



MASTERARBEIT | MASTER'S THESIS

Titel | Title

Optimization algorithms in deep-learning based solutions of the
Schrödinger equation

verfasst von | submitted by

Lucia Marchionne

angestrebter akademischer Grad | in partial fulfilment of the requirements for the degree of
Master of Science (MSc)

Wien | Vienna, 2024

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 066 821

Studienrichtung lt. Studienblatt | Degree
programme as it appears on the student
record sheet:

Masterstudium Mathematik

Betreut von | Supervisor:

Univ.-Prof. Dr. Philipp Grohs

Abstract

The Schrödinger equation plays a central role in the study of quantum mechanics and the ability to solve it is of uttermost importance. However, its complexity makes it particularly difficult to find exact solutions. Among the methods to find approximate solutions is the variational method, which essentially consists in minimizing an energy function. In this thesis, two different approaches will be compared to solve this optimization problem, the Adam minimizer and the Supervised Wave function Optimization (SWO) algorithm. The focus of this thesis centers on their efficacy in addressing the minimization problem within the context of the one-dimensional quantum harmonic oscillator. The SWO algorithm emerges as a better choice. It not only requires fewer epochs for convergence but also exhibits enhanced computational speed.

Kurzfassung

Die Schrödingergleichung spielt eine zentrale Rolle in der Quantenmechanik, und die Fähigkeit, sie zu lösen, ist von größter Bedeutung. Aufgrund ihrer Komplexität ist es jedoch besonders schwierig, exakte Lösungen zu finden. Zu den Methoden, um Näherungslösungen zu finden, gehört die Variationsmethode, die im Wesentlichen darin besteht, eine Energiefunktion zu minimieren. In dieser Arbeit werden zwei verschiedene Ansätze zur Lösung dieses Optimierungsproblems verglichen, der Adam-Algorithmus und der Supervised Wave function Optimization (SWO)-Algorithmus. Der Schwerpunkt dieser Arbeit liegt auf ihrer Wirksamkeit bei der Lösung des Minimierungsproblems im Zusammenhang mit dem eindimensionalen quantenmechanischen harmonischen Oszillator. Der SWO-Algorithmus erweist sich als die bessere Wahl. Er benötigt nicht nur weniger Schritte zur Konvergenz, sondern weist auch eine höhere Rechengeschwindigkeit auf.

Contents

1	Introduction	1
1.1	Background	1
1.2	Objective of the study	2
1.3	Structure of the thesis	2
1.4	Bra-ket notation	3
2	Quantum Mechanics and its mathematical background	5
2.1	Operators on Hilbert spaces	5
2.1.1	Self-adjoint operators	5
2.1.2	Spectrum of an operator	6
2.2	The Schrödinger Equation	9
2.2.1	The wave function and quantum observables	9
2.2.2	The Hamiltonian	11
2.2.3	The Schrödinger Equation	13
3	A brief introduction to Neural Networks	17
3.1	Shallow neural networks and the Universal Approximation Theorem . . .	17
3.1.1	Proof of the Universal Approximation Theorem	19
3.2	Training a Neural Network	22
3.2.1	Stochastic Gradient Descent	23
4	Approximating solutions to the Schrödinger Equation	27
4.1	The variational principle	27
4.2	The variational method	30
4.2.1	Variational Monte Carlo	30
4.2.2	The Metropolis-Hastings algorithm	31
5	Optimizers	33
5.1	Adam	33
5.2	SWO	33
5.2.1	Influence of the parameters	36
6	The Quantum Harmonic Oscillator	39
6.1	The problem	39
6.2	Analytical solutions	40
6.3	Approximate solutions	41
6.3.1	Adam vs. SWO	41

Contents

7 Concluding remarks	45
Literaturverzeichnis	46

1 Introduction

1.1 Background

Systems of particles on atomic and subatomic scale do not obey the laws of classical mechanics. Quantum mechanics was developed at the start of the 20th century to explain their peculiar behaviour. At the core of this new theory lies the Schrödinger equation

$$H\psi = E\psi, \quad (1.1)$$

named after Erwin Schrödinger who published it in 1926 [37]. It is an eigenvalue equation where H is a *Hamiltonian* operator on a Hilbert space $\mathcal{H}$, ψ is a *wave function*, a vector in $\mathcal{H}$, and E is a real number. The Schrödinger equation describes how the state of a quantum system, a system that obeys the laws of quantum mechanics, changes over time. Solving it also yields the system's possible energy levels, which is useful knowledge in many fields. These range from the prediction of molecular properties, useful in drug discovery [13], to quantum engineering and quantum computing [27].

However, the Schrödinger equation is not easy to solve. In fact, it can only be solved exactly for a handful of problems, such as the free particle, the hydrogen atom and the quantum harmonic oscillator [36]. In addition, many methods have been developed to find approximate solutions of more complex systems. For instance, perturbation theory [36], [41] is based on the assumption that the Hamiltonian of a certain system can be written as a sum of two. The first summand is a Hamiltonian H whose Schrödinger equation is easier to solve, and the second summand is small enough to be considered a perturbation of H . Another important technique in this context is the WKB method [36], [11]. The Hamiltonian is the sum of kinetic and potential energy functions, and the WKB method relies on the assumption that the potential energy function is nearly constant. Assuming further that ψ has exponential form, this leads to major simplifications of the Schrödinger equation.

In this thesis, a variational method to approximate solutions of the Schrödinger equation will be used [36], [39]. The variational method is based on the *variational principle*, which states that the lowest eigenvalue of H is always less than or equal the *local energy* at any trial wave function $\psi \in \mathcal{H}$

$$\langle \psi | H | \psi \rangle. \quad (1.2)$$

The local energy is a real quantity that can be numerically estimated quite easily and the variational method consists in minimizing it. Once the local energy is minimized, the resulting value of 1.2 and ψ will approximate E and ψ in 1.1, respectively. In general, the tentative solution ψ is represented by a function that depends on many parameters.

1 Introduction

The task of optimizing a function of many parameters is one that can be efficiently accomplished with the use of deep neural networks. Neural networks are very powerful mathematical objects, in the sense that they can be used to approximate a very large class of functions. This is done by efficiently tuning the (many) parameters on which the neural network depends. This property is leveraged in this context by using neural networks to represent the trial wave functions ψ [9], [28], [16], [15].

1.2 Objective of the study

The scope of this thesis is to compare the performance of two neural network training algorithms when implemented to solve the one-dimensional quantum harmonic oscillator problem. More specifically, the Adam [19] algorithm and the Supervised Wavefunction Optimization algorithm [21] will be implemented as part of a variational method and compared.

The implementation of the algorithms and of the solution of the problem was done in Python 3.8 using the haiku¹ and jax² libraries, which allowed for fast computations. Furthermore, the first half of this thesis will be dedicated to the presentation of introductory notions in quantum mechanics and deep learning.

1.3 Structure of the thesis

In the next two chapters, a brief overview on the topics of quantum mechanics and deep learning will be given. Only the concepts relevant to this thesis will be mentioned and explained.

In particular, the first part of Chapter 2 includes notions from functional analysis on operators on Hilbert spaces and their spectra. These are important to the understanding of the fundamental concepts in quantum mechanics that will be explained in the second part of Chapter 2, building up to the Schrödinger equation.

Chapter 3 includes the definition of a shallow neural network as well as the statement and a proof of the Universal Approximation Theorem. Furthermore, the neural network training algorithm of Stochastic Gradient Descent will be covered as a basis for the Adam algorithm that will be described and used in Chapters 5 and 6.

The focus of Chapter 4 revolves around variational methods to find approximate solutions to the Schrödinger equation. More specifically, the variational principle will be proved before stating the variational Monte Carlo algorithm that will be used in Chapter 6 to analytically estimate the solution of the ground state Schrödinger equation for the one-dimensional quantum harmonic oscillator. On the one hand, this will be done using the Adam algorithm and on the other hand, using the Supervised Wave function Optimisation algorithm, both of which are described in Chapter 5. In Chapter 6, after

¹<https://dm-haiku.readthedocs.io>

²<https://jax.readthedocs.io/en/latest/>

defining the problem of the quantum harmonic oscillator, the results of the computations with the two algorithms will be compared.

1.4 Bra-ket notation

Bra-ket notation was introduced by Dirac in 1939 [7] and is the standard notation in quantum mechanics literature. It will be used throughout this thesis as well. Assume that $\mathcal{H}$ is a complex Hilbert space with inner product $\langle \cdot, \cdot \rangle$.

Definition 1.1. A *ket* has the form $|v\rangle$ and denotes a vector in $v \in \mathcal{H}$.

Definition 1.2. A *bra* is a continuous linear functional $\mathcal{H} \rightarrow \mathbb{C}$, i.e., an element of $\mathcal{H}'$. For any $v \in \mathcal{H}$, the expression $\langle v|$ denotes the bra defined as

$$\begin{aligned} \langle v| : \mathcal{H} &\rightarrow \mathbb{C} \\ |w\rangle &\mapsto \langle v|w\rangle := \langle v, w \rangle. \end{aligned}$$

If H is an operator on $\mathcal{H}$, i.e. $H : \mathcal{H} \rightarrow \mathcal{H}$, then one can define the linear functional

$$\begin{aligned} \langle v|H : \mathcal{H} &\rightarrow \mathbb{C} \\ |w\rangle &\mapsto \langle v|H|w\rangle := \langle v|Hw\rangle. \end{aligned}$$

Using the notation $\langle v|H|w\rangle$ highlights the two-fold interpretation of this complex number: it is the bra $\langle v| \in \mathcal{H}'$ applied to the ket $H|w\rangle \in \mathcal{H}$ as well as the bra $\langle v|H \in \mathcal{H}'$ applied to the ket $|w\rangle \in \mathcal{H}$.

2 Quantum Mechanics and its mathematical background

This chapter will cover the key notions in quantum mechanics that are needed to understand the Schrödinger equation and the necessary mathematical background. This will include statements from functional analysis, in particular from the theory of operators. In shaping the content of this chapter, I have referenced [11], [36], [10], [22], [41], [31], [40].

2.1 Operators on Hilbert spaces

Operators are a central concept in the study of quantum mechanics, because in a certain sense, they play the role of functions in classical mechanics [42], [8].

In Section 2.1.1, the definitions of *bounded*, *unbounded* and *self-adjoint operators* will be given. Next, in Section 2.1.2, the *spectrum* will be discussed. These notions will become relevant in Section 2.2.1.

2.1.1 Self-adjoint operators

Let $\mathcal{H}$ be a Hilbert space with inner product $\langle \cdot, \cdot \rangle$ and corresponding norm $\|\cdot\|$. We start with the definition of an operator.

Definition 2.1. Let $\text{dom}(T) \subseteq \mathcal{H}$ be a subspace and consider a linear map

$$T : \text{dom}(T) \rightarrow \mathcal{H}.$$

Then T is called an *operator on $\mathcal{H}$ with domain $\text{dom}(T)$* . If $\text{dom}(T)$ is dense in $\mathcal{H}$, then T is called *densely defined*. If $\text{dom}(T) = \mathcal{H}$, then T is called an *operator on $\mathcal{H}$* .

The next definition is useful to understand the concept of self-adjointness, which is crucial in the study of operators in quantum mechanics.

Definition 2.2. An operator $T : \text{dom}(T) \rightarrow \mathcal{H}$ is called *symmetric*, if

$$\langle T\psi, \phi \rangle = \langle \psi, T\phi \rangle \text{ for all } \psi, \phi \in \text{dom}(T).$$

Definition 2.3. An operator $T : \text{dom}(T) \rightarrow \mathcal{H}$ is said to be *bounded* if there exists a constant $M > 0$ such that

$$\|T\psi\| \leq M \|\psi\| \text{ for all } \psi \in \text{dom}(T).$$

2 Quantum Mechanics and its mathematical background

Otherwise, T is called *unbounded*. The space of bounded linear operators on $\mathcal{H}$ is denoted by $L(\mathcal{H})$.

Self-adjointness can be defined in a fairly straightforward way in the case of bounded operators.

Definition 2.4. Let $T \in L(\mathcal{H})$. Then its *adjoint* $T^* \in L(\mathcal{H})$ can be characterized by the condition

$$\langle T\psi, \phi \rangle = \langle \psi, T^*\phi \rangle \text{ for all } \psi, \phi \in \mathcal{H}.$$

If $T = T^*$, the operator T is called *self-adjoint*. Note that in this case, T is self-adjoint if and only if it is symmetric.

In the case of unbounded operators, defining the adjoint is not as straightforward, because of domain-related issues.

Definition 2.5. Let $T : \text{dom}(T) \rightarrow \mathcal{H}$ be a densely defined operator on H , not necessarily bounded. We will define the *adjoint* $T^* : \text{dom}(T^*) \rightarrow \mathcal{H}$ of T . For a fixed $\phi \in \mathcal{H}$, define the map $f_\phi : \text{dom}(T) \rightarrow \mathbb{C}$ by $f_\phi(\psi) = \langle T\psi, \phi \rangle$. Furthermore, define the subspace

$$\text{dom}(T^*) := \{\phi \in \mathcal{H} \text{ s.t. } f_\phi \text{ is continuous}\}.$$

By the Hahn-Banach theorem (see also Theorem 3.4), if $\phi \in \text{dom}(T^*)$, the function f_ϕ has a unique continuous extension to a linear functional $F : \mathcal{H} \rightarrow \mathbb{C}$. Moreover, by the Riesz Representation theorem (see also Theorem 3.5), the functional F has a unique representation in the form $\psi \mapsto \langle \psi, z \rangle$ with $z \in \mathcal{H}$. Define $T^*\phi := z$. If $T = T^*$, the operator T is called *self-adjoint*. Note that by construction, it holds

$$\langle T\psi, \phi \rangle = \langle \psi, T^*\phi \rangle \text{ for all } \psi \in \text{dom}(T), \phi \in \text{dom}(T^*).$$

2.1.2 Spectrum of an operator

Definition 2.6. The *spectrum* of an operator T on a Hilbert space is defined by

$$\sigma(T) := \{\lambda \in \mathbb{C} \mid T - \lambda \text{ has no bounded inverse}\}$$

Definition 2.7. A vector $\psi \neq 0$ that satisfies the equation $T\psi = \lambda\psi$ for a (complex) number λ is said to be an *eigenvector* of T with corresponding *eigenvalue* λ . In this case, $T - \lambda$ fails to be injective and therefore cannot have bounded inverse. Therefore any such λ is in the spectrum of T .

