

How a representation-independent approach to quantum mechanics can aid student learning

James K. Freericks^{1*} 

¹ Department of Physics, Georgetown University, Washington, DC, USA

Received 09 December 2025 • Accepted 09 June 2026

Abstract

Quantum mechanics is traditionally taught in a specific representation—the position-space representation—because the time-independent Schrödinger equation is simplest to formulate there. But this leads to a number of conceptual issues both associated with language and with the formalism, especially if one starts more formally and then emphasizes representations for calculations. In this work, we describe a number of these issues and argue why restricting to a representation-independent formalism (that is, working with abstract operators and state vectors independent of any specific representation) fixes many of these issues. Then representations can be introduced later when needed for specific applications (but we find the need to do this can be sharply reduced from current paradigms).

Keywords: quantum mechanics, representation-independent formalism, precise language

INTRODUCTION

Back when I first learned classical mechanics in my freshman year in college, the instructor spent a lot of time focusing on the notion that the position vector we use to describe motion is an abstract vector that is independent of the choice of coordinates used to represent it. In fact, we had many choices we could make for coordinates. I was shocked when I was first shown a y-axis that pointed downward for the problem of a particle falling due to gravity. Eventually, I was able to accept and even believe that this was the right thing to do.

Then when I moved into quantum mechanics, we went back to working in a specific coordinate system—nearly everything was done in position space! Why? What happened to the use of abstract vectors? Why were we no longer working in a representation-independent fashion? The most common answer to this question is that because the kinetic energy is the same in all Hamiltonians, it is easier to work in position space, so that the representation in terms of derivatives all enter in the same form. Indeed, energy eigenvalue problems all become Sturm-Liouville problems in position space,

which have been heavily studied by mathematicians. If one tried to work in momentum space, then one has many different types of equations—while the harmonic oscillator takes the same form, the hydrogen problem becomes an integral equation. Trying to work in other representations is even worse. So, doesn't the necessity of how one solves the problems in practice force the position representation onto us?

Does it? Should we give up so easily?

Mathematicians have long known that working with a representation-independent formalism is superior to working on a specific concrete basis. A classic example is differential geometry. The early form of the theory was worked out in terms of the specific coordinate system being used and took the form of the following equations:

$$(\nabla_\mu X)^\nu = \partial_\mu X^\nu + \Gamma_{\mu\lambda}^\nu X^\lambda \text{ and} \quad (1)$$

$$R_{\sigma\mu\nu}^\rho = \partial_\mu \Gamma_{\nu\sigma}^\rho - \partial_\nu \Gamma_{\mu\sigma}^\rho + \Gamma_{\mu\lambda}^\rho \Gamma_{\nu\sigma}^\lambda - \Gamma_{\nu\lambda}^\rho \Gamma_{\mu\sigma}^\lambda,$$

where the Einstein summation convention for repeated indices is used. The first line is the covariant derivative and the second is the Riemann curvature tensor, with Γ the Christoffel symbol. The main issue with these formulas, aside from their being a bit lengthy, is that they do not show the invariance of these results under

A preliminary version of this paper was presented at the GIREP 2025 Conference.

Contribution to the literature

- This work discusses the pros and cons of precision in the mathematical formalism and in the language used to describe quantum mechanics in instruction.
- It argues that sticking to one formal approach reduces cognitive load for the students and can reduce the length of calculations.
- Precision in language also provides clarity for understanding concepts and how they inter-relate. Numerous examples are given to illustrate the different ideas.

coordinate transformations. Instead, in a representation-independent approach, one defines the covariant derivative as a linear map that satisfies the Leibniz rule and then the curvature tensor is defined in terms of an appropriate form of two covariant derivatives acting on the vector field. This latter form is manifestly invariant with respect to coordinate changes, emphasizing the structure of the theory without relying too heavily on how it is concretely implemented in a specific example. One does often have to get concrete and work on a specific basis when performing an actual calculation; we will discuss this further below with respect to quantum mechanics, where many calculations can actually be carried out in the representation-invariant formalism and do not need to be evaluated in a particular representation. Indeed, unlike what many believe, very few problems in quantum mechanics require one to use a concrete representation. A glimpse of this result can be seen in the common way the harmonic oscillator is worked with in instruction, where the operator-based method can be used to find energy eigenstates, eigenvalues, and compute matrix elements, all without requiring wavefunctions, as is needed if one solves the differential form of the Schrödinger equation for wavefunctions in position space. In quantum mechanics formalism, a representation-independent development of the material is when one uses only abstract operators and abstract states to perform calculations of states that have definite properties (such as energy) and to compute matrix elements and wavefunctions in the same fashion. One never uses differential equations when working in a representation-independent fashion. Manipulations are always algebraic in nature.

