psi'_1 + \psi'_2 = \psi_1 \boxplus \psi_2$ (because we know that $\psi_1 \boxplus \psi_2$ is numerically exact), whereby $\psi'_1 \neq \psi_1$ and $\psi'_2 \neq \psi_2$ (because $|\psi_1|^2 + |\psi_2|^2 \neq |\psi_1 \boxplus \psi_2|^2$ in the case of destructive interference $\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r}) = 0$ and there must exist $\psi'_1(\mathbf{r})$ and $\psi'_2(\mathbf{r})$, such that $|\psi'_1(\mathbf{r})|^2 + |\psi'_2(\mathbf{r})|^2 = |\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r})|^2$, because the electron must have traveled through one of the two slits). The Huyghens' principle must be physically flawed (even if numerically correct). There are two mathematical reasons to justify this: The first one is that summing spinors is just not mathematically defined. The second one is that Huyghens' principle itself is not generally valid as shown by the work of Hadamard and Kirchhoff. We will also discuss below (see Subsection 9.2) that Huyghens' principle is physically meaningless, even when it is numerically correct. This is already suggested by the phenomenon of destructive interference.

The best way to render this intelligible is pointing out that the probabilities of QM are *not absolute probabilities* $p(A)$ but *conditional probabilities* $p(A|B)$ for an event A to take place *provided* the event B has taken place. It is the choice of a set-up that serves here as the event B . The quantum mechanical probabilities are conditional probabilities for an event A *provided* the set-up is the one you are considering. This is very clear for the theoretical probability amplitude, as it is the wave function that you have to calculate by solving the wave equation for a given potential, i.e. a given set-up. Quantum mechanical probabilities are thus only valid within a unique context defined by a specific set-up: The wave equation is non-local and contextual. Changing the context from a double-slit experiment to a single-slit experiment implies even changing the axiomatics of the logic from $\mathcal{A}_2$ to $\mathcal{A}_1$. As long as we stay within a same context, we can tacitly assume $p(B) = 1$ and forget that the probabilities are conditional. But this dependence on the set-up (a tacit condition B) implies that we cannot transpose a probability from one set-up to another one. The consequence of this is that there does not exist a *common* probability distribution that could be used for the three experiments (with potentials V_1 , V_2 and V_3). In fact, we will see in Section 12 that the probability for a coherent local interaction at slit S_1 in a double-slit context is completely different from the probability for a coherent local interaction at slit S_1 in a single-slit context, despite the fact that the local contexts are strictly identical. The reason for this is that the probabilities must be defined with respect to a global context and that the global contexts are completely different. One global context is decidable while the other global context is undecidable.

It is here that the argument that we are not measuring electrons with the set-up but the set-up with electrons intervenes, because it permits to differentiate the probabilities that occur in the single-slit and the double-slit experiments based on their *global contexts*, while there is no way to differentiate these probabilities based on the *local interactions* (as very well summarized by Feynman's question: "How does the electron know if the other slit is closed or otherwise?"). The local interaction processes in the single-slit and double-slit experiments are in both cases the same and undecidable. This way of differentiating the probabilities coincides exactly to the way we treat the probabilities in QM and can thus be considered as a classical explanation for it. The set-ups are different and lead to different constraints on the definition of the probabilities, because the measured probability distribution must faithfully represent the complete information about the set-up (as illustrated by the *Gedankenexperiment* with the spies), which is not only locally undecidable (i.e. the interactions are coherent) but also globally undecidable (i.e. two slits are open). The situation is somewhat analogous to what happens in the paradox of Bertrand for geometrical probabilities, where the value of a probability depends on the procedure or protocol you use to define it. In QM, the probabilities are also geometrical and the set-up is the procedure or protocol you use to define them. Bertrand's paradox serves exactly to illustrate that probability calculus cannot always be intuitive. We might have considered Bertrand's paradox as purely academical nitpicking, but here it finds its way to a real-world physical application. The paradox is thus that it is reasonable to assume (allowing for the weak proviso expressed in Footnote 20) that all local interactions are rigorously the same but that we have to assign different probabilities to them in different contexts. It is thus a probability paradox. The paradox does not reside in the physics. The paradox is rooted in a profound difference between the ways we must calculate probabilities in binary and in ternary logic.

An accessory justification for the fact that we cannot consider the probabilities as independent from the context of the experimental device is that in the double-slit experiment with the potential V_3 there exist probabilities for an answer "do not know" to the question: "did the particle travel through slit S_j ?". Such probabilities do not exist in the two other experiments, where only one slit is open and the answers can only be "yes" or "no". Due to the absence of a common probability distribution for the three experiments, one cannot calculate the diffraction pattern of the double-slit experiment by averaging over the probability distributions of the single-slit experiments. This argument is further confirmed by what QM itself tells about probability distributions (see Subsubsection 7.3.2), as will be further discussed and elaborated in the very important Footnote 32, Section 12.

7.3.2 Corollary: Why Bell-type inequalities clash with QM

These considerations are also shedding new light on the Bell inequalities that are used to test QM like in the experiments of Aspect *et al.* They explain why we cannot assume that there would be a unique common hidden-variables distribution for all different set-ups that intervene in verifying a Bell inequality. It is the assumption used in the derivation of the Bell inequality that there would exist such a common probability distribution which renders this derivation wrong and makes it clash with QM (see Footnote 23). We have already used this argument of the absence of a common probability distribution in reference [1] to refute this derivation. We based it there on a completely different classical reasoning, such that we have now two different analyses leading to the same conclusion. In fact, in reference [1], we argued that there is no common distribution function that could describe all physical parameters needed to define the set-up. Here, we argue that there is no common probability distribution for the electrons that

interact with the various set-ups. A third argument (which formulates the second argument in a different language) is that certain operators in QM do not commute which entails that they do not have common eigenfunctions. The latter are just probability amplitude distributions. We can thus say that there is no common probability distribution for such non-commuting operators. Different contexts do not have common probability distributions. In the photon polarization experiments the number of possible different contexts is even infinite. We see thus that this argument, which comes from QM itself, corresponds exactly to our critique of the derivation of the CHSH inequality, which we derived from completely different ideas. There is thus absolutely no reason to pooh-pooh this critique. On the contrary, it permits to identify exactly where things go wrong. It permits even to see how we can construct inequalities that will be violated by experiments. When we can answer a question by yes or no within the probability distribution of an operator O_1 , that question will become sometimes undecidable within the probability distribution of an operator O_2 that does not commute with O_1 . In the CHSH inequality, there must exist separate probability distributions and eigenfunctions ψ_{jk} for the pairs of photons in the set-up used to measure the quantity $p(A_j \cap B_k)$, that are not common to those in the set-up used to measure the quantity $p(A_{j'} \cap B_{k'})$ with $(j' \neq j) \vee (k' \neq k)$. The inequality can thus not be obtained by integrating over a hypothetical *common* distribution function, because such a distribution function just does not exist. The experiments by Aspect *et al.* are thus very important and very beautiful because they are revealing something very deep and fundamental, even though they do not show that “Einstein was wrong”. Einstein thought to defeat the conclusions from $[A_j, A_k] \neq 0$ by considering two correlated photons, and indeed it is possible to measure then $p(A_j \cap B_k)$ as a surrogate for $p(A_j \cap A_k)$, but in the CHSH inequality the commutation relation resurfaces through $[A_j, A_{j'}] \neq 0 \vee [B_k, B_{k'}] \neq 0$ when we try to consider $p(A_j \cap B_k)$ and $p(A_{j'} \cap B_{k'})$ simultaneously. What the double-slit experiment shows is not that hidden variables do not exist but that their probability distributions in one set-up can become useless for the calculation of certain probabilities in an other set-up, because the latter are defined on an incompatible distribution. In a similar spirit, one cannot calculate Malus’ law $p(A_j \cap A_k)$ by classical logic from some $p(A_j)$, as any one who has tried may have found out in frustration. The reader may note that the *Gedankenexperiment* with the spies in Subsection 7.1 also settles the non-locality issue that seems to surround the experiments of Aspect *et al.* and has been dubbed entanglement.

7.3.3 Concluding remarks

We may feel that binary logic is different from ternary logic in that it is so obvious that it would not need any experimental evidence from physics to justify its rules. But the reason why we find it obvious might be that we have only been confronted in our human evolution with macroscopic physical evidence where binary logic prevails. Eventually, this paradox is not any more harsh than the paradox of Banach and Tarski in measure theory [21].

There is another way to present the problem of a contextual situation. We can think of checking that the history associated with ψ_1 and takes the particle through slit S_1 is truly decided by taking a path ψ^* backwards in time. If the context is truly decided, all backward paths must die on slit S_2 , such that $\psi_1\psi^* = \psi_1\psi_1^*$. If the context is truly undecided, the backward paths can thread through both S_1 and S_2 . Moreover, the phases must fulfill a compatibility condition. We have then $\psi_1\psi^* = \psi_1\psi_1^* + \psi_1\psi_2^*$. You can however, not consider this term in an isolated way in its own right, because the context is undecided. You must include the analogous reasoning for $\psi_2\psi^*$. The compatibility issue leads exactly to the wave functions we proposed in reference [1]. It also contacts truly nicely with Cramer’s handshake mechanism in his transactional interpretation of QM [22]. The waves that are traveling backwards in time are here only a mathematical expedient to check the global self-consistence of the wave function when it is non-local. They do not correspond to physical reality. They are just a way to deal with the non-locality of the wave function and the set-up. There is no true violation of causality in the wave function. Probabilities can only be approved as good probabilities for the non-local problem if all the handshakes have been carried out. This argument provides us with some intuition for the fact that in order to render a Huyghens’ principle for solving a Dirichlet problem exact, it may be necessary to take into account backward wave propagation (see Subsection 9.2).

8 Intermezzo: How to justify the quantum rules for the calculation of the probabilities of coherent and incoherent scattering

Whereas these remarks (and especially those in Footnote 32) solve the probability paradox conceptually, they do of course not explain the textbook rule that tells us how we must calculate probabilities in the coherent case. But that the paradox can be solved conceptually is already an important result. It must be clear from the analysis given above that we can have actually two types of sums $\psi_1^{(coh)} \boxplus \psi_2^{(coh)}$ and $\psi_1^{(incoh)} + \psi_2^{(incoh)}$. As explained, we can introduce the incoherent sum rule by mere definition and it corresponds to the prescriptions of QM when we apply the superposition principle. The coherent rule cannot be justified this way. It has to be done differently. We must eventually justify it by a different principle. In fact, as we already explained above, the rule $p = |\psi_1^{(coh)} \boxplus \psi_2^{(coh)}|^2$ is not exact. The exact rule is that for each set-up we must solve the corresponding wave equation in a mathematical rigorous way, and use the rule $p = |\psi|^2$. The justification of the textbook rule is then that $\psi \approx \psi_1^{(coh)} \boxplus \psi_2^{(coh)}$. We will justify the rule $p = |\psi|^2$ in Section 10, here we will try to explain why $\psi \approx \psi_1^{(coh)} \boxplus \psi_2^{(coh)}$ in the case (Q2) in the classification of Subsection 4.1.

Our attempt to derive the Born rule from the Born approximation given below may look a bit sloppy. One reason for this is that we cannot derive an equation for a point source (which leads to spherical-wave equation) from the equation for a spinning electron in its rest frame by global covariance. When a point source in $\mathbf{r}_0$ is emitting electrons, the electron wave function ought to be spherical, e.g. for a SU(2)-type wave function:

$$\psi(\mathbf{r}, t) = \frac{e^{-\frac{i}{\hbar}(Et - p|\mathbf{r} - \mathbf{r}_0|)}}{|\mathbf{r} - \mathbf{r}_0|} \begin{pmatrix} 1 \\ 0 \end{pmatrix}. \quad (17)$$

Here the factor $|\mathbf{r} - \mathbf{r}_0|^{-1}$ makes sure that $\psi^\dagger \psi \propto |\mathbf{r} - \mathbf{r}_0|^{-2}$. We assume from now on that $\mathbf{r}_0 = \mathbf{0}$ in order to simplify the notations. We have then $\mathbf{r} = r\mathbf{e}_r$ and $\mathbf{p} = p\mathbf{e}_r$. The vector $\mathbf{p}$ is thus not a constant and depends on $\mathbf{r}$. It is not possible to obtain such a wave function by simple covariance from the wave function for an electron at rest, because with a Lorentz transformation one can only obtain a constant value for $\mathbf{p}$. Let us now derive a Dirac-like equation for this wave function. We have:

$$-\frac{\hbar}{i} \frac{\partial}{\partial t} \psi = E\psi. \quad (18)$$

$$\frac{c\hbar}{i} \frac{\partial \psi}{\partial x} = \frac{c\hbar}{i} \left(\frac{1}{\hbar} p \frac{\partial r}{\partial x} \right) \psi - \frac{c\hbar}{i} \frac{1}{r^2} \frac{\partial r}{\partial x} e^{-\frac{i}{\hbar}(Et - p|\mathbf{r} - \mathbf{r}_0|)} \begin{pmatrix} 1 \\ 0 \end{pmatrix}. \quad (19)$$

Now $\frac{\partial r}{\partial x} = \frac{x}{r}$, such that we obtain:

$$\frac{c\hbar}{i} \frac{\partial \psi}{\partial x} = \frac{cp}{r} \psi - \frac{c\hbar}{i} \frac{x}{r^2} \psi. \quad (20)$$

This leads to:

$$\frac{c\hbar}{i} \nabla \psi = (c\mathbf{p} - \frac{c\hbar}{i} \frac{\mathbf{r}}{r^2}) \psi. \quad (21)$$

In reality we must write the vectors here under matrix form:

$$\frac{c\hbar}{i} [\nabla \cdot \boldsymbol{\sigma}] \psi = [(c\mathbf{p} - \frac{c\hbar}{i} \frac{\mathbf{r}}{r^2}) \cdot \boldsymbol{\sigma}] \psi. \quad (22)$$

The Dirac-like equation would thus be:

$$\left[-\frac{\hbar}{i} \frac{\partial}{\partial t} \mathbb{1} + \frac{c\hbar}{i} \left[\left(\nabla + \frac{\mathbf{r}}{r^2} \right) \cdot \boldsymbol{\sigma} \right] \right] \psi = m_0 c^2 \psi \quad (23)$$

This equivalent to introducing a vector potential $\mathbf{A}$ according to $-q\mathbf{A} = \frac{c\hbar}{i} \frac{\mathbf{r}}{r^2}$ or:

$$\mathbf{A} = i \frac{c\hbar}{q} \frac{\mathbf{r}}{r^2} = \frac{i}{\alpha} \frac{q\mathbf{r}}{r^2}. \quad (24)$$

This is purely formal, because the “vector potential” is imaginary. The imaginary vector potential must thus not obey a gauge condition like a real vector potential. By adding a real double-slit potential, we could this way obtain a better formalism to make the calculations that correspond to the double-slit experiment. We have not worked this out, because this would be a whole program in its own right. But this development shows why we have in general to abandon the wave equation approach for an S-matrix approach.