Not all operators on infinite dimensional vector spaces have eigenvalues.

Example 2.8. For $t \in [0, 1]$, define the multiplication operator

$$M : L^2[0, 1] \rightarrow L^2[0, 1], \quad (Mf)(t) = t \cdot f(t).$$

2.1 Operators on Hilbert spaces

The operator M is clearly bounded and also self-adjoint, because

$$\langle (Mf)(t), g(t) \rangle = \langle t \cdot f(t), g(t) \rangle = \langle f(t), t \cdot g(t) \rangle = \langle f(t), (Mg)(t) \rangle \quad \text{for all } t \in [0, 1].$$

Furthermore, the operator $M - \lambda$ possesses the continuous inverse

$$(M - \lambda)^{-1} : L^2[0, 1] \rightarrow L^2[0, 1], \quad ((M - \lambda)^{-1}f)(t) = \frac{f(t)}{t - \lambda}$$

if and only if $\lambda \in \mathbb{C} \setminus [0, 1]$. It will be shown later (Theorem 2.11) that the spectrum of a self-adjoint operator is a subset of the real numbers. Therefore, the spectrum of M is the interval $[0, 1]$. However, M does not have any eigenvalues. Indeed, assume the contrary. Then

$$Mf = \lambda f \Leftrightarrow t \cdot f(t) = \lambda \cdot f(t) \text{ for all } t \in [0, 1],$$

a contradiction, since $f \neq 0$.

This next definition and lemma are based on [32], Chapter 25 and will be used to prove the important Theorem 2.11 as well as later in Section 4.1.

Definition 2.9. An element $\lambda \in \mathbb{C}$ is called an *approximate eigenvalue* of an operator T if there exist a sequence $(\psi_n)_{n \geq 0} \subseteq \text{dom}(T)$ such that

$$\|(T - \lambda I)\psi_n\| \rightarrow 0 \text{ as } n \rightarrow \infty$$

Lemma 2.10. Let T be a self-adjoint operator and let $\lambda \in \sigma(T)$. Then λ is an approximate eigenvalue.

Proof. Let $\lambda \in \sigma(T)$. Then $(T - \lambda I)^{-1}$ is not bounded. Therefore, for all $n \geq 0$ there exists an $x_n \in \text{dom}((T - \lambda I)^{-1})$ such that

$$\|(T - \lambda I)^{-1}x_n\| > n.$$

Assume without loss of generality that $\|x_n\| = 1$, and define $\xi_n := (T - \lambda I)^{-1}x_n$. Note that $\xi_n \in \text{dom}(T)$ for all n . Then it holds

$$\left\| (T - \lambda I) \frac{\xi_n}{\|\xi_n\|} \right\| = \frac{\|x_n\|}{\|\xi_n\|} \rightarrow 0 \text{ as } n \rightarrow \infty.$$

□

Theorem 2.11. Let T be a self-adjoint operator. Then $\sigma(T) \subseteq \mathbb{R}$.

Proof. Let $\lambda \in \sigma(T)$. By the lemma above, we can find a (normalized) sequence $(\psi_n)_{n \geq 0} \subseteq \text{dom}(T)$ such that $\|(T - \lambda I)\psi_n\| \rightarrow 0$. By the Cauchy-Schwarz inequality,

$$|\langle (T - \lambda I)\psi_n, \psi_n \rangle| \leq \|(T - \lambda I)\psi_n\| \|\psi_n\| \rightarrow 0.$$

2 Quantum Mechanics and its mathematical background

At the same time, since $\|\psi_n\| = 1$, it holds for all $n \geq 0$

$$\langle (T - \lambda I)\psi_n, \psi_n \rangle = \langle T\psi_n, \psi_n \rangle - \lambda.$$

Combining these two facts, we obtain

$$\langle T\psi_n, \psi_n \rangle \rightarrow \lambda \quad \Rightarrow \quad \overline{\langle T\psi_n, \psi_n \rangle} \rightarrow \bar{\lambda}.$$

On the other hand, by conjugate symmetry of the inner product and self-adjointness of T , it holds

$$\overline{\langle T\psi_n, \psi_n \rangle} = \langle \psi_n, T\psi_n \rangle = \langle T\psi_n, \psi_n \rangle,$$

and therefore, by uniqueness of the limit, $\lambda = \bar{\lambda}$, which implies $\lambda \in \mathbb{R}$. $\square$

Remark. As a matter of fact, the opposite implication also holds: an operator is self-adjoint if and only if its spectrum is real.

We will now turn to the statement of the *Spectral Theorem*, whose proof can be found in [11] and [43]. The Spectral Theorem will be used in Section 4.2 to justify the variational method to approximate solutions to the Schrödinger equation. Before stating it, it is necessary to define *spectral measures*.

Definition 2.12. Let $\mathcal{B}(\mathbb{R})$ be the Borel σ -algebra on $\mathbb{R}$. Define

$$\begin{aligned} E : \mathcal{B}(\mathbb{R}) &\rightarrow L(\mathcal{H}) \\ A &\mapsto E_A. \end{aligned}$$

The map E is called a *spectral measure*, if

1. E_A is an orthogonal projection for every $A \in \mathcal{B}(\mathbb{R})$,
2. $E_\emptyset = 0$ and $E_\mathbb{R} = I$, where I is the identity operator,
3. if $\psi \in \mathcal{H}$ and $A_1, A_2, \dots \in \mathcal{B}(\mathbb{R})$ are pairwise disjoint with $A := \bigcup_i A_i$, it holds

$$\sum_i E_{A_i} \psi = E_A \psi.$$

Integration with respect to a spectral measure can be defined analogously to integration with respect to a non-spectral measure. Start by letting $h := \sum_{i=1}^n c_i \chi_{A_i}$ be a step function on $\mathbb{R}$, where χ_{A_i} is the indicator function on $A_i \subseteq \mathbb{R}$. Then define

$$\int h \, dE := \sum_{i=1}^n c_i E_{A_i}.$$

By density of step functions, this can be extended to measurable functions by passing onto the limit. Details of this construction can be found in [43]. In a similar manner,

for $\psi, \phi \in \mathcal{H}$, integration with respect to the measure

$$\begin{aligned} E : \mathcal{B}(\mathbb{R}) &\rightarrow L(\mathcal{H}) \\ A &\mapsto \langle E_A \psi, \phi \rangle \end{aligned}$$

can also be defined. Again, details can be found in [43].

Theorem 2.13. *[Spectral Theorem] Let $T : \text{dom}(T) \rightarrow \mathcal{H}$ be a self-adjoint operator. Then there is a unique spectral measure E such that*

$$\langle T\psi, \phi \rangle = \int_{\sigma(T)} t \, d\langle E\psi, \phi \rangle \text{ for all } \psi \in \text{dom}(T), \phi \in \mathcal{H}. \quad (2.1)$$

Furthermore, if $h : \mathbb{R} \rightarrow \mathbb{R}$ is measurable and $D_h := \{\psi \in \mathcal{H} \text{ s.t. } \int |h(\lambda)|^2 d\langle E\psi, \psi \rangle < \infty\}$, then the operator $h(T) : D_h \rightarrow \mathcal{H}$ defined by

$$\langle h(T)\psi, \phi \rangle = \int_{\sigma(T)} h(t) \, d\langle E\psi, \phi \rangle \quad (2.2)$$

is self-adjoint.

2.2 The Schrödinger Equation

The Schrödinger equation is a fundamental equation in quantum mechanics, central to the understanding of the behaviour of particles at quantum level, such as electrons in atoms and molecules. Before stating it, some key concepts in quantum mechanics, including wave functions and Hamiltonian operators, will be introduced. Some references on these topics include [11], [36], [39], [29].

2.2.1 The wave function and quantum observables

In classical mechanics, the state of a system is completely determined by its position and momentum. That is, each *observable*, meaning each quantity or property of the system that can be measured or detected directly or indirectly, can be expressed as a function of these two quantities. However, it was shown by Heisenberg in 1927 that position and momentum are canonically conjugate variables, meaning that their values cannot be simultaneously predicted to arbitrary accuracy [14]. Quantum mechanics was developed to address the limitations of classical physics in explaining this and other types of behaviour on atomic and subatomic scales.

Consider a system of a single particle in $\mathbb{R}^d$ and a measurable set $K \subseteq \mathbb{R}^d$. The fundamental new idea is that the state of such a system is completely determined by a *wave function*

$$\psi : \mathbb{R}^d \rightarrow \mathbb{C}.$$

2 Quantum Mechanics and its mathematical background

In this context, the integral

$$\int_K |\psi(x)|^2 dx$$

is interpreted as the probability of the particle being found in K , provided that ψ is normalized, i.e.

$$\int_{\mathbb{R}^d} |\psi(x)|^2 dx = 1,$$

because the particle must be found somewhere. In other words, $\psi \in L^2(\mathbb{R}^d)$ and $|\psi|^2$ is the probability density function for the position of the particle.

In more general terms, $L^2(\mathbb{R}^d)$ could be replaced by an appropriate Hilbert space $\mathcal{H}$. It is important to note that the wave function itself does not describe an observable. Instead, observables in quantum mechanics are represented by operators on $\mathcal{H}$.

Definition 2.14. An *observable* is a self-adjoint operator on a quantum Hilbert space.

In the Dirac-von Neumann formulation of quantum mechanics [8], [42], it is assumed that to each real-valued function f describing an observable on the classical phase space there is an associated self-adjoint, but not necessarily bounded, operator $\hat{f}$ on the quantum Hilbert space.

Theorem 2.11 provides the reason why observables are required to be self-adjoint. Indeed, the spectrum of an operator T represents the collection of potential results attainable when measuring the physical quantity related to T . For instance, the spectrum of the Hamiltonian operator for a certain quantum system is the set of potential outcomes attainable when measuring the system's total energy. In light of this interpretation, it is desirable for the spectrum of T to be real, which motivates the definition.

Definition 2.15. Let T be an observable on $\mathcal{H}$ and $\psi \in \mathcal{H}$ a state. Then the *expectation value of T at state ψ* is defined as

$$\langle T \rangle_\psi := \langle \psi | T | \psi \rangle$$

In case of normalized ψ , the expectation value of an observable T at state ψ is the probabilistic expected value of the outcome when measuring the physical quantity related to T at state ψ . For example, let H be the Hamiltonian operator for a certain quantum system. Then $\langle H \rangle_\psi$ is the expected value of the system's total energy when at state ψ .

Example 2.16. Consider a single particle in $\mathbb{R}^d$ and its wave function $\psi \in L^2(\mathbb{R}^d)$. Let $K \subseteq \mathbb{R}^d$ be measurable.

- The bounded operator $A : L^2(\mathbb{R}^d) \rightarrow L^2(\mathbb{R}^d)$ given by $\psi \mapsto \chi_K \psi$ is an observable. Its expectation value

$$\langle A \rangle_\psi = \int_{\mathbb{R}^d} \chi_K |\psi|^2 = \int_K |\psi|^2$$

represents the probability of the particle being found in K .

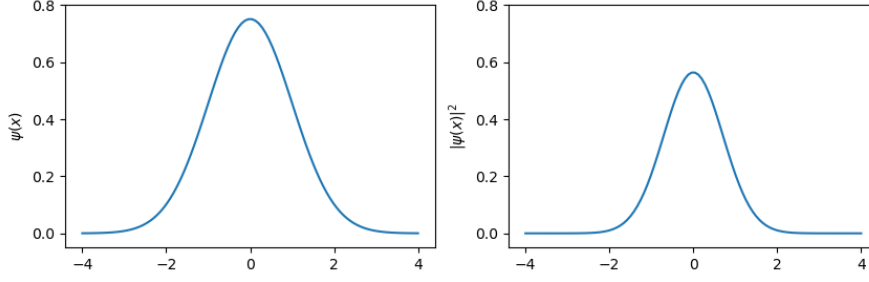


Figure 2.1: An example wave function $\psi(x) = \pi^{-1/4} \cdot e^{-x^2/2}$ (L) and the corresponding probability density function $|\psi|^2$ (R).

- The unbounded operator $X_j : \text{dom}(X_j) \rightarrow L^2(\mathbb{R}^d)$ given by $\psi \mapsto x_j \psi$ is an observable representing spatial (coordinate) position.

2.2.2 The Hamiltonian

Consider a system of N particles moving in a force field within $\mathbb{R}^n$. For $j \in \{1, \dots, N\}$, let $\mathbf{x}_j \in \mathbb{R}^n$ and $\mathbf{p}_j \in \mathbb{R}^n$ represent particle j 's position and momentum, respectively. Furthermore, let $V(\mathbf{x})$ be a potential energy function describing the force field.

In the Hamiltonian approach to classical mechanics, the energy of a system can be thought of as a real-valued function (the Hamiltonian) of position and momentum. Denote $\mathbf{x} := (\mathbf{x}_1, \dots, \mathbf{x}_N) \in \mathbb{R}^{n \times N}$ and $\mathbf{p} := (\mathbf{p}_1, \dots, \mathbf{p}_N) \in \mathbb{R}^{n \times N}$. The classical Hamiltonian is the sum of kinetic and potential energy, and therefore has the form

$$H(\mathbf{x}, \mathbf{p}) := \sum_{j=1}^N \frac{1}{2m_j} |\mathbf{p}_j|^2 + V(\mathbf{x}),$$

where $m_j \in \mathbb{R}$ denotes the mass of particle j .

