

Logical foundations of physics. Resolution of classical and quantum paradoxes in the finitistic paraconsistent logic NAFL

Radhakrishnan Srinivasan*
302/A4, Lokmilan CHS Ltd.
Chandivali Farm Road
Mumbai 400072, India
rk_srinivasan@yahoo.com
rk_srinivasan@outlook.com

Abstract

Non-Aristotelian finitary logic (NAFL) is a finitistic paraconsistent logic that redefines finitism and correctly captures the notion of a potential infinity. Classical infinitary reasoning is refuted in NAFL, with the consequent negative resolution of Hilbert's program. It is argued that the existence of nonstandard models of arithmetic is an artifact of infinitary classical semantics, which must be rejected by the finitist, for whom the meaning of "finite" is not negotiable. The main postulate of NAFL semantics defines formal truth as time-dependent axiomatic declarations of the human mind, a consequence of which is the following metatheorem. If the axioms of an NAFL theory T are pairwise consistent, then T is consistent. This metatheorem, which is the more restrictive counterpart of the compactness theorem of classical first-order logic, leads to the conclusion that T supports only constructive existence, and consequently, nonstandard models of T do not exist, which in turn implies that infinite sets cannot exist in consistent NAFL theories. It is shown that arithmetization of syntax, Godel's incompleteness theorems and Turing's undecidability of the halting problem, which lead classically to nonstandard models, cannot be formalized in NAFL theories. The NAFL theories of arithmetic and real numbers are defined. Several paradoxical phenomena in quantum mechanics, such as, quantum superposition, Wigner's friend paradox, entanglement, the quantum Zeno effect and wave-particle duality, are shown to be justifiable in NAFL, which provides a logical basis for the incompatibility of quantum mechanics and infinitary (by the NAFL yardstick) relativity theory. Finally, Zeno's dichotomy paradox and its variants, which lead to meta-inconsistencies in classical infinitary reasoning, are shown to be resolvable in NAFL.

*Permanent address (for correspondence): 1102 Tejas Heights, Plot 245, Mumbai Tamil Sangam Road, Sion East, Mumbai 400022, India.

1 The finitist's objection to infinite sets / classical model theory

We start with an argument for why the finitist cannot accept the existence of nonstandard models of arithmetic, and consequently, must reject classical model theory [1, 2] and the existence of infinite sets. Let $\mathbb{N} = \{0, 1, \dots\}$ be the set of natural numbers. That every natural number has been exceeded within $\mathbb{N}$ is expressed by the following sentence:

$$\forall n \exists m m > n. \quad (1)$$

Infinitely many natural numbers have not been exceeded within $\mathbb{N}$, or equivalently, there is no infinitely large natural number in $\mathbb{N}$, as seen from:

$$\neg \exists m \forall n m \geq n. \quad (2)$$

Eqs. (1) and (2) are theorems of first-order Peano arithmetic (**PA**) [3]. If we were to interchange the quantifiers in (1) and syntactically deduce the existence of an infinitely large natural number, we would be committing a fallacy. Instead let us proceed model-theoretically and consider a well-known method for deducing the existence of nonstandard models of arithmetic. Let c be a constant in the language of **PA**. Break up (1) into infinitely many axioms, as follows:

$$(\exists m m = c) \wedge c > n, \quad \text{where } n = 0, 1, \dots \quad (3)$$

Add these axioms to **PA**, to obtain the theory $\mathbf{PA}' = \mathbf{PA} + (3)$. It is a standard result that if **PA** is consistent, every finite subset of $\mathbf{PA}'$ is consistent. From the compactness theorem of first-order logic, it follows that if **PA** is consistent, then $\mathbf{PA}'$ is consistent. Hence if **PA** is consistent, there must exist a model of $\mathbf{PA}'$, which is also a model of **PA**, in which the constant c exceeds every standard natural number in $\mathbb{N}$. We have deduced the existence of a nonstandard model of **PA** in which c is dubbed as a *nonstandard finite* natural number.

Currently, **PRA** (Primitive recursive arithmetic), which is a weak subsystem of **PA**, is widely accepted as a theory that correctly captures the principles of finitism [4]. The existence of nonstandard models of arithmetic can also be proven from Gödel's first incompleteness theorem [5], which can be formalized in **PRA** and is therefore considered to be finitistically valid. So currently, the existence of nonstandard natural numbers is presumably considered to be acceptable to the finitist.

In this paper, we take the stand that the finitist must consider the number c in (3) as infinite. Note that any theory that proves the compactness theorem must recognize (3) as an *infinite* set of axioms. When the compactness theorem is applied to every *finite* subset of the *infinite* set of axioms in (3), *finite* and *infinite* have already been defined in the standard sense and the number c is infinite by this yardstick. The finitist recognizes c as an absolute (or “completed”) infinity. In other words, for the finitist, $c + 1 = c$ (which implies that c is a maximum for the natural numbers), in contradiction to (1) and (2), which

are theorems of **PA**; hence c cannot exist in a model of **PA**. Those who have strong faith in classical model theory, in the consistency of **PA** and in infinitary reasoning in general, will perhaps conclude from this contradiction that non-standard models of **PA** must exist. The finitist, however, ought to draw the straightforward conclusion that **PA**, and even **PRA**, are inconsistent theories. What can be the source of such an inconsistency? Consider (3) with the axioms for c removed:

$$\exists m \, m > n, \quad \text{where } n = 0, 1, \dots \quad (4)$$

Note that the compactness theorem guarantees that the infinitely many theorems of **PA** in (4) must have a model, with the consequence that infinitely many natural numbers n have been exceeded within $\mathbb{N}$, which is why we are able to postulate the existence of c in (3). But “infinitely many natural numbers have been exceeded within $\mathbb{N}$ ” means, even grammatically, that there exists an infinitely large natural number in $\mathbb{N}$, in contradiction to (2). The argument that c is “externally infinite” but “internally finite” in nonstandard theories/models ought not to be acceptable to the finitist, for whom the meaning of *finite* is not negotiable.

We stress that this inconsistency is deduced from the truth of infinitely many sentences (4) in the standard model of **PA**. This model-theoretic derivation does not have any syntactic equivalent because a proof in **PA** can contain only finitely many sentences. Within the theory **PA**, what we can deduce from (1) and (2) are, respectively, the finitistically acceptable conclusions that any given natural number has been exceeded within $\mathbb{N}$, but not infinitely many natural numbers. The finitist must therefore reject classical model theory and the existence of $\mathbb{N}$ as an infinite set. But this does not commit the finitist to ultrafinitism. In the ensuing sections, we will demonstrate that the finitist can indeed accept (1) and (2) via non-Aristotelian finitary logic (NAFL), [6] in which $\mathbb{N}$ is an infinite proper class and the model theory (which associates truth with provability in NAFL theories) is nonclassical. The nonconstructive existence of the constant c in (3) is not permitted in NAFL theories, as will be proven in the ensuing section. It is not possible to deduce the existence of nonstandard natural numbers either syntactically or model-theoretically in NAFL. The NAFL model of (1) correctly captures the theoretical / syntactic notion that “given any natural number n , there exists a natural number m greater than n ”, without actually enumerating infinitely many instances of (1), as does a classical model of (1). In other words, NAFL semantics correctly captures the notion of a potential, rather than an actual, infinity.

Remark 1 (Potential infinity). *In recent papers, [7, 8, 9] the authors have claimed that there exist notions of potential infinity that are justifiable in classical first-order logic (FOL). Linnebo and Shapiro [7] conclude that the Aristotelian notion of potential infinity, which requires modal arithmetic, is correctly captured in the FOL theory **PA**. Simpson [8] argues that potential infinity is justifiable in certain formal systems that are conservative extensions of the FOL theory **PRA**. These formal systems are powerful in the sense that they account for much of mathematics, including classical analysis, algebra and geometry.*

Similarly, Eberl [9] has developed a model theory that claims to justify potential infinity in FOL theories stronger than **PA**, which formalize much of mathematics.

*In the opinion of the author, all of the above claims concerning potential infinity are flawed, for the following reason. The FOL theories that these authors use, namely, **PRA**, **PA** and stronger theories, do not even capture the notion of “finite”, in the sense that they admit nonstandard models. But the notions of potential infinity that these authors embrace ought not to accept nonstandard natural numbers as finite. The potentialist ought to reject the existence of the nonstandard natural number c defined by the infinitely many axioms (3), because the definition of c requires the existence of an actual (or completed) infinity of standard natural numbers. In other words, for the potentialist, “finite” ought to mean “standard finite”. If indeed the potentialist rejects the existence of c , then, as argued in this section, **PRA** and **PA** are inconsistent theories. Therefore one is inevitably led to the logic NAFL for an acceptable definition of potential infinity.*

2 NAFL semantics

Our goal is to replace the infinitary semantics of classical first-order logic with finitary NAFL semantics, while retaining the finitistically unproblematic part of classical syntax. We will break with tradition and describe NAFL semantics first, as this will better motivate the modifications/restrictions that are required in classical syntax in order to obtain NAFL syntax.

Classically, truths are pre-existing and the purpose of logic is to discover these truths and describe them consistently. This, of course, is the philosophy of Platonism, which makes classical model theory and indeed, any nontrivial classical reasoning, unavoidably infinitary.

2.1 The main postulate of NAFL semantics

In the logic NAFL, formal truths exist only with respect to axiomatic theories, and there are no pre-existing truths in just the language of these theories. Let **T** be a consistent NAFL theory which resides temporarily in the human mind and let P be a sentence in the language of **T**. Let **T*** be either the theory **T** or a consistent extension of **T**, which the human mind creates by (possibly) adding axioms to **T**. Note that a sentence that is provable in **T** is also provable in **T***, and a sentence that is undecidable (*i.e.*, neither provable nor refutable) in **T*** is also undecidable in **T**. Here **T*** is defined as the *interpretation* of **T**, which basically generates an NAFL model of **T**, as follows. The provable sentences of **T*** are defined as true with respect to **T**. In particular, if P is provable (refutable) in **T***, then P is true (false) w.r. to **T**. The paraconsistency of NAFL enters via sentences that are undecidable in **T***, which are defined as “neither true nor false” w.r. to **T**. If P is undecidable in **T***, a nonclassical NAFL model of **T** is generated in which $P \wedge \neg P$ is true, *i.e.*, both P and $\neg P$

are true, where “ P ” is interpreted as “ $\neg P$ is not provable in $\mathbf{T}^*$ ” and “ $\neg P$ ” is interpreted as “ P is not provable in $\mathbf{T}^*$ ”. Clearly, when P is undecidable in $\mathbf{T}$, such a nonclassical model of $\mathbf{T}$ can exist if and only if the law of the excluded middle ($P \vee \neg P$) and the law of noncontradiction ($\neg(P \wedge \neg P)$) are not theorems of $\mathbf{T}$. Note that $\mathbf{T}$ does not *prove* $P \wedge \neg P$; any NAFL theory that proves a contradiction is inconsistent, just as in classical logic. However, the notion of consistency for an NAFL theory $\mathbf{T}$ is more restrictive, in the sense that $\mathbf{T}$ must admit nonclassical models for undecidable sentences. In this paper, the terminology “consistent” / “inconsistent” for an NAFL theory is always used in the NAFL sense; if we wish to refer to the weaker classical notion of consistency, we will explicitly mention that.

2.2 Comments on the main postulate

Formal truths, which are only created when the human mind constructs NAFL theories, are essentially axiomatic assertions of the human mind. There *are* pre-existing (Platonic) truths in NAFL, but these are truths *about* NAFL theories and cannot be formalized within them. E.g., the notion of provability is not formalizable in NAFL theories. That a sentence P is either decidable or undecidable in an NAFL theory $\mathbf{T}$ is taken as a Platonic truth about $\mathbf{T}$ that is independent of the human mind, irrespective of whether a proof of the said decidability/undecidability exists. Note that the interpretation $\mathbf{T}^*$ completely specifies the NAFL model of $\mathbf{T}$ via the finitary concept of provability, without collecting together all the truths in an infinite set, as in a classical model. The only classical truths in the NAFL model of $\mathbf{T}$ are the sentences decidable in $\mathbf{T}^*$, and all other sentences are in a nonclassical, superposed state of “neither true nor false”.

It should be emphasized that a given interpretation $\mathbf{T}^*$, and hence the truths in the NAFL model of $\mathbf{T}$, have only a temporary existence in the human mind. At different times, the same human being could consider different theories $\mathbf{T}^*$ as the interpretation of $\mathbf{T}$, and at a given time, different human beings could have different interpretations $\mathbf{T}^*$ in mind. Thus NAFL semantics admits time-dependent truths and $\mathbf{T}^*$ is chosen according to the free will of the human mind, which axiomatically asserts the classical $\mathbf{T}$ -truths via provability in the interpretation $\mathbf{T}^*$. The nonclassical $\mathbf{T}$ -truths have a different interpretation as noted above, and are not axiomatic assertions because $\mathbf{T}^*$ does not prove a contradiction. Note that an NAFL theory $\mathbf{T}$ could have a classically “false” axiom (e.g., “The sun rises in the west”), which is a *true* statement with respect to $\mathbf{T}$ via provability in $\mathbf{T}$. NAFL rejects Platonism, and there is no obligation to keep formal NAFL truths in conformity with any pre-existing “reality”.

To obtain an NAFL theory of physics, e.g., quantum mechanics ($\mathbf{QM}$), the human mind chooses the axioms of $\mathbf{QM}$ and its interpretation $\mathbf{QM}^*$ such that the theorems of $\mathbf{QM}^*$ agree with real-life *observations*, e.g., made via experiments. Thus the proposition “The Schrödinger cat is alive (dead)”, which is undecidable in $\mathbf{QM}$, could be an axiom of $\mathbf{QM}^*$ if the cat has been *observed* to be alive (dead) in a real-life experiment. Therefore, even in a physical NAFL

theory, formal truths are axiomatic in nature and the connection with reality comes indirectly, via the human mind.

2.3 Metatheorems of NAFL semantics

Important metatheorems, which capture the essence of the NAFL philosophy of finitism, are derivable from the main postulate of NAFL semantics, and these will guide us on how NAFL syntax must be formulated. In what follows, “ $\mathbf{T} \vdash P$ ” means “ $\mathbf{T}$ proves P ”, where P is a proposition or a sentence.

Metatheorem 1. *If $\mathbf{T}$ is a consistent NAFL theory and if $\mathbf{T} \vdash (P \vee Q)$, then either $\mathbf{T} \vdash P$ or $\mathbf{T} \vdash Q$, i.e., both P and Q cannot be undecidable in $\mathbf{T}$.*

Proof. Suppose $\mathbf{T} \vdash (P \vee Q)$ with P and Q undecidable in $\mathbf{T}$. Consider an NAFL model of $\mathbf{T}$ generated by the interpretation $\mathbf{T}^* = \mathbf{T}$. The main postulate of NAFL semantics requires that if $\mathbf{T}$ is consistent, such an NAFL model of $\mathbf{T}$ must exist. In this model $P \vee Q$ has the classical truth value of “true”, while each of P and Q has the nonclassical truth value of “neither true nor false”. But this contradicts the fact that $P \vee Q$ can be classically true if and only if at least one of the sentences P and Q has the classical truth value of “true”. $\square$

Corollary 1. *If $\mathbf{T}$ is a consistent NAFL theory and if $\mathbf{T} \vdash (P \rightarrow Q)$ or if $\mathbf{T} \vdash \neg(P \wedge Q)$, then both P and Q cannot be undecidable in $\mathbf{T}$.*

Proof. Follows from Metatheorem 1, because $(P \rightarrow Q) \leftrightarrow (\neg P \vee Q)$ and $\neg(P \wedge Q) \leftrightarrow (\neg P \vee \neg Q)$. $\square$

Remark 2. *It follows from Corollary 1 that the classical inference rule corresponding to conditional proof (see Sec. 3.1.1), namely,*

$$((\mathbf{T} + P) \vdash Q) \vdash (\mathbf{T} \vdash (P \rightarrow Q)),$$

fails in NAFL when P and Q are undecidable in $\mathbf{T}$. A classical proof of $P \rightarrow Q$ in the theory $\mathbf{T}$ would start with a hypothesis P that leads to the conclusion Q , following which the hypothesis P is discharged and $P \rightarrow Q$ is inferred. The logic NAFL, wherein truth is axiomatic in nature, requires that the hypothesis P cannot be discharged because it is fixed as an axiomatic declaration, and hence, as a theorem of $\mathbf{T} + P$. The conclusion Q is also a theorem of $\mathbf{T} + P$, and $P \rightarrow Q$, equivalent to $\neg P \vee Q$, follows trivially from Q and the inference rule corresponding to disjunction introduction (see Sec. 3.1.1).