9 Justifying the coherent sum rule $\psi = \psi_1 \boxplus \psi_2$

9.1 By using the Fourier transform and/or the Born approximation

We want now to discuss why $\psi = \psi_1 \boxplus \psi_2$ will often be a good approximation of the wave function. Here we will discuss how we can invoke the Born approximation and/or the Fourier transform to justify this idea. In the next section we will refine this by using Huyghens’ principle from optics. In Subsection 9.3, we will improve this further by specifying the Huyghens’ principle for QM.

The general form of a wave in the experiment is $\psi(\mathbf{r}) = a(\mathbf{r})e^{i\mathbf{k} \cdot \mathbf{r}}$, whereby we note $(x, y) = \mathbf{r}$, $(k_x, k_y) = \mathbf{k}$. The reason why we write $a(\mathbf{r})$ instead of just 1 is as follows. We will consider plane waves. This is an approximation as the source is in reality a small surface. But when the source is very far away from the slit, it can be treated as though it were a point source. The waves are then spherical. However the spherical wave front that arrives at the slit will locally, i.e. on the length scale of the slit, just be indistinguishable from the tangent plane to the spherical wave front. E.g. for plane waves propagating along the y-direction we

would thus have $k_x = 0$, $k_y = |\mathbf{k}| = k$. Taking the idea of plane wave literally, we must consider the source as a line $y = -d_G$. We will first assume that the phase of the wave is zero when the particle is emitted by the source. A particle that arrives in a point $(x, 0)$ of the plane of the set-up must then have left the source in $(x, -d_G)$. The amplitude of a plane wave in the Oxy -plane is just $a(\mathbf{r}) = 1, \forall \mathbf{r} \in \mathbb{R}^2$. The wave is then $\psi(\mathbf{r}) = e^{ik(y+d_G)}$. The values in the plane of the slit are thus $\psi(x, 0) = e^{id_G k}$. However, the potential imposes boundary conditions on the wave function. The wave function must vanish in all points where $V(\mathbf{r}) = \infty$. This does not need to be true when the particles are reflected by the device, but we are interested only in the particles that are transmitted through the device. The points $(x, 0) \parallel V(x, 0) = \infty$ do not transmit. In all other points the wave just “sees” free space, such that the local solution of the wave equation in these points will thus just be the initial plane wave. The boundary conditions are not compatible with the plane-wave *Ansatz* because the *Ansatz* just implies $a(\mathbf{r}) = 1$. Combined with the assumption that a particle that arrives at $(x, 0)$ is emitted at $(x, -d_G)$ this forces us to assume that not all points of the source on the line $y = -d_G$ are emitting. The incoming wave is then not $\psi(\mathbf{r}) = e^{ik(y+d_G)}$, but $\psi(\mathbf{r}) = C(x) e^{ik(y+d_G)}$, where:

$$\begin{aligned} C(x) &= 1, \quad \forall x \in [-D/2 - w/2, -D/2 + w/2] \cup [D/2 - w/2, D/2 + w/2], \\ C(x) &= 0, \quad \forall x \in \mathbb{R} \setminus ([-D/2 - w/2, -D/2 + w/2] \cup [D/2 - w/2, D/2 + w/2]). \end{aligned} \quad (25)$$

We are thus modulating the wave function with $C(x)$. We have then $\psi(x, 0) = e^{id_G k} C(x)$. In all these equations $k = \frac{2\pi}{\lambda}$ contains the information about the energy of the incoming particle. To get rid of the term $e^{id_G k}$, we will assume from now on that the particles leave the source with the phase $e^{-id_G k}$, such that for $y > 0$ the phase of the wave function is just $e^{ik_y y}$. We will further modify this. We will consider that the incoming wave has not the form $e^{ik_y y}$ but:

$$\psi(x, y) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{k_x^2}{2\sigma^2}} e^{i(k_x x + k_y y)}. \quad (26)$$

This implies that the distribution of the momentum $p_x = \hbar k_x$ is not a Dirac measure $\delta(p_x)$ positioned at $p_x = 0$, but a Gaussian distribution:

$$G(k_x) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{k_x^2}{2\sigma^2}}, \quad (27)$$

centered around this value with a width $\hbar\sigma$, which is certainly more realistic from a physical viewpoint.²⁷ The transmission amplitude $T(k_x)$ will become:

$$T(k_x) = \int_{-\infty}^{\infty} C(x) \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{k_x^2}{2\sigma^2}} e^{ik_x x} dx. \quad (28)$$

This expresses that only in the points where the potential is zero the particle is transmitted. The factor $\frac{1}{\sigma} e^{-\frac{k_x^2}{2\sigma^2}}$ can be put in front of the integral. The Fourier transform of $C(x)$ is given by:

$$\mathcal{F}(C)(k_x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^{\infty} C(x) e^{i(k_x x)} dx = \frac{1}{\sqrt{2\pi}} \int_{-D/2-w/2}^{-D/2+w/2} e^{ik_x x} dx + \frac{1}{\sqrt{2\pi}} \int_{D/2-w/2}^{D/2+w/2} e^{ik_x x} dx. \quad (29)$$

Here each term:

$$\frac{1}{\sqrt{2\pi}} \int_{\varsigma D/2-w/2}^{\varsigma D/2+w/2} e^{ik_x x} dx = \frac{i w}{\sqrt{2\pi}} J_0\left(\frac{k_x w}{2}\right) e^{i \varsigma k_x D/2}, \quad (30)$$

with $\varsigma \in \{-1; 1\}$, represents the contribution ψ_j of a single slit S_j . When both slits S_1 and S_2 are open, the term:

$$e^{ik_x D/2} + e^{-ik_x D/2} = 2 \cos \frac{k_x D}{2}, \quad (31)$$

will occur as a pre-factor of the terms $i \frac{w}{\sqrt{2\pi}} J_0(\frac{k_x w}{2})$ and represent the interference. With the pre-factor $\frac{1}{\sigma} e^{-\frac{k_x^2}{2\sigma^2}}$ we recover completely the expression of the initial width distribution $G(k_x)$ at the end of the calculation. The term $i w J_0(\frac{k_x w}{2})$ will fulfill the rôle of a hull function for the interference term. The final amplitude is thus $i 2w \cos \frac{k_x D}{2} J_0(\frac{k_x w}{2}) G(k_x)$. This calculation yields thus exactly the rule $\psi = \psi_1 \boxplus \psi_2$.

²⁷ We could qualify this as a “wave packet” but one should not identify the electron with a physically meaningful wave packet. The electron is and remains a point particle. We conceive the wave packet rather as a heuristic mathematical tool. It is a more appropriate tentative solution of the wave equation than the Dirac measure in view of the global character of the boundary conditions. It corresponds more to an *S*-matrix approach.

The final amplitude $\mathcal{F}(C)$ is related to the Fourier transform $\mathcal{F}(V)$ of the potential in a 1-1 fashion. The amplitude is not exactly the Fourier transform of the potential V which takes values ∞ in some points $\mathbf{r}$ and 0 in some points $\mathbf{r}'$. It is the Fourier transform of a related function C which takes values 0 in the very same points $\mathbf{r}$ and 1 in the very same points $\mathbf{r}'$. We see thus that the probability amplitudes $C(k_x)$ are the Fourier transform of the “potential” C . This is at least in its spirit in agreement with the Born approximation:

$$\frac{d\sigma}{d\Omega} = \frac{m_0}{4\pi^2} |\mathcal{F}(V_s)(\mathbf{q})|^2. \quad (32)$$

for the differential cross section that describes the scattering of a particle with mass m_0 by a scattering potential V_s , where $\mathbf{q}$ is the momentum transferred in the scattering. The result we obtained can be related to the result obtained in the Born approximation. The relation between V and C is that we must renormalize the infinite values that V takes to 1. This yields V_1 . And then we must take $C = 1 - V_1$. It is more rigorous to formulate this the other way around by considering that the idealized double-slit potential V is given by:

$$\forall x \in \mathbb{R} : V(x) = \lim_{V_0 \rightarrow \infty} V_0 V_1(x). \quad (33)$$

This can be related to a Born-type scattering potential V_s as follows:

$$\forall x \in \mathbb{R} : V_s(x) = V_0(1 - V_1(x)) = V_0 C(x), \quad C(x) = \lim_{V_0 \rightarrow \infty} \frac{V_s}{V_0} = 1 - V_1(x). \quad (34)$$

We see then that Born describes the scattering by a potential V_s that has the form of two identical mountains in an empty landscape. The relation is illustrated between C and V_s is illustrated in figure 2. The Fourier transform of 1 in $1 - V_1(x)$ leads just to a Dirac measure, that we did not reproduce in our geometrical calculation of Eqs. 26-31. In Born’s problem, a lot of electrons will not be scattered at all and yield a Dirac measure $\delta(p_x)$ that accounts for the difference between the result we found for C from the geometrical calculation in Eqs. 26-31 and Born’s result for $V_s = V_0 C$. The scattering problem for the potential V_s (when $V_0 \rightarrow \infty$) is the “negative” of the transmission measurement for the potential V . The scattering problem is physical, while the transmission experiment is geometrical. We can use one to calculate the other because the total transmission of the sum of the potentials V_s and V will be zero when $V_0 \rightarrow \infty$.

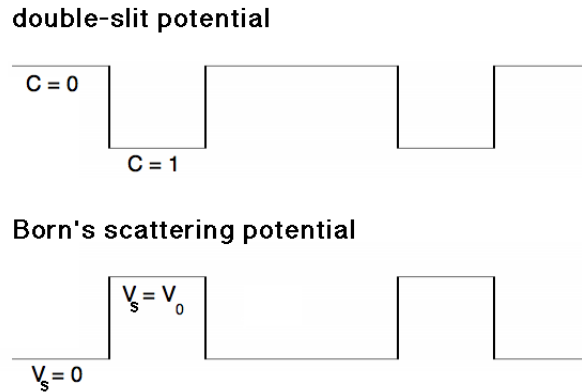


Fig. 2. Comparison of the double-slit potential V (with its values $V(x) \in \{0, \infty\}$ labeled in terms of $C(x) \in \{0, 1\}$) and Born’s scattering potential V_s that corresponds to it.

We can thus solve the Schrödinger equation for a double-slit experiment exactly without any use of the superposition principle, and calculate the scattering. In principle the Born approximation will yield a good approximation to this scattering. It is this Born approximation that will obey the pseudo “superposition principle” that we just must sum the two single-slit wave functions of the contributions of the two slits. The principle is not a real superposition principle but a Huyghens’ principle. For the more rigorous true solution, we cannot prove this sum rule. In fact, we have in our approach presented the experiment as a geometrical transmission experiment rather than as a physical scattering experiment, and without introducing the Gaussian width distribution, we could never have drawn in the Fourier transform into our calculation. The exact identity between the result of the geometrical calculation for C and the result of the physical calculation for the related potential V_s in the Born approximation can only underline that the result of the geometrical calculation is also an approximation (see also Footnote 25). This is in our opinion a way to justify that the sum rule yields a good, but not exact description of the wave function in a double-slit experiment.

The calculation also confirms the thesis announced in Section 6 that we are measuring the set-up with electrons, rather than measuring electrons with the set-up, because the wave function is the Fourier transform $\mathcal{F}(V)(\mathbf{q})$ of the potential $V(\mathbf{r})$.

The information contained in $\mathcal{F}(V)(\mathbf{q})$ is the same as the information contained in $V(\mathbf{r})$. This is especially obvious because the Fourier transform of $\mathcal{F}(V)(\mathbf{q})$ is again $V(\mathbf{r})$. The amplitude of the wave function contains thus nothing more and nothing less than the complete information about the experimental set-up. The electrons are providing so to say a photograph, not of the set-up, but of the Fourier transform of the set-up, which is an equivalent representation of the information. It is the thesis from Section 6 that renders the double-slit experiment intelligible. It provides a strong motivation for Bohr's interpretation that we must take into account the rôle of the measuring device in analyzing physical experiments. The result also provides evidence for the thesis that the Dirac equation is classical and that it is the way we solve it that renders it quantum mechanical.

Finally, it also can be used to clarify something that is really mysterious, viz. that to solve certain problems we must postulate that the wave function must be a (single-valued) function.²⁸ This postulate leads to a quantization condition in a straightforward manner. This is hard to justify if we think about QM as a set-up measuring electrons, because we are focussing then only on the local character of the interactions. It is much more easy to accept if we think about QM as electrons measuring a set-up, because we this protects us against glossing over the global character of the definition of the probabilities which are conditional. For a given set-up we must find a "probability amplitude function" (in the form of a Fourier transform) that describes it unambiguously and its is only reasonable that this should be a function, because the Fourier transform of a function is a function. The wave function must be global, and be assembled from its definition over local patches of its definition domain. We knit the patches together and in order to avoid any possible internal contradictions in the global assembly, we must make sure that calculations based on alternative paths through the definition domain lead to consistent results. The histories must be consistent. Therefore we use the Huyghens' principle and require the wave function to be a single-valued function. This way we can make sure that (e.g. in a Dirichlet problem) the solution proposed to satisfy the local boundary conditions in one place does not clash on a global level with the solution proposed to satisfy the local boundary conditions in another remote place. From this point of view, it is not at all obvious that the existence of a global solution would be granted, and there is then no surprise that this can lead to weeding out certain global inconsistencies through quantization.²⁹ The thesis advocated in Section 6, is thus a kind of paradigm shift as described by Kuhn [23], whereby all at once the drawing of a duck changes into the drawing of a rabbit.