As discussed in the preceeding section, given a quantum system, there exists a Hamiltonian operator

$$\hat{H} : \text{dom}(\hat{H}) \rightarrow \mathcal{H},$$

associated to the energy of the system, corresponding to the classical Hamiltonian. Here, $\text{dom}(\hat{H}) \subseteq \mathcal{H}$ merely denotes the domain of $\hat{H}$. By analogy with classical mechanics, the Hamiltonian operator is defined to have the form

$$\hat{H} := -\hbar^2 \sum_{j=1}^N \frac{1}{2m_j} \nabla_j^2 + \hat{V},$$

where $\hat{V}$ is some potential energy operator defined by $(\hat{V}\psi) = \hat{V}(x)\psi(x)$ for some $\hat{V} :$

2 Quantum Mechanics and its mathematical background

$\mathbb{R}^n \rightarrow \mathbb{R}$. Thus, for a wave function $\psi \in \text{dom}(\hat{H})$ and $x \in \mathbb{R}^n$, it holds

$$\hat{H}\psi(x) = -\hbar^2 \sum_j \frac{1}{2m_j} \frac{\partial^2 \psi}{\partial x_j^2} + \hat{V}(x)\psi(x).$$

Here,

$$-\hbar^2 \frac{1}{2m_j} \nabla_j^2 = \frac{\hat{\mathbf{p}}_j \cdot \hat{\mathbf{p}}_j}{2m_j}$$

is the kinetic energy operator of particle j , where $\hat{\mathbf{p}}_j := -i\hbar \nabla_j$ denotes the momentum operator.

The potential energy operator is self-adjoint if defined on

$$\text{dom}(\hat{V}) := \{\psi \in L^2(\mathbb{R}^n) \text{ s.t. } \hat{V}\psi \in L^2(\mathbb{R}^n)\},$$

provided that $\hat{V} : \mathbb{R}^n \rightarrow \mathbb{R}$ is a measurable function ([11], Section 9.8). The kinetic energy operator is also self-adjoint if defined on an appropriate domain ([11], Section 9.8). However, the Hamiltonian, as a sum of two self-adjoint operators, need not be self-adjoint ([11], Section 9.9).

The question of self-adjointness is a delicate one and goes beyond the scope of this thesis. Typically, given a (Hamiltonian) operator $\hat{H}$, its domain $\text{dom}(\hat{H})$ is chosen specifically to ensure it is self-adjoint. A detailed analysis of the technical aspects of this process may not necessarily be of interest: the operators that will be discussed in this thesis and the functions they act on are always physically meaningful and provide no relevant issues regarding their domains. A comprehensive study of this topic can be found in [11], [41], [31].

Example 2.17. The (quantum) harmonic oscillator is an important model in physics. It is one of the simplest quantum systems that can be solved analytically. Furthermore, it can be applied to describe a variety of physical phenomena in the field of solid state physics [20], as well as the vibrations of molecules ([2], Section 10.8). In the case of a single particle subject to a restoring force, the Hamiltonian has the form (using Hartree atomic units)

$$\hat{H} := -\frac{1}{2} \nabla^2 + \hat{V}, \text{ where } \hat{V}(x) = \frac{1}{2} x^2. \quad (2.3)$$

The Hartree atomic units are a system of units used in physics to simplify calculations. By definition, the constant $\hbar$ and the electron's mass m_e are formally set to 1. A more in-depth derivation of 2.3 can be found in Section 6.2.

The fairly simple problem of the one-dimensional quantum harmonic oscillator will be the focus of Chapter 6. Nevertheless, in the following example, the electronic Hamiltonian for bigger quantum systems - in particular, molecules - will be described.

Example 2.18. Consider a molecule as a quantum system of n electrons and M nuclei. Let the electrons be described by position vectors $r_i \in \mathbb{R}^3, i = 1, \dots, n$, and the nuclei

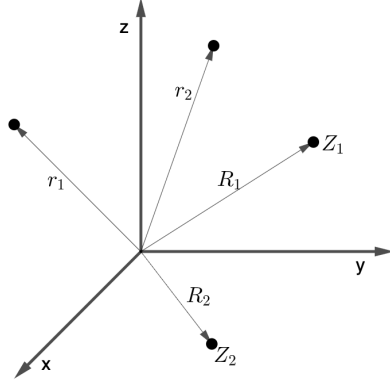


Figure 2.2: An example molecular coordinate system.

by position vectors $R_I \in \mathbb{R}^3, I = 1, \dots, M$ and atomic numbers $Z_I \in \mathbb{R}, I = 1, \dots, M$. Furthermore, denote by N_I the ratio of the mass of nucleus I to the one of an electron. Then, using Hartree atomic units and cartesian coordinates, the (molecular) Hamiltonian operator becomes

$$\hat{H} = -\sum_{j=1}^n \frac{1}{2} \nabla_j^2 - \sum_{I=1}^M \frac{1}{2N_I} \nabla_I^2 + \hat{V}, \quad (2.4)$$

$$\text{where } \hat{V} = \sum_{i>j} \frac{1}{|r_i - r_j|} + \sum_{I>J} \frac{Z_I Z_J}{|R_I - R_J|} - \sum_{i,I} \frac{Z_I}{|r_i - R_I|}. \quad (2.5)$$

The function $\hat{V}$ is the *Coulomb potential* and describes the interaction between the particles. In the expression above, the first and second terms represent the repulsion between electrons and between nuclei, respectively; the third term represents the attraction between electrons and nuclei.

This expression for $\hat{H}$ can be simplified by using the Born-Oppenheimer Approximation [5]. This approximation is based on the observation that electrons are much lighter than nuclei and therefore move much faster. Thus they can be considered as moving particles in a field of fixed nuclei. Under this assumption, the kinetic energy of the nuclei, the second term in 2.4, is zero, and the repulsion between the nuclei, the second term in 2.5, is constant.

2.2.3 The Schrödinger Equation

This section is based on [8] (Section 29) and [22]. Further details can also be found in [42] and [41].

Let $\hat{H}$ be a self-adjoint Hamiltonian operator on a Hilbert space $\mathcal{H}$ (see Section 2.2.2) describing the energy of a certain quantum system. The *time-dependent Schrödinger*

2 Quantum Mechanics and its mathematical background

equation has the form

$$i\hbar \frac{\partial}{\partial t} \psi(x, t) = \hat{H} \psi(x, t), \quad (2.6)$$

where i is the imaginary unit, $\hbar$ the reduced Planck's constant and the wave function $\psi \in \mathcal{H}$ depends on both position and time. It is a differential equation that was postulated by Erwin Schrödinger in 1925 and published in 1926 [37].

Consider a quantum system with constant energy E and, without loss of generality, consider a state that at time $t = 0$ is an eigenstate of $\hat{H}$ with eigenvalue E , i.e.

$$\hat{H} \psi(x, 0) = E \psi(x, 0). \quad (2.7)$$

In this case, the time-dependent Schrödinger equation 2.6 may be formally solved to obtain

$$\psi(x, t) = e^{-iEt/\hbar} \psi(x, 0). \quad (2.8)$$

Combining 2.6, 2.7, 2.8 it can be concluded that

$$\hat{H} \psi(x, t) = E \psi(x, t) \quad (2.9)$$

for all $t > 0$, i.e., that $\psi(x, t)$ is also an eigenstate of $\hat{H}$ with eigenvalue E for all $t > 0$. Equation 2.9 implies that when measuring the energy of such a state, the probability of a particular outcome does not depend on the time of measurement. A state with these properties is called a *stationary state*.

Therefore, for stationary states, Equation 2.9 reads

$$\hat{H} \psi(x) = E \psi(x). \quad (2.10)$$

Equation 2.10 is called *time-independent Schrödinger Equation*. It is a differential equation where $\hat{H}$ is given and $\psi \in \mathcal{H}$, $E \in \mathbb{R}$ are the unknowns. It is of particular interest to solve it for E_0 , the lowest eigenvalue, and its corresponding eigenfunction. This corresponds to the ground state, whose energy is the lowest possible energy the system can have. It is the most stable state and the system will eventually settle into it. Analysing the efficiency of different methods to determine it is the goal of this thesis.

Example 2.19. Consider a simple system of one electron, one proton and a force described by the Coulomb potential between them, such as the hydrogen atom. In this case, the Schrödinger equation can be solved analytically. The equation in Example 2.18 can be used to write down the (molecular) Hamiltonian in cartesian coordinates, where $n = M = 1$, $Z_1 = 1$ and N_1 is the proton-to-electron mass ratio μ ¹. To find the eigenvalues of the Hamiltonian analytically (as is done in, for example, [8]), it is useful to consider the nucleus as a fixed center of force and rewrite the Hamiltonian in radial form:

$$\hat{H} = -\frac{1}{2} \frac{d^2}{dr^2} - \frac{1}{r},$$

¹The exact value of μ can be found on <https://physics.nist.gov/cgi-bin/cuu/Value?mpsme>.

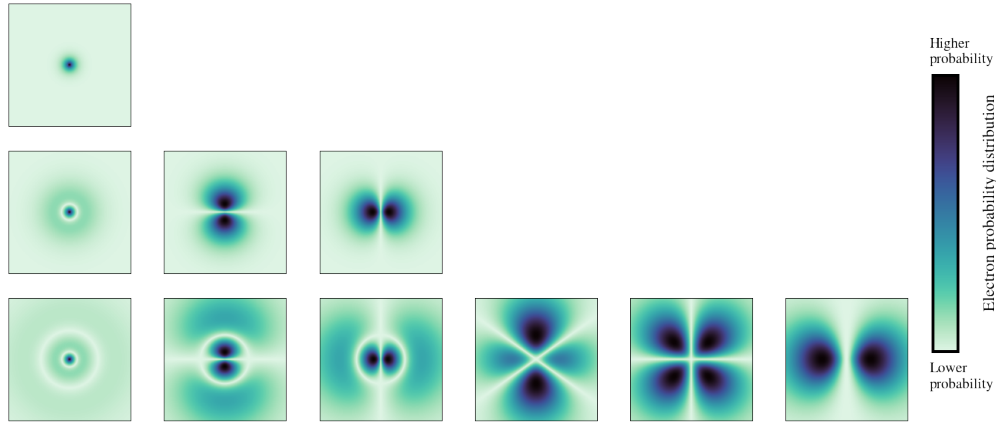


Figure 2.3: Plots of $|\psi_n|^2$ for $n = 1$ (first row), $n = 2$ (second row), $n = 3$ (third row). Here, the ψ_n are found as solutions of the Schrödinger equation for the hydrogen atom $\hat{H}\psi_n = E_n\psi_n$ (see Example 2.19). To each wave function corresponds a unique combination of values for the electron's energy, its total angular momentum, orientation of its angular momentum and spin. The plots were created partially using <https://github.com/sebastianmagns/hydrogen-wavefunctions>.

where r is the distance between the nucleus and the electron. This approximation is justified as the nucleus is much heavier than the electron.

The Schrödinger equation has a countable number of solutions in this case. These are the only allowed values of energy that the system can take, and are called *energy levels*. The general formula for the n -th, $n > 0$, energy level is

$$E_n = \frac{1}{2n^2},$$

in Hartree atomic units. In conventional units, $E_n = \frac{-13.6 \text{ eV}}{2n^2}$. For an energy level $n > 0$, E_n is the energy that the electron needs to jump to the next energy level. If $n > 1$, E_n is also the energy that it emits when jumping back to the previous energy level. The energy levels of the hydrogen atoms are *degenerate*, meaning that to each energy level there is one or more associated wave functions ψ_n (see Figure 2.3 for some examples).

3 A brief introduction to Neural Networks

Neural networks were first introduced by neurophysiologist Warren McCulloch and mathematician Walter Pitts in 1943 as a way to model the human nervous system [23]. After some initial promising developments [44], the limited available computational power at the time prevented further research in the field. Only in the late 1980's neural networks regained popularity and crucial results were shown. Most importantly for the scope of this thesis, it was proved that, under certain conditions, neural networks can approximate a large class of functions to any desired accuracy (Universal Approximation Theorem, [17], [6]). Leveraging this property and approximating the solution to the Schrödinger equation via a neural network has proved to be a successful approach ([9], [28], [16]).

In this chapter, I will define shallow neural networks and present a proof of the Universal Approximation Theorem. Then I will describe (Stochastic) Gradient Descent, a fundamental algorithm to train neural networks. Some references on these topics include [3], [38], [25], [26].

3.1 Shallow neural networks and the Universal Approximation Theorem

Neural networks come in all shapes and sizes. The simplest ones are *shallow neural networks* 3.1. In this section, if not otherwise specified, K denotes the n -dimensional unit cube and $C(K)$ the set of continuous functions on K .

Definition 3.1. Let $d, N \in \mathbb{N}$ and let $\sigma : \mathbb{R} \rightarrow \mathbb{R}$ be a non-linear function. A *shallow neural network* is defined by its *architecture* (d, N, σ) . Here, d denotes the network's *input dimension*, N the number of *hidden neurons* and σ the *activation function*. Given a shallow neural network, its corresponding *realization function* $\Phi : \mathbb{R}^d \rightarrow \mathbb{R}$ is given by

$$\Phi(x) = \sum_{i=1}^N w_i^{(2)} \sigma(\langle w_i^{(1)}, x \rangle + b_i^{(1)}) + b_i^{(2)},$$

where $w_i^{(1)} \in \mathbb{R}^d$ and $b_i^{(1)}, w_i^{(2)} \in \mathbb{R}$ for $i = 1, \dots, N$. The numbers $w_i^{(j)}$, $j = 1, 2$ are called the *weights* of the neural network and $b_i^{(j)}$, $j = 1, 2$ the *biases* of the neural network. Lastly, the set of all the parameters of the neural network will be denoted by θ and, if needed, the dependence of Φ on both x and θ will be highlighted by writing $\Phi_\theta(x)$.

3 A brief introduction to Neural Networks

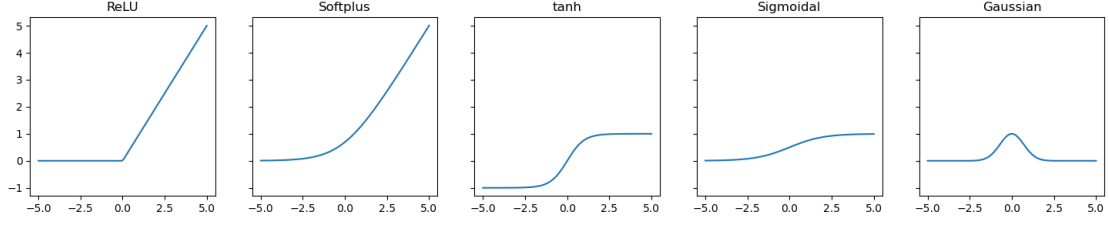


Figure 3.1: Examples of activation functions.

Some widely used activation functions include the *Rectified Linear Unit (ReLU)* (see Figure 3.1) and the *sigmoid* function (see Figure 3.1).