Because physics has chosen to emphasize the coordinate representation so much in the development of quantum mechanics, we find many concepts become conflated, leading students (and faculty) to misunderstandings. One of the most common incorrect statements found in essentially all books is with regards to precisely what the abstract momentum operator is.

If we work in an abstract Hilbert space, and we use Dirac notation, then operators denote with hats, act on abstract state vectors and map them to abstract state vectors. This is the representation-independent way of formulating their definitions. Then an abstract state of definite momentum is defined by the eigenvalue-eigenvector relation given by

$$\hat{p}|p_0\rangle = p_0|p_0\rangle, \quad (2)$$

where $\hat{p}$ is the abstract operator, p_0 is the eigenvalue, and $|p_0\rangle$ is the abstract state vector of definite momentum. Note how the abstract state vector is mapped into a scaling factor multiplying itself by the action of the momentum operator, indicating it is an eigenvector of the abstract momentum operator. Within this context, derivatives cannot act on state vectors. The Hilbert space is not constructed with a differentiable structure in terms of its state vectors. Instead, in a representation-independent approach, one works solely with states and operators and not their representations in position space, momentum space, or some other representation.

One can also determine this state of definite momentum within the representation-independent formalism—it follows directly from then canonical commutation relation $[\hat{x}, \hat{p}] = i\hbar$. We start by noting that there is a special state, the state at the origin of momentum space, which satisfies $\hat{p}|0_p\rangle = 0$ (because $p_0 = 0$), where the p subscript is used so as to not confuse this state with any other state labeled by 0. Then, the state $|p_0\rangle$ is constructed as

$$|p_0\rangle = e^{\frac{i}{\hbar}p_0\hat{x}}|0_p\rangle. \quad (3)$$

This is verified by acting $\hat{p}$ onto the state, using a “multiply by one” to create a Hadamard lemma, evaluating the Hadamard lemma (which truncates after the second term), and then using the defining property of the momentum state at the origin. We have

$$\begin{aligned} \hat{p}|p_0\rangle &= \hat{p}e^{\frac{i}{\hbar}p_0\hat{x}}|0_p\rangle = \\ &\underbrace{e^{\frac{i}{\hbar}p_0\hat{x}}e^{-\frac{i}{\hbar}p_0\hat{x}}}_{\text{multiply by 1}}\hat{p}e^{\frac{i}{\hbar}p_0\hat{x}}|0_p\rangle = \\ &\underbrace{e^{\frac{i}{\hbar}p_0\hat{x}}e^{-\frac{i}{\hbar}p_0\hat{x}}\hat{p}e^{\frac{i}{\hbar}p_0\hat{x}}}_{\text{Hadamard lemma}}|0_p\rangle = \\ &e^{\frac{i}{\hbar}p_0\hat{x}}\left(\hat{p} - \frac{i}{\hbar}p_0[\hat{x}, \hat{p}] + \dots\right)|0_p\rangle = \\ &e^{\frac{i}{\hbar}p_0\hat{x}}(\hat{p} + p_0)|0_p\rangle = p_0e^{\frac{i}{\hbar}p_0\hat{x}}|0_p\rangle = p_0|p_0\rangle. \end{aligned} \quad (4)$$

So, one can find the state of definite momentum in a representation-independent fashion, given the existence of a state of definite momentum with zero momentum; the existence of such a state can be established once squeezed states of the harmonic oscillator are discussed,

and one takes an extreme squeezing limit. Note that the Hadamard lemma is given by the infinite series

$$e^{\hat{A}} \hat{B} e^{-\hat{A}} = \hat{B} + [\hat{A}, \hat{B}] + \frac{1}{2!} [\hat{A}, [\hat{A}, \hat{B}]] + \dots, \quad (5)$$

with ever increasing nested commutators. In the case worked with here, the series truncates after the second term because $[\hat{x}, \hat{p}] = i\hbar$ is a constant, and so it commutes with $\hat{x}$. Note that only the abstract properties of the states and vectors were used in this derivation. This is what is meant by a representation-independent calculation.