The calculation also illustrates the idea how we one can generalize the Dirac equation from a path $\mathcal{P}$ to $\mathbb{R}^4$ in the presence of a potential, by just generalizing it first to $\mathbb{R}^4$ and then imposing the potential. However, the calculation we described here is based on the introduction of a packet of waves with central symmetry. On the local scale of the slits, we have replaced the plane wave which is the exact generalization from $\mathcal{P}$ to $\mathbb{R}^4$, by a radially propagating wave. This is thus a kind of cheat. (We have further modulated this radially propagating wave with a Gaussian, but fortunately this modulation does not intervene in the calculation above). Without this cheat, we would never have been able to make the Fourier analysis, because the term $k_x x$ would have been absent in the exponential (An S-matrix approach could justify all this). Furthermore, with a narrow incoming beam we will not reproduce the full span of the diffraction pattern behind the slits, which seems to bend around the corner. This is because we describe the experiment by pure physical transmission with a highly idealized potential. We can try to fiddle with the width of the Gaussian to obtain a reasonable result. But we can appreciate that this is all sketchy and intuitive rather than hard proof. We should actually not expect to be able to render the argument any more rigorous because the tentative expression $\psi_1 \boxplus \psi_2$ is not rigorous. Taking it literally would lead to insuperable conceptual difficulties (see Footnote ??). What is exact is the extrapolation of the wave equation from $\mathcal{P}$ to $\mathbb{R}^4$ leading to the boundary conditions expressed in Eq. 25, not the approximative solution for that extrapolated equation we tried to find here in terms of presumably real histories. We loose our grip on the detailed real histories in the approximation $\psi_1 \boxplus \psi_2$. We will see that the wave function can allow for a broad fan of incoming angles, if we do not consider these angles as physical and only occurring within a purely mathematical Huyghens' principle. First we will describe Kirchhoff's Huyghens' principle, then we will describe Feynman's improved Huyghens' principle in the form of his path integral method. With Kirchhoff's approach one cannot justify the generalization of the definition domain of the wave equation. It takes the generalized wave equation as its starting point, and the wave equation in question applies to photons rather than to electrons. It illustrates how the same general principles can show up in different situations, and lead to the same conclusions, while certain details can be actually different. But Feynman's approach can be used to justify the generalization of the definition domain as Feynman has shown that his Huyghens' principle leads to the Schrödinger equation. In both approaches we will see that it is crucial to have the electrons explore the full set-up, even if this must be considered as a purely mathematical description that does not need to correspond to the real physical situation. What counts in the physical situation is that we can state that the wave function contains the complete information about the set-up, which is expressed by (the non-physical) mathematical image that

²⁸ This intervenes in the solution of the wave equation for the spectrum of the hydrogen atom, as discussed in [1], p. 206, p. 281. It also intervenes in the discussion we made in reference [1] of the double-slit experiment (see pp. 327-335).

²⁹ We could object that we are only interested in the probabilities such that one could allow for inconsistencies in the phase. But the phase is also important because the spinning and orbital motion are coupled in QM. When a sub-atomic particle rotates faster, then its rest mass will be increased. This in turn will change its orbit. And due to Thomas precession effects, changing the orbit will change the rate of the spinning motion and the rest mass. Postulating that the wave function must be a function permits to treat this coupling by defining unambiguously the frequency of the spinning motion through the knowledge of the exact value of the phase. Electrons do have a phase, such that the phase is a quantity with a physical meaning. It is the rotation angle (modulo 4π) of the spinning electron. The phase is what makes a spinor different from a vector. Furthermore, if we dropped the phase from the wave function, the result $|\psi|$ would no longer constitute the full information about the set-up. That is, from $|\mathcal{F}[V(\mathbf{k})]|$ we can no longer recover $V(\mathbf{r})$ by the inverse Fourier transform, like we can from $\mathcal{F}[V(\mathbf{k})]$. This e.g. well-known in crystallography where it gives rise to a phase problem and the introduction of so-called Patterson functions.

the electrons visit all parts of the set-up even if this takes very counter-intuitive paths. These are the elaborations described in the following sections.

9.2 Photons in optics: By using the historical Huyghens' principle

If a wave equation allows for a Huyghens' principle we can argue that the solution $\psi = \psi_1 \boxplus \psi_2$ should be a very good approximation. For light waves we can e.g. use a principle due to St. Venant discussed in reference [24] (pp. 203-204). The idea is to put the wave function and its derivatives zero on the diffracting screen (exactly as we argued for the infinite potential). As Longhurst is pointing out, this leads to problems at the points that separate the regions where $V = \infty$ and $V = 0$ as in these points the continuity conditions are not satisfied. Let us however assume that we can neglect this problem. By using Venant's method we can derive then $\psi = \psi_1 \boxplus \psi_2$.

However, this derivation is not fully exact for another reason. In St. Venant's method the wave propagation must be strictly forward. But this is not true in the Huyghens principle for light rays. The derivation of Huyghens' principle for light rays by Kirchhoff is also discussed in reference [24]. In this Huyghens' principle there is an obliquity factor that reduces the weight of backwards traveling waves. Nevertheless, it does not completely rule out backwards traveling waves. St. Venant's principle will however only prove the rule $\psi = \psi_1 \boxplus \psi_2$ exactly if we exclude the possibility of backwards traveling waves all together. Else, we can consider for the potential V_3 (when the two slits are open) a path C_1PC_2Q where P is situated before the slits (such that on C_1P the wave is traveling backwards) and Q is situated behind the slits (such that on PC_2Q the wave is traveling forwards). Such a path does not exist for the potentials V_1 and V_2 , which clearly indicates that the solution $\psi = \psi_1 \boxplus \psi_2$ is not exact. Here the waves should not be considered as physically meaningful. They are just a mathematical tool that occurs in the Huyghens' principle for the wave equation. The Huyghens' principle itself is also just a mathematical tool because it contains features that cannot be justified physically, such as the obliquity factor, a weighting factor $\frac{1}{\lambda}$ and a phase difference of $\frac{\pi}{2}$ between primary waves and secondary waves. There is no dictionary that would enable to translate Kirchhoff's mathematics in a 1-1 fashion to meaningful physics for light rays. But we may present some wrong physical interpretation for it. This will exhibit counterintuitive aspects because it is wrong, but it could be helpful to present pseudo-physical pictures for the calculations that have great mnemonic value. This could be true for other Huyghens' principles as well. Feynman's path integral method can also be considered as a Huyghens' principle and must also be considered as a purely mathematical tool. It is in this respect that we can understand what Feynman reported e.g.: *"...there is also an amplitude for light to go faster (or slower) than the conventional speed of light. You found out in the last lecture that light doesn't only go in straight lines; now, you find out that it doesn't only go at the speed of light!"* [25], which is completely incomprehensible if taken literally.

9.3 The path integral method as the exact Huyghen's principle for quantum mechanics

We derive here a formula given by Dirac [26], that has been used by Feynman as the starting point in his development of the path integral method [27]. This equation is a form of the Huyghens' principle. In the rest frame of the electron $\psi(\rho, \tau) = e^{-\frac{i}{\hbar}m_0c^2\tau}$, such that:

$$\psi(\rho, \tau') = \psi(\rho, \tau) e^{-\frac{i}{\hbar}m_0c^2(\tau' - \tau)} \quad (35)$$

and

$$\psi(\rho', \tau') = \int_{\mathcal{B}} \psi(\rho, \tau) e^{-\frac{i}{\hbar}m_0c^2(\tau' - \tau)} \delta(\rho' - \rho) d\rho \quad (36)$$

for any set $\mathcal{B}$ that contains ρ' . Here ρ and ρ' are positions in the rest frame of the electron. With $\mathcal{B} = \mathbb{R}^3$ we can take the following form for the Dirac measure:

$$\delta(u) =_D \lim_{\alpha \downarrow 0} \frac{1}{2\sqrt{\pi\alpha}} e^{-u^2/4\alpha} \quad (37)$$

where the subscript D will serve to remind us that we are rather dealing with an equivalence in the sense of distributions. By putting:

$$u = \xi' - \xi; \quad \frac{i m_0}{2\hbar(\tau' - \tau)} = -\frac{1}{4\alpha} \quad (38)$$

where $\rho = (\xi, \eta, \zeta)$, we will obtain:

$$\delta(\xi' - \xi) =_D \lim_{\tau' - \tau \rightarrow 0} \sqrt{\frac{m_0}{2\pi\hbar i(\tau' - \tau)}} \times \exp \left[\frac{i m_0}{2\hbar} \left(\frac{\xi' - \xi}{\tau' - \tau} \right)^2 (\tau' - \tau) \right] \quad (39)$$

and in three dimensions:

$$\delta(\rho' - \rho) =_D \lim_{\tau' - \tau \rightarrow 0} \left(\frac{m_0}{2\pi\hbar i(\tau' - \tau)} \right)^{\frac{3}{2}} \exp \left[\frac{i m_0}{2\hbar} \left(\frac{\rho' - \rho}{\tau' - \tau} \right)^2 (\tau' - \tau) \right] \quad (40)$$

If we work non-relativistically, for a free particle the Lagrangian is given by $\mathcal{L}(\mathbf{r}, \mathbf{v}) = \frac{1}{2} m_0 v^2$. Hence $\frac{1}{2} m_0 \left[\frac{\rho' - \rho}{\tau' - \tau} \right]^2 \cdot (\tau' - \tau)$ can be written as $\mathcal{L} \left[\rho, \frac{\rho' - \rho}{\tau' - \tau} \right] \cdot (\tau' - \tau)$, such that

$$\delta(\rho' - \rho) e^{-\frac{i}{\hbar} m_0 c^2 (\tau' - \tau)} =_D \lim_{\tau' - \tau \rightarrow 0} \left(\frac{m_0}{2\pi\hbar i(\tau' - \tau)} \right)^{\frac{3}{2}} \exp \left[\frac{i}{\hbar} \mathcal{L} \left(\rho, \frac{\rho' - \rho}{\tau' - \tau} \right) \cdot (\tau' - \tau) \right] \quad (41)$$

which after substitution into Equation (36) yields:

$$\psi(\rho', \tau') = \lim_{\tau' - \tau \rightarrow 0} \int_{\mathbb{R}^3} \psi(\rho, \tau) \left(\frac{m_0}{2\pi\hbar i(\tau' - \tau)} \right)^{\frac{3}{2}} \exp \left[\frac{i}{\hbar} \mathcal{L} \left(\rho, \frac{\rho' - \rho}{\tau' - \tau} \right) \cdot (\tau' - \tau) \right] d\rho \quad (42)$$

Introducing an instantaneous Lorentz transformation that maps: $\tau' - \tau \mapsto \varepsilon = t' - t$, $(\tau, \tau') \mapsto (t, t')$, $(\rho, \rho') \mapsto (\mathbf{r}, \mathbf{r}')$, we obtain Dirac's formula:

$$\lim_{\varepsilon \rightarrow 0} \psi(\mathbf{r}', t + \varepsilon) = \lim_{\varepsilon \rightarrow 0} \left(\sqrt{\frac{m_0}{2\pi\hbar i \varepsilon}} \right)^3 \int_{\mathbb{R}^3} \psi(\mathbf{r}, t) \exp \left[\frac{i}{\hbar} \mathcal{L} \left(\mathbf{r}, \frac{\mathbf{r}' - \mathbf{r}}{\varepsilon} \right) \varepsilon \right] d\mathbf{r} \quad (43)$$

where we have used the Lorentz covariance of $\mathcal{L}$ which continues to express the proper time, such that we are allowed to extrapolate the formula to the motion of a particle in a potential. Introducing the potential this way changes the nature of the extrapolation from geometrical to physical. In fact, the Eq. 38 we started from is purely geometrical. Eq. 43 is in a sense an integral form of the equation $\frac{d}{dt} \psi = \frac{im_0 c^2}{\hbar} \psi$. The fact that the fully relativistic wave function should have four components as exemplified by the Dirac equation, shows why Eq. (43) can only have a limited domain of validity (In fact, there is no exact instantaneous Lorentz transformation of the type we postulated). Note also that our expression contains the right normalization constant and that in the classical Lagrangian the rest energy is ignored. Feynman had to introduce the normalization constant a *a posteriori* into Dirac's formula in order to obtain the correct expressions. Feynman has shown that Equation (43) leads directly to the Schrödinger equation [27]. The Schrödinger equation can also be derived from the Dirac equation. We should also not be too much surprised that in Feynman's method a particle seems to take all possible paths. As the information we obtain about the particle is rather time-like, there is very little information about the path the particle has taken: In free space, in the rest frame of the particle, there is none! In a hand-waving fashion we can argue that a way to reproduce total absence of information about the paths is to assume that all paths occur with the same weight. More rigorously it must be a consequence of a finding by Hadamard that for certain partial differential equations the solutions exhibit a Huyghens' principle [28].