Definition 3.2. A function $f : \mathbb{R} \rightarrow \mathbb{R}$ is called *sigmoidal* if

$$\lim_{x \rightarrow -\infty} f(x) = 0 \text{ and } \lim_{x \rightarrow +\infty} f(x) = 1.$$

Theorem 3.3 (Universal Approximation Theorem). *Let σ be a continuous sigmoidal function. Then the set of shallow neural networks with activation function σ is dense in $C(K)$ with respect to the supremum norm.*

In other words, given $\varepsilon > 0$ and any $f \in C(K)$, one can find a shallow neural network with sigmoidal activation function σ such that for its realization function Φ , it holds

$$\sup_{x \in K} |\Phi(x) - f(x)| < \varepsilon.$$

This first version of the Universal Approximation Theorem was proven by Cybenko [6] in 1989 and a similar proof will be presented in Section 3.1.1. Since then, many further results have been shown. In particular, it is known that the set of shallow neural networks with activation function σ is dense in $C(K)$ if and only if σ is nonpolynomial [30].

Shallow neural networks are fairly simple but theoretically powerful structures. However, in many applications, *deep* neural networks are typically used. These are built by successively stacking shallow neural networks in multiple layers. Several versions of the Universal Approximation Theorem have been proved for deep neural networks. One notable result is that the set of neural networks of arbitrary depth and bounded width with activation function σ is dense in $C(K)$ if σ is continuously differentiable at at least one point, with nonzero derivative at that point [18]. This includes ReLU (see Figure 3.1) neural networks.

On the one hand, these results are promising. On the other hand, different problems might arise in practice. For instance, it has been shown that the training of ReLU networks to high uniform accuracy is intractable [4].

3.1 Shallow neural networks and the Universal Approximation Theorem

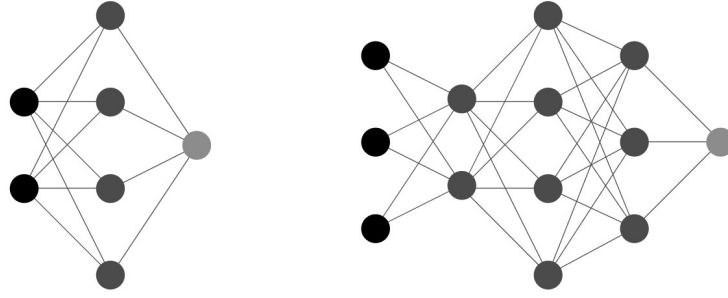


Figure 3.2: Schematic representation of a shallow neural network (L) and of a deep neural network with three inner layers (R).

3.1.1 Proof of the Universal Approximation Theorem

The Universal Approximation Theorem follows from some important theorems in Real, Functional and Fourier Analysis. These include the Hahn-Banach Theorem, the Riesz Representation Theorem and the Lebesgue Dominated Convergence Theorem.

Theorem 3.4 (Theorem of Hahn-Banach). *Let X be a normed space and $U \subseteq X$ a linear subspace. Let $u : U \rightarrow \mathbb{R}$ be a continuous linear functional. Then there exists a continuous linear functional $x : X \rightarrow \mathbb{R}$ such that*

$$x|_U = u \text{ and } \|x\| = \|u\|.$$

Proof. For the proof of this theorem, refer to Theorem III.1.5 in [43]. □

Theorem 3.5 (Riesz Representation Theorem). *Let K be a compact metric space and $x : C(K) \rightarrow \mathbb{R}$ be a nonzero functional. Then there exists a signed nonzero measure μ such that*

$$x(h) = \int_K h(t) d\mu(t)$$

for all $h \in C(K)$.

Proof. For the proof of this theorem, refer to Theorem II.2.5 in [43]. □

Theorem 3.6 (Dominated Convergence Theorem). *Suppose that $(f_n)_{n \in \mathbb{N}}$ is a sequence of measurable functions on a measure space (S, Σ, μ) such that*

$$f_n(x) \rightarrow f(x), \quad (x \in S)$$

as $n \rightarrow \infty$. Suppose that there exists a Lebesgue integrable function g on S such that

$$|f_n(x)| \leq g(x), \text{ for all } n \in \mathbb{N}, \quad (x \in S).$$

3 A brief introduction to Neural Networks

Then it holds

$$\lim_{n \rightarrow \infty} \int_S f_n d\mu = \int_S f d\mu.$$

Proof. For the proof of this theorem, refer to Theorem 11.32 in [33]. $\square$

Furthermore, the following corollary of the Hahn-Banach Theorem is needed.

Corollary 3.7. *[Of the Hahn-Banach Theorem] Let X be a normed space and $U \subseteq X$ a closed linear subspace. Let $t \in X \setminus U$. Then there exists a continuous linear functional $x : X \rightarrow \mathbb{R}$ such that*

$$x|_U = 0 \text{ and } x(t) \neq 0.$$

Proof. Let $q : X \rightarrow X/U$ be the quotient map. Then $q(u) = 0$ for all $u \in U$ and $q(t) \neq 0$. Next, consider $f : \text{span}\{q(t)\} \rightarrow \mathbb{R}$ defined by

$$f(\lambda q(t)) = \lambda \|q(t)\|.$$

By the Hahn-Banach Theorem (Theorem 3.4), there exist a continuous linear extension $F : X/U \rightarrow \mathbb{R}$ such that $F|_{\text{span}\{q(t)\}} = f$. Finally, define $x := F \circ q$. Then for all $u \in U$, it holds

$$x(u) = F(q(u)) = F(0) = 0,$$

and moreover

$$x(t) = F(q(t)) = \|q(t)\| \neq 0.$$

$\square$

In the following, a proof of the Universal Approximation Theorem based on the one in [6] will be presented.

Proof of Theorem 3.3. Step 1. It is easy to see that the set of shallow neural networks as in Definition 3.1 is a linear subspace of $C(K)$. Call it S and denote its closure by $\bar{S}$. Assume that $\bar{S}$ is a proper subset of $C(K)$, i.e., the set $C(K) \setminus \bar{S}$ is nonempty. By Corollary 3.7, there exists a nonzero continuous linear functional x on $C(K)$ such that $x|_{\bar{S}} = 0$. By the Riesz Representation Theorem, there exists a signed nonzero measure μ such that

$$x(h) = \int_K h(t) d\mu(t)$$

for all $h \in C(K)$.

Let $w \in \mathbb{R}^n$, $b \in \mathbb{R}$ and define $h := \sigma(\langle w, \cdot \rangle + b)$. Then $h \in \bar{S} \subset C(K)$ and therefore

$$0 = x(h) = \int_K \sigma(\langle w, t \rangle + b) d\mu(t) \text{ for all } w \in \mathbb{R}^n, b \in \mathbb{R}. \quad (3.1)$$

Step 2a. In order to prove the theorem, we will show that condition 3.1 implies that $\mu = 0$, a contradiction. To this end, define for $\alpha, \beta \in \mathbb{R}$ the functions

$$f_{\alpha, \beta}(t) := \sigma(\alpha(\langle w, t \rangle + b) + \beta) = \sigma(\alpha \langle w, t \rangle + \alpha b + \beta)$$

3.1 Shallow neural networks and the Universal Approximation Theorem

and note that, since σ is continuous and sigmoidal, for any t, w, b, β it holds

$$f_{\alpha, \beta}(t) \rightarrow \begin{cases} 0 & \text{if } \langle w, t \rangle + b < 0 \\ 1 & \text{if } \langle w, t \rangle + b > 0 \\ \sigma(\beta) & \text{if } \langle w, t \rangle + b = 0 \end{cases} =: X(t) \quad \text{as } \alpha \rightarrow \infty.$$

In other words, X is the pointwise limit of $f_{\alpha, \beta}$ as $\alpha \rightarrow \infty$. Since each $f_{\alpha, \beta}$ is clearly bounded by an integrable function, the Dominated Convergence Theorem implies

$$\lim_{\alpha \rightarrow \infty} \int_K f_{\alpha, \beta} d\mu = \int_K \lim_{\alpha \rightarrow \infty} f_{\alpha, \beta} d\mu = \int_K X d\mu. \quad (3.2)$$

On the other hand, by 3.1, we have

$$\begin{aligned} \lim_{\alpha \rightarrow \infty} \int_K f_{\alpha, \beta} d\mu &= \lim_{\alpha \rightarrow \infty} \int_K \sigma(\alpha(\langle w, t \rangle + b) + \beta) d\mu \\ &= \lim_{\alpha \rightarrow \infty} \int_K \sigma(\langle \alpha w, t \rangle + (\alpha b + \beta)) d\mu = 0. \end{aligned} \quad (3.3)$$

Combining 3.2 and 3.3, we get

$$0 = \int_K X d\mu = \mu(\Pi_1^{w, b}) + \sigma(\beta)\mu(\Pi_2^{w, b})$$

for all β, w, b , where $\Pi_1^{w, b} = \{t \in \mathbb{R}^n | \langle w, t \rangle + b > 0\}$ and $\Pi_2^{w, b} = \{t \in \mathbb{R}^n | \langle w, t \rangle + b = 0\}$. By taking the limit $\beta \rightarrow -\infty$, this implies

$$\mu(\Pi_1^{w, b}) = 0 \quad \text{for all } w \in \mathbb{R}^n, b \in \mathbb{R}. \quad (3.4)$$

Step 2b. To show that 3.4 implies $\mu = 0$ and conclude the proof, note that for any $b_1 < b_2$ we have

$$0 = \mu(\Pi_1^{w, b_1}) - \mu(\Pi_1^{w, b_2}) = \int_K \chi_{(b_1, b_2]}(\langle w, t \rangle) d\mu(t)$$

for all $w \in \mathbb{R}^n$ where $\chi_{(b_1, b_2]}$ is the indicator function of the interval $(b_1, b_2]$. The same relation extends to step functions by linearity and to continuous functions by density. Then we obtain

$$\int_K g(\langle w, t \rangle) d\mu(t) = 0 \quad \text{for all } w \in \mathbb{R}^n, g \in C(\mathbb{R}).$$

In particular, this applies to the function $g(t) = \cos(t) + i \sin(t)$ we obtain

$$0 = \int_K \cos(\langle w, t \rangle) + i \sin(\langle w, t \rangle) d\mu(t) = \int_K \exp(i \langle w, t \rangle) d\mu(t) \quad (3.5)$$

for all $w \in \mathbb{R}^n$. This implies that the Fourier transform of μ is zero, and we can conclude

3 A brief introduction to Neural Networks

that $\mu = 0$ [34]. □

Remark. A function $f : \mathbb{R}^n \rightarrow \mathbb{R}$ is called *discriminatory* on a compact set K , if for a signed Borel measure $\mu \in M(K)$, the following implication is true

$$\int_K f(\langle w, t \rangle + b) d\mu(t) = 0 \text{ for all } w \in \mathbb{R}^n, b \in \mathbb{R} \implies \mu = 0. \quad (3.6)$$

In other words, f is discriminatory on K if the only signed Borel measure for which the integral in 3.6 is zero for all $w \in \mathbb{R}^n, b \in \mathbb{R}$ is the zero measure. In Step 1 of the proof of Theorem 3.3, it was actually shown that the set of shallow neural networks with activation function σ is dense in $C(K)$ with respect to the supremum norm if σ is a discriminatory function. In Step 2, it was shown that any continuous sigmoidal function is discriminatory. In particular, some common activation functions such as tanh and sigmoid (Figure 3.1) are discriminatory.

Lemma 3.8. *The common activation function $ReLU : \mathbb{R} \rightarrow \mathbb{R}$ (Figure 3.1) is not sigmoidal but is discriminatory.*

Proof. That $ReLU$ is not sigmoidal is clear. Let $K \subseteq \mathbb{R}$ be a compact set and $\mu \in M(K)$ be a signed Borel measure. Assume that the following holds for all $w \in \mathbb{R}, b \in \mathbb{R}$:

$$\int_K ReLU(wt + b) d\mu(t) = 0.$$

We want to show that $\mu = 0$. Consider the following function $f : \mathbb{R} \rightarrow \mathbb{R}$:

$$f(t) = \begin{cases} 0 & \text{if } t < 0 \\ t & \text{if } t \in [0, 1] \\ 1 & \text{if } t > 1 \end{cases} = ReLU(t) - ReLU(t - 1).$$

The function f is continuous and sigmoidal, hence discriminatory. Furthermore, it holds

$$\int_K f(wt + b) d\mu(t) = \int_K ReLU(wt + b) d\mu(t) - \int_K ReLU(wt + (b - 1)) d\mu(t) = 0.$$

Since f is discriminatory, this implies $\mu = 0$. □

3.2 Training a Neural Network

The *training* of a neural network refers to the process of fine-tuning its parameters in order to achieve a good approximation of the target function. This is done by minimizing a loss function based on a set of data points.

Example 3.9. Let Φ_θ be the realization function of a given neural network with parameters θ and f be the target function. Given m data points $(r_i, f(r_i))$, $i = 1, \dots, m$,

one possible loss function could be

$$\mathcal{L}((r_i)_{i=1}^m, \theta) = \sum_{i=1}^m |f(r_i) - \Phi_\theta(r_i)|^2.$$

In this context, the function $\mathcal{L}(\cdot, \theta)$ is a stochastic function, in the sense that its output depends on which datapoints it is evaluated on. When training a neural network, often one is interested in minimizing $\mathbb{E}[\mathcal{L}(\cdot, \theta)]$ with respect to the network parameters θ . This is done by evaluating $\mathcal{L}(\cdot, \theta)$ at different timesteps on (possibly) different sets of datapoints and adjusting the parameters accordingly, for instance with (Stochastic) Gradient Descent (Section 3.2.1) steps. I will denote the realization of $\mathcal{L}$ at timestep k with $\mathcal{L}_k$.

The classical way to train a neural network is by using gradient-based methods, since computing the pointwise derivatives $\nabla_\theta \Phi_\theta(x)$ of the realization function (and therefore the ones of the loss function) can be done efficiently and accurately by means of automatic differentiation [35].

Stochastic Gradient Descent (SGD), a well-known algorithm for minimizing a differentiable (convex) function will first be described in this section. The loss function in the training of a neural network is usually differentiable, but more often than not, highly non-convex. However, SGD is often used anyway and usually yields good results.