The reader should reflect on this crucial difference between classical logic and NAFL. Does the classical claim that P is a temporary hypothesis (pertaining to a pre-existing truth), but not an axiom, make sense when P is undecidable in $\mathbf{T}$? NAFL correctly rejects the classical proof in $\mathbf{T}$, the point being that formally, P is unavoidably an axiom with respect to $\mathbf{T}$ even if we have in mind a pre-existing, metamathematical truth for the hypothesis P .

Metatheorem 2 (Pairwise consistency implies consistency). *Suppose the axioms of an NAFL theory $\mathbf{T}$ are pairwise consistent, in the sense that every pair of axioms constitutes a consistent NAFL theory. Then $\mathbf{T}$ is consistent.*

Proof. Suppose three arbitrarily chosen axioms of $\mathbf{T}$, namely, P , Q and R , are inconsistent in the classical sense, *i.e.*, a contradiction is provable from $(P \wedge Q \wedge R)$. Consider an NAFL theory $\mathbf{T}'$ with the single axiom P . Observe that $\mathbf{T}'$ is consistent, by assumption of pairwise consistency and the requirement that a subtheory of a consistent theory is consistent. Clearly, $\mathbf{T}' \vdash P$ and $\mathbf{T}' \vdash (P \rightarrow \neg(Q \wedge R))$. By the modus ponens inference rule (see Sec. 3.1.1), $\mathbf{T}' \vdash \neg(Q \wedge R)$. It follows from Corollary 1 that either $\mathbf{T}' \vdash \neg Q$ or $\mathbf{T}' \vdash \neg R$. But then it follows that the axioms P , Q and R are not pairwise consistent, as required. This contradiction implies that P , Q and R are consistent in the classical sense, *i.e.*, no contradiction can be deduced from $(P \wedge Q \wedge R)$.

Consider an NAFL theory $\mathbf{T}_1$ with the pair of axioms P and Q . Note that $\mathbf{T}_1$ is consistent by assumption and define its interpretation $\mathbf{T}_1^* = \mathbf{T}_1 + R$. From the main postulate of NAFL semantics, we conclude that $\mathbf{T}_1^*$ generates an NAFL model $\mathcal{M}$ of $\mathbf{T}_1$ in which all sentences provable in $\mathbf{T}_1^*$ (including its axioms P , Q and R) are true w.r. to $\mathbf{T}_1$ and all $\mathbf{T}_1^*$ -undecidable sentences have the nonclassical truth value of “neither true nor false” w.r. to $\mathbf{T}_1$. But then it follows that $\mathcal{M}$ is also a model of $\mathbf{T}_1^*$ generated by the interpretation $\mathbf{T}_1^{**} = \mathbf{T}_1^*$. We have proved that $\mathbf{T}_1^*$ is a consistent NAFL theory. We conclude that every triple of the axioms of $\mathbf{T}$ is consistent.

Now let P , Q , R and S be arbitrarily chosen axioms of $\mathbf{T}$ and consider the theory $\mathbf{T}_2$ with the axioms $(P \wedge Q)$, R and S . The preceding arguments show that these axioms of $\mathbf{T}_2$ are pairwise consistent and hence consistent. Therefore every quadruple of the axioms of $\mathbf{T}$ is consistent.

Proceeding by induction, we conclude that every finite subset of the axioms of $\mathbf{T}$ is consistent. Since any inconsistency in $\mathbf{T}$ must be derivable from a finite subset of its axioms, it follows that $\mathbf{T}$ is consistent. $\square$

Corollary 2 (Constructive existence). *Metatheorem 2 implies that, in general, classical nonconstructive existence is not permitted, and hence arbitrary constants do not exist, in consistent NAFL theories.*

Proof. In the language of an NAFL theory $\mathbf{T}$ with equality, let c_1 , c_2 and c_3 be constant symbols, and let $\mathbf{T}$ include the following three axioms:

$$c_1 = c_2, \quad c_2 = c_3, \quad c_3 \neq c_1. \quad (5)$$

Clearly these three axioms are inconsistent, but pairwise consistent in the classical sense. This is so because classical logic permits c_1 , c_2 and c_3 to be specified nonconstructively, *i.e.*, they are arbitrary constants. However, we know from Metatheorem 2 that an inconsistent set of axioms of an NAFL theory cannot be pairwise consistent. It follows that consistent NAFL theories do not permit the (nonconstructive) existence of arbitrary constants. $\square$

See Metatheorem 6 (Sec. 3.2.3) for a direct proof of the above result.

Remark 3. *If $\mathbf{T}$ is a consistent NAFL theory of arithmetic that proves the existence of the constant c , then Corollary 2 requires that $\mathbf{T}$ must provide a construction for c , which necessarily implies that c is a standard natural number.*

It follows that the axioms in (5) are not pairwise consistent in the NAFL sense, because c_1 , c_2 and c_3 are arbitrary constants. It is easy to see that if the axioms in (5) are specified with constructions for c_1 , c_2 and c_3 (e.g. $(c_1, c_2, c_3) = (1, 1, 2)$), then they will not be pairwise consistent even in the classical sense.

Metatheorem 3. *A consistent NAFL theory of arithmetic does not admit nonstandard models, i.e., nonstandard natural numbers do not exist in NAFL.*

Proof. Consider the infinitely many axioms in (3), which classically prove the existence of a nonstandard natural number c . These axioms are classically pairwise consistent because c has been specified nonconstructively. However, from Corollary 2, it follows that any consistent NAFL theory of arithmetic that proves the existence of a constant c must also prove that c is a standard natural number with a specific construction, e.g., $c = 100$ (see Remark 3). We conclude that by the NAFL yardstick, the axioms (3) are inconsistent and not pairwise consistent, because the axiom corresponding to $n = c$ would be the contradiction $c > c$. $\square$

Remark 4. *In NAFL, truth for the existence of a constant c is established constructively via provability in a theory $\mathbf{T}$, whereas classically, such a truth can be pre-existing and imposed nonconstructively on $\mathbf{T}$ from outside. This is the main reason why NAFL, unlike classical logic, does not permit c to be a nonstandard natural number.*

Remark 5. *It follows from Metatheorem 3 that Gödel's incompleteness theorems, which imply the existence of nonstandard models of arithmetic, cannot be formalized in consistent NAFL theories. The same applies to Turing's proof of the unsolvability of the halting problem. We will explain this in detail while considering the syntax of NAFL theories.*

Remark 6. *Note that an NAFL theory of arithmetic will prove infinitely many instances of (1), as in (4). The main postulate of NAFL semantics, wherein truth is equated with the finitary concept of provability, ensures that it is not possible to deduce a contradiction (e.g., the existence of nonstandard natural numbers) from the infinitely many sentences in (4).*

Metatheorem 4. *Infinite sets do not exist in consistent NAFL theories.*

Proof. In set theory, define the infinite sets E_n as follows:

$$(\forall n \in \mathbb{N}) E_n = \{n, n + 1, \dots\},$$

where $\mathbb{N} = \{0, 1, \dots\}$. It is provable that every finite intersection of these sets is nonempty:

$$(\forall n \in E_1) \bigcap_{j=0}^n E_j = E_n \neq \emptyset. \quad (6)$$

It is also provable that the infinite intersection is empty:

$$\bigcap_{j=0}^{\infty} E_j = \emptyset. \quad (7)$$

The finite intersections in (6) essentially state that removing the elements $(0, 1, \dots, n-1)$ from the set $\mathbb{N}$ leaves the residue E_n . Therefore we may interpret (6) as “Each finite set $\{0, 1, \dots, n-1\}$ has been exceeded within $\mathbb{N}$ by the infinitely many natural numbers in E_n ”. Similarly, we may interpret (7) as “Infinitely many natural numbers have not been exceeded within $\mathbb{N}$ ” (or equivalently, “ $\mathbb{N}$ does not contain an infinitely large natural number”). In the logic NAFL, wherein nonstandard models of arithmetic do not exist, these are contradictory interpretations, because (6) essentially expresses the fact that infinitely many natural numbers have been exceeded within $\mathbb{N}$.

This contradiction occurs because the language of a theory which admits infinite sets has the expressive power to state infinitely many instances of (1) (as given by (4)) in a single provable sentence, namely, (6). In the absence of nonstandard models of arithmetic (see Metatheorem 3), the only possible interpretation of (6) in NAFL is the contradiction that infinitely many natural numbers have been exceeded within $\mathbb{N}$. As noted in Remark 6, NAFL semantics avoids this contradiction in arithmetic, wherein infinitely many provable sentences of (4) are required to obtain the equivalent of (6). Therefore, by the NAFL yardstick, (6) and (7) are illegal (infinitary) sentences that contradict each other. It follows that infinite sets cannot exist in consistent NAFL theories. $\square$

Remark 7. *Later, in Metatheorem 8, we will give a syntactic proof of the nonexistence of infinite sets in theories formulated in NFOL (the NAFL version of first-order logic). The very definition of infinite sets requires infinitely many instances of (1) and hence it is not surprising that nonstandard models of infinite sets must necessarily exist. Further, the existence of nonstandard models of arithmetic can be formalized and proven in suitable set theories that admit infinite sets. Therefore Metatheorem 3 directly implies Metatheorem 4. We will later see that NAFL theories do admit infinite (proper) classes, but to avoid sentences like (6) and (7), quantification over infinite classes must be banned. Further, there are no variables that range over infinite classes (i.e., arbitrary infinite classes do not exist) in the language of NAFL theories. Infinite classes are constants that must be defined constructively.*

Remark 8 (The conflict between NAFL and classical logic). *Classical model theory requires infinite sets to exist. Working backwards, one can see that all the metatheorems of NAFL semantics are violated by the existence of infinite sets and hence the main postulate of NAFL semantics is classically false. The conflict between NAFL and classical logic is then reduced to whether one ought to accept finitism via the main postulate or accept the classical Platonic existence of infinite sets.*

3 The syntax of NAFL theories

An NAFL theory **NT** has two levels of syntax, namely, the proof syntax (p-syntax) and the theory syntax (t-syntax). The language, well-formed formulas

and rules of inference of $\mathbf{NT}$ are, in general, the same as those of classical logic, with possibly some additional restrictions. The p-syntax determines all the theorems and the undecidable sentences that follow from the axioms of a classical theory $\mathbf{T}$. The t-syntax, which specifies the axioms, theorems and undecidable sentences of the corresponding NAFL theory $\mathbf{NT}$, contains only a proper subset of the theorems of $\mathbf{T}$.

For example, if P and Q are legitimate sentences in the t-syntax of $\mathbf{NT}$ that are undecidable in $\mathbf{T}$ (and hence, undecidable in $\mathbf{NT}$) and if the p-syntax contains a proof of $P \vee Q$, then the t-syntax will exclude $P \vee Q$ (*i.e.*, $P \vee Q$ is not a legitimate sentence in the t-syntax), because Metatheorem 1 requires that $P \vee Q$ cannot be a theorem of $\mathbf{NT}$. In particular, the law of the excluded middle $P \vee \neg P$ and equivalently, the law of noncontradiction $\neg(P \wedge \neg P)$, are not theorems of $\mathbf{NT}$ when P is undecidable in $\mathbf{NT}$. The basic idea is that sentences like $P \vee Q$, which are interpreted classically in the p-syntax, may occur in proofs of the theorems of $\mathbf{NT}$, but the reverse implication does not go through, *i.e.*, $\mathbf{NT}$ does not imply $P \vee Q$, because $\mathbf{NT}$ interprets P and Q in a nonclassical sense, as noted in the main postulate of NAFL semantics.

At first sight, it may seem strange, even contradictory, that an $\mathbf{NT}$ -undecidable sentence like P is treated classically in the p-syntax and nonclassically in the t-syntax. In particular, if Q is a legitimate sentence in the t-syntax of $\mathbf{NT}$ and if $Q \rightarrow (P \wedge \neg P)$ is provable in the p-syntax of $\mathbf{NT}$, then that constitutes a proof by contradiction of $\neg Q$ in $\mathbf{NT}$, *despite* the failure of the law of noncontradiction for P in the t-syntax. This is explained as follows. Note that the main postulate of NAFL semantics requires that formal truths exist only with respect to axiomatic theories. There is nothing intrinsically nonclassical about P , which may have a classical truth value with respect to another theory that decides P . To deduce the nonclassical nature of P with respect to the theory $\mathbf{NT}$, one first needs to determine that P is undecidable in $\mathbf{NT}$, and to avoid circularity, the p-syntax of $\mathbf{NT}$ cannot possibly include the notion of $\mathbf{NT}$ -undecidability of P and the consequent failure of the law of noncontradiction. It is only *after* P is determined to be undecidable in $\mathbf{NT}$ that P loses its classical meaning with respect to $\mathbf{NT}$.

3.1 Propositional logic

Let NPL denote the NAFL version of propositional logic that uses a natural deduction system. Let $\mathbf{T}$ be a theory in classical propositional logic and let $\mathbf{NT}$ be the corresponding NPL theory. Here we define a theory as a set of nonlogical axioms.

3.1.1 The p-syntax of the NPL theory $\mathbf{NT}$

The p-syntax of $\mathbf{NT}$ determines all the theorems that can be deduced from the axioms of the classical theory $\mathbf{T}$. We start with the description of classical

propositional logic [3], which is a formal system $\mathcal{L} = \mathcal{L}(A, \Omega, Z, I)$, where A , Ω , Z and I are defined as follows, for a natural deduction system.