In fact, one can use the representation-independent approach to also find the wave function for the state of definite momentum, since it is equal to

$$\begin{aligned} \psi_{p_0}(x) &= \langle x | p_0 \rangle = \langle x | e^{\frac{i}{\hbar} p_0 \hat{x}} | 0_p \rangle = \\ e^{\frac{i}{\hbar} p_0 x} \langle x | 0_p \rangle &= e^{\frac{i}{\hbar} p_0 x} \left\langle 0_x \left| e^{\frac{i}{\hbar} x \hat{p}} \right| 0_p \right\rangle = \\ e^{\frac{i}{\hbar} p_0 x} \langle 0_x | 0_p \rangle &= \frac{1}{\sqrt{2\pi\hbar}} e^{\frac{i}{\hbar} p_0 x}, \end{aligned} \quad (6)$$

where the wavefunction of the state of definite momentum at $x = 0$ is equal to its normalization constant; that is $\langle 0_x | 0_p \rangle = 1/\sqrt{2\pi\hbar}$. Here, we used the property that the state $|x\rangle$ is a state of definite position, which satisfies $\hat{x}|x\rangle = x|x\rangle$. Note that this wavefunction is identical to the one usually found from employing differential equations.

But hold it, isn't the momentum operator given by $-i\hbar d/dx$? You may be surprised that the answer is no! It is not! And this means that essentially every quantum mechanics textbook is using improper language, including all of the most popular textbooks, even your favorite.

Let us show the correct way to determine this matrix representation of the momentum operator. The treatment we give here follows closely the discussion in Cohen-Tannoudji et al. (1977). Consider a state $|\psi\rangle$ and it's associated wavefunction $\langle x|\psi\rangle$ in the position-space basis. Then to compute the action of the momentum operator on this basis, we want to compute $\langle x|\hat{p}|\psi\rangle$, which can be thought of as the wavefunction in position space corresponding to the abstract state vector $\hat{p}|\psi\rangle$, determined after the abstract momentum operator acts on the state. This is done by inserting a complete set of states in momentum space, using the fact that $\langle x|p\rangle = e^{\frac{i}{\hbar}xp}/\sqrt{2\pi\hbar}$, and the Feynman trick of differentiating under the integral sign. We have

$$\begin{aligned} \langle x|\hat{p}|\psi\rangle &= \int_{-\infty}^{\infty} dp \langle x|\hat{p}|p\rangle \langle p|\psi\rangle = \\ \int_{-\infty}^{\infty} dp \langle x|p\rangle p \langle p|\psi\rangle &= \int_{-\infty}^{\infty} dp \frac{e^{\frac{i}{\hbar}xp}}{\sqrt{2\pi\hbar}} p \langle p|\psi\rangle = \\ -i\hbar \frac{d}{dx} \int_{-\infty}^{\infty} dp \frac{e^{\frac{i}{\hbar}xp}}{\sqrt{2\pi\hbar}} \langle p|\psi\rangle &= \\ -i\hbar \frac{d}{dx} \int_{-\infty}^{\infty} dp \langle x|p\rangle \langle p|\psi\rangle &= -i\hbar \frac{d}{dx} \langle x|\psi\rangle = \\ -i\hbar \frac{d}{dx} \psi(x). \end{aligned} \quad (7)$$

The derivative acts only on a wavefunction, not on an abstract state vector. Officially speaking it is a continuous diagonal matrix, so it should be written as

$$\hat{p} \leftrightarrow -i\hbar \frac{d}{dx} \delta(x - x'), \quad (8)$$

so that the continuous matrix-vector multiplication becomes

$$\hat{O}|\psi\rangle \leftrightarrow \int_{-\infty}^{\infty} dx' O(x, x') \psi(x'), \quad (9)$$

where the double arrows denote the equivalence between the abstract objects on the left and their concrete representations on the right. Then the proper expression for the form of the continuous matrix that represents the momentum operator acting on a state becomes

$$\langle x|\hat{p}|\psi\rangle = -i\hbar \frac{d}{dx} \int dx' \delta(x - x') \psi(x'), \quad (10)$$

but because this is always just equal to $-i\hbar d\psi(x)/dx$, the matrix structure is usually suppressed. *Note that the derivative is not an abstract operator!* Abstract operators only act on state vectors. It is instead the matrix representation of the momentum operator in position space. This is simply because we require evaluating the operation of the momentum operator onto a state vector in the position basis to determine the derivative form. This matrix representation language may be a "long-winded" expression to say, but it is the right one to use. Interestingly, while this derivation appears in the Cohen-Tannoudji et al.'s (1977) book, the language used in that book conflates these two ideas, just as almost all the others do; we do find that Sakurai and Napolitano (2020) and Shankar (1994) each do treat this material correctly in symbols, but they still do not call this representation of the momentum operator a matrix.