3.2.1 Stochastic Gradient Descent

In this section, I will describe Gradient Descent and its variant Stochastic Gradient Descent in the context of minimizing any differentiable convex function $\mathcal{L} : \mathbb{R}^n \rightarrow \mathbb{R}$. This can be applied to the training of neural networks by taking $\mathcal{L}$ to be the desired loss function, for instance as in Example 3.9.

Gradient Descent

Gradient Descent is a classical optimization algorithm for minimizing a differentiable convex function $\mathcal{L} : \mathbb{R}^n \rightarrow \mathbb{R}$.

Definition 3.10. Let $\mathcal{L} : \mathbb{R}^n \rightarrow \mathbb{R}$ be a differentiable function. Then the *gradient* of $\mathcal{L}$ is the vector of partial derivatives, i.e.

$$\nabla_\theta \mathcal{L} : \mathbb{R}^n \rightarrow \mathbb{R}^n \tag{3.7}$$

$$\theta \mapsto \left(\frac{\partial \mathcal{L}(\theta)}{\partial \theta_1}, \dots, \frac{\partial \mathcal{L}(\theta)}{\partial \theta_n} \right). \tag{3.8}$$

Note that the gradient of $\mathcal{L}$ at a point $\theta \in \mathbb{R}^n$ is the direction of steepest ascent. At each iteration, GD takes a step in the opposite direction, effectively decreasing the value of $\mathcal{L}$. This works well in case of a convex function, but might result in stopping at a local minimum or plateau in other cases.

Algorithm 1 Gradient Descent

Require: Starting point $\theta^0 \in \mathbb{R}^n$, number of iterations $K > 0$ and sequence of step sizes $\eta^k > 0$, $k = 1, \dots, K$

- 1: **for** $k = 0, 1, \dots, K$ **do**
- 2: Set $d^k := -\nabla_{\theta} \mathcal{L}(\theta^k)$
- 3: Set $\theta^{k+1} := \theta^k + \eta^k d^k$
- 4: **end for**
- 5: **return** θ^K

Stochastic Gradient Descent

In large-scale problems, calculating the full gradient can often be expensive. In Stochastic Gradient Descent, this is replaced by a random vector d^k whose expected value at each iteration is the gradient. How this vector is chosen depends on the application, see for instance Example 3.11.

Algorithm 2 Stochastic Gradient Descent

Require: Starting point $\theta^0 \in \mathbb{R}^n$, number of iterations $K > 0$ and sequence of step sizes $\eta^k > 0$, $k = 1, \dots, K$

- 1: **for** $k = 0, 1, \dots, K$ **do**
- 2: Choose $d^k \in \mathbb{R}^n$ at random from a distribution such that $\mathbb{E}[d^k | \theta^k] = \nabla_{\theta} \mathcal{L}(\theta^k)$
- 3: Set $\theta^{k+1} := \theta^k + \eta^k d^k$
- 4: **end for**
- 5: **return** θ^K

Example 3.11. In the setting of Example 3.9, one might choose a subset $S \subseteq \{1, \dots, m\}$ of size $1 \leq |S| < m$ and set

$$d^k := \nabla_{\theta} \mathcal{L}((r_i)_{i \in S}, \theta) = \nabla_{\theta} \mathcal{L}_k(\theta). \quad (3.9)$$

Concretely, this is done by shuffling the given data points and dividing them in smaller subsets, the so-called *batches*. At each step of SGD, a different batch is chosen to calculate d^k as in 3.9.

In practice (as in Example 3.11), one needs to compute $\nabla_{\theta} \mathcal{L}((r_i)_{i \in S}, \theta)$ for a certain set of data points $(r_i)_{i \in S}$, where S is an index set. This can be done efficiently using *backpropagation*, an algorithm introduced in 1989 by Rumelhart et al [35].

3.2 Training a Neural Network

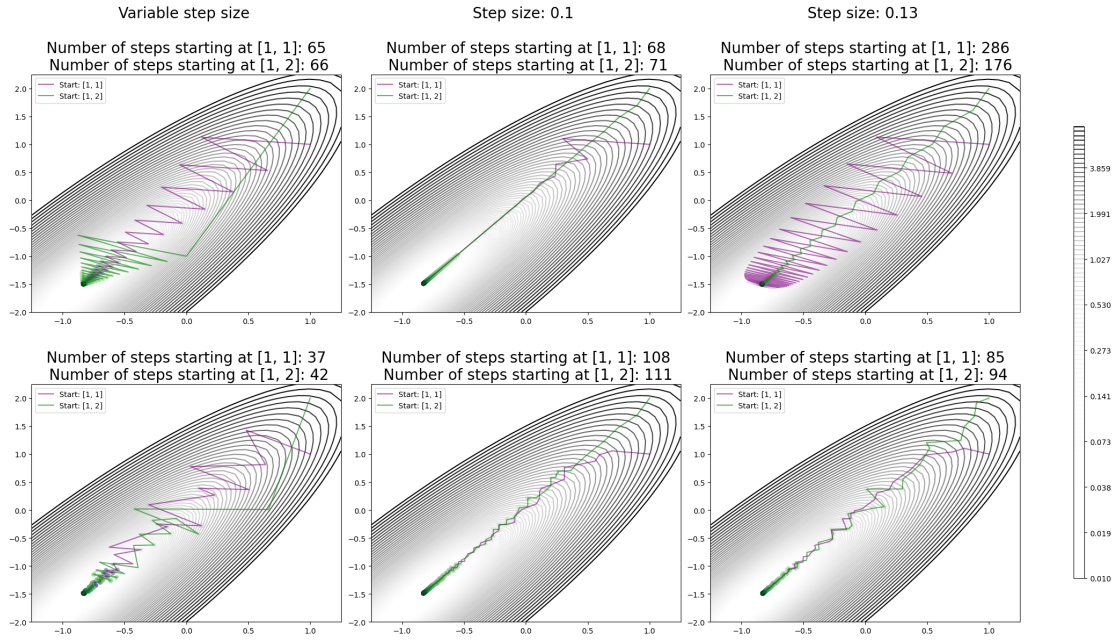


Figure 3.3: Steps of Gradient Descent (top) and Stochastic Gradient Descent (bottom) to minimize the function $f(x, y) = 6x^2 - 6xy + 2y^2 + x + y + 1$ starting from the points (1, 1) (purple) and (1, 2) (green) and using different step sizes (Armijo step size [1] (L), 0.1 (C), 0.13 (R)). The algorithms stopped when $|\nabla f| < 10^{-2}$. It is clear that the success of the algorithms depends on the choice of starting point and step size.

4 Approximating solutions to the Schrödinger Equation

The complex nature of most quantum system makes solving their Schrödinger equation a challenging task. It is only possible to solve the Schrödinger equation exactly for very small systems. Approximate methods to solve it include the WKB method [11] and perturbation theory [41], [31]. In this thesis, a variational method will be used. Variational methods are approximation methods that involve choosing a trial wavefunction dependent on a set of parameters and optimizing these parameters to minimize an energy expectation value [41], [31].

Variational methods are based on the so-called *variational principle*. In the first section of this chapter, the variational principle will be stated and proved for self-adjoint operators on finite dimensional Hilbert spaces. A corresponding theorem for the infinite dimensional case will be stated and its proof will be sketched. Next, the variational method will be explained, with focus on variational Monte Carlo methods.

4.1 The variational principle

If an operator for a quantum system is self-adjoint on an appropriate finite dimensional Hilbert space, its eigenvalues are real and its eigenfunctions form an orthonormal basis of $\mathcal{H}$. This fact is used in the proof of the following theorem.

Theorem 4.1 (Variational Principle, finite dimensional case). *Let H be a self-adjoint operator on an n -dimensional Hilbert space $\mathcal{H}$ and $\psi \in \mathcal{H}$ a function such that $\langle \psi | \psi \rangle = 1$. Let E_0 be the lowest eigenvalue of H . Then it holds*

$$\langle \psi | H | \psi \rangle \geq E_0. \quad (4.1)$$

Proof. Let $\{\phi_i\}$ be the set (normalized) eigenfunctions of H . Since H is self-adjoint, the set $\{\phi_i\}$ forms an orthonormal basis of $\mathcal{H}$. Hence, we have the following representation of ψ

$$\psi = \sum_{i=0}^n \alpha_i \phi_i. \quad (4.2)$$

For the coefficients α_i , it holds

$$\sum_{i=0}^n \alpha_i^* \alpha_i = \sum_{i,j} \alpha_j^* \alpha_i \langle \phi_j | \phi_i \rangle = \langle \psi | \psi \rangle = 1.$$

4 Approximating solutions to the Schrödinger Equation

Then, the following holds

$$\begin{aligned}\langle \psi | H | \psi \rangle &= \left\langle \sum_{j=0}^n \alpha_j \phi_j \middle| H \middle| \sum_{i=0}^n \alpha_i \phi_i \right\rangle = \left\langle \sum_{j=0}^n \alpha_j \phi_j \middle| \sum_{i=0}^n \alpha_i H \phi_i \right\rangle = \left\langle \sum_{j=0}^n \alpha_j \phi_j \middle| \sum_{i=0}^n \alpha_i E_i \phi_i \right\rangle \\ &= \sum_{i,j} \alpha_j^* \alpha_i \langle \phi_j | E_i \phi_i \rangle = \sum_{i,j} \alpha_j^* \alpha_i E_i \langle \phi_j | \phi_i \rangle \geq E_0 \sum_{i=0}^n \alpha_i^* \alpha_i = E_0.\end{aligned}$$

□

Unfortunately, the proof of Theorem 4.1 is in general not valid if $\mathcal{H}$ is infinite dimensional. This is because self-adjoint operators on infinite dimensional Hilbert spaces need not have any eigenvalues, see example 2.8, and therefore no representation like 4.2.

However, a similar statement to Theorem 4.1 does hold in the infinite dimensional case as well. In the remaining of this section, it will be presented and proved in a similar way as in [40] and [45]. We start with the following statement.

Theorem 4.2. *Let H be self-adjoint. Then*

$$\inf \sigma(H) = \inf \{ \langle \psi | H | \psi \rangle \text{ s.t. } \psi \in \text{dom}(H) \text{ and } \|\psi\| = 1 \}. \quad (4.3)$$

Proof. Define c to be the right-hand-side of Equation 4.3. If $c = -\infty$, it trivially holds $\inf \sigma(H) \geq c$. The same inequality is true also if $c \neq -\infty$. Indeed, let $\lambda \in \sigma(H) \subseteq \mathbb{R}$. It holds by the Cauchy-Schwarz inequality and by self-adjointness of H

$$\begin{aligned}\|(H - \lambda I)\psi\| \|\psi\| &\geq |\langle (H - \lambda I)\psi | \psi \rangle| \geq \langle (H - \lambda I)\psi | \psi \rangle \\ &= \langle H\psi | \psi \rangle - \lambda \langle \psi | \psi \rangle = \langle \psi | H | \psi \rangle - \lambda \langle \psi | \psi \rangle \\ &\geq c \|\psi\|^2 - \lambda \|\psi\|^2 = (c - \lambda) \|\psi\|^2,\end{aligned}$$

which implies for all ψ

$$\|(H - \lambda I)\psi\| \geq (c - \lambda) \|\psi\|.$$

Since $H - \lambda I$ is invertible, it must hold $\text{Ker}(H - \lambda I) = \{0\}$ and therefore $c - \lambda \geq 0$. This implies $\sigma(H) \subseteq [c, \infty)$ and thus $\inf \sigma(H) \geq c$.

To prove the reverse inequality, we use the Spectral Theorem. Recall that by Theorem 2.13, it holds

$$\begin{aligned}\langle \psi | H | \psi \rangle &= \langle H\psi | \psi \rangle = \int_{\inf \sigma(H)}^{\sup \sigma(H)} t \, d\langle E\psi | \psi \rangle \\ &\geq \inf \sigma(H) \int_{\inf \sigma(H)}^{\sup \sigma(H)} d\langle E\psi | \psi \rangle = \inf \sigma(H) \langle \psi | \psi \rangle,\end{aligned}$$

where the last equality follows from 2.2 with $h(t) = 1, t \in \mathbb{R}$. This implies $\inf \sigma(H) \leq c$, thus proving the claim. □

4.1 The variational principle

If the smallest eigenvalue E_0 of a self-adjoint operator H lies at the bottom of the spectrum of H , then by Theorem 4.2 it holds

$$\langle \psi | H | \psi \rangle \geq E_0 \quad (4.4)$$

for all normalized $\psi \in \text{dom}(H)$.

Similar bounds can also be obtained for larger eigenvalues, as long as they lie below the essential spectrum of H . The *essential spectrum* is the set of values in the spectrum of H that are *not* isolated eigenvalues with finite multiplicity.

Theorem 4.3. *Let H be self-adjoint. Let $E_0 \leq E_1 \leq \dots \leq E_N$ be the eigenvalues of H below the essential spectrum. Define*

$$\begin{aligned} O(\phi_0, \dots, \phi_n) &= \{\psi \in \text{dom}(H) \text{ s.t. } \|\psi\| = 1 \text{ and } \psi \in \text{span}\{\phi_0, \dots, \phi_n\}^\perp\}, \\ S(\phi_0, \dots, \phi_n) &= \{\psi \in \text{dom}(H) \text{ s.t. } \|\psi\| = 1 \text{ and } \psi \in \text{span}\{\phi_0, \dots, \phi_n\}\}. \end{aligned}$$

Then, for $n = 0, \dots, N-1$, it holds

$$E_{n+1} = \sup_{\phi_0, \dots, \phi_n} \inf_{\psi \in O(\phi_0, \dots, \phi_n)} \langle \psi | H | \psi \rangle \quad (4.5)$$

$$= \inf_{\phi_0, \dots, \phi_n} \sup_{\psi \in S(\phi_0, \dots, \phi_n)} \langle \psi | H | \psi \rangle. \quad (4.6)$$

The idea of the proof of statement 4.5 will be sketched. Details can be found in [45], [40], [31], [41].