- The set A is a countably infinite set of symbols that denote atomic formulas (propositional variables), e.g.:

$$A = \{p_1, q_1, r_1, p_2, q_2, r_2, \dots\}.$$

- The set of logical connectives is defined as $\Omega = \Omega_1 \cup \Omega_2$, where

$$\Omega_1 = \{\neg\}, \quad \Omega_2 = \{\vee, \wedge, \rightarrow, \leftrightarrow\}.$$

- The set of logical axioms is empty: $I = \emptyset$.
- The set Z is defined by the following eleven inference rules, where $p, q, r, \dots$ denote formulas:

$$\{p \rightarrow q, p \rightarrow \neg q\} \vdash \neg p \quad (\text{negation introduction}),$$

$$\neg p \vdash (p \rightarrow r) \quad (\text{negation elimination}),$$

$$\neg \neg p \vdash p \quad (\text{double negation elimination}),$$

$$\{p, q\} \vdash (p \wedge q) \quad (\text{conjunction introduction}),$$

$$(p \wedge q) \vdash p \text{ and } (p \wedge q) \vdash q \quad (\text{conjunction elimination}),$$

$$p \vdash (p \vee q) \text{ and } q \vdash (p \vee q) \quad (\text{disjunction introduction}),$$

$$\{p \vee q, p \rightarrow r, q \rightarrow r\} \vdash r \quad (\text{disjunction elimination}),$$

$$\{p \rightarrow q, q \rightarrow p\} \vdash (p \leftrightarrow q) \quad (\text{biconditional introduction}),$$

$$(p \leftrightarrow q) \vdash (p \rightarrow q) \text{ and } (p \leftrightarrow q) \vdash (q \rightarrow p) \quad (\text{biconditional elimination}),$$

$$\{p, p \rightarrow q\} \vdash q \quad (\text{modus ponens}),$$

$$(p \vdash q) \vdash (p \rightarrow q) \quad (\text{conditional proof}).$$

- In addition to the logical rules of inference given above, there may be nonlogical rules of inference [10] that are specific to a given NPL theory **NT**. See Secs. 3.1.2 and 3.1.3.

The well-formed formulas of the language $\mathcal{L}$ are inductively defined as follows:

- Any atomic formula belonging to the set A is a formula of $\mathcal{L}$.
- If p is a formula, then $\neg p$ is a formula.
- if p and q are formulas, then $(p \rightarrow q)$, $(p \leftrightarrow q)$, $(p \vee q)$ and $(p \wedge q)$ are formulas.
- There are no other formulas.

One can use the above inference rules to deduce all the theorems of **T** from the set of nonlogical axioms (which can also be empty) of a classical theory **T** in propositional logic. This completes the definition of the p-syntax of the corresponding NPL theory **NT**.

3.1.2 The t-syntax of the NPL theory **NT**

The t-syntax of **NT** defines the axioms, theorems and undecidable propositions of **NT**. The theorems of **NT** are a proper subset of the classical theorems of **T** deduced in the p-syntax. The following metatheorem provides important guidance on the definition of the t-syntax of NPL theories.

Metatheorem 5. *Let P and Q be well-formed formulas of the language $\mathcal{L}$ of propositional logic. A proposition of the form $P \vee Q$, including propositions like $P \rightarrow Q$ and $\neg(P \wedge Q)$, is not a legitimate axiom in the t-syntax of consistent NPL theories.*

Proof. Consider an NPL theory **NT** with the single axiom $P \vee Q$ in its t-syntax. Here P and Q are undecidable in **NT** and from Corollary 1, it follows that **NT** is an inconsistent theory. Further, any NPL theory containing $P \vee Q$ as one of its axioms is inconsistent, even if either P or Q gets decided by the other axioms of this theory. This is so because of our imposed requirement that subtheories of consistent NAFL theories should be consistent, which is violated by the inconsistency of the theory **NT**. $\square$

Remark 9. *Consider an NPL theory with the following 3 axioms in its t-syntax:*

$$P \rightarrow Q, \quad Q \rightarrow \neg P, \quad P.$$

Note that these axioms are inconsistent and pairwise consistent in the classical sense. Metatheorem 2 requires that these axioms cannot be pairwise consistent in the NAFL sense, despite the fact that each pair of axioms decides either P or Q , so that Corollary 1 is not violated. Metatheorem 5 ensures that this requirement is met. Note that the proof of Metatheorem 2 uses the requirement that subtheories of consistent NAFL theories should be consistent, which is also the rationale for Metatheorem 5.

In order to define the t-syntax of **NT**, it is convenient to convert all well-formed formulas of $\mathcal{L}$ to the disjunctive normal form (DNF), which is a formula of the form:

$$\bigvee_{i=1}^n \bigwedge_{j=1}^{m_i} L_{ij},$$

where L_{ij} is a literal, defined as an atomic formula (*i.e.*, a propositional variable) or its negation. For any propositional formula A it is possible to construct an equivalent DNF B containing the same variables as A . [11] The DNF, which is a disjunction of one or more conjunctions of one or more literals (including the cases where there is a single disjunct and / or a single conjunct), is generated as follows.

- Eliminate all occurrences of $\rightarrow$ and $\leftrightarrow$ from the formula in question. Use logical equivalences, *e.g.*, the following, for this purpose:

$$P \rightarrow Q \equiv \neg P \vee Q$$

$$P \leftrightarrow Q \equiv (\neg P \vee Q) \wedge (P \vee \neg Q)$$

$$P \leftrightarrow Q \equiv (P \wedge Q) \vee (\neg P \wedge \neg Q)$$

- Move all negations inward, so that finally, negations appear only as negations of propositional variables. Use logical equivalences, e.g., the following, for this purpose:

$$\neg\neg P \equiv P.$$

$$\neg(P \wedge Q) \equiv \neg P \vee \neg Q.$$

$$\neg(P \vee Q) \equiv \neg P \wedge \neg Q.$$

- Use the distributive laws, e.g.

$$(P \wedge (Q \vee R)) \equiv ((P \wedge Q) \vee (P \wedge R))$$

$$(P \vee (Q \wedge R)) \equiv ((P \vee Q) \wedge (P \vee R))$$

The rules for constructing the t-syntax of **NT**, with all well-formed formulas assumed to be in DNF, are as follows:

- From Metatheorem 5, it follows that an axiom of **NT** cannot contain the disjunction symbol $\vee$. Therefore the only legitimate axioms of **NT** are those that are conjunctions of one or more literals.
- The theorems of **NT** are those theorems of **T** for which at least one of the disjuncts is provable in **T**. This follows from Metatheorem 1. In particular, theorems of **T** that are conjunctions of one or more literals are theorems of **NT**. If P is a well-formed formula for which each of the disjuncts is refutable in **T**, then $\neg P$ will qualify as a theorem of **NT**.
- If there are theorems of **T** that contain the disjunction symbol $\vee$ and for which at least two of the disjuncts are undecidable in **T** and other disjuncts, if present, are either undecidable or refutable in **T**, then such theorems of **T** are not in the t-syntax of **NT**, *i.e.*, they are neither theorems nor undecidable formulas of **NT**. That they are not theorems of **NT** follows from Metatheorem 1 and that they are not undecidable formulas of **NT** follows from the main postulate of NAFL semantics (see Sec. 2), as seen from the following example.
- Consider the law of the excluded middle $P \vee \neg P$, which is neither a theorem (as follows from Metatheorem 1) nor an undecidable formula of **NT** when P is a well-formed formula that is undecidable in **T**. From the main postulate of NAFL semantics, it follows that there must exist models of **NT** in which undecidable formulas are true (via provability in the interpretation **NT***), or false (via provability in **NT***), or neither true nor false (via undecidability in **NT***). Clearly, there cannot exist a model of **NT** in which $P \vee \neg P$ is classically false (via provability of $P \wedge \neg P$ in **NT***) because consistent NAFL theories, like **NT***, cannot prove contradictions. Hence $P \vee \neg P$ is not an undecidable formula of **NT** and is therefore not in the t-syntax of **NT**.

- For each nonlogical axiom of $\mathbf{T}$ that is not present and / or is not a legitimate axiom in the t-syntax of $\mathbf{NT}$, an appropriate nonlogical rule of inference [10], from which the axiom can be classically inferred, should be included in the p-syntax of $\mathbf{NT}$. For example, if P and Q are undecidable in $\mathbf{T}$ and if $P \rightarrow Q$ is an axiom of $\mathbf{T}$, then $P \rightarrow Q$ is not in the t-syntax of $\mathbf{NT}$. In this case the nonlogical rule of inference $P \vdash Q$ (“From P , infer Q ”), from which $P \rightarrow Q$ can be classically inferred, should be included in the p-syntax of $\mathbf{NT}$, so that the theory $\mathbf{NT} + P$ will prove Q . The notion of truth as provability with respect to theories in NAFL semantics (see Sec. 2) implies that “From P , infer Q ” is to be interpreted as “From a proof of P , infer Q ”, or, equivalently, “From an axiomatic assertion of P , infer Q ”. It is important to note that this nonlogical rule of inference does *not* lead to a proof of $P \rightarrow Q$ in the theory $\mathbf{NT}$, as is evident from Remark 2. Hence the nonlogical rule of inference is *not* equivalent to an assertion of the illegal axiom $P \rightarrow Q$ in the t-syntax of $\mathbf{NT}$. It is this fact which permits the existence of a nonclassical model of $\mathbf{NT}$ in which $P \wedge \neg P$ and $Q \wedge \neg Q$ hold, *i.e.*, both P and Q are nonclassically “neither true nor false”, as required by the main postulate of NAFL semantics. As noted in the proof of Metatheorem 1, the existence of such a nonclassical model is not possible if $P \rightarrow Q$ were to be asserted as an axiom of $\mathbf{NT}$.

Remark 10. *The above argument explains why NPL has a natural deduction system (see Sec. 3.1.1) in which there are no logical axioms, which, if present, would be violated by the requirement of the existence of nonclassical models of NPL theories.*

- The axioms of $\mathbf{NT}$ are those axioms of $\mathbf{T}$ that are legitimate in the t-syntax of $\mathbf{NT}$, namely, the axioms that are conjunctions of one or more literals. The theorems of $\mathbf{NT}$ may be deduced from these axioms and the nonlogical rules of inference that will replace any axioms of $\mathbf{T}$ that are not conjunctions of one or more literals.
- Consider a well-formed formula P of $\mathcal{L}$ that is not a theorem of $\mathbf{T}$, and for which at least one disjunct is undecidable in $\mathbf{T}$ and other disjuncts, if present, are either undecidable or refutable in $\mathbf{T}$. Then P is an undecidable formula of $\mathbf{NT}$. In particular, if P is a conjunction of one or more literals and if P is undecidable in $\mathbf{T}$, then P is undecidable in $\mathbf{NT}$.

3.1.3 Simple examples of NPL theories

Consider a classical theory $\mathbf{T}$ in propositional logic whose axioms are p and $p \rightarrow q$, where p and q are propositional variables. Using modus ponens, we conclude that q is a theorem of $\mathbf{T}$.

- Here $p \rightarrow q$ is not a legitimate axiom in the t-syntax of $\mathbf{NT}$. Hence a nonlogical rule of inference $p \vdash q$ (“From p , infer q ”) should be included in the p-syntax of $\mathbf{NT}$. The t-syntax would have the single axiom p , from

which q may be deduced via this inference rule. Note that $p \rightarrow q$ (or equivalently, $\neg p \vee q$) is also a theorem of $\mathbf{NT}$. From Remark 2, it follows that the above inference rule does *not* directly lead to a proof of $p \rightarrow q$ in $\mathbf{NT}$, and hence is not equivalent to an axiomatic assertion of $p \rightarrow q$. Rather, the inference rule together with a proof of p leads to a proof of q , from which we infer $p \rightarrow q$ as a theorem (but not a legitimate axiom) of $\mathbf{NT}$, via a different inference rule (disjunction introduction).

- The propositions $p \vee \neg p$ and $q \vee \neg q$, which are theorems of $\mathbf{T}$, are also theorems of $\mathbf{NT}$ that can be deduced from p and q respectively.
- The proposition $r \vee \neg r$, where r is a propositional variable, is a theorem of $\mathbf{T}$ but is not in the t-syntax of $\mathbf{NT}$ because each of the disjuncts r and $\neg r$ is undecidable in $\mathbf{T}$.
- The propositions $p \vee r$ and $q \vee r$ are theorems, but not legitimate as axioms, of $\mathbf{NT}$.
- The formula $r \vee s$, where s is a propositional variable, is an undecidable proposition of $\mathbf{NT}$. Note that $r \vee s$ cannot be legitimately added as an axiom to $\mathbf{NT}$, but addition of either r or s as an axiom to $\mathbf{NT}$ would make $r \vee s$ provable in the extended theory.

Consider the case where $\mathbf{T} = \emptyset$ is the null set of axioms. The theorems of $\mathbf{T}$ are the classical tautologies, which are not in the t-syntax of $\mathbf{NT}$. Hence the corresponding NPL theory $\mathbf{NT} = \emptyset$ has no theorems. The main postulate of NAFL semantics requires that the interpretation $\mathbf{NT}^* = \mathbf{NT}$ (see Sec. 2) will generate a model of $\mathbf{NT}$ in which $P \wedge \neg P$ is the case (meaning both P and $\neg P$ are true and hence nonclassically “neither true nor false”) for every proposition P in the t-syntax of $\mathbf{NT}$.

Suppose $\mathbf{T}$ has the single axiom $p \rightarrow q$, where p and q are propositional variables. Then again $\mathbf{NT} = \emptyset$. The axiom $p \rightarrow q$ (equivalent to $\neg p \vee q$) in the p-syntax of $\mathbf{NT}$ does not lead to any theorems in its t-syntax. This is so because both p and q are undecidable in $\mathbf{NT}$ and hence $\neg p \vee q$, which is provable in $\mathbf{T}$, is not in the t-syntax of $\mathbf{NT}$ either as a theorem or as an undecidable proposition. Here $p \vdash q$ should be included as a nonlogical rule of inference in the p-syntax of $\mathbf{NT}$, so that the theory $\mathbf{NT} + p$ would prove q .

3.2 First-order logic

Let NFOL denote the NAFL version of classical first-order logic (FOL) [3] that uses a natural deduction system. As in propositional logic, NFOL theories have a proof syntax (p-syntax) and a theory syntax (t-syntax).

3.2.1 The p-syntax of NFOL theories

The p-syntax of NFOL theories is classical, but unlike propositional logic, there are significant restrictions demanded by the finitary reasoning of NAFL, as will be described in Secs. 3.2.2 and 4.5.1.

Alphabet

The alphabet of the language $\mathcal{L}$ of an NFOL theory consists of the logical and the non-logical symbols.

Logical symbols

- The universal quantifier $\forall$ and the existential quantifier $\exists$.
- The logical connectives $\wedge$ (conjunction), $\vee$ (disjunction), $\rightarrow$ (implication), $\leftrightarrow$ (biconditional) and $\neg$ (negation).
- Parentheses $()$.
- Infinitely many variables $x, y, z, \dots$ or $x_0, x_1, x_2, \dots$.
- The equality symbol $=$.
- The truth constants $\top$ (for “true”) and $\perp$ (for “false”).

Non-logical symbols

- For each natural number n , the n -place predicate symbols $P^n, Q^n, R^n, \dots$ or $P_1^n, P_2^n, P_3^n, \dots$. The superscripts are often omitted. Zero-place predicate symbols are identified with propositional variables and are sometimes called sentence letters.
- For each natural number n , the n -place function symbols $f^n, g^n, h^n, \dots$ or $f_1^n, f_2^n, f_3^n, \dots$. The superscripts are often omitted. Zero-place function symbols are constants, denoted by $a, b, c, \dots$, or $a_0, a_1, a_2, \dots$.

Terms, formulas

The terms and formulas (also called well-formed formulas), as well as the free and bound occurrences of variables in a formula, are defined inductively as in FOL. The atomic formulas are of the form $P(t_1, t_2, \dots, t_n)$ and $t_1 = t_2$, where P is an n -place predicate symbol and $t_1, t_2, \dots, t_n$ are terms. A sentence / closed term is a formula / term with no free variable occurrences.

Deductive system

A natural deduction system, as defined in Sec. 3.1.1, is used, where the symbols $\{p, q, r, P, Q, \dots\}$ denote sentences rather than propositions. The inference rules for quantifiers are universal generalization, universal instantiation, existential generalization and existential instantiation, defined as in FOL. The following axiom schemas for equality complete the deductive system:

- Reflexivity: For each variable x , $x = x$.