Now wait, you might say, we are working in the L_2 Hilbert space of square-integrable functions and derivatives are operators that act on those functions. Yes, this is correct. But it is also the old-fashioned "concrete representation in a particular basis" way of discussing things (such as using the coordinate approach to differential geometry with Christoffel symbols). The representation-independent way is always the preferred way, at least by mathematicians, and using it properly allows one to make more precise statements and hopefully to make the discussions less confusing for the students. It allows you to see more clearly how the theory is developed and how the different parts inter-relate.

One can always discuss the L_2 representation if desired, but one should always use the proper language of wavefunctions and matrices when describing it, since it corresponds to choosing a concrete basis.

Aren't there textbooks that already do this? For example, Sakurai and Napolitano (2020) (or Shankar, 1994) focus on working with abstract states, don't they? Yes and no. It is true that much of the formal development in these texts is done in a representation-independent way, but a significant amount of the

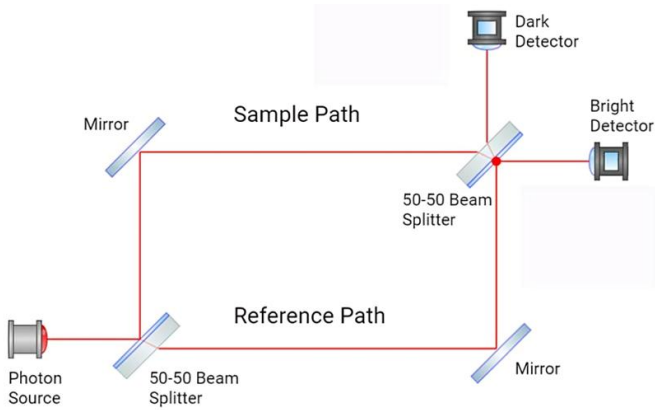


Figure 1. Mach-Zehnder interferometer with a source to the left, a beam splitter, the sample and reference paths, a second beam splitter and two detectors labelled bright and dark (note that the beam splitter's mirrored surface is flipped between the two splitters) (the author's own elaboration)

calculations default to the position representation, especially when solving the energy eigenvalue problems. So, while it is moving in the right direction, it does not go far enough. Once you start switching from abstract states to their representations, it brings in confusion for the students as they need to know whether they are working abstractly or concretely and when do they switch from one to another? We believe it is much simpler for students to stick with the abstract approach and only bring in representations in special cases where they really are necessary.

THE IMPORTANCE OF LANGUAGE

Physicists have known for a long time the importance of using precise language, especially in quantum mechanics (Brookes & Ektina, 2007). To a neophyte, it seems like a bunch of jargon being spewed out, but when used properly, the precision of the language is key to conveying the precision of the ideas. But we do not do this very well in quantum mechanics. We now give some examples of just how we fail in being precise, going beyond just conflating matrices with abstract operators.

Consider the notion of superposition. Philosophers have been insistent that we should say when a system is in superposition, that its state is indeterminate (Maudlin, 2019). Physicists will often become less clear on this notion and will often say things such as the quantum is at two different locations at the same time or goes through both slits of a two-slit experiment, and so forth. Insisting on using the words indeterminate state only helps clarify what we know and what we do not know—for we have never seen an experiment showing a quantum that splits into two instantiations in a superposition.

Entanglement is similar, but perhaps not as problematic. The key language here is to relate which-way information to the formation of an entangled state,

as the two usually go hand in hand. The entanglement, just like the which-way information, often suppresses interference effects and ultimately leads to ideas such as decoherence.