Equation 4.4 yields a very useful bound for E_0 and a method to approximate its eigenfunction ψ_0 (see also Section 4.2). To compute an approximation to the second smallest eigenvalue, assume that the eigenfunction ψ_0 corresponding to E_0 has already been found. Then one could restrict H to $\text{span}\{\psi_0\}^\perp$, so that E_1 is the smallest eigenvalue. Proceeding as before will yield an approximation of E_1 , since now it holds (without explicitly stating the - still very important - assumptions that $\psi \in \text{dom}(H)$ and $\|\psi\| = 1$)

$$E_1 = \inf_{\psi \in \text{span}\{\psi_0\}^\perp} \langle \psi | H | \psi \rangle.$$

Continuing in this manner will yield higher eigenvalues as well. However, an issue arises when $\psi_0, \psi_1, \dots$ are not found exactly and approximations are used to restrict the domain of H . Assume $\phi_0 \in \text{dom}(H)$ is a normalized approximation for ψ_0 . Then even when restricting H to $\text{span}\{\phi_0\}^\perp$, a component in the direction of ψ_0 will still be present. Therefore it holds

$$E_1 \geq \inf_{\psi \in \text{span}\{\phi_0\}^\perp} \langle \psi | H | \psi \rangle.$$

Equality holds when $\phi_0 = \psi_0$, and in this case the value of E_1 is the greatest of all expectations taken over all $\phi_0 \in \text{dom}(H)$:

$$E_1 = \sup_{\phi_0 \in \text{dom}(H)} \inf_{\psi \in \text{span}\{\phi_0\}^\perp} \langle \psi | H | \psi \rangle.$$

4 Approximating solutions to the Schrödinger Equation

Proceeding in this manner yields analogous equalities for higher eigenvalues as stated in Theorem 4.3.

Remark. In case ψ is not normalized, then the proofs of 4.1 and 4.2 can be easily adapted to obtain

$$\frac{\langle \psi | H | \psi \rangle}{\langle \psi | \psi \rangle} \geq E_0. \quad (4.7)$$

4.2 The variational method

By the proof of Theorem 4.1 and the discussion at the end of the previous section, it is clear that equality in (4.1) and (4.7) is achieved if and only if $\psi = \phi_0$. That is, if ψ is the ground state wave function or, in other words, the solution of the ground state Schrödinger Equation

$$H\psi = E_0\psi,$$

where $E_0 \in \mathbb{R}$ is the smallest possible eigenvalue of H . This observation, together with Theorem 4.1, motivates the *variational method* for solving the Schrödinger Equation for a self-adjoint operator H . Starting with a *trial wave function* ψ_θ parametrized by a set of parameters θ , one solves the minimization problem

$$\min_{\theta} \langle \psi_\theta | H | \psi_\theta \rangle. \quad (4.8)$$

The solution to this problem is an approximation to the solution of the Schrödinger equation.

This is the method that will be used in this thesis to find an approximate solution to the Schrödinger Equation. Here, ψ_θ will be assumed to be a neural network with parameters θ (see Chapter 3). The parameters θ of the neural network will be optimized to solve 4.8 using two different algorithms (see Chapter 5). In Chapter 6, the results obtained with the two algorithms will be compared.

4.2.1 Variational Monte Carlo

The name *variational Monte Carlo* refers to the use of Monte Carlo algorithms to implement the variational method. These algorithms are designed to solve potentially deterministic problems using randomness. They do so by relying on iterative random sampling.

Define the local (as in, defined for a normalized quantum state ψ) energy of an operator $H : \mathcal{H} \rightarrow \mathcal{H}$ as the expectation value of H at state ψ :

$$\langle E \rangle_\psi := \langle \psi | H | \psi \rangle = \int \psi^*(r) H \psi(r) dr. \quad (4.9)$$

The variational principle 4.1 can now be restated as to say that the local energy of an approximation of the solution of the Schrödinger Equation is always higher than the one

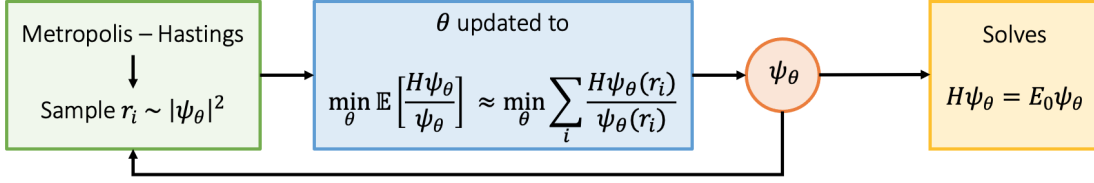


Figure 4.1: A schematic representation of the method used in this thesis to solve the Schrödinger Equation is depicted. A neural network ψ_θ is initialized and random datapoints r_i are sampled according to $|\psi_\theta|^2$ by means of the Metropolis-Hastings algorithm. These datapoints are used to train ψ_θ to minimize the loss function $\sum_i \frac{H\psi_\theta(r_i)}{\psi_\theta(r_i)}$, an approximation of the local energy $\langle E \rangle_\psi$. In this thesis, this minimization is done on the one hand with the Adam algorithm [19] and on the other hand with the SWO algorithm [21]. The newly obtained ψ_θ is used to sample new datapoints until convergence is reached.

of the true solution. Therefore, the idea of variational method is merely to minimize the local energy.

In variational Monte Carlo, Markov-chain Monte Carlo methods are used to numerically estimate the local energy $\langle E \rangle_\psi$. By the properties of the expected value, it holds

$$\begin{aligned} \langle E \rangle_\psi &= \int \psi^2(r) \frac{H\psi(r)}{\psi(r)} dr \\ &= \mathbb{E}_{r \sim \psi^2} \left[\frac{H\psi(r)}{\psi(r)} \right]. \end{aligned}$$

In order to evaluate this expression, n values $r_1, \dots, r_n$ will be numerically sampled from ψ^2 with the Metropolis-Hastings algorithm (see Section 4.2.2) and then the following expression will be used

$$\langle E \rangle_\psi \approx \frac{1}{n} \sum_{i=1}^n \frac{H\psi(r_i)}{\psi(r_i)} = \frac{1}{n} \sum_{i=1}^n E_L(r_i),$$

where $E_L(x) := \psi^{-1}(x)H\psi(x)$.

4.2.2 The Metropolis-Hastings algorithm

Metropolis-Hastings [24], [12] is an algorithm to numerically find a sequence of points distributed according to a given probability density function. In this thesis, it has been used to sample data points $r_i \sim |\psi_\theta|^2, i = 1, \dots, n$, where $\psi_\theta : \mathbb{R}^d \rightarrow \mathbb{R}$ is a trial wave function. See Figure 4.2 and Algorithm 3.

Algorithm 3 Metropolis-Hastings

Require: probability distribution $|\psi_\theta|^2$, initial samples $r_i \in \mathbb{R}^d$, $i = 1, \dots, n$

- 1: **while** not converged **do**
 - 2: Generate normally distributed samples $s_i \in \mathbb{R}^d$, $i = 1, \dots, n$
 - 3: Set $s_i \leftarrow r_i + s_i$, $i = 1, \dots, n$
 - 4: Calculate ratios $R_i = |\psi_\theta|^2(s_i)/|\psi_\theta|^2(r_i)$, $i = 1, \dots, n$
 - 5: Generate uniformly distributed values $u_i \in [0, 1]$, $i = 1, \dots, n$
 - 6: **if** $R_i > u_i$ **then** Set $r_i \leftarrow s_i$, $i = 1, \dots, n$
 - 7: **end if**
 - 8: **end while**
 - 9: **return** r_i , $i = 1, \dots, n$
-

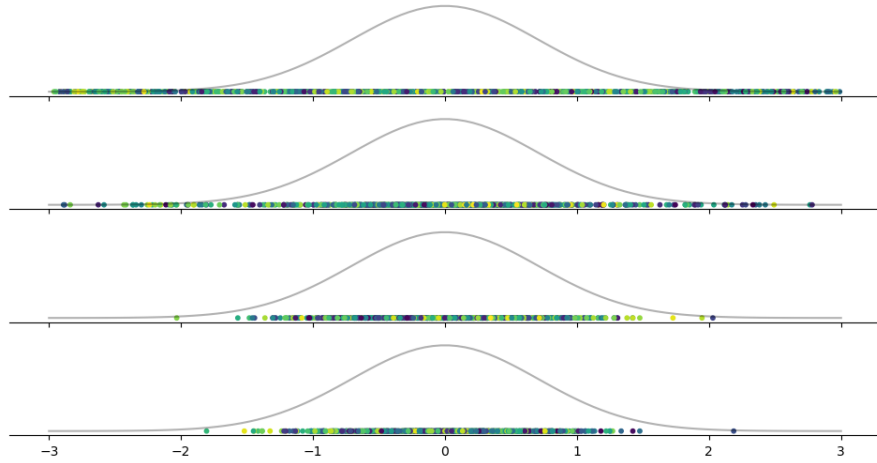


Figure 4.2: Four steps of the Metropolis-Hastings algorithm.

5 Optimizers

In this section, a more detailed overview of the used algorithms will be given. Minimizing a loss function during the training a neural network can in general be done using Stochastic Gradient Descent (Section 3.2.1), an algorithm that was first presented in the mid-twentieth century. Since then, many variants have been introduced, most notably Adam by Kingma and Ba [19].

Moreover, to solve the specific minimization problem

$$\min_{\theta} \langle \psi_{\theta} | H | \psi_{\theta} \rangle. \quad (5.1)$$

that arises in the context of Variational Monte Carlo, the Supervised Wave function Optimization (SWO) algorithm was introduced by Kochkov and Clark in 2018 [21].

5.1 Adam

In this section, I will describe Adam as it was first introduced in [19].

Let $\mathcal{L}(\theta)$ be a stochastic loss function with parameters θ (see for instance Examples 3.9 and 3.11). As in Example 3.9, call $\mathcal{L}_k(\theta)$ the realization of $\mathcal{L}$ at timestep k . Adam is an algorithm to solve

$$\min_{\theta} \mathbb{E}[\mathcal{L}(\theta)].$$

It does so by extending Stochastic Gradient Descent. During training, Adam maintains the moving averages of gradients (first moment) and square gradients (second moment) for each parameter in the network. This smooths out gradient noise and speeds up convergence. Furthermore, it also helps with escaping from local minima. This is particularly useful when the objective function $\mathcal{L}$ is not convex, as it often is the case in the training of neural networks.

5.2 SWO

The Supervised Wave function Optimization (SWO) algorithm was introduced to solve the energy optimization problem 5.1. The SWO algorithm works over multiple epochs. At each epoch, it minimizes the difference between the current solution ψ_{θ} and a target wave function ψ_T :

$$\min_{\theta} \mathcal{L}(\psi_{\theta}, \psi_T), \quad (5.2)$$

where $\mathcal{L}$ is a loss function. The target wave function ψ_T is epoch-specific. Furthermore, it is improved at each step, meaning that after each epoch, the new ψ_T will be closer

Algorithm 4 Adam (as in [19])

Require: Starting point $\theta^0 \in \mathbb{R}^n$, number of iterations $K > 0$ and step size $\eta > 0$, first moment vector $M^0 = 0$, second moment vector $V^0 = 0$ and exponential decay rates for the moment estimates $\beta_1, \beta_2 \in [0, 1)$

- 1: **for** $k = 0, 1, \dots, K$ **do**
- 2: Compute (for instance with backpropagation) $d^k := \nabla_{\theta} \mathcal{L}_k(\theta^k)$
- 3: Compute bias-corrected first moment estimate:

$$M^{k+1} \leftarrow \frac{\beta_1 M^k + (1 - \beta_1) d^k}{(1 - (\beta_1)^2)}$$

- 4: Compute bias-corrected second raw moment estimate:

$$V^{k+1} \leftarrow \frac{\beta_2 V^k + (1 - \beta_2) (d^k)^2}{(1 - (\beta_2)^2)}$$

- 5: Set $\theta^{k+1} := \theta^k + \eta M^{k+1} / (\sqrt{V^{k+1}} + \varepsilon)$
 - 6: **end for**
 - 7: **return** θ^K
-

to the true solution. There are different ways to define the improved ψ_T through the current trial wave function ψ_{θ} , as mentioned in [21]. For the implementation specific to this thesis ψ_T will be defined by means of *imaginary time evolution*.

Remark (Imaginary time evolution). The time-dependent Schrödinger equation 2.6 with Hamiltonian H may be formally solved to obtain

$$\psi(x, t) = e^{-iHt/\hbar} \psi(x, 0). \quad (5.3)$$

In order to take the exponential of an (unbounded) operator, a continuous functional calculus is required [43]. Usually, $e^{-iHt/\hbar}$ is defined by the exponential power series in H and is itself an operator, called the *time propagator*. Equations 2.6 and 5.3 govern the evolution in time of the wave function. Substituting $\tau = it/\hbar$ in 5.3, we obtain the *imaginary time propagator* $e^{-\tau H}$.

Furthermore, using the representation of ψ_{θ} as a linear combination of the eigenfunctions of H where the eigenvalues E_i are in ascending order, one obtains the following:

$$e^{-\tau H} \psi_{\theta} = \sum_i \alpha_i e^{-\tau E_i} \psi_i \quad (5.4)$$

$$= \sum_i \alpha_i e^{-\tau E_i} \psi_i \approx \psi_0 \text{ as } \tau \rightarrow \infty \quad (5.5)$$

Let ψ_{θ} be the solution of the minimization problem 5.2 at an epoch. In the implementation of SWO used in this thesis, the next epoch's target wave function ψ_T is defined

as

$$\psi_T = \psi_\theta - \beta H \psi_\theta,$$

where $\beta > 0$ is a fixed parameter.

Formally approximating $e^{-\beta H}$ using the exponential power series, it holds

$$\psi_T \approx e^{-\beta H} \psi_\theta, \quad (5.6)$$

the imaginary time evolution of ψ_θ (compare with Remark 5.2 and Equation 5.3).

Approximation 5.6 is indeed meaningful for our purposes: a repeated application of the imaginary time propagator on ψ_θ is equivalent to letting β go to infinity, which brings the current wave function closer to the ground state one (see Remark 5.2).