- Substitution for functions: For all variables x and y , and any function symbol f ,

$$x = y \rightarrow f(\dots, x, \dots) = f(\dots, y, \dots). \quad (8)$$

- Substitution for formulas: For any variables x and y and any formula $\phi(x)$, if ϕ' is obtained by replacing any number of free occurrences of x in ϕ by y , such that these remain free occurrences of y ,

$$x = y \rightarrow (\phi \rightarrow \phi'). \quad (9)$$

In addition to the logical rules of inference given above, there may be nonlogical rules of inference [10] that are specific to a given NFOL theory $\mathbf{NT}$ (see Sec. 3.2.4). The p-syntax of an NFOL theory $\mathbf{NT}$ is completed by specifying a set of nonlogical axioms, corresponding to the classical FOL theory $\mathbf{T}$, satisfying the relevant NAFL restrictions specified in Sec. 3.2.2 below.

3.2.2 NAFL restrictions in the p-syntax of NFOL theories

Here we consider NFOL theories with a countably infinite universe. For the NFOL theory of real numbers, see Sec. 4.3. According to the main postulate of NAFL semantics (see Sec. 2), models of NFOL theories are generated by provability in *interpretations*, which are also NFOL theories. Therefore semantics for various infinite entities, such as, an infinite universe, functions and predicates, are to be provided via provability in NFOL theories. Classical semantics treats these infinite entities as pre-existing infinite sets in a Platonic universe, which makes truth distinct from provability. This classical distinction between syntax and semantics gets blurred in NAFL, which rejects Platonism and its infinitary consequences.

Sort for n -tuples

A separate sort is required for the treatment of n -tuples and infinite classes of n -tuples in the p-syntax of NFOL theories that admit a countably infinite universe. The complete set of rules / NAFL restrictions for this sort are as follows.

- The sort for n -tuples in NFOL theories accepts the natural numbers, rather than sets, as primitives. The class $\mathbb{N} = \{0, 1, 2, \dots\}$ of all natural numbers exists in this sort.
- Every NFOL theory proves the existence of n -tuples, where, for $n \geq 1$, an n -tuple is a finite sequence (or a list) defined by

$$(a_1, a_2, \dots, a_n) = (f(1), f(2), \dots, f(n)),$$

where f is a unary function defined on the domain $\{1, 2, \dots, n\}$, with the codomain as the universe of the NFOL theory. Here the a_j are constant symbols that map to objects in the universe. Note that we have defined

n -tuples as primitives rather than as sets. These n -tuples do not belong to the universe of objects (over which the quantifiers range) and may be considered as a separate sort.

- Infinite sets do not exist in consistent NFOL theories (see Metatheorem 4, Remark 7 and also Metatheorem 8 below). Infinite classes of n -tuples must necessarily exist in NFOL theories that prove the existence of infinitely many (finite) objects and these are proper classes that do not belong to the universe of objects or n -tuples.
- Infinite classes of n -tuples are constants, in the sense that there are no class variables that range over infinite classes, and hence there are no class quantifiers. Arbitrary infinite classes do not exist and no formula of an NFOL theory with a countably infinite universe can either directly or indirectly (e.g. via coding) refer to or quantify over infinitely many infinite classes. Since functions and predicates, when defined on an infinite domain or universe, are infinite classes of n -tuples (see Remark 16 below), it follows that there are no arbitrary functions or arbitrary predicates in NOFL theories.
- Every NFOL theory proves the existence of infinite sequences of n -tuples (where $n \geq 1$) as mappings from the domain $\mathbb{N}$, whenever the elements of the infinite sequence can be defined constructively, in the following sense. Given any natural number i , there must exist an algorithm (effective procedure) that constructs the i -th element of the sequence. For example, the infinite sequence of all prime numbers must exist in an NFOL theory of arithmetic.
- **Existence axiom for infinite classes:** As noted in Corollary 2, NAFL theories do not admit nonconstructive existence. Infinite classes of n -tuples that exist in NFOL theories must be countable and constructively specified. The class existence axiom schema may be stated as follows:

For each natural number $n \geq 1$, let $\phi(x_1, \dots, x_n)$ be a formula in the language of a consistent NFOL theory $\mathbf{T}$ (with a countably infinite universe) with no free variables other than $(x_1, \dots, x_n)$. Define $\tau_n = (a_1, \dots, a_n)$, where the a_j are objects of the universe whose existence is provable in $\mathbf{T}$. If there are only finitely many n -tuples τ_n that satisfy ϕ , then $\mathbf{T}$ proves the constructive existence of a finite list of such n -tuples. If there are infinitely many n -tuples τ_n that satisfy ϕ , then $\mathbf{T}$ proves the constructive existence of an infinite sequence of distinct n -tuples $(\tau_{ni})_{i \in \mathbb{N}}$ such that each element of this sequence satisfies ϕ , and $\mathbf{T}$ also proves the existence of the corresponding infinite class of n -tuples $\{\tau_{ni}\}_{i \in \mathbb{N}}$.

Note carefully the restriction that the formula ϕ should not contain any free variables other than $(x_1, \dots, x_n)$. The infinite class that we construct from $(x_1, \dots, x_n)$ cannot depend on other free variables because that would

violate the requirements that infinite classes are constants, and that no formula of an NFOL theory can refer to infinitely many infinite classes. The only legitimate formulas in NFOL theories are those which admit constructive existence, in the above strict sense, of the finite list (which could be empty), or infinite class, of n -tuples that satisfy these formulas. In particular, corresponding to every n -place function symbol f , there must exist an infinite sequence, and an infinite class, of distinct $(n + 1)$ -tuples $(x_1, \dots, x_n, y) = (a_1, \dots, a_{n+1})$ satisfying the relation $f(x_1, \dots, x_n) = y$.

- **Axiom of extensionality for infinite classes:** *If $\mathbb{A}$ and $\mathbb{B}$ are infinite classes of n -tuples (where $n \geq 1$) and if for all n -tuples x , $x \in \mathbb{A} \leftrightarrow x \in \mathbb{B}$, then $\mathbb{A} = \mathbb{B}$.*
- The universal class $\mathbb{U} = \{x : x = x\}$ must provably exist in every NFOL theory that proves the existence of a countably many finite objects. As required by the class existence axiom, the universal class must be enumerable in an infinite sequence. For example, in an NFOL theory of arithmetic, $\mathbb{U} = \mathbb{N} = \{0, 1, 2, \dots\}$.
- Constant symbols are nullary functions and Corollary 2 establishes semantically that there are no arbitrary constants in NAFL theories. See also Metatheorem 6 below for another proof of this result.
- **Examples:** In an NFOL theory of arithmetic, the prime numbers can be constructed as an infinite sequence and therefore there exists the corresponding infinite class of primes. The function $f(x) = x^2$ generates an infinite class of ordered pairs of natural numbers, namely, $\{(0, 0), (1, 1), (2, 4), \dots\}$. If $\mathbf{T}$ is an NFOL theory of finite sets, then $\mathbf{T}$ proves the existence of the infinite class $\{\emptyset, \{\emptyset\}, \{\{\emptyset\}\}, \dots\}$ provided $\mathbf{T}$ proves the existence of each element of this class, where $\emptyset$ denotes the empty set.

3.2.3 Metatheorems pertaining to the p-syntax of NFOL theories

This section also applies to NFOL theories with a countably infinite universe. These metatheorems characterize the finitistic restrictions imposed by NAFL on the p-syntax of NFOL theories.

Metatheorem 6. *Consistent NFOL theories do not admit arbitrary constants. If $\mathbf{T}$ is a consistent NFOL theory and if $\mathbf{T} \vdash \exists x x = c$, where c is a constant symbol, then $\mathbf{T} \vdash c = a$, where $a \in \mathbb{U}$ and $\mathbb{U}$ is the universal class whose existence is provable in $\mathbf{T}$.*

Proof. Suppose, to obtain a contradiction, $\mathbf{T} \vdash \exists x x = c$ and c is an arbitrary constant. Then $\mathbf{T}$ admits models in which $c = a_1$ and $c = a_2$, where $a_1 \in \mathbb{U}$, $a_2 \in \mathbb{U}$ and $a_1 \neq a_2$. The main postulate of NAFL semantics (see Sec. 2) implies that there must exist a nonclassical model of $\mathbf{T}$ in which $(c = a_1 \wedge c = a_2)$ holds. But such a nonclassical model cannot exist, because $(c = a_1 \wedge c = a_2)$ violates (8), which is the axiom of equality corresponding to substitution for functions.

Indeed, the existence of such a nonclassical model would imply (from (8)) that $(f(c) = f(a_1) \wedge f(c) = f(a_2))$ for any given function f . Since $f(c)$ must be uniquely defined (by the definition of a function), this yields $f(a_1) = f(a_2)$ for any given f , which implies $a_1 = a_2$, contrary to our assumption. $\square$

Metatheorem 7. *Let $\mathbb{A}$ be an infinite class whose existence is provable in a consistent NFOL theory $\mathbf{T}$, and let $\mathbf{T} \vdash \exists x (x = a)$. Then the proposition $a \in \mathbb{A}$ must be decidable in $\mathbf{T}$.*

Proof. Assume that $a \in \mathbb{A}$ is undecidable in $\mathbf{T}$. Then, upon choosing the interpretation $\mathbf{T}^* = \mathbf{T}$ (see Sec. 2), the main postulate of NAFL semantics requires that there must exist a nonclassical model of $\mathbf{T}$ in which $(a \in \mathbb{A} \wedge a \notin \mathbb{A})$ is true. But $(a \in \mathbb{A} \wedge a \notin \mathbb{A})$ violates the axiom of extensionality for infinite classes, which is a theorem of $\mathbf{T}$ and therefore the said nonclassical model of $\mathbf{T}$ cannot exist. $\square$

Remark 11. *The axiom of extensionality ensures the uniqueness of infinite classes. In NAFL, truth is associated with provability in theories via the main postulate of NAFL semantics. Therefore Metatheorem 7 requires that uniqueness of infinite classes must hold with respect to NFOL theories, whereas classically, this uniqueness holds in a Platonic universe. In particular, Metatheorem 7 implies that an infinite class must be constructively specified in an NFOL theory $\mathbf{T}$, and the construction must be provable in $\mathbf{T}$, i.e., it must hold in every model of $\mathbf{T}$.*

Metatheorem 8. *Infinite sets do not exist in consistent NFOL theories.*

Proof. Metatheorem 7 implies that NFOL theories with a countably infinite universe must specify a unique construction for the universal class $\mathbb{U}$. In particular, it follows from Metatheorem 7 that the proposition that a given infinite class $\mathbb{A} \in \mathbb{U}$ (i.e., $\mathbb{A}$ is a set) cannot be undecidable in a consistent NFOL theory $\mathbf{T}$. Since there exists a model of $\mathbf{T}$ in which $\mathbb{A} \notin \mathbb{U}$, we may conclude that $\mathbf{T} \vdash (\mathbb{A} \notin \mathbb{U})$, i.e., $\mathbb{A}$ is a proper class. $\square$

Remark 12. *Essentially the same proof, which uses the uniqueness of the construction for the universal class $\mathbb{U}$, also shows that consistent NFOL theories do not admit nonstandard models.*

Metatheorem 9. *Let x be a free variable in a formula of an NFOL theory $\mathbf{T}$. Let $\mathbb{U} = \{a_j\}_{j \in \mathbb{N}}$ be the universal class of objects whose existence is provable in $\mathbf{T}$, i.e., when universally quantified, x ranges over the values $a_j, j \in \mathbb{N}$. Then x (when free) must assume an infinite superposition of all possible values, i.e.,*

$$(x = a_0) \wedge (x = a_1) \wedge (x = a_2) \dots \quad (10)$$

Proof. As noted in Sec. 3.1.3, an NPL theory $\mathbf{NT} = \emptyset$ (with the null set of axioms) must have a model $\mathcal{M}$ in which every proposition in the t-syntax of $\mathbf{NT}$ has a nonclassical truth value of “neither true nor false”, i.e., each proposition and its negation are both assigned the truth value “true”. To model the

semantics of a free variable of an NFOL theory, let the propositions in the t-syntax of $\mathbf{NT}$ include the atomic propositions $x = a_j$ for each $j \in \mathbb{N}$. When x is a quantified variable, possibly infinitely many interpretations $\mathbf{NT}^*$ of $\mathbf{NT}$ (see Sec. 2), each containing the single axiom $x = a_j$, would be required. When x is a free variable, we set the interpretation $\mathbf{NT}^* = \mathbf{NT}$, and it follows that each of the propositions $x = a_j$ must be assigned the truth value “true” in the model $\mathcal{M}$. Then it is clear that (10) must hold in $\mathcal{M}$. What the superposition expresses is that there is no proof in $\mathbf{NT}^*$ that $x = a_j$, for each $j \in \mathbb{N}$. $\square$

Remark 13 (Variables versus constants). *A variable is allowed to have multiple values within a theory, unlike a constant, whose value is fixed. This is why in the above proof the nonclassical model $\mathcal{M}$ exists in which a variable assumes a superposition of all possible values. The axiom of equality (8) prohibits the existence of such a nonclassical model for constants, as seen in the proof of Metatheorem 6.*

Remark 14 (Free variables as infinite classes). *Note that the infinite superposition of values of x in (10) cannot be expressed explicitly (because formulas have to be of finite length), unlike a finite superposition, e.g. “ $x = a_1 \wedge x = a_2$ ”. We overcome this limitation by postulating $x = \mathbb{U}$, in a nonclassical model of an NFOL theory, for a free variable x of a formula. Thus if $\mathbb{U} = \{a_0, a_1, \dots\}$, then “ $x = \mathbb{U}$ ” is to be interpreted as a code for the infinite superposition in (10). The key point to note is that in NAFL, a free variable is an infinitary object (not belonging to the universe $\mathbb{U}$ of finite objects), whereas in classical logic, a free variable x is interpreted as an arbitrary object belonging to the universe (i.e., satisfying the infinite disjunction $x = a_0 \vee x = a_1 \dots$) and is therefore considered finite. On the other hand, in an NFOL theory $\mathbf{T}$, $x \in \mathbb{U}$ if and only if one of the infinitely many disjuncts ($x = a_j$) is provable in $\mathbf{T}$, and this is true only for a quantified variable x when it is assigned each of the values a_j (universal quantifier) or specific values a_j (existential quantifier).*

Remark 15. *It follows from Remark 14 that only quantified formulas are meaningful in NFOL theories. If an NFOL theory is specified with formulas containing free variables, such variables are assumed to be universally quantified. An interpretation $\mathbf{T}^*$ of an NFOL theory $\mathbf{T}$ may be thought of as consisting of infinitely many sub-interpretations, wherein the quantified variables of formulas are assigned values.*

Remark 16 (Functions and predicates as infinite classes). *When an n -tuple $(x_1, \dots, x_n)$ is assigned a value $(a_1, \dots, a_n)$, the function $f(x_1, \dots, x_n) = y$ is an $(n+1)$ -tuple in a model of an NFOL theory, i.e., $f = \tau_n = (a_1, \dots, a_n, a_{n+1})$, where the a_j belong to the universal class $\mathbb{U}$ and $y = a_{n+1}$. When the n -tuple $(x_1, \dots, x_n)$ is not assigned a value, the main postulate of NAFL semantics requires that the function $f(x_1, \dots, x_n) = y$ be in an infinite superposition of all possible values τ_n in a nonclassical model of an NFOL theory. In this case f is postulated to be an infinite class of $(n+1)$ -tuples, i.e., $f = \{\tau_{ni}\}_{i \in \mathbb{N}}$. Similar considerations apply to a predicate $P(x_1, \dots, x_n)$, which may be interpreted as a function which maps the n -tuple $(x_1, \dots, x_n)$ to a truth constant ($\top$ or $\perp$).*

Remark 17. *Observe that the nonclassical postulations of free variables and functions as infinite classes in Remarks 14 and 16 respectively must hold in every model of an NFOL theory. Therefore these postulations must be provable in the sort for n -tuples of a consistent NFOL theory $\mathbf{T}$, i.e.,*

$$\mathbf{T} \vdash (x = \mathbb{U}) \quad \& \quad \mathbf{T} \vdash (f = \{\tau_{ni}\}_{i \in \mathbb{N}}), \quad (11)$$

whenever x is a free variable in a formula and the argument $(x_1, \dots, x_n)$ of the n -place function f is unspecified. Note that, as demanded by consistency, $\mathbf{T}$ does not formally prove a contradiction in (11) despite the fact that the nonclassical interpretation of (11) in a model of $\mathbf{T}$ involves infinite superpositions of values for x and f , which are indeed contradictions. We believe that these infinite superpositions are appropriately encoded by infinite classes in (11) because an infinite class is indeed a contradictory object in NAFL. But NAFL semantics blocks the deduction of a contradiction from infinite classes, for which infinitely many sentences are required (see Sec. 2, in particular, Remarks 6 and 7).