Let us illustrate these ideas with the Mach-Zehnder interferometer using the simplified Feynman (1985) path-integral approach to describe the experiments, as developed in his *QED* book. The Mach-Zehnder interferometer is illustrated in Figure 1. It has a source that inputs light onto a beam splitter, or onto a polarizing beam splitter, then the light that is transmitted goes on the reference path and the light that is reflected goes on the sample path forming an indeterminate state. After reflection off a mirror on each path, they converge again on a second beam splitter and travel either to the bright detector or the dark detector. The probabilities in each detector are determined by the precise difference in length between the two paths, which gives rise to the phase difference $\Delta\phi$ between the two paths.

Feynman (1985) uses an arrow (or, equivalently a complex number probability amplitude) to describe each alternative way an event can occur. Events are described by the initial and final conditions of an experiment, which are that a photon leaves the source and that a detector detects the photon. When traveling through air, the arrow rotates at a speed given by the frequency of the light. After entering the beam splitter the arrow is renormalized by a factor of $1/\sqrt{2}$ for a 50-50 beam splitter and it is multiplied by -1 if the reflection occurs from air. For a mirror, the arrow is multiplied by -1 due to the reflection from air.

Let us use these rules to describe the Mach-Zehnder interferometer with ordinary beam splitters. There are two events. The first is that the photon leaves the source and is detected in the bright detector, the second is that it leaves the source and is detected in the dark detector. Each event has two alternative ways it can occur—by a photon traveling on the sample or the reference path. Note that this delineation of two paths is not saying the photon travels on both, it is just the strategy used in a path-integral approach where all possible alternative ways for an event to occur must be considered. Here, the simplified method uses just two alternate ways. Note that it is a calculation recipe, not an ontology for the physical system.

For the bright detector, we have two ways (reference and sample). For the reference path, the arrow is shrunk by $(1/\sqrt{2})^2 = 1/2$ for the two beam splitters and the phase is $\phi_{source \rightarrow BS} + \phi_{BS \rightarrow mirror}^R + \pi + \phi_{mirror \rightarrow BS}^R + \pi + \phi_{BS \rightarrow bright}$ and for the sample path we have the same shrink factor and the phase is $\phi_{source \rightarrow BS} + \pi + \phi_{BS \rightarrow mirror}^S + \pi + \phi_{mirror \rightarrow BS}^S + \phi_{BS \rightarrow bright}$. If the reference and sample path lengths are the same, the phases $\phi_{BS \rightarrow mirror}^R + \phi_{mirror \rightarrow BS}^R$ and $\phi_{BS \rightarrow mirror}^S + \phi_{mirror \rightarrow BS}^S$ are the same. So the two arrows of length 0.5 lie in the same direction. We hook them head to tail to

get the final arrow, which is length 1 and take its squared length to determine the probability, which is also 1. So all of the photons go to the bright detector (and hence its name). As the length along the sample path is changed, the sample path arrow will rotate, while the reference path arrow remains fixed. As the sample path rotates all the way from an additional angle of 0 to an additional angle of 2π , we will see an interference pattern for the probability that behaves as $\cos^2(\Delta\phi/2)$, with

$$\Delta\phi = \phi_{BS \rightarrow \text{mirror}}^S + \phi_{\text{mirror} \rightarrow BS}^S - (\phi_{BS \rightarrow \text{mirror}}^R + \phi_{\text{mirror} \rightarrow BS}^R). \quad (11)$$

The dark detector is similar. There are two alternative ways from source to detector. The phase for the reference path is $\phi_{\text{source} \rightarrow BS} + \phi_{BS \rightarrow \text{mirror}}^R + \pi + \phi_{\text{mirror} \rightarrow BS}^R + \phi_{BS \rightarrow \text{dark}}$, because it transmits at the second beam splitter and for the sample path it is $\phi_{\text{source} \rightarrow BS} + \pi + \phi_{BS \rightarrow \text{mirror}}^S + \pi + \phi_{\text{mirror} \rightarrow BS}^S + \phi_{BS \rightarrow \text{dark}}$ because the last beam splitter reflects from glass, not air. One can now see that the phase difference between the two paths is $\Delta\phi + \pi$. With no phase difference ($\Delta\phi = 0$), the arrows are equal length but point opposite and cancel, while for other phase differences, the probability to be detected in the dark detector is $\sin^2(\Delta\phi/2)$. Of course the total probability to be measured in both detectors is always equal to 1, regardless of the phase difference on the two paths.