9.4 Summary

The three approaches in Subsections 9.1-9.3 are not rigorous in all their details, but they all converge to the same general idea that one can generalize the Dirac equation in the presence of a double-slit potential from a single path $\mathcal{P}$ to $\mathbb{R}^4$. The wave function over this generalized definition domain is obtained by considering other alternative paths through this extended domain, which are consistent histories. These alternative paths cover the entire domain. The solution of the wave equation over this extended domain is a wave function which contains the complete information about the experimental set-up (e.g. in the form of the Fourier transform of the potential) in a one-to-one correspondence. It is this simultaneous extension of the equation and its solution by a Huyghens' principle that permits to understand the coherent sum rule and the so-called "particle-wave duality". The latter does not imply that the electron would sometimes behave like a particle and sometimes like a wave. An electron always behaves like a particle. What behaves like a wave is the wave function, i.e. the full information about the set-up collected by a statistical ensemble of electrons used to study the set-up.³⁰ That the information collected about the set-up is not presented directly but

³⁰ It may be noted in this respect that even in a classical water wave the individual water molecules do not move in wavelike manner. They only display some local circular motion (see e.g. Figure 6 from reference [29]). The wave behaviour is also here a property of an ensemble of a large number of molecules rather than of the individual molecules. The analogy stops here because the water molecules are in mutual interaction, while in a carefully designed double-slit experiment there may be only one particle present at the time. Moreover, electrons do move over large distances. They do have a phase due to their spinning motion and mathematically this does look like a wave on $\mathcal{P}$. However, the spinor ψ that describes the spinning motion is a temporal wave, not a wave that truly propagates in space (as indicated by its phase velocity $c^2/v > c$, see also reference [1], pp. 199-205). The extrapolation of ψ from $\mathcal{P}$ to $\mathbb{R}^4$ carries the local wavelike behaviour on $\mathcal{P}$ over to a global wavelike behaviour on $\mathbb{R}^4$. It is the fact that the wavelike appearance just expresses the spinning motion which explains that the probability waves can propagate in vacuum. The propagation in vacuum was historically considered as a riddle for electromagnetic waves.

e.g. under the form of a Fourier transform of the potential, is due to the fact that an electron has a phase associated with its spinning motion. It is this spinning motion that makes the representation of the information collected about the set-up look like a wave. It turns out that this wave is a probability amplitude (see Section 10). The procedure followed here to generalize the definition domain is case-specific. The generalizations must be discussed on a case-by-case basis. E.g. the generalization of the wave function for the energy spectrum of the hydrogen problem must be made in a completely different way. It is also lengthy and involved. It is based on the result announced in Footnote 8 and on Kepler's area theorem (see reference [1]). Tunneling requires yet another approach.

10 A note on the probability densities

The argument developed in this paper is not hundred percent rigorous, due to the difficulty of the problem. The small loopholes makes one doubt. One can then try to inspect the various parts of the argument to see if we cannot find a weak spot. This is a frustrating exercise in a complex problem whereby the various parts are interconnected because it makes you wander from one part to the other and feel like running in circles all over the place. Born's rule that the probabilities can be obtained by "squaring the wave function" plays a major rôle in the paradox of the double-slit experiment. This leaves one wondering if the origin of the paradox may not lie hidden somewhere within this rule. Is this rule really completely exact? Without a perfect understanding of the origin of this rule we cannot foster any hope of sorting out this problem. We will therefore try to justify this rule here as well as we can. We will see that it is all about counting particles. A first justification is that a spinor:

$$\psi = \begin{pmatrix} \xi_0 \\ \xi_1 \end{pmatrix}, \quad \psi^\dagger \psi = \xi_0^* \xi_0 + \xi_1^* \xi_1 = 1, \quad (44)$$

of SU(2) is normalized to one by definition, as expressed in the second part of Eq. 44. A triad of basis vectors is in one-to-one correspondence with the rotation that is necessary to obtain it from a reference triad. To describe the spinning motion of an electron we attach such a triad to it such that it rotates with it. Each electron has such a triad attached to it permitting to describe its spin, such that counting triads becomes equivalent to counting electrons. A quantity $p(\mathbf{r}) = \psi^\dagger(\mathbf{r})\psi(\mathbf{r}) = |c|^2$ becomes this way a measure for a number of electrons in $\mathbf{r}$. We must carry this idea with us all along the derivation of the Dirac equation starting from the Rodrigues equation.

Group theory shows that vectors (or four-vectors) are covariant quadratic expressions in terms of spinors. In the theory of relativity probabilities must be considered as components of a four-vector, viz. the probability charge-current four-vector. The quadratic quantity $\psi^\dagger \psi$ behaves like a probability charge-current density four-vector in the Lorentz group and satisfies a continuity equation. As a matter of fact, for every quantum mechanical equation in a textbook, the introduction of the wave equation is followed by a derivation of a continuity equation for the probability charge-current four-vector from it. These derivations of the expressions for the probability charge-current density four-vector $j_\mu \equiv (c\rho, \mathbf{j})$ from the free-space Dirac equation or the Schrödinger equation are well known (see e.g. reference [30]). As we have mentioned, the wave functions that occur in the Pauli and Dirac equations are sums of spinors, because the eigenfunctions ψ of a spin operator $\hat{S}$ are always sums of spinors. The reason for this is that the spin operator is a reflection operator and reflections do not have eigenfunctions on the group. They only have eigenfunctions on the group ring. We can consider the case of the Schrödinger equation as derived from the Dirac equation, such that also for this equation the wave function is such a sum. This special sum rule is not at all a result of the group theory, because the group theory does not even provide a meaning for such linear combinations. As far as the group theory is concerned, reflection operators have no eigenfunctions on the group and the result we obtain by carrying out the algebra mindlessly is meaningless, because in doing so we fail to acknowledge for the fact that the group is not a vector space but a manifold. The quadratic link between four-vectors and spinors can only be derived from the group theory for true pure spinors describing group elements, not for linear combinations of pure spinors. However the fact that we can derive a continuity equation from the wave equation shows then that such a quadratic link also exists for some special linear combinations of spinors, viz. those that are eigenvectors of the spin operator. The results obtained in Section 9 show that the approximate solution $\psi = \psi_1 \boxplus \psi_2$ of the wave equation for the double-slit experiment is also such a special linear combination of spinors. For these special sums that represent spin operator eigenfunctions we can justify the rule for coherent summing because they lead to the Pauli and Dirac equations, and for these equations the continuity equations derived from them show that that $\psi^* \psi$ and $\psi^\dagger \psi$ behave as probabilities.

In reference [1] (Eq. C12, p. 356), we show that for an SL(2, $\mathbb{C}$) matrix:

$$\begin{pmatrix} a & b \\ c & d \end{pmatrix}, \quad \text{with} \quad ad - bc = 1, \quad (45)$$

This conundrum disappears if electromagnetic waves are also probability amplitudes, whereby it is now the polarization of the photon which displays the oscillatory behaviour. We grasp our intuition about physical waves from the example of water waves, sound waves or waves propagating along a rope, where the wave transports energy and the energy transport is due to mutual interaction. But probability amplitudes are not physical waves in this intuitive sense. It is only their mathematical expression that makes the probability amplitudes look like waves. Because these probabilities are proportional to numbers of particles and particles represent energy, the probability amplitudes also "transport energy", but the transport mechanism is not based on mutual interaction, just on individual-particle motion.

$aa^* + bb^* + cc^* + dd^* = 2\gamma$, and that for the special sum of spinors used in the Dirac equation $\psi^\dagger\psi = aa^* + bb^* + cc^* + dd^* = 2\gamma$ (see reference [1], p. 166) such that $\psi^\dagger\psi$ accounts correctly for the electron probability density in the probability charge-current density four-vector because it is proportional to it. (This emphasizes that relativistic probability densities do not comply fully with our intuition about probability densities. They are no longer pure scalars but components of a four-vector. This somewhat analogous with the fact that we need oriented surface elements in $\mathbb{R}^3$, e.g. in order to take into account exposure to sunlight.)

When we try to justify this rule for photons by the same technique we run into some problems. Using the same methodology for the Klein-Gordon equation we find a continuity equation for the charge-current-like four-vector:

$$j_\mu = \iota(\psi^* \partial_\mu \psi - \psi \partial_\mu \psi^*), \quad (46)$$

such that:

$$c\rho = \iota(\psi^* \frac{\partial \psi}{\partial ct} - \psi \frac{\partial \psi^*}{\partial ct}). \quad (47)$$

As $\psi^* \frac{\partial \psi}{\partial ct} - \psi \frac{\partial \psi^*}{\partial ct} = [\psi^* \frac{\partial \psi}{\partial ct}] - [\psi^* \frac{\partial \psi}{\partial ct}]^*$, $c\rho \in \mathbb{R}$ is real. However, $c\rho$ can be negative, such that it cannot be defined as a probability density. The reason for this failure is the fact that the Klein-Gordon equation is of second order. To know its solution we must know both ψ and $\frac{\partial \psi}{\partial t}$ at some moment in time. The fact that $\frac{\partial \psi}{\partial t}$ helps in defining the solution makes then its way into the Eq. 46.

Another approach is thus necessary. The equation for electromagnetic waves propagating in free space is a special case of the Klein-Gordon equation for $m_0 = 0$. One might consider to resolve the problem of deriving an expression for the probability density for electromagnetic waves by taking the square root of the equation for electromagnetic waves in free space. This would lead to a Dirac-like equation whereby the rest mass m_0 of the particle is zero. From this equation we could then derive an expression for the probability density just like for the Dirac equation. But this runs contrary to the idea that photons must have scalar wave functions. It follows that the rule that $p = |\psi|^2$ must be derived in a completely different way for light waves in free space. The derivation can be found in the Feynman Lectures [31] (paragraphs 27.2 -27.4). As the energy of the electromagnetic field is a measure for the number of photons (which for our purposes we assume here to be all of the same energy), we can consider this energy to be proportional to the probability density. The Poynting vector is then a measure for the energy flow. We can use thus the Poynting vector to count photons. This is similar to the way the normalization $\psi^\dagger\psi = \xi_0^*\xi_0 + \xi_1^*\xi_1 = 1$ in SU(2) attaches the number 1 to a spinor (which is in turn attached to one electron) as described above.

We mentioned that the equation for electromagnetic waves propagating in free space is a special case of the Klein-Gordon equation whereby $m_0 = 0$. The continuity equation Eq. 46 is therefore an entirely correct result for light waves. However, we have no physical meaning for it. The waves are only mathematical tools to describe the probability distribution defined by a specific set-up. (In the Dirac equation, the wave has originally physical meaning because it describes the spinning motion of the electron over a physical path, but this meaning does not apply to the extrapolation of the wave function from this physical path to a definition domain that corresponds to the whole of $\mathbb{R}^4$. It is this extrapolated function we use when we solve the equation). The quantity j_μ corresponds thus to some mathematical flow that occurs in the wave function ψ but has no physical meaning in terms of physical quantities like energy or mass. It resembles in this sense the mathematical flow that occurs in Huygens' principle which also has no true physical meaning. We may note that the term $\nabla \cdot \mathbf{j}$ (where $\mathbf{j}$ is given by Eq. 46) that occurs in the continuity equation derived from the Klein-Gordon equation is obtained from the expression $\psi^* \Delta \psi - \psi \Delta \psi^*$ and that the structure of this expression exhibits a lot of similarity with the structure of the quantity $\psi_1 \Delta \psi_2 - \psi_2 \Delta \psi_1$ Kirchhoff started from in the derivation of his Huyghens' principle. Quantities of this type can thus be used to define mathematical flows that do not need to have a physical meaning, but that may intervene in checking the global consistency of the solution of a Dirichlet problem with its boundary conditions.

From the wave equations $\square \psi_1 = 0$ and $\square \psi_2 = 0$ one can also establish a continuity equation for quantities:

$$j_\mu = \iota(\psi_1^* \partial_\mu \psi_2 - \psi_2 \partial_\mu \psi_1^*), \quad (48)$$

For $\psi_1 = e^{i(Et - \mathbf{p} \cdot \mathbf{r})}$ and $\psi_2 = e^{-i(Et - \mathbf{p} \cdot \mathbf{r})}$, we have then:

$$c\rho = \iota(\psi_1^* \frac{\partial \psi_2}{\partial ct} - \psi_2 \frac{\partial \psi_1^*}{\partial ct}) = E(\psi_1^* \psi_2 + \psi_1 \psi_2^*). \quad (49)$$

The quantity $\psi_1^* \psi_2 + \psi_1 \psi_2^*$ intervenes in $|\psi_1 \boxplus \psi_2|^2$ in addition to the positive quantity $\psi_1^* \psi_1 + \psi_2 \psi_2^*$, and it can take both positive and negative values leading to constructive and destructive interference. The idea to combine $\psi_1 = e^{i(Et - \mathbf{p} \cdot \mathbf{r})}$ and $\psi_2 = e^{-i(Et - \mathbf{p} \cdot \mathbf{r})}$ in this argument can be found back in the idea of the paths C_1P and PC_2Q building up the path C_1PC_2Q described above, but it would be more general because it could apply for paths whereby P and Q are both behind the slits. This is certainly necessary if we want to explain interference. In that case the paths C_1B and BC_2 belong to the type of paths that could be considered. The idea of using paths C_1P and PC_2Q only serves to illustrate that $\psi = \psi_1 \boxplus \psi_2$ cannot be an exact rule. The term $\psi_1^* \psi_1 + \psi_2 \psi_2^*$ which only occurs when there is interference must thus be due to the backward action of Huyghens' principle, as the presence of complex conjugation shows. This does not imply backward action in physics, because Huyghens' principle is only a purely mathematical tool. But we can relate Huyghens' principle to the fact that we obtain the Dirac equation by extrapolation as pointed

out in Footnote 18. Due to this extrapolation, we obtain a Dirac equation for the double-slit experiment. In this equation a path p_1 through slit S_1 must be consistent with a path p_2 through slit S_2 . This consistency criterion can be expressed by calculating the phase difference over paths GC_1B and GC_2B , which must be a multiple of 2π . This is exactly the argument we developed in reference [1]. But we can also calculate the phase differences by considering GC_1BC_2G as a loop. Over this loop the phase must be a multiple of 2π . This approach introduces traveling backwards in time. We can also express the consistency criterion by using Huyghens' principle using waves that are traveling backwards in time. In fact, all paths between two points must lead to the same phase difference for global consistency of the wave function. This is thus the reason why we can consider all possible paths as Feynman did. All these approaches ensure that the wave function will be a function. One can question the condition that the wave function should be a function. We have given a first justification for it in Subsection 9.1. But we can also justify it is a heuristic method to insure that we treat correctly changes of mass due to e.g. Thomas precession. Precession changes the mass of a particle. Changing the mass will result in a change of orbit, and changing the orbit will change the precession. Therefore we must carefully monitor the precession and the way to do it is to postulate that the wave function must be a function. In certain cases, one even has to replace $\mathbb{R}^3$ by a Riemann manifold (containing multiple copies of $\mathbb{R}^3$) in order to recover a wave function that is really a function. Introducing harmonic polynomials allows then to recover a definition domain that corresponds to $\mathbb{R}^3$ [1].

The rôle played by Eq. 49 in the previous discussion shows that the mathematical flows are not always quantities devoid of any interest. Wave equations should therefore not be dismissed because they do not permit to derive easily an expression for a positive probability density, as was the case for the Klein-Gordon equation. In fact, squaring the free-space Dirac equation automatically yields a free-space Klein-Gordon equation, which has therefore to be correct. The mathematical quantity in Eq. 49 is thus also valid for the Dirac equation.