Metatheorem 10. *Arithmetization of syntax and Gödel's incompleteness theorems cannot be formalized in consistent NFOL theories.*

Proof. The symbols of an uninterpreted classical theory, such as, function / predicate symbols, are treated as finite objects and encoded by Gödel numbers, even if they are identified with infinite sets in an interpretation. This classical separation of syntax and semantics does not go through in NAFL. The interpretation $\mathbf{T}^*$ of an NAFL theory $\mathbf{T}$ (see Sec. 2) necessarily exists. If the human mind does not explicitly specify $\mathbf{T}^*$, then $\mathbf{T}^* = \mathbf{T}$, i.e., no axioms have been added to $\mathbf{T}$. Hence uninterpreted NAFL theories do not exist and at least some symbols of a specified NFOL theory, such as variables, functions and predicates, must necessarily be treated as infinite classes (see Remarks 14, 16 and 17). In other words, the appropriate codes for these symbols of specified NFOL theories are infinite classes and not Gödel numbers, which can only be used to encode finite objects. By the rules stated in Sec. 3.2.2, infinite classes are constants and no formula of an NFOL theory can either directly or indirectly (e.g. via coding) quantify over or refer to infinitely many infinite classes. Hence infinitely many of these symbols cannot be specified as objects of an NFOL theory with a countably infinite universe, i.e., there are no arbitrary symbols and there is no universe of symbols, to quantify over.

Suppose the human mind does not specify any NFOL theory, but merely considers the infinitely many symbols used in NFOL theories. In this case, these uninterpreted symbols are to be treated as (approximate) geometric shapes, and must be encoded by specific geometric constructions, rather than by Gödel numbers. For example, the symbol “ o ” may be encoded by a circle of a given radius, which, from the point of view of Euclidean geometry, is a finite construction. But clearly the algebraic construction of geometric objects, such as, circles, requires a continuum of real numbers and is therefore not finite. We will later see that in the NAFL version of real analysis, the objects of Euclidean geometry are superclasses of real numbers and are encoded by infinite sequences of rational

numbers. Hence uninterpreted symbols are infinite objects and infinitely many of these cannot be (encoded as) objects of the universe of a consistent NFOL theory with a countably infinite universe, for example, via encoding by Gödel numbers. $\square$

Remark 18. *Classically, symbols are introduced as primitives and treated as finite objects in formal systems. In NAFL, this procedure is valid only for finitely many symbols. When infinitely many symbols are under consideration, their infinite structure as geometric objects (for uninterpreted symbols) or as infinite classes (e.g. for symbols interpreted as functions, predicates and variables) cannot be ignored. Hence infinitely many symbols cannot be introduced as primitives in consistent NFOL theories.*

Metatheorem 11. *The halting problem of classical computability theory (and Turing’s proof of the unsolvability of the halting problem) cannot be formalized in consistent NFOL theories.*

Proof. The halting problem is the problem of determining whether an arbitrary computer program will halt on an arbitrary input. A computer program performs the role of a function P that maps infinitely many possible inputs x to either “halts” ($P(x) = 0$) or “not-halts” ($P(x) = 1$). In NFOL theories, such a function is an infinite class (see Remark 16) and as noted in Sec. 3.2.2, arbitrary infinite classes do not exist. Hence arbitrary functions, and in particular, arbitrary computer programs, do not exist in consistent NFOL theories, and quantification over infinitely many computer programs is illegal. $\square$

Metatheorem 12. *Various forms of the diagonal argument (for example, as used in Gödel’s incompleteness theorems, Turing’s proof of the unsolvability of the halting problem, Cantor’s proof that there are uncountably many real numbers, the diagonal lemma) cannot be formalized in consistent NFOL theories.*

Proof. The common feature of all forms of the diagonal argument is that they either require arbitrary infinite classes (such as, arbitrary functions or arbitrary formulas with one or more free variables) or quantification over infinitely many infinite classes, which are illegal in NFOL theories, as noted in Sec. 3.2.2. For example, one form of Cantor’s diagonal argument requires an infinite list of real numbers, each of which is represented by an infinite sequence of rationals. $\square$

Remark 19. *As noted in Remark 5, Gödel’s incompleteness theorems and Turing’s proof of the unsolvability of the halting problem are infinitary results by the NAFL yardstick. Indeed, the unsolvability of the halting problem implies that there must exist a program P for which it is true, but unprovable, that P does not halt. In the NAFL version of finitism, the assertion that P does not halt is meaningful if and only if there exists a proof that P does not halt. If such a proof does not exist, then the claim that P does not halt implies that P passes through infinitely many states sequentially, which in turn implies that all the natural numbers in $\mathbb{N}$ can be counted, one at a time. This tacit assumption of classical computability theory enables Turing’s proof, which is rejected in NAFL as infinitary.*

3.2.4 The t-syntax of NFOL theories with an infinite universe

All the theorems in the sort for n -tuples, including those pertaining to infinite classes of n -tuples, carry over from the p-syntax to the t-syntax of NFOL theories with an infinite universe. In what follows, “t-syntax” is an abbreviation for the main sort of the t-syntax. The building blocks of the t-syntax of NFOL theories are *basic sentences*, defined as follows:

$$Q_1x_1Q_2x_2\dots Q_nx_nL. \quad (12)$$

Here each Q_j is a quantifier, *i.e.*, either $\forall$ or $\exists$, and L is a literal (an atomic formula of FOL or its negation), containing either no free variables or all of and only the free variables $(x_1, x_2, \dots, x_n)$. See Sec. 3.2.1 for the definition of atomic formulas. If L does not contain any free variables (for example, if all the free variables of a literal are substituted with closed terms) then the quantifiers $Q_1, \dots, Q_n$ may be removed from (12).

Observe that the negation of a basic sentence in (12) is also a basic sentence; the negation symbol can be moved inwards to the literal when $\forall x_j$ and $\exists x_j$ are replaced by $\neg\exists x_j\neg$ and $\neg\forall x_j\neg$ respectively and double negations are eliminated. The well-formed formulas of the t-syntax of NFOL theories are sentences built up inductively from (12) using propositional logic, as follows.

- Any basic sentence is a formula of the t-syntax of NFOL theories.
- If p is a formula, then $\neg p$ is a formula.
- if p and q are formulas, then $(p \rightarrow q)$, $(p \leftrightarrow q)$, $(p \vee q)$ and $(p \wedge q)$ are formulas.
- There are no other well-formed formulas in the t-syntax of NFOL theories.

Let $\mathbf{NT}$ be an NFOL theory and let the p-syntax of $\mathbf{NT}$ correspond to the classical theory $\mathbf{T}$ which satisfies the NAFL restrictions specified in Sec. 3.2.2 (when the universe of $\mathbf{T}$ is countable) or Sec. 4.5.1 (when the universe of $\mathbf{T}$ is the real line). The t-syntax of $\mathbf{NT}$ may be defined by following a procedure analogous to that described in Sec. 3.1.2. The well-formed formulas of the t-syntax can be converted to a disjunctive normal form (DNF), as follows:

$$\bigvee_{i=1}^n \bigwedge_{j=1}^{m_i} B_{ij}, \quad (13)$$

where B_{ij} is a basic sentence. The rules for constructing the t-syntax of $\mathbf{NT}$, with all well-formed formulas assumed to be in DNF, are as follows:

- From Metatheorem 5, it follows that an axiom of $\mathbf{NT}$ cannot contain the disjunction symbol $\vee$. Therefore the only legitimate axioms of $\mathbf{NT}$ are those that are conjunctions of one or more basic sentences.

- The theorems of **NT** are those theorems of **T** for which at least one of the disjuncts is provable in **T**. This follows from Metatheorem 1. In particular, theorems of **T** that are conjunctions of one or more basic sentences are theorems of **NT**. If P is a well-formed formula for which each of the disjuncts is refutable in **T**, then $\neg P$ will qualify as a theorem of **NT**.
- If there are theorems of **T** that contain the disjunction symbol $\vee$ and for which at least two of the disjuncts are undecidable in **T** and other disjuncts, if present, are either undecidable or refutable in **T**, then such theorems of **T** are not in the t-syntax of **NT**, *i.e.*, they are neither theorems nor undecidable formulas of **NT**. For example, the law of the excluded middle $P \vee \neg P$ is neither a theorem nor an undecidable formula of **NT** when P is a well-formed formula that is undecidable in **T**.
- For each nonlogical axiom of **T** that is not present and / or is not a legitimate axiom in the t-syntax of **NT**, an appropriate nonlogical rule of inference [10], from which the axiom can be classically inferred, should be included in the p-syntax of **NT**. See Remarks 20 and 21.
- The axioms of **NT** are those axioms of **T** that are legitimate in the t-syntax of **NT**, namely, the axioms that are conjunctions of one or more basic sentences. The theorems of **NT** may be deduced from these axioms and the nonlogical rules of inference that will replace any axioms of **T** that are not conjunctions of one or more basic sentences.
- Consider a well-formed formula P of the t-syntax of **NT** that is not a theorem of **T**, and for which at least one disjunct is undecidable in **T** and other disjuncts, if present, are either undecidable or refutable in **T**. Then P is an undecidable formula of **NT**. In particular, if P is a conjunction of one or more basic sentences and if P is undecidable in **T**, then P is undecidable in **NT**.

Remark 20. *Note that a sentence like $\forall x (P(x) \vee Q(x))$, where P and Q are unary predicate symbols, is legitimate in the p-syntax of **NT**, but is not a well-formed formula of the t-syntax. Consider the following simple illustrative example. Let **T** include the axioms*

$$\forall x (P(x) \vee Q(x)), \quad (14)$$

$$\forall x (P(x) \vee Q(x)) \rightarrow \forall x \exists y R(x, y), \quad (15)$$

where R is a two-place predicate symbol. By *modus ponens*, **T** proves

$$\forall x \exists y R(x, y). \quad (16)$$

Note that (16), being a basic sentence, is a well-formed formula of the t-syntax and is therefore a theorem of **NT**. But (14) and (15) are not legitimate sentences in the t-syntax of **NT**. Let $\mathbb{U}$ denote the universal class whose constructive existence is required to be provable in **T**. Infinitely many instances of (14), of the

form $P(a_j) \vee Q(a_j)$, where $a_j \in \mathbb{U}$, are theorems of $\mathbf{T}$ and, being basic sentences, are well-formed formulas of the t -syntax of $\mathbf{NT}$. Whether each of these instances is a theorem of $\mathbf{NT}$ or is excluded from its t -syntax must be determined by the rules given above. The point of this example is that axioms like (14) and (15), which contain infinitely many disjunctions, are excluded from the t -syntax of $\mathbf{NT}$. Therefore $\mathbf{NT}$ does not have any axioms. The theorems of $\mathbf{NT}$ may be deduced from the following nonlogical rules of inference corresponding to (14) and (15):

$$\begin{aligned} & \forall x (\neg P(x) \vdash Q(x)), \\ & \forall x (\neg P(x) \vdash Q(x)) \vdash \forall x \exists y R(x, y). \end{aligned}$$

3.2.5 First-order Peano arithmetic

Let $\mathbf{NPA}$ be the NFOL version of classical first-order Peano arithmetic ($\mathbf{PA}$). [3] The signature of $\mathbf{PA}$ has symbols for the constant 0, addition (+), multiplication ($\times$) and the successor function ($S(x)$). The p -syntax of $\mathbf{NPA}$ corresponds to the theory $\mathbf{PA}$, together with the NAFL restrictions given in Sec. 3.2.2. In particular, from Metatheorem 10, functions cannot be defined in $\mathbf{PA}$ via arithmetization of syntax, and hence the signature of $\mathbf{PA}$ must include function symbols other than $S(x)$ as required. The p -syntax of $\mathbf{NPA}$ contains the nonlogical axioms of $\mathbf{PA}$, which are specified as follows:

$$\forall x \neg(0 = S(x)), \quad (17)$$

$$\forall x \forall y (S(x) = S(y) \rightarrow x = y), \quad (18)$$

$$\forall x (x + 0 = x), \quad (19)$$

$$\forall x \forall y (x + S(y) = S(x + y)), \quad (20)$$

$$\forall x (x \times 0 = 0), \quad (21)$$

$$\forall x \forall y (x \times S(y) = (x \times y) + x), \quad (22)$$

$$\forall y_1 \dots \forall y_k ((\phi(0, \bar{y}) \wedge \forall x (\phi(x, \bar{y}) \rightarrow \phi(S(x), \bar{y}))) \rightarrow \forall x \phi(x, \bar{y})). \quad (23)$$

where (23) is the induction axiom schema for each formula $\phi(x, \bar{y})$ in the language of $\mathbf{PA}$ and $\bar{y}$ stands for $y_1, \dots, y_k$. The axioms (17) and (19) – (22) are basic sentences as defined in (12) and hence are theorems in the t -syntax of $\mathbf{NPA}$. The axiom (18) is not a well-formed formula of the t -syntax, but leads to infinitely many theorems (which are basic sentences) in the t -syntax, as follows:

$$\forall x \neg(S^{(n)}(x) = x), \quad \text{for } n = 1, 2, \dots, \quad (24)$$

where, for $n \geq 1$, $S^{(n)}(x)$ is the successor function applied n times to x , i.e., $S^{(1)}(x) = S(x)$, $S^{(2)}(x) = S(S(x))$, etc. Consider the induction axiom schema (23), which is also not a well-formed formula of the t -syntax. From the premise

$$\phi(0, \bar{y}) \wedge \forall x (\phi(x, \bar{y}) \rightarrow \phi(S(x), \bar{y}))$$

one may conclude

$$\phi(0, \bar{y}), \phi(S(0), \bar{y}), \dots, \phi(S^{(n)}(0), \bar{y})$$

for any natural number $n \geq 1$ and the reverse implication also holds. We conclude that a sentence equivalent to (23) is as follows:

$$\forall y_1 \dots \forall y_k ((\phi(0, \bar{y}) \wedge \forall x \phi(S(x), \bar{y})) \rightarrow \forall x \phi(x, \bar{y})). \quad (25)$$

Drop the quantifiers for $y_1, \dots, y_k$ and treat these as free variables:

$$(\phi(0, \bar{y}) \wedge \forall x \phi(S(x), \bar{y})) \rightarrow \forall x \phi(x, \bar{y}). \quad (26)$$

Upon eliminating the implication symbol from (26), we obtain the following equivalent of (23):

$$\neg \phi(0, \bar{y}) \vee \exists x \neg \phi(S(x), \bar{y}) \vee \forall x \phi(x, \bar{y}). \quad (27)$$

With $\phi(x, \bar{y})$ restricted to be a well-formed formula of the t-syntax, (27) can be put in the standard form (13) and each of the infinitely many instances of (27) (for every possible value of $\bar{y}$) is a theorem in the t-syntax of **NPA**.