You can see that the abstract states $|\psi_{\text{bright}}\rangle$ and $|\psi_{\text{dark}}\rangle$ correspond to the abstract set of rules we would use to evaluate the two arrows for the superposition of the two alternate ways the photon can go from source to detector, while the wavefunctions $\langle x_{\text{bright}} | \psi_{\text{bright}} \rangle$ and $\langle x_{\text{dark}} | \psi_{\text{dark}} \rangle$ are the concrete probability amplitudes we find after executing those rules at each detector using the concrete phase difference $\Delta\phi$ for the specific case corresponding to the current run of the experiment. The state is completely abstract and has no numbers associated with it, just rules, while the wavefunction is composed of two complex numbers here because we have a discrete-valued wavefunction with only two amplitudes, corresponding to each detector. We apply the rules to the specified cases and thereby determine the corresponding numbers.

These are two very different concepts and they should never be conflated. Yet many quantum texts, and even research articles, will improperly refer to a “ket” as a “wavefunction.” This will only lead to confusion as the two concepts are really quite different, as our example shows.

If we substitute a polarizing beam splitter for the first beam splitter, we have an extra degree of freedom corresponding to the polarization that we need to take care of. This allows us to then formulate how which-way information and entanglement are the same idea. In this realm, our language is actually pretty good, as the usage of entanglement or which-way information is applied

consistently in most textbooks, so we will not go through this example in detail.

Finally, there is one more use of language which can be improved in quantum mechanics that we will discuss here. In classical mechanics an object has definite position, momentum, and any other quality that describes its behavior, while in quantum mechanics, unless the characteristics come from commuting operators, the quantum particle can have only one property that is definite. When we discuss these “states of definite values,” we usually call them eigenstates and the definite value is called the eigenvalue. But this is more confusing than using a clearer language such as a state of definite position, or a state of definite momentum, or a state of definite energy. Using this alternate language reinforces the concept that we are choosing to find states with one and only one definite property, whereas the eigenvalue-eigenvector language tends to sound like a mathematical property that is disconnected from the underlying physics. In fact, we believe that just changing language in this one way will ease much of the confusion that students have with the eigenvalue problem in the context of quantum mechanics.

In summary, language in quantum mechanics is important. While some language is the choice of the instructor (such as states of definite values versus eigenstates, or which-way information versus entangled states) others are more serious and can cause deep-seeded confusion (such as conflating abstract state vectors with wavefunctions or operators with their matrix representations). The words really do matter. We owe it to our students to ensure we use words properly. Working with a representation-independent formalism makes this easier to do because objects discussed in quantum mechanics have well-defined differences in how they are used and of what they are.

HOW TO PERFORM CALCULATIONS IN A REPRESENTATION-INDEPENDENT FORMALISM

Working in a representation-independent formalism may sound intimidating. You may not know how to proceed. We will discuss a few examples here to make clear how it works, and we hope you will see that the representation-independent approach is often simpler and clearer than alternative methods.

We start with the time-independent Schrödinger equation used to find states of definite energy. In a representation-independent formalism, the equation looks like

$$\left(\frac{\hat{p}^2}{2m} + V(\hat{x}) \right) |\psi\rangle = E |\psi\rangle, \quad (12)$$

where it is clear we are just finding a state of definite total energy by adding the kinetic and potential energies

together. Contrast this to the position-space equation, which is

$$\left(-\frac{\hbar^2}{2m} \frac{d^2}{dx^2} + V(x)\right) \psi(x) = E\psi(x), \quad (13)$$

where we have the more intimidating derivatives appearing, which have no simple conceptual connection to kinetic energy.

Which version is easier to explain to a student? Which version needs to be derived? Which can just be stated?

In fact, the representation-independent way to solve quantum problems is actually well-known and is commonly used for two problems in conventional quantum mechanics instruction. These solutions are often viewed as the most elegant solutions for states with definite values of the operators. The two problems are the operator approach for the simple harmonic oscillator and the operator approach for finding states of definite total and z-component of angular momentum. Most texts describe how to find these states just by developing the arguments from how the operators act on the states. The eigenvalue problems are solved without every needing to find the wave functions. If needed, they are found later in a more complicated calculation.