What all this also shows is that the extrapolation of the wave function from a path to the whole of $\mathbb{R}^4$ we described is all but innocent. As we stated in reference [1]: The Dirac equation is classical. It is the way we solve it that turns it into QM! The present analysis confirms this. We can justify the extrapolation procedure by stipulating that we want to obtain an equation for a probability distribution. This emphasizes that the double-slit paradox is a probability paradox. We can perhaps not claim that this analysis is watertight, but it gives a good conceptual grasp on the problem. It must be noted that the justifications we considered for the rule $p(\mathbf{r}) = |\psi(\mathbf{r})|^2$ are derived within the context of a given potential, i.e. a given set-up. Extreme vigilance is needed in providing the complete and exact formulation of the set-up, as it defines all constraints on the wave function. This caveat is not sufficiently clearly spelled out in the rather casual formulation of Born's rule. In a state of lowered vigilance one may overlook the need to specify all constraints. This will then lead to paradoxes, as a wrong calculation whereby constraints are overlooked and not all conditions on the wave function are fulfilled will yield a result that is different from the result one obtains by a correct calculation that accounts for all constraints.

11 Strong and weak orthogonality as a criterion for absence or presence of interference

We encountered a few cases where we had $\psi(\mathbf{r}) = \sum_j c_j \psi_j(\mathbf{r})$, whereby $\forall \mathbf{r} \in \mathbb{R}^3 : \psi_j(\mathbf{r}) \psi_k(\mathbf{r}) = \delta_{jk} [\psi_j(\mathbf{r})]^2$. This is a much stronger condition than the orthogonality condition we normally use for wave functions and which is $\int_{\mathbb{R}^3} \psi_j^*(\mathbf{r}) \psi_k(\mathbf{r}) d\mathbf{r} = \delta_{jk}$. We could call it therefore strong orthogonality. We have then $|\psi(\mathbf{r})|^2 = \sum_j |c_j|^2 |\psi_j(\mathbf{r})|^2$. This corresponds then to the rule we derived in special cases for the superposition of pure states. But we have seen that for the superposition principle the strong orthogonality does not need to be justified. Strong orthogonality is only needed to avoid interference when we apply Huyghens' principle. When the condition $\forall \mathbf{r} \in \mathbb{R}^3 : \psi_j(\mathbf{r}) \psi_k(\mathbf{r}) = \delta_{jk} [\psi_j(\mathbf{r})]^2$ is not satisfied, we run then into the phenomenon of interference. Strong orthogonality implies that the interference term $\psi_j^*(\mathbf{r}) \psi_k(\mathbf{r}) + \psi_j(\mathbf{r}) \psi_k^*(\mathbf{r})$ becomes zero over the whole of $\mathbb{R}^3$. From this we can see that the condition $(\forall \mathbf{r} \in \mathbb{R}^3) (\psi_j(\mathbf{r}) \psi_k(\mathbf{r}) = \delta_{jk} [\psi_j(\mathbf{r})]^2)$ as we formulated it is actually too strong, as a strong orthogonality criterion $(\forall \mathbf{r} \in \mathbb{R}^3) (\psi_j(\mathbf{r})^* \psi_k(\mathbf{r}) + \psi_j(\mathbf{r}) \psi_k^*(\mathbf{r}) = 2\delta_{jk} \psi_j(\mathbf{r})^* \psi_j(\mathbf{r}))$ would suffice. The latter corresponds to $(\forall \mathbf{r} \in \mathbb{R}^3) (\Re[\psi_j^*(\mathbf{r}) \psi_k(\mathbf{r})] = \delta_{jk} \psi_j^*(\mathbf{r}) \psi_j(\mathbf{r}))$. In the Dirac theory this would become $(\forall \mathbf{r} \in \mathbb{R}^3) (\psi_j^\dagger(\mathbf{r}) \psi_k(\mathbf{r}) + \psi_j(\mathbf{r}) \psi_k^\dagger(\mathbf{r}) = 2\delta_{jk} \psi_j^\dagger(\mathbf{r}) \psi_j(\mathbf{r}))$. This reminds of the orthogonality condition $\Psi_j(\mathbf{r}) \Psi_k(\mathbf{r}) + \Psi_j(\mathbf{r}) \Psi_k(\mathbf{r})$ for 4×4 representation matrices of four-vectors in the representation theory of the Lorentz group, especially as the Dirac matrices $\gamma_x, \gamma_y, \gamma_z$ are anti-hermitian, and we are considering here an orthogonality property over $\mathbb{R}^3$ rather than over $\mathbb{R}^4$. Strong orthogonality can in this case easily be obtained, e.g.

$$\begin{pmatrix} \psi_1 \\ 0 \\ 0 \\ 0 \end{pmatrix} \quad \text{and} \quad \begin{pmatrix} 0 \\ \psi_2 \\ 0 \\ 0 \end{pmatrix} \quad (50)$$

would be strongly orthogonal, whatever the values $\psi_1(\mathbf{r})$ and $\psi_2(\mathbf{r})$ the functions ψ_1 and ψ_2 may take. The functions ψ_1 and ψ_2 themselves do not need to be orthogonal in any sense whatsoever. Strong orthogonality would then easily be obtained in the spinor space and not require orthogonality in function space $F(\mathbb{R}^3, \mathbb{R})$.

This shows the importance of considering wave functions like $\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r})$ as not being a superposition of states. One circumstance that is pointing in this sense is that the result $\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r})$ was shown not to be exact. When we use the weak

orthogonality condition $\int_{\mathbb{R}^3} \psi_j(\mathbf{r})^* \psi_k(\mathbf{r}) d\mathbf{r} = \delta_{jk}$, the quantities ψ_1 and ψ_2 that occur in $\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r})$ may turn out to be fortuitously orthogonal. However, for true probability amplitudes ψ_j we should be able to select any patch $V \subset \mathbb{R}^3$ and the functions ψ_j should behave as probability densities over it. As on such subsets V (imposed on the double-slit experiment) the probabilities would correspond to a true physical problem, we could expect that the wave functions should be orthogonal over V . By considering all possible subsets V , this leads to the idea that the wave functions should satisfy the strong orthogonality condition $\forall \mathbf{r} \in \mathbb{R}^3 : \psi_j(\mathbf{r}) \psi_k(\mathbf{r}) = \delta_{jk} [\psi_j(\mathbf{r})]^2$ to avoid interference when we apply Huyghens' principle. This leads then to the idea that the isolated functions ψ_1 and ψ_2 in $\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r})$ are not probability densities for the double-slit experiment because they are not orthogonal (unless they are fortuitously orthogonal). This would justify that we are not allowed to calculate $|\psi(\mathbf{r})|^2$ according to $|\psi(\mathbf{r})|^2 = |c_1|^2 |\psi_1(\mathbf{r})|^2 + |c_2|^2 |\psi_2(\mathbf{r})|^2$, because $\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r})$ is not a superposition and explain the origin for the difference between coherent and incoherent summing of probabilities.

12 Also the undecidability does not apply to the single electrons but to the set-up

12.1 Logically binary and and experimentally ternary logic

We must now gather all our results. We have seen that the wave function contains the full information about the experimental set-up. The paths we may explore according to the Huyghens' principle are not the real paths. The principle does therefore not give access to the real paths, such that we cannot determine them by simulation. The most shocking result is undoubtedly that we do not reproduce the coherent sum rule in the double-slit experiment when we average statistically following binary logic. It would be worth investigating if we could reproduce the coherent sum rule by following ternary logic, because the fundamental paradox can be reproduced in the maths, such that you can scrutinize them to find the answers. The answers are wrapped up within the subtlety of the Dirichlet problem.

Does nature follow ternary logic? And why does nature switch back to binary logic when we put a light source behind the slits as suggested by Feynman? The answer must be that what your experimental results tell you is not that nature follows sometimes ternary logic and sometimes binary logic. The undecidability does not reside in nature itself. What the experimental results tell is if the set-up you designed follows binary logic or ternary logic. It is a verdict about the set-up not about the electrons. The electrons always go through one of the two slits, following binary logic. They have no other options. This is just logical. But the information collected with the experimental set-up may not follow this binary logic because it may hide us a part of the information. It is the logic of the design of the experimental set-up that determines how much you will find out about the true path of the electron. If the set-up is made in such a way that you cannot tell which way the electron went because the interactions have not left the slightest trace of the passage of the electron through the set-up, the logic of your set-up is ternary. And if the set-up is made in such a way that one can tell, then its logic is binary. You can switch between binary and ternary logic at will by modifying the set up. It is the fact that you can switch from one outcome to the other by changing the set-up that strongly suggests that the changes are due to a difference between binary and ternary logic of the set-up, in agreement with Feynman's observations (see Footnote 17). And it is the ease with which you can switch that proves that the electron itself follows binary logic.³¹

12.2 Is the Moon is out there when we are not watching?

If we had divine powers that would permit you to watch the electron without interacting with it, we would have seen and we would know through which slit the electron has traveled. And that would be true all the time. But the problem is that there do not exist set-ups that correspond to this ideal *Gedankenexperiment* of observing the electron without interaction. The experiments are not a realization of a divine *Gedankenexperiment*. They are done with real-world set-ups, where facts are created by interactions of matter with matter and not by pure thought.

The undecidability criterion corresponds to a global constraint that has a spectacular impact on the definition of the probabilities. The probabilities are *conditional* and *not absolute*. They are *physically* defined by the physical information gathered from the interactions with the set-up, *not absolutely* by some absolute divine knowledge about the path the electron has taken.

³¹ Adding elements to the set-up K_1 that swap its logic from binary to ternary or *vice versa* implies that we have modified it to another set-up K_2 . Consequently, we will also modify the wave function, which just describes the set-up, from χ_1 to χ_2 . We can e.g. assume that the wave functions χ_1 and χ_2 are Fourier transforms of potentials W_1 and W_2 . In fact, any change of a set-up K is a change of its potential W and thus of its wave function χ . Set-ups that follow different logics do not have a common distribution function (see Footnote 32). As they correspond to different logics, χ_2 has no bearing on any analysis one may carry out on the results obtained with K_1 , and χ_1 has no bearing on any analysis one may carry out on the results obtained with K_2 . It is especially useless to consider χ_2 as a temporal continuation of χ_1 because that would lead to the wrong over-interpretation that modifying K_1 to K_2 could alter χ_1 by an action that reaches out backwards in time. This is the paradox of the delayed-choice experiment or of the quantum eraser [32], which is due to the choice of a wrong angle of perspective on the problem, viz. by focusing on the idea that the set-up studies the electrons. The geometry of the set-up is a global and non-local property, which we explore with a large number of electrons to obtain its Fourier transform.

The set-up biases the information we can obtain about that divine knowledge by withholding a part of the information about it. Einstein is perfectly right that the Moon is still out there when we are not watching. But we cannot find out that the Moon is there if we do not register any of its interactions with its environment, even if it is there. If we do not register any information about the existence of the Moon, then the information contained in our experimental results must be biased in such a way that everything looks as though the Moon were not there. Therefore, in QM the undecidability must affect the definition of the probabilities and bias them, such that $p'_j \neq p_j$, for $j = 1, 2$. The experimental probabilities must reflect the undecidability. That the undecidability influences the definition of the probabilities is perhaps astounding but the mathematics already tell it does.^{32 33}

12.3 A topological argument

In a rigorous formulation, the undecidability becomes a consequence of the fact that the wave function must be a function, because it is the integral transform of the potential, which must represent all the information about the set-up and its built-in undecidability. As the phase of the wave function corresponds to the spin angle of the electron, even this angle is thus uniquely defined.

This idea is worked out in reference [1], pp. 329-333, and depends critically on the fact that the space traversed by the electrons that end up in the detector is not simply connected. It is based on the simplifying ansatz that the way the electron travels through the set-up from a point $\mathbf{r}_1$ to a point $\mathbf{r}_2$ has no incidence whatsoever on the phase difference of the wave function between $\mathbf{r}_1$ and $\mathbf{r}_2$. The idea is based on an Aharonov-Bohm type of argument: For two alternative paths Γ_1 and Γ_2 between $\mathbf{r}_1$ and $\mathbf{r}_2$, we have $[\int_{\Gamma_1} E dt - \mathbf{p} \cdot d\mathbf{r}] - [\int_{\Gamma_2} E dt - \mathbf{p} \cdot d\mathbf{r}] = 2\pi n$, where $n \in \mathbb{Z}$. The union of the two paths defines a loop. In a single-slit experiment this loop can be shrunk continuously to point which can be used to prove that $n = 0$. In a double-slit experiment the loop cannot be shrunk to a point when Γ_1 and Γ_2 are threading through different slits, such that $n \neq 0$ becomes then possible. Each interference fringe corresponds to one value of $n \in \mathbb{Z}$. A phase difference of $2\pi n$ occurs also in the textbook approach where one argues that to obtain constructive interference the difference in path lengths behind the slits must yield a phase difference $2\pi n$. But this resemblance does not run deep and is superficial. The textbook approach deals with phase differences between ψ_1 and ψ_2 in special points $\mathbf{r}_2$, while our approach deals with different phases built up over paths Γ_1 and Γ_2 within $\psi_3 = \psi_1 \boxplus \psi_2$ for all points $\mathbf{r}_2$.