Remark 21. *The axioms in the t-syntax of **NPA** may be chosen as (17) and (19) – (22). The following nonlogical rules of inference, from which (18) and (23) can be classically inferred, must be included in the p-syntax:*

$$\forall x \forall y (S(x) = S(y) \vdash x = y),$$

$$\forall y_1 \dots \forall y_k ((\phi(0, \bar{y}) \wedge \forall x (\phi(x, \bar{y}) \vdash \phi(S(x), \bar{y}))) \vdash \forall x \phi(x, \bar{y})).$$

*These axioms / nonlogical rules of inference are all needed, and **NPA** is the weakest theory (in the unrestricted language of arithmetic) that is possible in NAFL, if we reject ultrafinitism. Dropping any of these axioms / inference rules will lead to nonstandard models of arithmetic, not permitted in NAFL (see Metatheorem 3 and Remark 12). To make the subtheories of **NPA** consistent, restrictions are needed in the language of these theories so as to exclude undecidable sentences that will lead to nonstandard models. This is justified in NAFL, wherein truths exist only with respect to axiomatic theories, and the undecidable sentences in question become meaningless in weaker theories. In other words, in NAFL, sentences must derive their meaning only from the axiomatic theories in which they are formulated. On the other hand, in classical logic, these sentences are required to have an intrinsic truth value that is independent of theories, which is why they end up having nonstandard models.*

Remark 22. *The sort for n-tuples in the NFOL theory **NPA** accepts the natural numbers and the infinite proper class $\mathbb{N} = \{0, 1, \dots\}$ as primitives (see Sec. 3.2.2). The axioms (17) – (23) in the p-syntax prove the existence of the universal class $\mathbb{U} = \{0, S(0), S(S(0)), \dots\}$ and may be viewed as defining the (pre-existing) natural numbers in terms of numerals. This is the minimum level of Platonism that one must accept when considering theories that prove the existence of infinitely many objects. Given that nonstandard models of **NPA** do not exist (see Metatheorem 3 and Remark 12), the induction axiom schema (23) (or equivalently, (25)) may be stated in the form*

$$\forall y_1 \dots \forall y_k (\forall x \in \mathbb{U} \phi(x, \bar{y}) \rightarrow \forall x \phi(x, \bar{y})). \quad (28)$$

Note that (28) follows trivially from the definition of universal quantification. Similarly, all the other axioms in (17) – (22) may also be proven once the existence of the universal class $\mathbb{U}$ has been established, wherein $\mathbb{U}$ is identified with the pre-existing $\mathbb{N}$. But to make such an identification, all the axioms / nonlogical rules of inference are needed, as noted in Remark 21.

Remark 23. The classical theory Robinson arithmetic (**RA**) [3] has the axioms (17) – (22) and the following additional axiom:

$$\forall x (\neg(x = 0) \rightarrow \exists y (S(y) = x)). \quad (29)$$

Classically, the induction axiom (23) is required to prove (29). The theory **NPA** proves (29) in the p -syntax and the infinitely many instances of (29) in the t -syntax, which follow trivially from (17) and (24). The sentence (29) is not a well-formed formula in the t -syntax of **NPA**.

Remark 24. To the extent that undecidable sentences lead to nonstandard models of arithmetic, consistency of the NFOL theory **NPA** rules out undecidable sentences in its t -syntax. Note that this is diametrically opposite to the conclusion that follows from Gödel’s incompleteness theorems, e.g., for the FOL theory **PA**. From Metatheorem 10, it follows that Gödel’s incompleteness theorems cannot be formalized in **NPA**. Given the complex structure of the Gödel sentence, it is unlikely that it will even be a well-formed formula (of the form (13)) in the t -syntax of **NPA**. However, Gödel’s proof that this sentence is true in the standard model of **PA** implies that infinitely many instances of this sentence will be provable in **NPA**.

Remark 25. Interestingly, consistency of **NPA** requires that Fermat’s last theorem (FLT), which is a legitimate sentence in the t -syntax of **NPA** (with the exponentiation function added to its signature), must be either provable or refutable in **NPA**. Hence NAFL requires that FLT, if true, must necessarily have an elementary proof in **NPA**. Wiles’s proof of FLT [12, 13] assumes the consistency of classical theories stronger than **PA** and hence cannot be formalized in **NPA**.

4 Real analysis in NAFL

NAFL theories do not admit infinite sets, arbitrary infinite classes or quantification over infinitely many infinite classes. We will demonstrate how the restrictions on infinite classes may be relaxed in order to carry out a limited version of real analysis in NAFL.

4.1 Integers

The integers may be defined as the infinite class

$$\mathbb{Z} = \{\dots, -2, -1, 0, 1, 2, \dots\},$$

where each integer is represented by an ordered pair of natural numbers (m, n) , which intuitively stands for $m - n$. Thus $-1 = (k, k + 1)$, $0 = (k, k)$, $1 = (k + 1, k)$, etc., for any given natural number k . The abstract definition of integers as equivalence classes of ordered pairs of natural numbers requires set theory and is not possible in NAFL (note that these equivalence classes are infinite (proper) classes in NAFL and quantification over infinitely many of these is not allowed). For our purposes, we only need to be concerned with defining integers constructively in terms of natural numbers, as unique canonical representatives of the corresponding equivalence classes. We make the following definitions of equality and the arithmetical operations on the integers [14]:

$$\begin{aligned}
(m, n) &= (i, j) \leftrightarrow m + j = n + i, \\
(m, n) + (i, j) &= (m + i, n + j), \\
-(m, n) &= (n, m), \\
(m, n) \times (i, j) &= (mi + nj, mj + ni), \\
(m, n) < (i, j) &\leftrightarrow m + j < n + i. \\
0 &= (0, 0), \quad 1 = (1, 0).
\end{aligned}$$

The NFOL theory **NPA** (see Sec. 3.2.5) can be suitably extended to the theory **NPAZ** by adding the above defining axioms for the integers to the p-syntax of **NPA** in a separate sort. The t-syntax of **NPAZ** may be defined in a separate sort, according to the rules described in Sec. 3.2.4. Thus **NPAZ** has two main sorts, one for the natural numbers and one for the integers, in both the p-syntax and the t-syntax.

We prefer not to transfer the above defining axioms for the integers in terms of the natural numbers from the p-syntax of **NPAZ** to its t-syntax. Thus the sort for integers in the t-syntax of **NPAZ** will contain only statements about integers, without any reference to natural numbers. In particular, if nonlogical axioms that refer only to the integers are added to the p-syntax of **NPAZ**, they may be transferred to its t-syntax. The proofs in the p-syntax of **NPAZ** may be considered as being executed via a translation of statements about integers in **NPAZ** to statements about natural numbers in **NPA**.

4.2 Rationals

Rational numbers may be similarly defined as ordered pairs of integers (a, b) such that $b > 0$. Intuitively, $(a, b) = \frac{a}{b}$ and we may require a and b to be coprime in order to obtain the infinite class $\mathbb{Q}$ of unique canonical representatives of rational numbers. We make the following definitions of equality and the arithmetical

operations on the rationals [14]:

$$\begin{aligned}
(a, b) &= (c, d) \leftrightarrow ad = bc, \\
(a, b) + (c, d) &= (ad + bc, bd), \\
-(a, b) &= (-a, b), \\
(a, b) \times (c, d) &= (ac, bd), \\
0 &= (0, 1), \quad 1 = (1, 1), \\
(a, b) < (c, d) &\leftrightarrow ad < bc.
\end{aligned}$$

The above definitions for the rationals may be added to the p-syntax of the theory **NPAZ** in a separate sort to obtain the theory **NPAQ**. The t-syntax of **NPAQ** may be defined in a separate sort, as described in Sec. 3.2.4. Thus **NPAQ** has three main sorts, one each for the natural numbers, the integers and the rationals, in both the p-syntax and the t-syntax.

We prefer not to transfer the above defining axioms for the rationals in terms of the integers from the p-syntax of **NPAQ** to its t-syntax. Thus the sort for rationals in the t-syntax of **NPAQ** will contain only statements about rational numbers, without any reference to natural numbers or integers. In particular, if nonlogical axioms that refer only to the rationals are added to the p-syntax of **NPAQ**, they may be transferred to its t-syntax. The proofs in the p-syntax of **NPAQ** may be considered as being executed via a translation of statements about rationals in **NPAQ** to statements about integers in **NPAZ**.

4.3 Real numbers

The theory **NPAQ** permits the construction of infinite sequences of rationals (see Sec. 3.2.2), which we denote as $(q_n)_{n \in \mathbb{N}}$, or just (q_n) . We view the real numbers as originating from Euclidean geometry, *i.e.*, as points on the real line. A real number may be represented (or encoded) as a Cauchy sequence of rationals [14], which is a sequence of rationals (q_n) such that

$$\forall \epsilon (\epsilon > 0 \rightarrow \exists m \forall n (m < n \rightarrow |q_m - q_n| < \epsilon)).$$

Here ϵ ranges over $\mathbb{Q}$. If $x = (q_n)$ and $y = (q'_n)$ are real numbers, we define $x = y$ to mean that $\lim_{n \rightarrow \infty} |q_n - q'_n| = 0$, *i.e.*,

$$\forall \epsilon (\epsilon > 0 \rightarrow \exists m \forall n (m < n \rightarrow |q_n - q'_n| < \epsilon)),$$

and we define $x < y$ to mean that

$$\exists \epsilon (\epsilon > 0 \wedge \exists m \forall n (m < n \rightarrow q_n + \epsilon < q'_n)).$$

Also define

$$x + y = (q_n + q'_n), \quad x \times y = (q_n \times q'_n), \quad -x = (-q_n), \quad 0 = (0), \quad 1 = (1),$$

where (0) and (1) are infinite rational sequences of 0 and 1 respectively.

4.4 Superclasses of real numbers

The above axioms constructing the real numbers as Cauchy sequences of rationals cannot be immediately added to the p-syntax of the theory **NPAQ**, which treats infinite sequences (which are infinite classes of ordered pairs) as constants. There are no variables in **NPAQ** that range over infinite sequences. The question then would be how we can characterize collections of real numbers, over which variables can range. Note that a class of real numbers does not exist, because each real number is an infinite proper class.

At first sight this looks like an intractable problem and it is tempting to conclude that real analysis requires infinitary reasoning that is not compatible with NAFL. But there is a surprisingly simple solution. The finitist should surely accept Euclidean geometry and therefore there ought to be a finitistically acceptable way to characterize the collection of real numbers in a geometric object, say, a line segment, in the theory **NPAQ**. The following notion of a superclass of real numbers [15] achieves this objective.

Definition of superclass

A nonempty superclass S is a collection of real numbers that exists if and only if there exists a sequence $(q_n)_{n \in \mathbb{N}}$ of rationals whose embedding in the extended real line $\bar{\mathbb{R}}$ [16] has all its limit points as members of S and every member of S is a limit point of the embedding of $(q_n)_{n \in \mathbb{N}}$ in $\bar{\mathbb{R}}$. We may set $S = (q_n)_{n \in \mathbb{N}}$, i.e., every member of S is identified with a Cauchy subsequence of $(q_n)_{n \in \mathbb{N}}$. The empty superclass of reals $\emptyset$ is postulated to exist.

Remark 26. *If S is a superclass of real numbers, then S is closed, i.e., every limit point of S is a member of S . This follows from the definition of S and the classical result that the set of limit points of a set is closed.*

Remark 27. *As an example, consider the real interval $[0, 1]$, which geometrically represents a line segment. We conclude that $[0, 1]$ is a superclass because an enumeration of all the rationals in the rational interval $[0, 1]$, when embedded in $\bar{\mathbb{R}}$, has precisely the real interval $[0, 1]$ as its limit points. The open / semi-open intervals of reals $(0, 1)$, $[0, 1)$ and $(0, 1]$ do not have any geometric equivalents (in the diagrammatic sense) and are not superclasses. Note that all these intervals are of unit length, and diagrammatically, it is impossible to extend a line segment to unit length without including the end points. Thus our encoding of reals as superclasses is strictly faithful to Euclidean geometry.*

Remark 28. *The standard real line $\mathbb{R} = (-\infty, \infty)$ is not a superclass. But the extended real line $\bar{\mathbb{R}} = [-\infty, \infty]$ is a superclass. This follows from the fact that any sequence of embedded rationals that has all the standard real numbers in $\mathbb{R}$ as its limit points (e.g., a sequence that enumerates all the rationals in $\mathbb{Q}$) also has $\pm\infty$ as its limit points in $\bar{\mathbb{R}}$. In other words, NAFL only supports the extended real number system. For example, the set $\{\frac{1}{x} : x > 0\}$ in the standard real number system is not a superclass because it excludes the limit points 0 and*

$+\infty$. But $\{\frac{1}{x} : x \geq 0\}$ in the extended real number system is a superclass. Here we define $\frac{1}{0} = \lim_{x \rightarrow 0^+} \frac{1}{x} = +\infty$ (see Sec. 4.5.3).

Remark 29. Every object in three-dimensional Euclidean geometry is a superclass that can be encoded as a sequence of rational 3-tuples in the theory **NPAQ**.

4.5 The theory **NPAR** of real numbers

NPAR may be defined as an NFOL theory that extends **NPAQ** and has all the machinery of NFOL defined in Secs. 3.2.1 and 3.2.4. However, Sec. 3.2.2 does not apply and needs to be modified, as in Sec. 4.5.1 below, because the objects of **NPAR** are the real numbers, each of which is an infinite sequence of rationals.

The axioms for the real numbers in the p-syntax of **NPAR** are those stated in Sec. 4.3. These axioms are added to the p-syntax of **NPAQ** in a separate sort, and are not transferred to the t-syntax of **NPAR**. Thus **NPAR** has four main sorts, one each for the naturals, integers, rationals and reals, in both its p-syntax and t-syntax. Note that **NPAQ** does not have variables that range over infinite classes (in particular, infinite sequences), which are treated as constants. The proofs in the p-syntax of **NPAR** may be considered as being executed via a translation of statements about reals to statements about Cauchy sequences of rationals in an extended version of **NPAQ**. This extended theory contains class quantifiers and variables ranging over the Cauchy sequences of rationals that represent reals.

4.5.1 NAFL restrictions in the p-syntax of **NPAR**

For the extension of **NPAQ** mentioned above to be finitistically acceptable, it is necessary that all collections of reals whose existence is implied by **NPAR** be restricted to be superclasses, which are represented by infinite sequences of rationals in **NPAQ** (see Sec. 4.4). These restrictions may be spelled out as follows.

- **NPAR** admits only constructive reasoning, and every statement about real numbers should be translatable to statements about Cauchy sequences of rationals in an extended version of **NPAQ**. Arbitrary predicates, arbitrary functions and arbitrary constants (see Metatheorem 6) do not exist in **NPAR**.
- The variables of **NPAR** range over the universal superclass of reals, which is the extended real line $\mathbb{R} = [-\infty, \infty]$. Note that $\mathbb{R}$ may be represented in **NPAQ** by any infinite sequence that enumerates all the rationals in $\mathbb{Q}$.
- Nonstandard models of **NPAR** do not exist because nonstandard models of **NPAQ** do not exist (see Metatheorem 3 and Remark 12).