By Stone's Theorem [11], the imaginary time propagator is not unitary. Hence, ψ_T as defined in 5.6 will not be normalized. A normalization factor N is introduced and updated during training. Using bra-ket notation and the fact that H is self-adjoint, the following holds

$$\begin{aligned} \langle N\psi_T | N\psi_T \rangle &= N^2 \langle \psi_T | \psi_T \rangle \\ &= N^2 (\langle \psi_\theta | - \beta \langle \psi_\theta | H) (| \psi_\theta \rangle - \beta H | \psi_\theta \rangle) \\ &= N^2 (\langle \psi_\theta | \psi_\theta \rangle - 2\beta \langle \psi_\theta | H | \psi_\theta \rangle + \beta^2 \langle \psi_\theta | H^2 | \psi_\theta \rangle). \end{aligned}$$

Setting the last expression equal to 1 and solving for N , we obtain

$$N(\theta) = (\langle \psi_\theta | \psi_\theta \rangle - 2\beta \langle E \rangle_{\psi_\theta} + \beta^2 \langle E^2 \rangle_{\psi_\theta})^{-0.5}, \quad (5.7)$$

where

$$\langle E \rangle_{\psi_\theta} = \langle \psi_\theta | H | \psi_\theta \rangle \text{ and } \langle E^2 \rangle_{\psi_\theta} = \langle \psi_\theta | H^2 | \psi_\theta \rangle.$$

In order to more easily avoid numerical instabilities, an exponential moving average of the normalization factors is kept during training. Furthermore, as discussed in Section 4, the following approximations are used:

$$\langle E \rangle_{\psi_\theta} \approx \frac{1}{n} \sum_{i=1}^n \frac{H\psi(r_i)}{\psi(r_i)} \text{ and } \langle E^2 \rangle_{\psi_\theta} \approx \frac{1}{n} \sum_{i=1}^n \left(\frac{H\psi(r_i)}{\psi(r_i)} \right)^2,$$

where $r_1, \dots, r_n$ are numerically sampled from ψ^2 using the Metropolis-Hastings algorithm.

Having defined ψ_T , one can use standard Gradient Descent techniques like Adam to find improved parameters

$$\theta = \arg \min_{\Theta} \mathcal{L}(\psi_\Theta, \psi_T)$$

(compare with 5.2). As a loss function during this inner minimization, a stochastic estimate L_2 of the $\mathcal{L}_2$ loss function is implemented. For each step of the inner minimization, a set of configurations $\mathbf{r} = \{r_j \in \mathbb{R}^n | j = 1, \dots, m\}$ is generated using the

5 Optimizers

Metropolis-Hastings algorithm such that

$$\mathbb{P}(r_j) \propto |\psi_{\theta_i}|^2 \quad \text{for all } j = 1, \dots, m.$$

Then, it holds

$$\mathcal{L}_2(\mathbf{r}, \psi_{\Theta}, \psi_T) \approx \sum_{j=1}^m \frac{(\psi_T(r_j) - \psi_{\Theta}(r_j))^2}{\mathbb{P}(r_j)} =: L_2(\mathbf{r}, \psi_{\Theta}, \psi_T).$$

Algorithm 5 Supervised Wave function Optimization (as in [21])

Require: Parameter $\beta > 0$, number of outer iterations $K > 0$, number of inner iterations $L > 0$, initial neural network parameters θ , initial set of configurations $\mathbf{r}$

- 1: Set $\alpha := 2/(L - 1)$
 - 2: Set $N^{**} = N(\theta)$ ▷ Initialize normalization factor with 5.7
 - 3: **for** $k = 0, 1, \dots, K$ **do**
 - 4: Set $\psi_T \leftarrow \psi_{\theta} - \beta H \psi_{\theta}$ ▷ Compute target wave function
 - 5: Set $N \leftarrow N^{**}$ ▷ Update normalization factor
 - 6: **for** $l = 0, 1, \dots, L$ **do**
 - 7: $\theta \leftarrow \arg \min_{\Theta} L_2(\mathbf{r}, \psi_{\Theta}, N \psi_T)$ ▷ Solve inner minimization
 - 8: Set $N^* \leftarrow N(\theta)$
 - 9: Set $N^{**} \leftarrow \alpha N^* + (1 - \alpha) N^{**}$ ▷ EMA of normalization factor
 - 10: Set $\mathbf{r} \leftarrow \text{M-H}(\mathbf{r}, \psi_{\theta})$ ▷ Compute new set of configurations
 - 11: **end for**
 - 12: **end for**
 - 13: **return** θ
-

5.2.1 Influence of the parameters

The influence of the parameters β and L has been analysed by looking at the performance of SWO in solving the one-dimensional quantum harmonic oscillator problem as described in Chapter 6.

The number of epochs needed to reach convergence does not decrease significantly as β increases, and the time needed for each epoch to run seems to slightly increase. This would suggest that a choice of a very high β is not necessarily a better one.

Increasing the values of L reduces the number K of epochs needed to reach convergence. At the same time however, each epoch takes much longer to complete. Setting $L = 1$ reduces SWO to the algorithm used to solve the inner minimization (Step 7 in Algorithm 5), in this case, Adam. I have found a value of L between 10 and 30 to be a good compromise between performance and runtime.

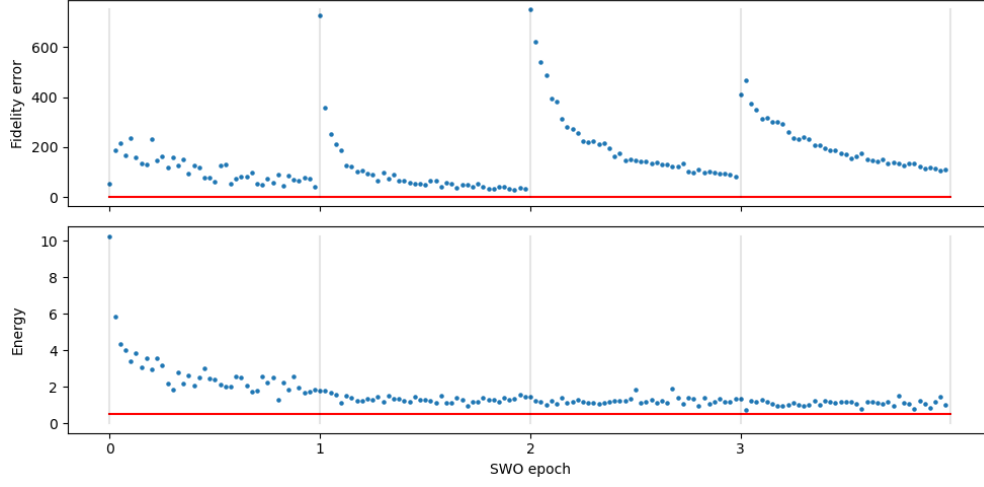


Figure 5.1: Fidelity error $L_2(\mathbf{r}, \psi_\Theta, N\psi_T)$ and local energy $\langle E \rangle_{\psi_\theta}$ during the first four epochs ($k = 0, 1, 2, 3$) of one application of SWO to the quantum harmonic oscillator problem (Section 6.2). The parameters used are $\beta = 0.12$ and $L = 40$. The inner minimization is solved using Adam [19]. The fidelity error converges to zero in each epoch, whereas the energy converges to the solution $E_0 = 0.5$ overall.

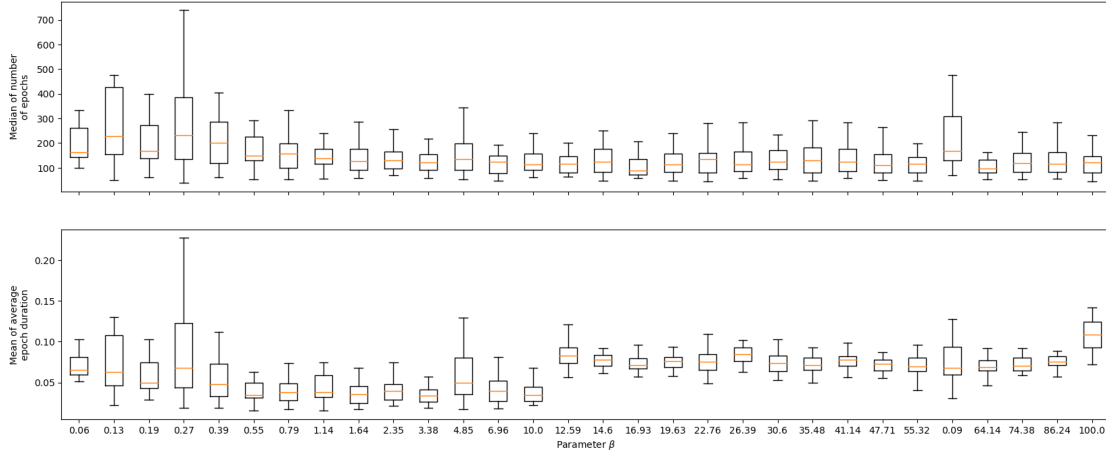


Figure 5.2: For 30 values of β (notice the scaling of the x-Axis), I have run 30 experiments to solve the quantum harmonic oscillator problem with the SWO algorithm. The boxplots represent the number of epochs needed to reach convergence as described in Chapter 6 (T) and the average duration of one epoch during the experiment (B). The few outliers have been omitted for clarity.

6 The Quantum Harmonic Oscillator

The one-dimensional quantum harmonic oscillator is a fairly simple problem for which there is an exact solution. This makes it a good first toy problem to test the performance of the two algorithms in question.

In the first part of this chapter, the Schrödinger equation for the quantum harmonic oscillator will be derived heuristically. Then, this chapter will cover both analytical solutions of the one-dimensional quantum harmonic oscillator as well as the approximate solutions obtained using the variational method and the algorithms described in Chapter 5.

6.1 The problem

The classical harmonic oscillator is a model describing a system subject to a restoring force F proportional to the displacement x from the system's equilibrium position, i.e.

$$F = -kx, \text{ for some } k > 0. \quad (6.1)$$

The simplest example of such a system is that of an object of mass m connected to a spring with force constant k (see Figure 6.1). In this case, 6.1 holds by Hooke's law.

Furthermore, by 6.1, the (parabolic) potential is given by

$$V(x) = \frac{1}{2}kx^2 = \frac{1}{2}m\omega^2x^2, \quad (6.2)$$

where $\omega = \sqrt{k/m}$ is the angular frequency of oscillation.

In quantum mechanics, the harmonic oscillator serves as a fundamental model that can be applied to a wide array of systems. In the following, I will consider the one-dimensional quantum harmonic oscillator. Therefore, let $x \in \mathbb{R}$. Recall (Section 2.2.2)

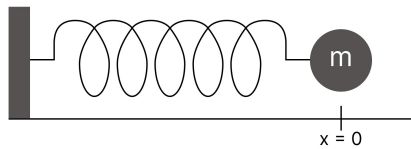


Figure 6.1: The constant k describes the stiffness of the spring.

that in this case, using 6.2 it holds

$$\begin{aligned}\hat{H}\psi(x) &= -\frac{\hbar}{2m}\psi''(x) + V(x)\psi(x) \\ &= -\frac{\hbar}{2m}\psi''(x) + \frac{1}{2}m\omega^2x^2\psi(x).\end{aligned}$$

This equation can be simplified using nondimensionalization. Using Hartree atomic units (i.e., measuring the energy in units of $\hbar\omega$ and the distance in units of $\sqrt{\hbar/m\omega}$), the equation above becomes

$$\hat{H}\psi(x) = -\frac{1}{2}\psi''(x) + \frac{1}{2}x^2\psi(x). \quad (6.3)$$

6.2 Analytical solutions

It can be easily verified that the ground state solution to the this problem's Schrödinger Equation

$$\hat{H}\psi = E_0\psi,$$

is the pair

$$E_0 = 0.5, \quad \psi(x) = \frac{1}{\pi^{1/4}} \cdot e^{-\frac{x^2}{2}},$$

where $1/\pi^{1/4}$ is a normalization constant (see Figure 2.2.1).

What about excited states? The following derivation is based on [45]. Define the so-called *ladder operators*

$$a := \frac{1}{\sqrt{2}} \left(\frac{d}{dx} + x \right) \text{ and } a^* := \frac{1}{\sqrt{2}} \left(-\frac{d}{dx} + x \right).$$

Note that for $\psi \in \text{dom}(\hat{H})$, it holds

$$\begin{aligned}a^*(a\psi) + \frac{1}{2}\psi &= \frac{1}{\sqrt{2}} \left(-\frac{1}{\sqrt{2}} (\psi'' + \psi + x\psi') + \frac{1}{\sqrt{2}} (x\psi' + x^2\psi) \right) + \frac{1}{2}\psi \\ &= -\frac{1}{2}\psi''(x) + \frac{1}{2}x^2\psi = \hat{H}\psi,\end{aligned}$$

and therefore

$$\hat{H} = a^*a + \frac{1}{2} \quad (6.4)$$

This calculation assumes that ψ is in the domain of a and a^* as well as in the domain of $\hat{H}$. This is not a trivial assumption, but it can be shown that the calculation above is indeed valid ([11], Section 11.4). With a similar calculation, one also obtains

$$aa^* = a^*a + 1. \quad (6.5)$$

Observation 6.4 is key to the proof of the following theorem, which provides the remaining

eigenvalues and eigenfunctions for $\hat{H}$.

Theorem 6.1. *Let $n \geq 0$. The n -th eigenfunction of the Hamiltonian operator 6.3 is given by*

$$\psi_0(x) = \frac{1}{\pi^{1/4}} \cdot e^{-\frac{x^2}{2}} \text{ and } \psi_n(x) = \frac{1}{\sqrt{n!}} a^* \psi_{n-1}(x). \quad (6.6)$$

Furthermore, the corresponding eigenvalue is $E_n = n + 0.5$.

Proof. To prove the theorem, it is enough by 6.4 to prove that

$$a^* a \psi_n = n \psi_n. \quad (6.7)$$

We will prove 6.7 by induction using Equation 6.5. The case $n = 0$ holds because $a \psi_0 = 0$. Assume 6.7 holds for $n - 1$. Then

$$\begin{aligned} a^* a \psi_n &= \frac{1}{\sqrt{n!}} a^* a a^* \psi_{n-1} = \frac{1}{\sqrt{n!}} a^* (a^* a \psi_{n-1} + \psi_{n-1}) \\ &= \frac{1}{\sqrt{n!}} a^* ((n-1) \psi_{n-1} + \psi_{n-1}) = \frac{n}{\sqrt{n!}} a^* \psi_{n-1} \\ &= n \psi_n. \end{aligned}$$

□

6.3 Approximate solutions

In the following, I will use the variational method as introduced in Section 4 to approximate this solution. As a trial wave function, I will take

$$\psi_\theta(x) = N_\theta \cdot e^{-x^2},$$

where N_θ is a neural network with one inner layer of depth 5 and activation function $\tanh(x)$. Other activation functions, such as the ones depicted in Figure 3.1, did not yield satisfying results. This ansatz has been chosen with prior knowledge of the solution to speed up computation. Furthermore, multiplying the neural network by the exponential function localized the outputs and increased stability. To find the solution of

$$\min_{\theta} \langle \psi_\theta | H | \psi_\theta \rangle,$$

I will use the Adam optimizer (see Section 5.1) and the Supervised Wave Optimization (SWO) algorithm (see Section 5.2). I will compare the results in the next section.