It turns out all of the problems that we solve in quantum mechanics can be solved this way using the Schrödinger factorization method. This includes problems like particle in a box, harmonic oscillator, isotropic oscillator, hydrogen, Morse potential, and more. So, one can solve states of definite energy using just abstract quantum states and operators and never needing to find the wave functions. It is here that current textbooks fail to work in the representation-independent formalism and instead go to the position representation and differential equations.

Even within the representation-independent formalism, the choice of how objects are constructed matters. The conventional way to write a time ordered product, given by

$$\hat{U}(t) = T_t \exp\left(-\frac{i}{\hbar} \int_0^t dt' \hat{H}(t')\right), \quad (14)$$

is rather opaque. The time-ordered nature is hard to see, being buried in the symbol T_t , and is often confused by students, who may choose to evaluate the integral directly and then exponentiate it instead of working with it in the proper time-ordered way. An alternative form that is more convenient to express the time-evolution operator is the Trotter product formula

$$\hat{U}(t) = e^{-\frac{i}{\hbar} H(N\Delta t)\Delta t} e^{-\frac{i}{\hbar} H((N-1)\Delta t)\Delta t} \dots e^{-\frac{i}{\hbar} H(2\Delta t)\Delta t} e^{-\frac{i}{\hbar} H(1\Delta t)\Delta t}, \quad (15)$$

because integrating an abstract operator is a complicated mathematical procedure. Here, the time ordering is explicit. The fact that one needs to take a limit as $\Delta t \rightarrow 0$ with $N\Delta t = t$ held constant is not so difficult to discuss and carry out with students. Any problem that can be solved in the conventional way can also be solved this

way and one does not need to develop any integral equations to formulate these solutions, making the math load lighter as well.

RELATIONSHIP OF THESE IDEAS TO PHYSICS EDUCATION RESEARCH

While there is not a huge amount of physics education research (PER) performed on teaching in the representation-independent formulation of quantum mechanics, there has been some work done on how students are exposed to different formalisms, how they view the different formalisms (and what they mean), and how they can learn to switch between the abstract and concrete formalisms. Due to space restrictions, we consider only two of these works in detail here. The first examines how students are exposed to different formulations of the eigenvalue problem by their instructor, in a quantum mechanics setting (Wawro & Serbin, 2025). In the lectures, students were exposed to the eigenvalue problem as an abstract statement of operators and vectors and as a concrete representation in terms of matrices and vectors in a specific basis. The lecturer was careful not to equate the two but spoke of them as equivalent approaches. The conclusion was that as the instructor changed how form was linked to function in the course of teaching the materials, students were exposed to how one works with these ideas at an intermediate level. The second work, focused on ensuring students understand the difference between abstract states and wavefunctions, with a tutorial (called a QuILT) prepared to enhance learning and sensemaking (Marshman et al., 2024). In this work, students who have had conventional instruction show many types of confusion about operators, states, and wavefunctions, and struggle to get all of this diverse material correct. By using the additional QuILT tutorial, students achieve much higher success.

Our opinion here is a different one. We do not disagree with this work, we instead feel it can be interpreted differently. Rather than the belief that understanding stems from the ability to smoothly transition from abstract objects to their concrete representations, we argue that early and consistent restriction to a representation-independent formalism using well-chosen language provides a clearer ontology for states and operators, enforces linguistic and symbolic precision, and reduces ambiguity introduced by representational shifts that often occur “on the fly” during instruction. In this view, representations are not parallel conceptualizations to be coordinated with their abstract counterparts but instead are coordinate-dependent re-expressions of abstract objects, introduced only when required for specific calculations. We believe that by restricting to a representation-independent formalism from the start and not introducing additional ideas such as matrices or wavefunctions until they are

needed, one can avoid the confusion students often have between these different formal tools. Indeed, from a pedagogical perspective, it seems to be a much lighter cognitive load to focus only on one formal strategy than multiple ones at the same time. We have not yet investigated this claim via PER, but, as long as one can develop materials so that students are comfortable with increasing levels of abstraction, it seems clear that reducing the formal elements should be easier for students on a number of different levels. We do note that this is what mathematicians have been doing for a long time, as described above with the differential geometry example.

We are actively engaged in research that examines student attitudes to conventional versus representation-independent approaches. Student teams of two were asked to examine some common quantum problems with bound states of the hydrogen atom. We asked them to consider the highest angular momentum states for a given n value ($l = n-1$). They were asked to compute the expectation value of $1/\hat{r}$ and of $\hat{p}_r$, using both the operator method and the wave functions in position space. The students always chose to work with the position-space expressions first, because they were more comfortable with the calculus-based expressions, but after solving the problems both ways, they acknowledged that the operator-based method was much simpler. This study is ongoing and we plan to complete it soon.