12.4 The contributions of the slits S_1 and S_2 to the double-slit wave function

We can approach this somewhat differently. We will show that the textbook quantum mechanics prescription $\psi_3 = \psi_1 \boxplus \psi_2$ belongs to “pre-geometry” in the analogy we discussed above. We accept that the question through which slit the electron has traveled is undecidable and accept ternary logic. We should then play the game and not attempt in any instance to reason about the question which way the electron has traveled, because this information is not available. But we can also add a new axiom, the axiom of

³² We explained in Subsection 7.3 that this can be understood by pointing out that there does not exist a *common* probability distribution that could be used for the three experiments (with potentials V_1 , V_2 and V_3). In fact, the definition domain of the wave function for the double-slit experiment (i.e. the potential V_3) when the scattering is coherent does not contain subsets S_j on which $\psi(\mathbf{r}, t) \neq 0$ and on which one could decide that the electron has traveled through one of the slits S_j . These sets S_j are just empty and the whole domain of the coherent wave function where $\psi(\mathbf{r}, t) \neq 0$ is undecidable as can be seen from the tentative interpretation of the Fourier transform given immediately hereafter in the main text. In a transition regime between purely classical and purely quantum mechanical the wave function would be $\psi = c_1\psi_1^{(inc)} + c_2\psi_2^{(inc)} + c_3\psi_3^{(coh)}$, with $|c_1|^2 + |c_2|^2 + |c_3|^2 = 1$, where $\psi_1^{(inc)}$, $\psi_2^{(inc)}$, and $\psi_3^{(coh)} \approx \psi_1^{(coh)} \boxplus \psi_2^{(coh)}$ are three different wave functions. They correspond to the answers “yes”, “no”, “do not know” in ternary logic to the question: “did the electron travel through slit S_1 ?”. The wave function ψ is then a true superposition and we must sum incoherently to obtain the probability: $p = |c_1|^2|\psi_1^{(inc)}|^2 + |c_2|^2|\psi_2^{(inc)}|^2 + |c_3|^2|\psi_3^{(coh)}|^2$. In fact, histories with coherent scattering and with incoherent scattering are incompatible, such that they must be relegated to different wave functions. Here $\psi_3^{(coh)}$ corresponds to coherent scattering, while $\psi_1^{(inc)}$ and $\psi_2^{(inc)}$ correspond to incoherent scattering. In the purely quantum mechanical regime we must put $c_1 = c_2 = 0$ and $c_3 = 1$. In the purely classical regime we must put $c_1 = c_2 = \frac{1}{\sqrt{2}}$, $c_3 = 0$. In the intermediate case we can see very clearly that the coherent probability $|c_3\psi_3^{(coh)}|^2 \approx |c_3(\psi_1^{(coh)} \boxplus \psi_2^{(coh)})|^2$ should not be associated with the incoherent probabilities $|c_1|^2|\psi_1^{(inc)}|^2$ and $|c_2|^2|\psi_2^{(inc)}|^2$. Binary coherent probabilities $|c_1|^2|\psi_1^{(coh)}|^2$ and $|c_2|^2|\psi_2^{(coh)}|^2$ from single-slit set-ups are just not present in the expression $p = |c_1|^2|\psi_1^{(inc)}|^2 + |c_2|^2|\psi_2^{(inc)}|^2 + |c_3|^2|\psi_3^{(coh)}|^2$ for the probabilities in the double-slit set-up. It is only that numerically $\psi_3^{(coh)} \approx \psi_1^{(coh)} \boxplus \psi_2^{(coh)}$ accidentally happens to be a good approximation.

³³ The solution in Footnote 32 in terms of an absence of a *common* probability distribution is very similar to our critique of the derivation of the CHSH inequality in Footnote 23, Subsubsection 7.3.2 and reference [1], p. 278, which shows that this Bell-type inequality cannot be used to analyze the experiments of Aspect *et al.* [19] as has been done to draw the conclusion that QM cannot be a local-variable theory (See also reference [33]). In this critique we also argued that there is no proof that there exists a *common* probability distribution for the various measurements that must be carried out to test the CHSH inequality. In absence of such a proof, it is impossible to draw any conclusion from the experiments. For the double-slit experiment, one could adopt a hard line and argue that the number j of the slit S_j through which the particle has traveled is a hidden variable and that it does not exist. But this mixes up the concepts of determinism and decidability as explained further in the main text, and illustrates how a dogma condemning hidden-variable theories could really thwart any further progress in understanding physics in a very detrimental way. We may note that the CHSH inequality is also based on a logic of yes or no answers.

the divine knowledge, rendering the question decidable for a divine observer. We must then also play the game and accept the fact that the probabilities we will discuss can no longer be measured, such that the conclusions we draw will now no longer be compelled by experimental evidence but by the pure binary logic imposed by the addition of the axiom.

We take the exact solution ψ'_3 of the double-slit experiment and try to determine the parts ψ'_1 and ψ'_2 of it that stem from slits S_1 and S_2 . We can mentally imagine such a partition without making a logical error because each electron must go through one of the slits, even if we will never know which one. We have thus $\psi'_3 = \psi'_1 + \psi'_2$. We expect that $\psi'_1(\mathbf{r})$ must vanish on slit S_2 and $\psi'_2(\mathbf{r})$ on slit S_1 , but we must refrain here from jumping to conclusions by deciding that $\psi'_1 = \psi_1$ and $\psi'_2 = \psi_2$. We can only attribute probabilities $|\psi'_1|^2$ and $|\psi'_2|^2$ to the slits, based on the lack of experimental knowledge. We will use $\psi'_3 = \psi_3 = \psi_1 \boxplus \psi_2$ in our calculations, as we know it is a good numerical approximation. Let us call the part of $\mathbb{R}^3$ behind the slits V . Following the idea that $\psi'_1(\mathbf{r})$ would have to vanish on slit S_2 and $\psi'_2(\mathbf{r})$ on slit S_1 , we subdivide V in a region Z_1 where $\psi'_1(\mathbf{r}) = 0$ and a region N_1 where $\psi'_1(\mathbf{r}) \neq 0$. We define Z_2 and N_2 similarly. In the region $N_1 \cap Z_2$ we can multiply ψ'_1 by an arbitrary phase $e^{i\chi_1}$ without changing $|\psi'_1(\mathbf{r})|^2$. In the region Z_1 this is true as well as $|\psi'_1(\mathbf{r})|^2 = 0$. Similarly $\psi'_2(\mathbf{r})$ can be multiplied by an arbitrary phase $e^{i\chi_2}$ in the regions Z_2 and $Z_1 \cap N_2$. Let us now address the region $N_1 \cap N_2$. We must certainly have $|\psi'_1(\mathbf{r})|^2 + |\psi'_2(\mathbf{r})|^2 = |\psi_3(\mathbf{r})|^2$, because the probabilities for going through slit S_1 and for going through slit S_2 are mutually exclusive and must add up to the total probability of transmission. We might have started to doubt about the correctness of this idea, due to the way textbooks present the problem, but we should never have doubted. Let us thus put $\psi'_1(\mathbf{r}) = |\psi_3(\mathbf{r})| \cos \alpha e^{i\alpha_1}$, $\psi'_2(\mathbf{r}) = |\psi_3(\mathbf{r})| \sin \alpha e^{i\alpha_2}$, $\forall \mathbf{r} \in N_1 \cap N_2 = W$. In fact, if ψ'_j is a partial solution for the slit S_j , $\psi'_j e^{i\alpha_j}$ will also be a partial solution for the slit S_j . We must take here α , α_1 and α_2 as constants. If we took a solution whereby α , α_1 and α_2 were functions of $\mathbf{r}$, the result obtained would no longer be a solution of the Schrödinger equation in free space, due to the terms containing the spatial derivatives of α , α_1 and α_2 which are not zero. In first instance this argument shows also that we must take $\chi_1 = \alpha_1$ and $\chi_2 = \alpha_2$. We do not have to care about the phases χ_1 and χ_2 in the single-slit experiments, but this changes in the double-slit experiment, because we must make things work out globally. Over $N_1 \cap Z_2$, we must have $\psi'_1(\mathbf{r}) = \psi_3(\mathbf{r}) = \psi_1(\mathbf{r})$, as $\psi_2(\mathbf{r}) = 0$. Similarly, $(\forall \mathbf{r} \in N_2 \cap Z_1) (\psi'_2(\mathbf{r}) = \psi_3(\mathbf{r}) = \psi_2(\mathbf{r}))$ as $\psi_1(\mathbf{r}) = 0$. Then $\int_W |\psi'_1(\mathbf{r})|^2 d\mathbf{r} = \cos^2 \alpha \int_W |\psi_3(\mathbf{r})|^2 d\mathbf{r}$ and $\int_W |\psi'_2(\mathbf{r})|^2 d\mathbf{r} = \sin^2 \alpha \int_W |\psi_3(\mathbf{r})|^2 d\mathbf{r}$. Due to the symmetry, we must have $\int_W |\psi'_2(\mathbf{r})|^2 d\mathbf{r} = \int_W |\psi'_1(\mathbf{r})|^2 d\mathbf{r}$, such that $\alpha = \frac{\pi}{4}$. We see from this that not only $\int_W |\psi'_2(\mathbf{r})|^2 d\mathbf{r} = \int_W |\psi'_1(\mathbf{r})|^2 d\mathbf{r} = \frac{1}{2} \int_W |\psi_3(\mathbf{r})|^2 d\mathbf{r}$, but also $|\psi'_2(\mathbf{r})|^2 = |\psi'_1(\mathbf{r})|^2 = \frac{1}{2} |\psi_3(\mathbf{r})|^2$. In each point $\mathbf{r} \in W$ the probability that the electron has traveled through a slit to get to $\mathbf{r}$ is equal to the probability that it has traveled through the other slit. This is due to the undecidability. The choices $\alpha_1 \neq 0$, $\alpha_2 \neq 0$ we have to impose on the phases, are embodying here the idea that a solution of a Schrödinger equation with potential V_j , for $j = 1, 2$ cannot be considered as a solution of a Schrödinger equation with potential V_3 . The conditions we have to impose on α_1 and α_2 are thus a kind of disguised boundary conditions. They are not true boundary boundaries, but a supplementary condition (a logical constraint) that we want ψ'_1 and ψ'_2 to obey “divine” binary logic.

We can summarize these results as $\psi'_1 = \frac{1}{\sqrt{2}} |\psi_1 \boxplus \psi_2| e^{i\alpha_1}$ and $\psi'_2 = \frac{1}{\sqrt{2}} |\psi_1 \boxplus \psi_2| e^{i\alpha_2}$. Let us write $\psi_1 \boxplus \psi_2 = |\psi_1 \boxplus \psi_2| e^{i\chi}$. We can now calculate α_1 and α_2 by identification. This yields on W :

$$\begin{aligned} \psi'_1 &= \frac{|\psi_1 \boxplus \psi_2| e^{i(\chi + \frac{\pi}{4})}}{\sqrt{2}} = \frac{1}{\sqrt{2}} (\psi_1 \boxplus \psi_2) e^{+i\frac{\pi}{4}} \neq \psi_1 \\ \psi'_2 &= \frac{|\psi_1 \boxplus \psi_2| e^{i(\chi - \frac{\pi}{4})}}{\sqrt{2}} = \frac{1}{\sqrt{2}} (\psi_1 \boxplus \psi_2) e^{-i\frac{\pi}{4}} \neq \psi_2 \end{aligned} \quad (51)$$

What this shows is that the rule $\psi_3 = \psi_1 \boxplus \psi_2$ is logically flawed, because the correct expression is $\psi'_3 = \psi'_1 + \psi'_2$. We were able to get an inkling of this loophole by noticing that $\psi_3 = \psi_1 \boxplus \psi_2$ is not rigorously exact, even if it is an excellent approximation. The differences between ψ'_j and ψ_j , for $j = 1, 2$, are not negligible. The phases of ψ'_1 and ψ'_2 always differ by $\frac{\pi}{2}$ such that they are fully correlated. The difference between the phases of ψ_1 and ψ_2 can be anything. They can be opposite (destructive interference) or identical (constructive interference). In fact, in contrast to ψ_1 and ψ_2 , ψ'_1 and ψ'_2 reproduce the oscillations of the interference pattern. In this respect, the fact that α_1 and α_2 are different by a fixed amount is crucial. It permits to make up for the normalization factor $\frac{1}{\sqrt{2}}$ and end up with the correct numerical result of the flawed calculation $\psi_1 \boxplus \psi_2$. The phases of ψ'_1 and ψ'_2 conspire to render $\psi'_1 + \psi'_2$ equal to $\psi_1 \boxplus \psi_2$.

However, at the boundaries of $N_1 \cap Z_2$ and $N_2 \cap Z_1$ with W there are awkward discontinuities. In $N_1 \cap Z_2$, we must have $\psi'_1(\mathbf{r}) = \psi_1(\mathbf{r})$, while in $N_1 \cap N_2$, we have $\psi'_1(\mathbf{r}) = \frac{1}{\sqrt{2}} (\psi_1(\mathbf{r}) \boxplus \psi_2(\mathbf{r})) e^{+i\frac{\pi}{4}}$. The difference between the solutions is $\frac{\psi_2(\mathbf{r})}{\sqrt{2}} e^{i(\chi - \frac{\pi}{4})}$. This is the value $\psi'_2(\mathbf{r})$ would take over $N_1 \cap Z_2$ if we extrapolated it from W to $N_1 \cap Z_2$. We can consider that we can accept this discontinuity at the boundary, because over $N_1 \cap Z_2$, the question through which slit the electron has traveled is decidable, while over W it is undecidable, such that there is an abrupt change of logical regime at this boundary. In reality, the boundary between W and $N_1 \cap Z_2$ could be more diffuse and be the result of an integration over the slits, such that the description is schematic and the abruptness not real. The main aim of our calculation is to obtain a qualitative understanding rather than a completely rigorous solution. The same arguments can be repeated at the boundary between $N_2 \cap Z_1$ and W . If we accept this solution, then $\psi'_1(\mathbf{r})$ will vanish on slit S_2 and $\psi'_2(\mathbf{r})$ will vanish on slit S_1 .³⁴ We can also consider these discontinuities as a serious issue. We could then

³⁴ One could also try to smooth out the discontinuity mathematically, e.g. by taking something in the spirit of $(1 - e^{-|\mathbf{K} \cdot (\mathbf{r} - \mathbf{r}_0)|}) \psi_3 + e^{-|\mathbf{K} \cdot (\mathbf{r} - \mathbf{r}_0)|} \psi_1 + e^{-|\mathbf{K} \cdot (\mathbf{r} - \mathbf{r}_0)|} \psi_2$, on lines $\mathbf{r} = \mathbf{r}_0 + \mathbf{K}u$, $u \in \mathbb{R}$ through each point $\mathbf{r}_0 \in \partial W$ of the boundary of W , where $|\mathbf{K}|$ is very large and where on each line, the

postulate that we must assume that $N_1 \cap Z_2 = \emptyset$ & $N_1 \cap Z_2 = \emptyset$, in order to avoid the discontinuities. The fact that we have to choose $N_1 \cap Z_2 = \emptyset$ & $N_1 \cap Z_2 = \emptyset$ would then be a poignant illustration of the possible consequences of undecidability. Contrary to intuition, the value we have to attribute in a point of slit S_1 , to the probability that the particle has traveled through slit S_2 is now not zero as we might have expected but $\frac{1}{2}|\psi'_3(\mathbf{r})|^2$. The experimental undecidability biases thus the probabilities such that they are no longer the “divine probabilities”. The two different approaches correspond to Einstein-like and Bohr-like viewpoints. Both approaches are logically tenable when the detector screen is completely in the zone W , because the quantities in $N_1 \cap Z_2$ and $N_2 \cap Z_1$ are then not measured quantities. Of course we could try to measure them by putting the detector screen close to the slits, but this would be a different experiment, leading to different probabilities, as Feynman has pointed out in his analysis.