- The class $\mathbb{N}$ of all natural numbers does not exist in **NP****AR**, which only admits the superclass $\mathbb{N}_{\mathbb{R}} = \{0, 1, \dots, \infty\}$. In other words, when the naturals in $\mathbb{N}$ are embedded in $\mathbb{R}$, the real number $+\infty$ cannot be excluded and hence a countable infinity does not exist in **NP****AR**.
- The theory **NP****AR** proves the existence of every superclass, finite or infinite, of n -tuples of reals (for $n \geq 1$) that can be represented constructively by an infinite sequence of n -tuples of rationals in **NPAQ**. Here n -tuples are defined as in Sec. 3.2.2 and an infinite sequence of n -tuples of rationals is a sequence of the form $(q_{1i}, \dots, q_{ni})_{i \in \mathbb{N}}$. Canonical Cauchy subsequences of $(q_{1i}, \dots, q_{ni})_{i \in \mathbb{N}}$ represent the n -tuples of reals. There are no variables ranging over superclasses of n -tuples of reals, which are treated as constants in a separate sort.
- Let $\phi(x_1, \dots, x_n)$ be a formula with no free variables other than $(x_1, \dots, x_n)$. The only legitimate formulas ϕ are those that satisfy the following restriction. If there exist n -tuples $(x_1, \dots, x_n)$ that satisfy ϕ , then **NP****AR** must prove the constructive existence of the superclass $\{(x_1, \dots, x_n) : \phi\}$. Here the construction is as defined in Sec. 4.4, and ϕ expresses a predicate that may be interpreted as the (possibly empty) superclass that constructs it (see Remark 16, with “ $\{\tau_{ni}\}_{i \in \mathbb{N}}$ ” replaced by the appropriate superclass). All predicates that occur in formulas must be interpretable as superclasses in the above sense. The formula $y = f(x_1, \dots, x_n)$, where f is a continuous real-valued function, must also satisfy the above restrictions. The domain and range of f must necessarily be superclasses. If the function f has discontinuities at isolated points, then the above restrictions apply to every subdomain (separated by the discontinuities) on which f is continuous. At points of discontinuity, f will have multiple values, which could be either real numbers or superclasses of real numbers (see Sec. 4.5.2). More general discontinuous functions will be considered in future work. For example, consider the sentence

$$\forall x (x > 0 \rightarrow f(x) = \frac{1}{x}).$$

Here the formula $\phi(x)$, defined as “ $(x > 0 \rightarrow f(x) = \frac{1}{x})$ ”, is illegal in **NP****AR**, *despite* the fact that $\{x : \phi(x)\} = \mathbb{R}$, because it uses an illegitimate predicate “ $x > 0$ ” (note that $\{x : x > 0\}$ is not a superclass; see Remark 26).

- Only convergent real sequences can be defined, on the domain $\mathbb{N}_{\mathbb{R}}$, because $\mathbb{N}$ does not exist in **NP****AR**. A convergent sequence of n -tuples of reals $(x_{1j}, \dots, x_{nj})_{j \in \mathbb{N}_{\mathbb{R}}}$ is identified with the superclass of $(n + 1)$ -tuples of reals $\{j, x_{1j}, \dots, x_{nj}\}_{j \in \mathbb{N}_{\mathbb{R}}}$, which is represented in **NP****AR** by a sequence of $(n + 1)$ -tuples of rationals. The limit points of the embedded rational sequence in $\mathbb{R}$ are precisely the members of $\{j, x_{1j}, \dots, x_{nj}\}_{j \in \mathbb{N}_{\mathbb{R}}}$.

4.5.2 Multivalued functions

NPAR admits multivalued functions [17, 18], where the multiple values occur at points of discontinuity of the function. For example, consider the following classical definition of a step function:

$$\begin{aligned} f(x) &= -1, & x < 0, \\ &= 0, & x = 0, \\ &= 1, & x > 0. \end{aligned}$$

This definition is illegal in **NPAR** because the domains $x < 0$ and $x > 0$ are not superclasses (see Remark 26). The definition acceptable in **NPAR** is as follows.

$$f(x) = -1, \quad x \leq 0, \quad (30)$$

$$= \{-1, 1\}, \quad x = 0, \quad \text{dom}(f) \text{ not specified}, \quad (31)$$

$$= 1, \quad x \geq 0. \quad (32)$$

Here $\text{dom}(f)$ is the domain of the relevant branch of the function f and $\{-1, 1\}$ is a superclass. Note that $f(0)$ has three distinct values:

$$f(0) = \lim_{x \rightarrow 0^-} f(x) = -1, \quad \text{dom}(f) = [-\infty, 0], \quad (33)$$

$$= \{-1, 1\}, \quad \text{dom}(f) \text{ not specified}, \quad (34)$$

$$= \lim_{x \rightarrow 0^+} f(x) = 1, \quad \text{dom}(f) = [0, \infty]. \quad (35)$$

In (30), the domain $x \leq 0$ implies that x is constrained to approach 0 from below. But whenever $x \rightarrow 0^-$ via a superclass containing a sequence of reals whose limit point is 0, the superclass must also contain the limit point (see Remark 26). The limiting value of $f(x)$ as $x \rightarrow 0^-$ must necessarily be a value of $f(0)$, because of the requirement that the range of a continuous function $f(x)$ must be a superclass. This is the rationale for (33). A similar explanation applies for (35).

Whenever the axiom $f(0) = -1$ ($f(0) = 1$) is added to **NPAR** to obtain the interpretation **NPAR*** (see Sec. 2), it is understood that the domain of $f(x)$ in the extended theory is restricted to be $x \leq 0$ ($x \geq 0$), in accordance with (30) and (32). Thus the contradiction ($f(0) = -1 \wedge f(0) = 1$) is not deducible in either of the extended theories. In (31) and (34) no direction of approach to $x = 0$ is specified. Here “ $f(0) = \{-1, 1\}$ ” is essentially a code for the superposed state “ $f(0) = -1 \wedge f(0) = 1$ ”, which results from the main postulate of NAFL semantics when the value of $f(0)$ is not specified by the human mind that interprets **NPAR**. This superposed state is interpreted as “Neither $f(0) = -1$ nor $f(0) = 1$ is provable in **NPAR***”. In conclusion, a discontinuity breaks up a multivalued function into its different branches, which are chosen axiomatically by the human mind, in accordance with the main postulate of NAFL semantics.

Remark 30. *The above definition of a step function is strictly faithful to Euclidean geometry (see Remark 27) and is best understood when x is a variable*

representing physical time. In this case there is a clear intuition behind the multiple values of $f(x)$ at $x = 0$, namely, that $f(x)$ jumps from -1 to $+1$ in zero time as x passes from 0^- to 0^+ ; so it is only logical to expect that both values of f are present at $x = 0$. At any given time x , we may axiomatically specify either direction of approach to $x = 0$, in which case $f(0) = \pm 1$, with the domain for f restricted as in (30) or (32). If no direction of approach to $x = 0$ is specified at any given time x , the superposed state $f(0) = \{-1, 1\}$ essentially implies that there is no physically meaningful definition of $f(0)$, which has a purely logical interpretation as noted above. This time-dependent definition of $f(0)$ fits in perfectly with NAFL semantics, which admits time-dependent truths as axiomatic declarations of the human mind (see Sec. 2.2). Examples of such time-dependent truths that are relevant to physics are given in Sec. 5.

Remark 31. Consider the example in which a rod of unit length is rotating at uniform angular velocity from $\theta = 0$ to $\theta = \pi$, in a polar coordinate system (r, θ) . At $\theta = \frac{\pi}{2}$, the slope of the rod ($\tan \theta$) is classically undefined and has an infinite discontinuity. In **NPAR**, we may define this slope as follows:

$$\begin{aligned} \tan(\pi/2) &= \infty, & \text{dom}(\tan) &= [0, \pi/2], \\ &= \{-\infty, \infty\}, & \text{dom}(\tan) &\text{not specified}, \\ &= -\infty, & \text{dom}(\tan) &= [\pi/2, \pi]. \end{aligned}$$

Here $\text{dom}(\tan)$ is the domain of the relevant branch of the $\tan$ function. This is an example of a multivalued function in which an infinite discontinuity is correctly handled by the extended real number system of **NPAR**.

4.5.3 Division by zero

If $q = (a, b)$ and $q' = (c, d)$ are rational numbers represented by ordered pairs of integers (see Sec. 4.2) such that $b > 0$, $d > 0$ and $c \neq 0$, then we define division for rationals as follows:

$$q/q' = (a, b)/(c, d) = \text{sgn}(c)(ad, b|c|),$$

where $\text{sgn}(c) = +1(-1)$ if $c > 0$ ($c < 0$). If $y = (q_n)_{n \in \mathbb{N}}$ and $x = (q'_n)_{n \in \mathbb{N}}$ are real numbers represented by Cauchy sequences of rationals (see Sec. 4.3), where $q'_n \neq 0$, then we define division for real numbers as follows:

$$y/x = (q_n/q'_n)_{n \in \mathbb{N}}. \quad (36)$$

In general, y/x is a superclass and we identify canonical Cauchy subsequence(s) of (q_n/q'_n) with the value(s) of y/x . In particular, if $y > 0$ and $x = 0$, then y/x has multiple values:

$$\begin{aligned} (y/x) &= \infty, & x &\rightarrow 0^+, \\ &= \{-\infty, \infty\}, & x &= 0, \\ &= -\infty, & x &\rightarrow 0^-. \end{aligned}$$

The above notation is explained as follows. If x is represented by a Cauchy sequence of positive (negative) rationals, then $x \rightarrow 0^+$ ($x \rightarrow 0^-$). If the Cauchy sequence representing x approaches the limit $x = 0$ from both directions, then y/x is the superclass $\{-\infty, \infty\}$. Thus we see that the value of y/x depends on how x is specified. For example, consider the superclass $\{\frac{1}{x} : x \geq 0\}$. As noted in Remark 28, we define $\frac{1}{0} = \lim_{x \rightarrow 0^+} \frac{1}{x} = \infty$ in this context, because when $x \geq 0$, x is constrained to approach 0 from above, *i.e.*, $x \rightarrow 0^+$. Similarly, when considering the superclass $\{\frac{1}{x} : x \leq 0\}$, we may define $\frac{1}{0} = \lim_{x \rightarrow 0^-} \frac{1}{x} = -\infty$.

4.5.4 0/0 and dy/dx

If $y = x = 0$ in (36), then the value of y/x could be any real number (including $\pm\infty$) or a superclass of real numbers, depending on how the sequences (q_n) and (q'_n) are specified. This is in confirmation with the observation that $0 \times r = 0$, where r is any real number (including possibly $\pm\infty$ in specific contexts).

Remark 32. Note that $0 \times \infty$ and ∞/∞ can be put in the form $0/0$. In the theory **NPAR**, the value r obtained for $0/0$ in any specific context (*i.e.*, when specific Cauchy sequences are specified for the zeroes) can be viewed as the axiomatic assertion $0/0 = r$. There is no contradiction in making multiple such assertions because, as noted above, we do not expect $0/0$ to be a uniquely defined real number. If no Cauchy sequences are specified for the zeroes in $0/0$, then from the main postulate of NAFL semantics (see Sec. 2), $0/0$ is in a superposed state of all possible values, which may be encoded in **NPAR** as $0/0 = \bar{\mathbb{R}}$. Here $\bar{\mathbb{R}}$ is the superclass corresponding to the extended real line $[-\infty, \infty]$ and may be represented in **NPAR** by a sequence that enumerates all the rationals in $\mathbb{Q}$ (see Remark 28).

Remark 33. If $y = f(x)$, where $f(x)$ is a differentiable function, then by definition

$$f'(x) = \frac{dy}{dx} = \lim_{\Delta x \rightarrow 0} \frac{f(x + \Delta x) - f(x)}{\Delta x}. \quad (37)$$

To evaluate the limit in (37), one substitutes a sequence S of real numbers for Δx , where S has a single limit point 0. This results in a sequence S' of approximations for $f'(x)$ and the limit point of S' is the value of $f'(x)$ at a given value of x . In the theory **NPAR**, x is represented by a Cauchy sequence of rationals, and S and S' are represented by superclasses which are sequences of rationals (see Secs. 4.4 and 4.5.1). From Remark 26, it follows that S and S' must contain their limit points, namely, 0 and $f'(x)$ respectively. We conclude that in **NPAR**, $f'(x)$ is indeed $0/0$, evaluated according to a specific procedure that amounts to an axiomatic assertion (see Remark 32). Note that the classical Weierstrass $\epsilon - \delta$ argument for the evaluation of $f'(x)$ in (37) does not hold in **NPAR** because it requires the existence of open intervals of reals, which are not superclasses.

4.5.5 The t-syntax of NPAR and its extensions

As noted in Sec. 4.5, the defining axioms for the real numbers in terms of Cauchy sequences of rationals are not transferred to the t-syntax of **NPAR**, which will only contain theorems that are statements about real numbers. Because open intervals and a countable infinity of real numbers do not exist in **NPAR** (see Secs. 4.4 and 4.5.1), many of the proofs of classical real analysis and number theory cannot be formalized in **NPAR**. In particular, Cantor’s diagonal argument for the uncountability of the set of real numbers cannot be formalized in **NPAR**. Fermat’s last theorem cannot be formalized in **NPAR** because it would be an illegal statement about embedded natural numbers x, y, z and n that are restricted to be finite, *i.e.*, $x \neq \infty$, etc.

At first sight, **NPAR** appears to be severely restricted, but because multivalued functions and division by zero are legal, it is still possible to recover many of the results that are needed for physical theories. In future work, we will consider the detailed development of real analysis in **NPAR**. We will also consider extensions of **NPAR** that formalize Euclidean geometry, Newtonian mechanics and quantum mechanics. In particular, we will demonstrate that non-Euclidean geometries and relativity theory are infinitary theories by the NAFL yardstick and cannot be formalized in NAFL.

5 Resolution of classical and quantum paradoxes in NAFL

In the rest of this paper, we will demonstrate how the logic NAFL resolves, or has the potential to resolve, longstanding issues and paradoxes in classical / quantum physics.

5.1 Negative resolution of Hilbert’s program in NAFL

The main goal of Hilbert’s program was to provide a finitary consistency proof of classical mathematics, as seen from the following quote [19]:

In the early 1920s, the German mathematician David Hilbert (1862–1943) put forward a new proposal for the foundation of classical mathematics which has come to be known as Hilbert’s Program. It calls for a formalization of all of mathematics in axiomatic form, together with a proof that this axiomatization of mathematics is consistent. The consistency proof itself was to be carried out using only what Hilbert called “finitary” methods. [...] It was also a great influence on Kurt Gödel, whose work on the incompleteness theorems were motivated by Hilbert’s Program. Gödel’s work is generally taken to show that Hilbert’s Program cannot be carried out.

According to the NAFL definition of finitism, Gödel’s incompleteness theorems are themselves infinitary results; further, NAFL *refutes* classical infinitary reasoning (see, for example, Metatheorems 8 and 10). Therefore large parts of classical mathematics and in particular, the classical theories Peano arithmetic and Primitive recursive arithmetic, are inconsistent by the NAFL yardstick (see Secs. 1 and 2). In this sense, NAFL provides a negative resolution of Hilbert’s program.