6.3.1 Adam vs. SWO

I ran 100 experiments using the Adam optimizer and 100 experiments using the SWO algorithm. The algorithms ran for a maximum of 1000 iterations or until convergence

6 The Quantum Harmonic Oscillator

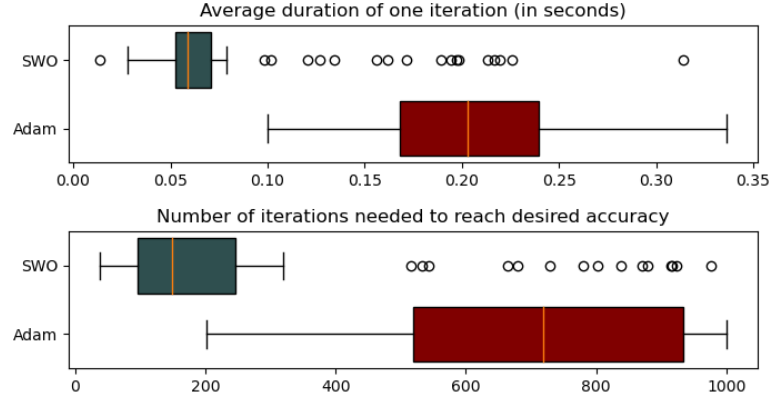


Figure 6.2: Top: Average duration of epochs over the aforementioned 100 experiments of Adam and SWO. Bottom: Number of epochs needed to converge over the same 100 experiments of Adam and SWO. Based on these parameters, SWO yields better results than Adam.

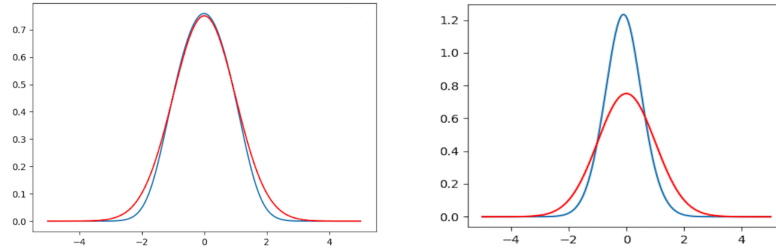


Figure 6.3: Blue: solutions found by SWO (L) and Adam (R). Red: true normalized solution (note the different scalings of the axes).

was reached. As a stopping criterion, the mean of the latest 10 values of the local energy should be less than $0.5 + \varepsilon$, where $\varepsilon = 0.001$.

In 77 out of the 100 experiments, Adam reached a solution to the desired accuracy in less than 1000 iterations. In this case, the average number of iterations needed to reach the desired accuracy and stop was 623. Each iteration took on average 0.2 seconds to complete. See Figure 6.2.

In all the experiments, SWO reached a solution to the desired accuracy in less than 1000 iterations. The average number of iterations needed to reach desired accuracy was 240. Each iteration took on average 0.08 seconds to complete. See Figure 6.2.

Remark. It is important to note that due to the introduction of a normalization factor (Equation 5.7), the solution found by using SWO is normalized. This is not the case when using the Adam minimizer (see Figure 6.3).

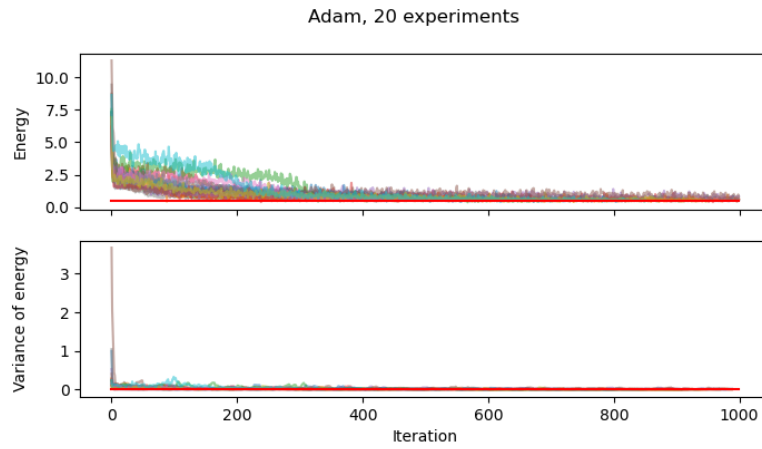


Figure 6.4: Value of the local energy (T) and of its variance (B) for 20 experiments using the Adam minimizer. The horizontal red lines represent the target energy and variance (0.5 and 0, respectively).

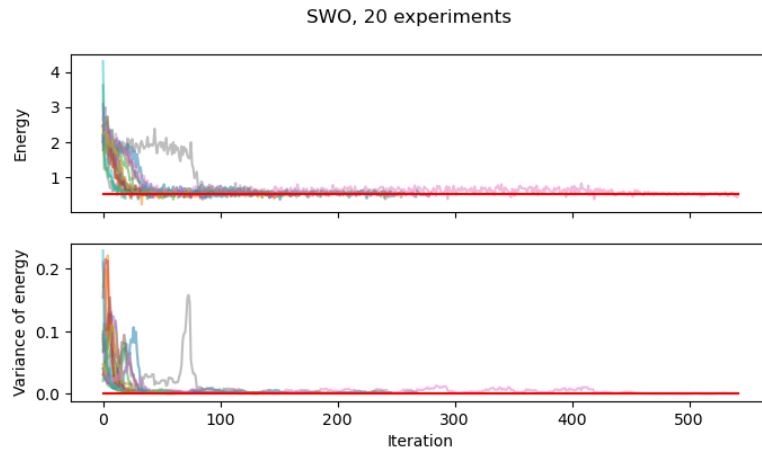


Figure 6.5: Value of the local energy (T) and of its variance (B) for 20 experiments using the SWO algorithm. The horizontal red lines represent the target energy and variance (0.5 and 0, respectively). Note the different scaling of the x-Axis compared to Figure 6.4.

7 Concluding remarks

The aim of this thesis was to investigate the performance of the SWO algorithm when implemented in the context of a deep-learning based variational Monte Carlo method to solve the Schrödinger equation. When applied to solving the simple problem of the one-dimensional quantum harmonic oscillator, the SWO algorithm performs much better than the Adam optimizer. It requires a lower number of iterations to reach convergence and each iteration is faster. Furthermore, the results obtained using the SWO algorithm are normalized, whereas the ones obtained with Adam are not.

These findings instill optimism regarding SWO's potential for success in addressing more complex problems. When doing so, its parameters should be adjusted, and it is reasonable to assume that each problem will require a separate analysis and appropriate tuning.

Bibliography

- [1] Larry Armijo. Minimization of functions having Lipschitz continuous first partial derivatives. *Pacific Journal of Mathematics*, 16(1):1 – 3, 1966.
- [2] P. W. Atkins and Ronald Friedman. *Molecular quantum mechanics*. Oxford University Press, New York, 4th ed edition, 2005.
- [3] Julius Berner, Philipp Grohs, Gitta Kutyniok, and Philipp Petersen. The modern mathematics of deep learning. In Philipp Grohs and Gitta Kutyniok, editors, *Mathematical Aspects of Deep Learning*, page 1–111. Cambridge University Press, 1 edition, 2022.
- [4] Julius Berner, Philipp Grohs, and Felix Voigtlaender. Learning relu networks to high uniform accuracy is intractable, 2023.
- [5] Max Born and Robert Oppenheimer. Zur quantentheorie der molekeln. *Annalen der Physik*, 389:457–484, 1927.
- [6] George Cybenko. Approximation by superpositions of a sigmoidal function. *Mathematics of Control, Signals, and Systems*, 1989.
- [7] P. A. M. Dirac. A new notation for quantum mechanics. *Mathematical Proceedings of the Cambridge Philosophical Society*, 35(3):416–418, July 1939.
- [8] Paul Adrien Maurice Dirac. *The Principles of Quantum Mechanics*. Clarendon Press, Oxford,, 1930.
- [9] Leon Gerard, Michael Scherbela, Philipp Marquetand, and Philipp Grohs. Gold-standard solutions to the schrödinger equation using deep learning: How much physics do we need?, 2022.
- [10] S.J. Gustafson and I.M. Sigal. *Mathematical Concepts of Quantum Mechanics*. Universitext (Berlin. Print). Springer, 2003.
- [11] Brian C. Hall. *Quantum Theory for Mathematicians*, volume 267 of *Graduate Texts in Mathematics*. Springer New York, New York, NY, 2013.
- [12] W. K. Hastings. Monte carlo sampling methods using markov chains and their applications. *Biometrika*, 57(1):97–109, 1970.
- [13] Alexander Heifetz, editor. *Quantum Mechanics in Drug Discovery*, volume 2114 of *Methods in Molecular Biology*. Springer US, New York, NY, 2020.

Bibliography

- [14] W. Heisenberg. Über den anschaulichen Inhalt der quantentheoretischen Kinematik und Mechanik. *Zeitschrift für Physik*, 43(3-4):172–198, March 1927.
- [15] Jan Hermann, Zeno Schätzle, and Frank Noé. Deep-neural-network solution of the electronic schrödinger equation. *Nature Chemistry*, 12(10):891–897, 2020.
- [16] Jan Hermann, James Spencer, Kenny Choo, Antonio Mezzacapo, W. M. C. Foulkes, David Pfau, Giuseppe Carleo, and Frank Noé. Ab-initio quantum chemistry with neural-network wavefunctions, 2022.
- [17] Kurt Hornik, Maxwell Stinchcombe, and Halber White. Multilayer feedforward networks are universal approximators. *Neural Networks*, 2:359–366, 1989.
- [18] Patrick Kidger and Terry Lyons. Universal approximation with deep narrow networks, 2020.
- [19] Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization, 2017.
- [20] Charles Kittel. *Introduction to solid state physics*. Wiley, Hoboken, NJ, 8th ed edition, 2005.
- [21] Dmitrii Kochkov and Bryan K. Clark. Variational optimization in the ai era: Computational graph states and supervised wave-function optimization, 2018.
- [22] Ajit Kumar. *Fundamentals of Quantum Mechanics*. Cambridge University Press, India, 2018.
- [23] Warren S. McCulloch and Walter Pitts. A logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematical Biophysics*, 5:115–133, 1943.
- [24] Nicholas Metropolis, Arianna W. Rosenbluth, Marshall N. Rosenbluth, Augusta H. Teller, and Edward Teller. Equation of state calculations by fast computing machines. *The Journal of Chemical Physics*, 21(6):1087–1092, 1953.
- [25] Mehryar Mohri, Afshin Rostamizadeh, and Ameet Talwalkar. *Foundations of machine learning*. Adaptive computation and machine learning. The MIT Press, Cambridge, Massachusetts, second edition edition, 2018.
- [26] Kevin P. Murphy. *Probabilistic machine learning: an introduction*. Adaptive computation and machine learning series. The MIT Press, Cambridge, Massachusetts, 2022.
- [27] Michael A. Nielsen and Isaac L. Chuang. *Quantum Computation and Quantum Information*. Cambridge University Press, 2000.
- [28] David Pfau, James S. Spencer, Alexander G. de G. Matthews, and W. M. C. Foulkes. Ab-initio solution of the many-electron schrödinger equation with deep neural networks. *Physical Review Research*, 2(3):033429, 2020.

- [29] Lucjan Piel. *Ideas of Quantum Chemistry*. Elsevier Science, 1st edition, 2007.
- [30] Allan Pinkus. Approximation theory of the mlp model in neural networks. *Acta Numerica*, 8:143–195, Jan 1999.
- [31] M. Reed and B. Simon. *Methods of Modern Mathematical Physics. IV Analysis of Operators*. Academic Press, New York, 1978.
- [32] James C. Robinson. *An Introduction to Functional Analysis*. Cambridge University Press, 2020.
- [33] W. Rudin. *Principles of Mathematical Analysis*. International series in pure and applied mathematics. McGraw-Hill, 1976.
- [34] W. Rudin. *Functional Analysis*. International series in pure and applied mathematics. McGraw-Hill, 1991.
- [35] David E. Rumelhart, Geoffrey E. Hinton, and Ronald J. Williams. Learning representations by back-propagating errors. *Nature*, 323:533–536, 1986.
- [36] Leonard I. Schiff. *Quantum Mechanics*. McGraw-Hill Book Company, 3rd edition, 1968.
- [37] Erwin Schrödinger. Quantisierung als eigenwertproblem. *Annalen der Physik*, 384:361–376, 1926.
- [38] Shai Shalev-Shwartz and Shai Ben-David. *Understanding Machine Learning: From Theory to Algorithms*. Cambridge University Press, 1st edition, 2014.
- [39] Attila Szabo and Neil S. Ostlund. *Modern Quantum Chemistry: Introduction to Advanced Electronic Structure Theory*. Dover Publications, 1989.
- [40] Gerald Teschl. *Mathematical Methods in Quantum Mechanics*. American Mathematical Society, 2009.
- [41] W. Thirring. *Quantum Mathematical Physics*. Physics and astronomy online library. Springer, 2002.
- [42] John von Neumann. *Mathematical Foundations of Quantum Mechanics*. Princeton University Press, new edition, 1955.
- [43] Dirk Werner. *Funktionalanalysis*. Springer-Lehrbuch. Springer Berlin Heidelberg, Berlin, Heidelberg, 2018.
- [44] Bernard Widrow and Marcian E. Hoff. Adaptive switching circuits. In *1960 IRE WESCON Convention Record, Part 4*, pages 96–104, New York, 1960. IRE.
- [45] Harry Yserentant. *Regularity and Approximability of Electronic Wave Functions*, volume 2000 of *Lecture Notes in Mathematics*. Springer Berlin Heidelberg, Berlin, Heidelberg, 2010.