In analyzing the path integrals, one should recover in principle the same results. However, the pitfall is here that one might too quickly conclude that $\psi'_j = \psi_j$ like in the textbook presentations, which leads us straight into the paradox. We see thus that the Huyghens’ principle is a purely numerical recipe that is physically meaningless, because it searches for a correct global solution without caring about the correctness of the partial solutions. It follows the experimental ternary logic and therefore is allowed to mistreat the phase difference that exists between the partial solutions ψ'_1 and ψ'_2 which always have the same phase difference, such that they cannot interfere destructively. The rule $\psi_3 = \psi_1 \boxplus \psi_2$ is perfectly right in ternary logic where we decide that we do not bother which way the particle has travelled, because that question is empirically undecidable. It is then empirically meaningless to separate ψ'_3 into two parts. This corresponds to Bohr’s viewpoint. It is tenable because it will not be contradicted by experiment. This changes if one wants to impose also binary logic on the wave function, arguing that conceptually the question about the slits should be decidable from the perspective of a divine observer. We do need then a correct decomposition $\psi'_3 = \psi'_1 + \psi'_2$. We find then out that $\psi'_j \neq \psi_j$ and we can attribute this change between the single-slit and the double-slit probabilities to the difference between the ways we must define probabilities in both types of logic. If we were able by divine knowledge to assign to each electron impact on the detector the corresponding number of the slit through which the electron has traveled, we would recover the experimental frequencies $|\psi'_j|^2$. This is Einstein’s viewpoint. Having made this clear, everybody is free to decide for himself if he prefers to study geometry in binary logic or pre-geometry in ternary logic. But refusing Einstein’s binary logic based on the argument that ψ'_1 and ψ'_2 cannot be measured appears to us a stronger and more frustrating *Ansatz* than the one that consists in introducing variables that cannot be measured. The refusal is of course in direct line with Heisenberg’s initial program of removing all quantities that cannot be measured from the theory. It is Heisenberg’s minimalism which preserves the experimental undecidability and ternary logic within the theory. To this we can add a supplementary logical constraint which enforces binary logic. As Eq. 51 shows, the difference between the fake partial ternary solutions ψ_j and the correct partial ternary solutions ψ'_j (where $j = 1, 2$) is much larger than we ever might have expected on the basis of the logical loophole that $\psi_3 = \psi_1 \boxplus \psi_2$ is not rigorously exact. In the approach with the additional binary constraint, whereby we follow the spirit of Bohr, we even do not reproduce $\psi'_1(\mathbf{r}) = 0$ on slit S_2 and $\psi'_2(\mathbf{r}) = 0$ on slit S_1 , because these quantities are not measured if we assume that they are only measured far behind the slits. In the Bohr-like approach, the conditions $\psi'_1(\mathbf{r}) = 0$ on slit S_2 and $\psi'_2(\mathbf{r}) = 0$ on slit S_1 are thus not correct boundary conditions for a measurement far behind the slits, because they violate the *Ansatz* of experimental undecidability. The bias in the experimentally measured probabilities due to the undecidability cannot be removed by the divine knowledge about the history. If we want a set-up with the boundary conditions that correspond to the unbiased case whereby $\psi'_1(\mathbf{r}) = 0$ on slit S_2 and $\psi'_2(\mathbf{r}) = 0$ on slit S_1 , we must assume that the detector screen is put immediately behind the slits, and the interference pattern can then not be measured, while everything becomes experimentally decidable. Otherwise, we must assume that $\psi'_1(\mathbf{r})$ and $\psi'_2(\mathbf{r})$ are not measured on the slits such that they can satisfy the undecidable solution in Eq. 51. The partial probabilities given by $\psi'_1(\mathbf{r})$ and $\psi'_2(\mathbf{r})$ are thus extrapolated quantities, and they are only valid within one set-up with a well-defined position of the detector screen. Even in the pure Heisenberg approach whereby one postulates that we are not allowed to ask through which slit the electron has traveled, the wave function contains extrapolated quantities that are not measured, despite the original agenda of that approach. We see also that there is *always* an additional logical constraint that must be added in order to account for the fact that the options of traveling through the slits S_1 or S_2 are mutually exclusive. This has been systematically overlooked, with the consequence that one obtains the result $p \neq p_1 + p_2$, which is impossible to make sense of. It is certainly not justified to use $p \neq p_1 + p_2$ as a starting basis for raising philosophical issues. Moreover, imposing the boundary condition that $\psi'_1(\mathbf{r}) = 0$ on slit S_2 and $\psi'_2(\mathbf{r}) = 0$ on slit S_1 remains a matter of choice, depending on which vision one wants to follow. If we do not clearly point out this choice, then confusion can enter the scene and lead to paradoxes, because adding this boundary condition amounts to adding information and information biases the definition of the probabilities.

In summary, $\psi_3 = \psi_1 \boxplus \psi_2$ is wrong if we cheat by wanting to satisfy also binary logic in the analysis of an experiment that follows ternary logic by attributing meaning to ψ_1 and ψ_2 . But it yields the correct numerical result for the total wave function if we play the game and respect the empirical undecidability by not asking which way the particle has traveled. We can thus only uphold that the textbook rule $\psi_3 = \psi_1 \boxplus \psi_2$ is correct if we accept that the double-slit experiment experimentally follows ternary logic. Within binary logic, the agreement of the numerical result with the experimental data is misleading, as such an agreement

vector $\mathbf{K}$ is taken perpendicular to the boundary of W . The domain W has at least two boundaries (there could be more due to the diffraction), but in each transition region only one of the families of lines will give rise to a contribution $e^{-i\mathbf{K} \cdot (\mathbf{r} - \mathbf{r}_0)}$ that is not negligible. This is elaborate but could be feasible because certain terms cancel one another. However, such a purely mathematical approach is highly artificial and contains a lot of arbitrariness, which is difficult to remove on the basis of physical arguments. Accepting the discontinuity of the schematic description seems in this respect more physical.

does not provide a watertight proof for the correctness of a theory. If a theory contains logical and mathematical flaws, then it must be wrong despite its agreement with experimental data.³⁵

Eq. 51 shows that interference does not exist, because the phase factors $e^{i\frac{\pi}{4}}$ and $e^{-i\frac{\pi}{4}}$ of ψ'_1 and ψ'_2 always add up to $\sqrt{2}$. We may note in this respect that χ itself is only determined up to an arbitrary constant within the experiment. The wave function can thus not become zero due to phase differences between ψ'_1 and ψ'_2 like happens with ψ_1 and ψ_2 in $\psi_1 \boxplus \psi_2$. When ψ'_3 is zero, both ψ'_1 and ψ'_2 are zero. Interference thus only exists within the purely numerical, virtual reality of the Huyghen's principle, which is not a narrative of the real world. We must thus not only dispose of the particle-wave duality, but also be very wary of the wave pictures we build based on the intuition we gain from experiments in water tanks. These pictures are apt to conjure up a very misleading imagery that leads to fake conceptual problems and stirs a lot of confusion. The phase of the wave function has physical meaning, and spinors can only be added meaningfully if we get their phases right.

12.5 Why is there an interference pattern?

This way we have solved the probability paradox that $\psi_3 = \psi_1 \boxplus \psi_2$ & $|\psi_3|^2 \neq |\psi_1|^2 + |\psi_2|^2$. But we have not yet explained the other part of the paradox, which is why we obtain an interference pattern. The general basic idea is here that the boundary conditions which describe the global context and define ψ_1 will change when we open the second slit, such that the definition of the probabilities changes (a fact we could qualify as a paradox of Bertrand). Just as the wave function must globally satisfy in a self-consistent way all the boundary conditions of the Dirichlet problem formulated by the wave equation (such that we must use Huyghens' principle to construct it), also the probabilities must be globally defined to make sure that they are consistent with the context. The requirement of global self-consistence is e.g. the reason we have to consider the topological connectivity of the definition domain in the Aharanov-Bohm experiment (and in our treatment of the double-slit experiment in reference [1]). A quantitative explanation is of course not sufficient.

The general idea points out the logic, but does not provide very much intuition about the reason why we have an interference pattern. On the other hand, the Born rule provides a good explanation, but it still looks abstract. We may still ask for a better understanding of the reasons why the interference pattern occurs and if it is really due to the ternary logic. In fact, up to now, we have only discussed the phases. We must also discuss the amplitudes of the wave functions. Let us observe in this respect that $\psi'_1(\mathbf{r})\psi'_2(\mathbf{r}) + \psi'_1(\mathbf{r})\psi'^*_2(\mathbf{r}) = 0$, $\forall \mathbf{r} \in V$, such that $\psi'_1(\mathbf{r})$ and $\psi'_2(\mathbf{r})$ are everywhere in V orthogonal with respect to the Hermitian norm. This result actually ensures that $|\psi'_1 + \psi'_2|^2 = |\psi'_1|^2 + |\psi'_2|^2$ such ψ'_1 and ψ'_2 are describing mutually exclusive probabilities. To obtain this orthogonality condition we must have actually that $\psi'_1 = \psi'_2 e^{\pm i\frac{\pi}{2}}$. When this condition is fulfilled, the functions ψ'_1 and ψ'_2 become exact wave functions for the double-slit experiment that in addition satisfy the boundary conditions that ψ'_1 is zero in slit S_2 and ψ'_2 is zero in slit S_1 . They can then actually be summed according to the superposition principle. Let us now assume that there exists a wave function ζ for the double-slit experiment. We do not assume here that we know that $\zeta = \psi_1 \boxplus \psi_2$ must be true. But it must be possible to decompose it into unknown functions ζ_1 and ζ_2 according to binary logic as described for ψ'_1 and ψ'_2 above. We can consider a continuous sweep over the detector screen from the left to the right, whereby we are visiting the points P . Let us call the source G , and the centres of the slits C_1 and C_2 and reduce the widths of the slits such that only C_1 and C_2 are open. The phase differences over the paths GC_1P through ψ_1 and GC_2P through ψ_2 will then continuously vary along this sweep. To be undecidable and obey ternary logic, the total wave function would have to be completely symmetrical with respect to ζ_1 and ζ_2 , such that it would have to be $\zeta_1 + \zeta_2$. But simultaneously, ζ_1 and ζ_2 would have to be mutually exclusive and satisfy binary logic because "God would know". All points where the phase difference is not $\pm \frac{\pi}{2} \pmod{2\pi}$ should therefore have zero amplitude and be excluded from the domain of ζ . From this point of view, an interference pattern (with its quantized momenta $\mathbf{p} = \hbar \mathbf{q}$) appears then as the only solution for the wave function that can satisfy simultaneously the requirements imposed by the binary and the ternary logic. The real experiment with non-zero slit widths would not yield the Dirac comb but a blurred result due to integration over the finite widths of the slits. A different point $D_2 \neq C_2$ might indeed provide an alternative path GD_2P that leads to the correct phase difference $\frac{\pi}{2}$ with GC_1P . The smaller the wavelength of the electron, the easier it will be to find such points D_2 , such that when we steadily increase the energy of the electron we will eventually end up in the classical tennis ball regime. Finally note that the phase difference of $\pm \frac{\pi}{2} \pmod{2\pi}$ over the paths GS_1P through ψ'_1 and GS_2P through ψ'_2 is not in contradiction with the phase difference $2\pi n$ over the paths GS_1P and GS_2P through ψ_3 we discussed in Subsection 12.3, because $\psi_3 = \psi'_1 + \psi'_2$.

The true reason why we can calculate $\psi'_3 = \psi'_1 + \psi'_2$ as $\psi_3 = \psi_1 \boxplus \psi_2$ can within the Born approximation also be explained by the linearity of the Fourier transform used in Eq. 32 [2], which is a better argument than invoking the linearity of the wave equation. The path integral result is just a more refined equivalent of this, based on a more refined integral transform. The reason for the presence of the Fourier transform in the formalism is the fact that the electron spins [2, 1]. One can derive the whole wave formalism purely classically, just from the assumption that the electron spins. Eq. 32 hinges also crucially on the Born rule

³⁵ This is a recurrent theme in QM [1, 2]. Many conceptual problems of QM trace back to the poor understanding of the geometrical meaning of the spinors in the group theory. The algebra is used as a black box, and in its philosophy to "shut up and calculate" it violates the meaning of the geometry while yielding numerically correct results. But these geometrical errors may manifest themselves by crystallizing under the form of some conceptual difficulties. In binary logic, the presence of destructive interference $\psi_1(\mathbf{r}) + \psi_2(\mathbf{r}) = 0$ in the theory leads to insuperable conceptual difficulties, which should have raised a red flag.

$p = |\psi|^2$. There is no rigorous proof for this rule but there exists ample justification for it [2]. The undecidability is completely due to the properties of the potential which defines both the local interactions and the global symmetry.