PER research documents well that students and instructors routinely conflate abstract objects with their representations. They also have shown that this can eventually lead to appropriate sensemaking, but it does require the instructor to present the materials properly each time, so that students eventually learn the correct way to do this. This current work argues instead that the issue is not just one of student difficulty, but rather a structural feature of how quantum mechanics is commonly taught, and this feature can be safely omitted to avoid unnecessary confusion, which may turn out to be even more beneficial to students.

SUMMARY

In this work, we have argued that there is a pedagogical benefit to working with a representation-independent formulation of quantum mechanics. When this is done, many aspects of the theory become easier to see, such as the time-independent Schrödinger equation. In general, representation-independent formulations provide a clearer view of the structure of the theory and how the different ideas fit together. Quantum mechanics is no exception to this. Many calculations can also be performed in a representation-independent fashion. This lightens the load for eigenvalue problems, where eigenvalues can be found without simultaneously determining wavefunctions. Usually this provides a

significant simplification to the analysis, as seen in the example for the operator method applied to the simple harmonic oscillator.

We have also stressed the importance of language—how one must carefully distinguish between an abstract state of the Hilbert space and the wavefunction (which represents the state in a particular basis). One has a similar issue with operators, which often are conflated with their matrix-representation counterparts. Finally, we think that using more precise language such as a state of definite energy provides more meaning than simply saying an energy eigenstate.

The important question to address is whether these ideas actually improve student learning. This requires careful research into student thinking about these ideas.

Finally, we need additional educational resources to be able to develop the theory along this representation-independent framework. A forthcoming book *Quantum mechanics done right: The shortest path from novice to researcher* (Freericks, 2026) develops quantum-mechanical ideas in this fashion. Hopefully, it will assist others in teaching their students with a more modern and scientifically precise approach. And even more hopefully, it will improve student learning and reduce their confusion in quantum mechanics.

Funding: This study was supported by the Air Force Office of Scientific Research under grant number FA9550-24-1-0292. It was also supported by the McDevitt bequest at Georgetown University.

Ethical statement: This work did not perform any research on human subjects and so no ethical approvals were required for it.

AI statement: Generative AI (Gemini-pro) was used to assist with replying to referee reports in the preparation of the final manuscript.

Declaration of interest: No conflict of interest is declared by the author.

Data sharing statement: No data was generated in this manuscript.

REFERENCES

- Brookes, D. T., & Ektina, E. (2007). Using conceptual metaphor and functional grammar to explore how language used in physics affects student learning. *Physical Review: Physics Education Research*, 3, Article 010105. <https://doi.org/10.1103/PhysRevSTPER.3.010105>
- Cohen-Tannoudji, C., Diu, B., & Laloë, F. (1977). *Quantum mechanics*. Wiley.
- Feynman, R. P. (1985). *QED: The strange theory of light and matter*. Princeton University Press.
- Freericks, J. K. (2026). *Quantum mechanics done right: The shortest path from novice to researcher*. Springer.
- Marshman, E., Maries, A., & Singh, C. (2024). Using multiple representations to improve student understanding of quantum states. *Physical Review: Physics Education Research*, 20, Article 020152. <https://doi.org/10.1103/PhysRevPhysEducRes.20.020152>

- Maudlin, T. (2019). *Philosophy of physics: Quantum mechanics*. Princeton University Press. <https://doi.org/10.1515/9780691190679>
- Sakurai, J. J., & Napolitano, J. (2020). *Modern quantum mechanics* (3rd ed.). Cambridge University Press. <https://doi.org/10.1017/9781108587280>
- Shankar, R. (1994). *Principles of quantum mechanics* (2nd ed.). Plenum Press. <https://doi.org/10.1007/978-1-4757-0576-8>
- Wawro, M., & Serbin, K. S. (2025). "What makes it eigenesque-ish?": A form-function analysis of the development of eigentheory concepts in a quantum mechanics course. *Educational Studies in Mathematics*, 119, 287-310. <https://doi.org/10.1007/s10649-025-10390-4>

<https://www.ejmste.com>