12.6 Making a clear distinction between determinism and experimental decidability

This scenario is not artificial as it is wired into the formalism of QM. We must point out that the path of the electron is determined, but that the fact that you cannot see and collect the corresponding information renders the question which way the electron travelled undecided. One must thus make a clear distinction between determinism and decidability in order to avoid the apparent contradiction that there exist situations where we could call the paths both “determined” and “not determined”, by pointing out and defining that these paths are “determined” but effectively “undecidable” within the context of the set-up (rather than “not determined”). This disambiguation will then remove the confusion. Determinism is about the “absolute truth” (Einstein), decidability is about what the set-up of an experiment can decide about that truth (Bohr). There are thus two levels of “truth”, an absolute one and an experimental one. The error you make by reasoning with binary logic and the incoherent rule on a set-up that follows ternary logic and a coherent rule is given by the real number $\psi_1^* \psi_2 + \psi_2^* \psi_1$, which expresses the overlap between ψ_1 and ψ_2 . The sign of that number can go either way, such that it models the alternating pattern of decidable and undecidable questions. The error can disappear on average over $\mathbb{R}^3$ after integration (weak orthogonality) and then still show up on a more detailed local scale (absence of strong orthogonality).³⁶ It is the Fourier transform which is responsible for the uncertainty relations, but these reflect only one possible type of undecidability. In the double-slit experiment, it is the rule of thumb $D \lesssim \lambda$ which can be used to predict undecidability of the question through which slit the electron may have traveled.

We have seen that the description of the set-up is idealized and cannot account for certain microscopic details. Fig. 1 is a purely geometrical description without any details about atomic positions and interaction potentials and can therefore not possibly be a complete description of the experimental set-up. It is therefore logical to assume that we would not only need hidden variables to deterministically describe the electrons but also to accurately describe the microscopic details of the set-up.

13 Synthetic Overview and Conclusion

The flow chart below summarizes the main ideas developed in this paper (with their rather complex interconnections). We explained how the double-slit paradox can be understood much better by considering it as an experiment whereby one uses electrons to study the set-up rather than an experiment whereby we use a set-up to study the behaviour of electrons. We have also shown that Heisenberg’s uncertainty principle is related to Gödel’s concept of undecidability and how this can be used in an intuitive way to make sense of the paradox of Young’s double-slit experiment. Taking this into account leads naturally to the rules that dictate how one should calculate the probabilities for coherent and incoherent scattering in QM. The calculations of QM look very specific because they are based on the notion of spin. It is this notion of spin which is responsible for the introduction of the Fourier transform and the probability amplitudes. The formalism looks therefore a very specific one-time shot. We may wish a broader perspective. This may require developing new mathematics, *viz.* a set theory and a probability theory based on ternary logic, in order to obtain a more general context for the calculations. However, the natural way to treat undecidability seems to be using symmetry arguments, and the corresponding group-theoretical arguments will presumably also introduce complex-valued eigenstates and probability amplitudes.

We could ask: “We are not so stupid that we would not be able to reason without contradictions about logic, are we?” Not a single defeatist, not a single quantum priest, not a single editor, not a single referee, not a single moderator has the right to decide the answer to that question single-handedly and secretly on behalf of whole mankind. Keeping the words of Dieudonné in our mind [12], it is our duty to be proud and confident that the answer will always be “no, we are not so stupid”. We can appreciate that sometimes it will really take a big fight, because the ambiguity between determinism and decidability was a really hard nut to crack. But even when we have the impression that there is no light, we shall never, ever give up! *“Pour l’honneur de l’esprit humain!”*

Acknowledgement. I thank my director, Kees van de Beek his continuous interest in this work and the CEA for the possibility to carry out this research. I also thank Andre Gontijo Campos for fruitful discussions.

³⁶ There is certainly information about the actual path of the electron written into the phase difference $\varphi_B - \varphi_G$ an electron builds up over a trajectory GB between source and detector. Within the patch corresponding to χ_j (discussed in Subsection 4.2) the phase difference $\varphi_1 - \varphi_2$ between the two alternative paths GC_1B and GC_2B is $2\pi j$ with $j \in \mathbb{Z}$. The traveling time from source to detector for particles that go through slit S_1 is thus different than that for particles that go through slit S_2 . Measuring this time in order to determine through which slit the electron has gone is also subject to an uncertainty relation, and will thus also involve considerations about the interaction of the electrons with parts of the set-up. This problem corresponds to a different set-up with a different wave function ($\psi(\mathbf{r}, t)$ instead of $\psi(\mathbf{r})$).

❶ Justification of Born's rule: $p = |\psi|^2$

- $\frac{\partial}{\partial t} \Rightarrow$ derivation of a continuity equation from the wave equation
- $\frac{\partial^2}{\partial t^2} \Rightarrow$ (photons): see Feynman
- probability charge-current four-vector must be quadratic expression in terms of spinors
- In SU(2) the quantity $\psi\psi^\dagger = \xi_0\xi_0^* + \xi_1\xi_1^* = 1$ for a single spinor. On the group ring $\psi\psi^\dagger$ counts thus spinors (electrons)

Section 10

❷ Motion of spinning electron on $\mathcal{P} \Rightarrow$ Dirac-like equation: $\psi = e^{-\frac{i}{\hbar}(Et - \mathbf{p} \cdot \mathbf{r})} \psi(0) \rightarrow$ ③ Two extrapolations:

Superposition principle: Dirac-like $\Rightarrow$ Dirac	Huyghens' principle: from $\mathcal{P}$ to $\mathbb{R}^4$
introduction of sets, statistical ensembles	potentialities, consistent histories
$\psi = c_1\psi_1 + c_2\psi_2$ $\textcircled{1} \rightarrow p = \sum_j c_j ^2 \psi_j ^2$	$\psi = c_1\psi_1 \boxplus c_2\psi_2$ $\textcircled{1} + \textcircled{3} \rightarrow p = \psi ^2$
Solution of the paradox of Schrödinger's cat	Transition: classical mechanics $\rightarrow$ QM
Section 2	Sections 2 - 3, Subsection 7.2, carried out in Section 9 ← ③

Section 1

❸ Extrapolation from $\mathcal{P}$ to $\mathbb{R}^4$ for a simple double-slit potential V (transition CM $\rightarrow$ QM)

- Fourier transform of V due to spin expressed by $\psi = e^{-\frac{i}{\hbar}(Et - \mathbf{p} \cdot \mathbf{r})} \psi(0) \leftarrow$ ② (compare with Born approximation)
- proof of Huyghens' principle: $\psi \approx \psi_1 \boxplus \psi_2 \rightarrow$ ②
- wave function must be a function $\Rightarrow$ quantization
- ψ contains all the information about the set-up $\rightarrow$ ④
- Feynman's path integral is a Huyghens' principle, the paths are purely mathematical and not physical

Section 9

❹ End of particle-wave duality

- electrons are particles
- probability amplitudes ψ for statistical ensembles of particles are waves (is literally what QM says) $\leftarrow$ ①
- ψ is non-local, the electrons interact only locally
- ψ contains all the information about the set-up by Fourier transform as $\mathcal{F}$ is bijective Section 9 $\leftarrow$ ③
- Useful paradigm shift: we are not measuring electrons with a set-up but a set-up with electrons Subsection 6.1

Section 7.1

❺ Undecidability in double-slit experiment Subsection 6.2

- Feynman's analysis hinting at a third possibility (undecided) Sections 4-5
- decided histories $\Rightarrow$ superposition principle: incoherent summing
- undecided histories $\Rightarrow$ Huyghens' principle: coherent summing
- Heisenberg's uncertainty relation is one example of a rule of thumb for undecidability
- applying binary averaging to a distribution that follows ternary logic is wrong Subsection 7.3, Section 12
 - $\rightarrow$ This is due to the absence of a common probability distribution:
 - $\rightarrow$ is literally what QM says for $p(A_1)$ and $p(A_2)$ when $[A_1, A_2] \neq 0$
 - $\rightarrow$ compare with criticism of Bell inequalities (derived from commutation relation or differently (see [1], p. 278))
- Another explanation of the interference pattern is possible (see [1], pp. 327-335)

Subsection 7.3

- ⑥ A symmetry description does not provide clues as to the mechanism (see [1], p. 46, pp. 337-340)
 → agreement with experiment by good fortune
 → necessity of a case-by-case study (H atom, tunneling)

Various remarks in the text

- ⑦ Strongly orthogonal parts of a wave function ψ do not produce interference

Section 11

References

1. G. Coddens, in *From Spinors to Quantum Mechanics*, Imperial College Press, (2015).
2. G. Coddens, CEA-01269569 (2016).
3. G. Farmelo, in *The strangest Man. The hidden Life of Paul Dirac: Quantum Genius*. Faber and Faber (2009), p. 430.
4. R.P. Feynman, R. Leighton, and M. Sands, in *The Feynman Lectures on Physics, Vol. 3*, Addison-Wesley, Reading, MA.
5. See e.g. B.E. Sagan in *The Symmetric Group. Representations, Combinatorial Algorithms, and Symmetric Functions*, Springer Graduate Texts in Mathematics, **203**, Second Edition, Springer, New York, (2001), who states clearly (on p.7) that the linear combinations are purely formal. Instead of a group ring ([7]), Sagan calls the set of purely formal linear combinations a *G*-module, and the corresponding calculus the group algebra.
6. H.S.M. Coxeter in *Regular Polytopes*, Dover (1973).
7. B.L. van der Waerden in *Group Theory and Quantum Mechanics*, Springer, Berlin-Heidelberg, (1974), translated from *Die gruppentheoretische Methode in der Quantenmechanik*, Springer, Berlin (1932).
8. R.B. Griffiths, in *Consistent Quantum Theory*, Cambridge University Press, Cambridge, (2002).
9. A. Sinha, A.H. Vijay and U. Sinha, Scientific Reports **5**, 10304 (2015); R. Sawant, J. Samuel, A. Sinha, S. Sinha and U. Sinha, Phys. Rev. Lett. **113**, 120406 (2014).
10. K. Gödel, in “Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme, I”, Monatshefte für Mathematik und Physik, **38**, 1, pp. 173-198 (1931); translated in *On Formally Undecidable Propositions of Principia Mathematica and Related Systems*, Dover, New York, (1992); D.R. Hofstadter in *Gödel, Escher, Bach: An Eternal Golden Braid*, Vintage Books, New York (1989); R.M. Smullyan, in *Gödel's Incompleteness Theorems*, Oxford University Press, New York, (1992).
11. P. Cohen, Proceedings of the National Academy of Sciences of the United States of America **50** (6), 1143-1148 (1963).
12. J. Dieudonné, in *Pour l'honneur de l'esprit humain - les mathématiques aujourd'hui*, Hachette, Paris (1987). Poincaré's half-plane model for hyperbolic geometry reveals how the information contained in Euclid's first four postulates is not sufficient to identify straight lines as really straight. The model shows how they can still be interpreted in terms of half circles in the half plane.
13. F. Hund, Z. Phys. **117**, 1 (1941); A. Hansen and F. Ravndal, Phys. Scr. **23**, 1033 (1981); P.J.M. Bongaarts and S.N.M. Ruijsenaars, Annals Phys. **101**, 289 (1976).
14. D. Bohm, in *The Undivided Universe: An ontological interpretation of quantum theory*, with B.J. Hiley, Routledge, (1993); D. Bohm, Phys. Rev. **85**, 166-193, (1952).
15. J.A. Wheeler in *Mathematical Foundations of Quantum Theory*, A.R. Marlow editor, Academic Press (1978), pp. 9-48; J.A. Wheeler and W.H. Zurek, in *Quantum Theory and Measurement* (Princeton Series in Physics).
16. A. Tonomura, J. Endo, T. Matsuda and T. Kawasaki, Am. J. Phys. **57**, 117 (1989); see also M.P. Silverman in *More than one Mystery, Explorations in Quantum Interference*, Springer, p. 3 (1995).
17. A.M. Gleason, A. M. (1957). Indiana University Mathematics Journal, **6**, 885 (1957).
18. A. Shimony in *The New Physics*, Paul Davies ed., Cambridge University Press, Cambridge (1983), pp. 373-395.
19. A. Aspect, J. Dalibard, and G. Roger, Phys. Rev. Lett. **49**, 91, 1804 (1982); J.F. Clauser and A. Shimony, Rep. Progr. Phys. **41**, 1881 (1978); S.J. Freeman and J.F. Clauser, Phys. Rev. Lett. **28**, 938 (1972); J.F. Clauser, M.A. Horne, A. Shimony and R.A. Holt, Phys. Rev. Lett. **23**, 880 (1969); J.F. Clauser and M.A. Horne, Phys. Rev. D **10**, 526 (1974).
20. G. Malinowski, in *Handbook of the History of Logic Volume 8. The Many-Valued and Non-monotonic Turn in Logic*, D.M. Gabbay, J. Woods (eds.), Elsevier, (2009).
21. S. Wagon in *The Banach-Tarski Paradox*, Cambridge University Press, Cambridge, (1985).
22. J. Cramer, Reviews of Modern Physics **58**, pp. 647-688, (1986).
23. T.S. Kuhn, in *The Structure of Scientific Revolutions*, University of Chicago Press, Chicago, (1962), see p.114 of 3rd edn.
24. R.S. Longhurst, in *Geometrical and Physical Optics*, Longman, London, (1967).
25. R.P. Feynman, *QED, The Strange Theory of Light and Matter*, Princeton University Press (1988).
26. P.A.M. Dirac, Physikalische Zeitschrift der Sowjetunion, Band 3, Heft 1 (1933).
27. R.P. Feynman in *The Character of Physical Law*, MIT Press (1967).
28. J.J. Duistermaat in *Huygens' Principle 1690-1990: Theory and Applications*, H. Blok, H. Ferwerds and H.K. Kuiken eds., (Elsevier, 1992), p. 273; J. Hadamard in *Le problème de Cauchy et les Equations aux Dérivées partielles linéaires hyperboliques*, reprint of the 1923 publication, (Dover, New York, 1952).
29. R.L. Wiegand and J.W. Johnson, in *Elements of wave theory*, Proceedings 1st International Conference on Coastal Engineering, Long Beach, California, ASCE, (1950) pp. 5-21.

30. P. Marmier and E. Sheldon, in *Physics of Nuclei and Particles*, Academic Press, New York, (1969). The derivation of a continuity equation from the wave equation is only a part of Born's very thorough original discussion (M. Born, *Zeitschrift für Physik* **38**, 803-827 (1926)).
31. R.P. Feynman, in *The Feynman Lectures on Physics, Vol. 2*, Addison-Wesley, Reading, MA (1970).
32. Y.-H. Kim, R. Yu, S.P. Kulik, Y. Shih et M.O. Scully, *Phys. Rev. Lett.* **84**, 1-5 (2000).
33. Th.M. Nieuwenhuizen, *Foundations of Physics* **41**, 580 (2011).