5.2 The logical incompatibility of quantum mechanics and special relativity theory

Quantum mechanics (**QM**) is a spectacularly successful theory, but its logical foundations are highly controversial. One of the most intractable problems of theoretical physics is the incompatibility between **QM** and special relativity theory (**SR**), which is perhaps most acute in quantum field theory. Freeborn *et al.*, [20] while considering the conclusions of Earman and Fraser, [21] make the following assessment:

As such, Haag’s theorem is perhaps the first of many signs that SR and QM are, in the end, irreconcilable.

The import of the arguments of Mamone-Capria [22] is that **SR** is not a “true” theory and needs to be replaced in order to match the requirements of **QM**.

Most of the classically minded theoreticians tend to believe that **SR**, which can be formalized in classical first-order logic, [23] is a logically consistent theory. In a previous preprint, [24] we have pointed out the existence of a meta-inconsistency (*i.e.*, an inconsistency at the metamathematical level) in **SR**, even from the point of view of classical logic. In the logic NAFL, this meta-inconsistency translates to a straightforward logical inconsistency in **SR**, as follows. The Lorentz transformations of **SR** require that the velocity (v) of a material object with respect to an inertial frame of reference can attain all values less than c (the velocity of light in vacuum), but not $v = c$. But the domain $0 \leq v < c$ is not a superclass (see Remark 26) and hence **SR** cannot be formulated as a consistent NAFL theory. This is a fatal objection because there is no obvious way in which **SR** can be extended to hold at the singular point $v = c$. On the other hand, NAFL successfully formalizes and justifies several of the “weird” phenomena of **QM**, such as, quantum superposition, entanglement and wave-particle duality, as outlined in the ensuing sections. Therefore NAFL supports **QM** and rejects **SR** as an infinitary theory. In this sense, NAFL points to a fundamental logical incompatibility between **QM** and **SR**. It is our belief that **QM**, unlike **SR**, does not rely on classical infinitary reasoning in an essential way, and in future work, we will consider how **QM** can be formalized in NAFL.

5.3 Quantum superposition, Wigner’s friend paradox and the collapse of the wave function

Consider the following famous thought experiment of Schrödinger [25] in quantum mechanics. A cat is placed inside a sealed box along with a radioactive source and a flask of poison. If an internally placed Geiger counter detects radioactivity caused by the decay of a single atom, a mechanical device shatters the flask, thus releasing the poison and killing the cat. There are many possible interpretations of this experiment. According to the Copenhagen interpretation, [26] while the box is closed, the system (*i.e.*, the cat along with the radioactive atom) exists in a superposition of the states $\langle \text{decayed atom} + \text{dead cat} \rangle$ and $\langle \text{undecayed atom} + \text{living cat} \rangle$, and when the box is opened and an observation is made, the wave function of the system collapses into one of the two states, and hence the superposed state can never be observed.

There are several questions that arise from this interpretation, for example, what constitutes a measurement or an observation? Does the Geiger counter perform a measurement? What does it mean for the cat to be in a superposed state of “alive and dead”? Is the collapse of the wave function an objective occurrence independent of the human mind? If the cat is observed to be alive when the box is opened, can we not conclude from a precise measurement of the cat’s age that the cat was always alive, in which case how could the cat have been in a superposed state while inside the sealed box? Similarly, if the cat is observed to be dead when the box is opened, a precise measurement of the time of death is in principle possible, say, via an autopsy, which may lead to the apparent contradiction that the cat was both dead and in a superposed state while inside the sealed box. An observer inside the sealed box would always observe the cat to be either alive or dead, so how can the superposed state exist for an observer outside the box? This is the paradox of Wigner’s friend, which is discussed in Sec. 5.3.2.

All of the above questions are answered elegantly and in a logically consistent manner in what we denote as the NAFL interpretation, [6] outlined below in Sec. 5.3.1. The NAFL interpretation can be thought of as the correct logical framework for the Copenhagen interpretation, which was always intended to be non-realist.

5.3.1 The NAFL interpretation of quantum superposition

Let **NQM** be the theory that formalizes quantum mechanics in NAFL. Let the Schrödinger cat experiment be set up at time $t = 0$, when the box is sealed. Let the box be opened at $t = 1$, when an observation is made of the cat’s state. Let P ($\neg P$) denote the proposition “The cat is alive (dead)”. Note that P is undecidable in **NQM**. According to the main postulate of NAFL semantics (see Sec. 2), P ($\neg P$) is true with respect to **NQM** if and only if the interpretation **NQM*** of **NQM** proves P ($\neg P$). Thus truth for P is axiomatic in nature and **NQM*** is chosen according to the free will of the human mind, *i.e.*, NAFL does not recognize any external reality for the state of the cat.

However, the connection with reality is made when the human mind (in this case, the observer) decides to keep the axiomatic declarations of $\mathbf{NQM}^*$ in tune with the observations made in real life.

Therefore, according to this convention, if the cat is observed to be alive (dead), the observer may choose $\mathbf{NQM}^* = \mathbf{NQM} + P$ ($\mathbf{NQM} + \neg P$) at time $t = 1$. For $0 \leq t < 1$, the observer makes no observations (and hence, no axiomatic declarations) on the state of the cat, *i.e.*, $\mathbf{NQM}^* = \mathbf{NQM}$. In this case $\mathbf{NQM}^*$ generates a nonclassical model $\mathcal{M}$ of $\mathbf{NQM}$ in which the cat is in a superposed state of “neither alive nor dead”, *i.e.*, $P \wedge \neg P$ is the case. In $\mathcal{M}$, “ P ” (“ $\neg P$ ”) is interpreted as “ $\mathbf{NQM}^*$ does not prove $\neg P$ (P)”. Further, when $t \geq 1$, the observer may deduce (in a suitable theory $\mathbf{NQM}^*$) that the cat was alive during $0 \leq t \leq 1$ or that the cat died before $t = 1$, say, at $t = 0.5$. These are time-dependent truths that only apply for $t \geq 1$, and hence do not clash with the superposed state that applied when $0 \leq t < 1$.

The alert reader may have noticed that the above informal description makes use of the domain $0 \leq t < 1$, which is not a superclass, for the superposed state of the cat. The precise state of the cat can be formally defined within the theory $\mathbf{NQM}^*$ only when $t \geq 1$ (*after* the box has been opened) because when the box is sealed, the superposed state of the cat ($P \wedge \neg P$) is not formally provable in $\mathbf{NQM}^*$. This is so because NAFL theories do not prove contradictions and the superposed state only exists in the model $\mathcal{M}$ of $\mathbf{NQM}$ generated by the interpretation $\mathbf{NQM}^*$. Thus $\mathbf{NQM}^*$ only makes use of the domain $t \geq 1$, which is a superclass.

The precise definition of the cat’s state within $\mathbf{NQM}^*$ is as follows. Let the state of the cat be denoted by the function $f(t)$ defined on the domain $0 \leq t \leq 1$ and with the range $\{0, 1\}$, where $f(t) = 1$ ($f(t) = 0$) denotes that the cat is alive (dead) at time t . Suppose the cat was observed to be alive (dead) at $t = 1$. Then the theory $\mathbf{NQM}^*$ will prove, from a precise determination of the cat’s age (time of death) that

$$f(t) = 1, \quad 0 \leq t \leq 1 \quad (f(t) = 0, \quad 0.5 \leq t \leq 1). \quad (38)$$

Note that (38) determines, *at* $t = 1$, that the cat was alive (dead at $t = 0.5$) while inside the sealed box. This is a time-dependent truth that only applies when $t \geq 1$ and hence there is no clash with the superposed state that existed earlier, when the box was sealed. In other words, the observer makes the axiomatic declaration at $t = 1$ that the cat was alive (dead) during $0 \leq t \leq 1$ ($0.5 \leq t \leq 1$), whereas the superposed state indicates that *when the box is sealed*, *i.e.*, when $t \in [0, 1)$, there is no proof in $\mathbf{NQM}^*$ that the cat is either alive or dead, leading to the nonclassical truth value of “neither alive nor dead” for the state of the cat. Thus, at $t = 1$, the observer moves from a state of ignorance, namely, “The cat was in the superposed state when the box was sealed”, to a state of knowledge, namely, “The cat was alive (dead at $t = 0.5$) when the box was sealed”, and there is no contradiction here because the NAFL interpretation does not ascribe any physical reality to the superposed state.

Let us now turn to the issue of how the superposed state of the cat can be represented consistently in $\mathbf{NQM}^*$. As noted earlier, there is a temptation to

deduce the following apparent contradiction. Does the fact that the cat was observed to be alive (dead) at $t = 1$ imply that **NQM*** generates a model $\mathcal{M}$ of **NQM** in which the superposed state existed for the time period $0 \leq t < 1$, which is not a superclass (see Remark 26), and does this make **NQM** an inconsistent NAFL theory? To answer this question in the negative, we would need a separate theory of the model $\mathcal{M}$ that makes use of only superclasses in the definition of the superposed state. Note that while NAFL theories like **NQM*** do not prove the existence of the superposed state, we may *encode* the superposed state, within **NQM*** and without any contradiction, as follows.

Let the state of the cat be denoted by the function $f(\tau, t)$, where $\tau \geq 0$, $t \geq 0$, and τ is the physical time, *i.e.*, $t = \tau$ represents the present and $t < \tau$ ($t > \tau$) represents the past (future). Let $f = 0$ ($f = 1$) denote that “The cat is dead (alive)” and let $f = 2$ stand for the superposed state of the cat, *i.e.*, “The cat is neither alive nor dead”. Note that we cannot define $f = \{0, 1\}$ to denote the superposed state because superclasses (like $\{0, 1\}$) are constants that can only occur as values of f at isolated points of discontinuity in the range of a function (see Sec. 4.5.1). Consider the case when the cat is observed to be alive at $t = 1$, and assume that the cat lives up to some finite time $c > 1$. The full time evolution of the state of the cat in the model $\mathcal{M}$ is encoded by the following multivalued function f defined in **NQM*** (see Sec. 4.5.2):

$$f(\tau, t) = 2, \quad \tau \in [0, 1], \quad t \in [0, 1], \quad (39)$$

$$= \{1, 2\}, \quad \tau = t = 1, \quad \text{dom}(f) \text{ not specified}, \quad (40)$$

$$= 1, \quad \tau \in [1, c], \quad t \in [0, c]. \quad (41)$$

Here $\text{dom}(f)$ is the domain of the relevant branch of f . Eq. (39) states that when $\tau \in [0, 1]$, the cat is in the superposed state, even as a prediction for future times up to $t = 1$, or in the past, up to $t = 0$. In (40), $f = \{1, 2\}$ expresses the fact that when $\tau = t = 1$, the (theory of the) model $\mathcal{M}$ does not prove that the cat is either alive or in the superposed state, because the direction of approach to $t = 1$ (*i.e.*, from the past or the future) is not specified. When $\tau \in [1, c]$, *i.e.*, after the box has been opened, we see from (41) that the observer has axiomatically declared that the cat is alive for all $t \in [0, c]$. Note that the state of the cat has multiple values when $t \in [0, 1]$, depending on whether t is in the present, past or future.

In the case when the cat is observed to be dead at $t = 1$ and is proven to die at $t = 0.5$, the following multivalued function of **NQM***, with a similar interpretation as above, represents the time evolution of the state of the cat in the model $\mathcal{M}$:

$$f(\tau, t) = 2, \quad \tau \in [0, 1], \quad t \in [0, 1],$$

$$= \{0, 2\}, \quad \tau = t = 1, \quad \text{dom}(f) \text{ not specified},$$

$$= 0, \quad \tau \geq 1, \quad t \geq 0.5,$$

$$= 1, \quad \tau \geq 1, \quad t \in [0, 0.5],$$

$$= 2, \quad \tau \geq 1, \quad t = 0.5.$$

5.3.2 Wigner’s friend paradox and the NAFL resolution

George Musser’s news article [27] describes the recent findings of Bong *et al* [28] on the Wigner’s friend paradox as follows:

Nearly 60 years ago, the Nobel Prize-winning physicist Eugene Wigner captured one of the many oddities of quantum mechanics in a thought experiment. He imagined a friend of his, sealed in a lab, measuring a particle such as an atom while Wigner stood outside. Quantum mechanics famously allows particles to occupy many locations at once – a so-called superposition – but the friend’s observation “collapses” the particle to just one spot. Yet for Wigner, the superposition remains: The collapse occurs only when he makes a measurement sometime later. Worse, Wigner also sees the friend in a superposition. Their experiences directly conflict.

Now, researchers in Australia and Taiwan offer perhaps the sharpest demonstration that Wigner’s paradox is real. In a study published this week in *Nature Physics*, [28] they transform the thought experiment into a mathematical theorem that confirms the irreconcilable contradiction at the heart of the scenario. The team also tests the theorem with an experiment, using photons as proxies for the humans. Whereas Wigner believed resolving the paradox requires quantum mechanics to break down for large systems such as human observers, some of the new study’s authors believe something just as fundamental is on thin ice: objectivity. It could mean there is no such thing as an absolute fact, one that is as true for me as it is for you.

Recent work on the Wigner’s friend paradox [29, 30, 28, 31], including rigorous no-go theorems and experimental confirmation, has been described in a comment by Brukner [32], who concludes:

It appears as if different observers live in different ‘bubbles’, with the notion of objectivity of events being limited to a single bubble. The terms ‘system’, ‘measurement device’, ‘unitary evolution’, ‘measurement’, ‘outcome’ and ‘facts’ are all defined only relative to a ‘bubble’.

Does this mean that quantum state assignment based on observations and measurement outcomes are subjective and expresses personal beliefs of individual observers? Perhaps, but the results of the no-go theorems do not force us to this conclusion. A single bubble may contain many observers, and they can all agree on the common facts arriving at an ‘objectivity within the bubble’. But across the boundaries of the bubbles there can be no agreement. The facts from other bubbles are not only inaccessible, but fundamentally incompatible: to regard the facts from other bubbles as objective, like one’s own, contradicts the predictions of quantum theory. Oddly

enough, this is true even if we know that the observers from other bubbles observe some facts.

The above quotes, while raising serious questions on the current realist interpretations of quantum mechanics, provide a spectacular confirmation of the subjective and axiomatic nature of NAFL truth, as defined in Sec. 2. In the NAFL interpretation of quantum superposition, Wigner’s friend paradox has a perfectly logical explanation. Note that an observer present inside the sealed box would observe the Schrödinger cat to be either alive or dead, *i.e.*, such an observer would not detect the superposed state of the cat. Nevertheless, for an external observer, the superposed state of the cat still exists, according to the NAFL interpretation. The internal observer has observed the classical state of the cat and has added the appropriate axiomatic declaration to his / her interpretation, whereas, for the external observer, the superposed (nonclassical) state of the cat has the logical meaning that there is no proof that the cat is either alive or dead. Clearly, there is no contradiction between these two states because the NAFL interpretation does not ascribe any physical reality to the superposed state.

It is important to note that according to the NAFL interpretation, an observer must be an intelligent living creature, capable of creating axiomatic theories and converting the real-world observations into axiomatic declarations. This calls into question the use of photons as observers in the experimental confirmation [28, 31] of the no-go theorems. A tacit assumption has been made here that photons do have the above attributes of an observer. While such an assumption is controversial, it certainly cannot be disproved. Perhaps a less controversial internal observer would be the Schrödinger cat itself, which would observe itself to be either alive or dead, while for the external observer, the cat is in a superposed state. According to the NAFL interpretation, the observations / axiomatic declarations of the internal observer are not part of the system of the external observer and vice versa, whereas other observable attributes of the internal observer, such as, the “alive” and “dead” states of the cat, are indeed part of the system of the external observer. This is in contrast to many current realist interpretations of quantum mechanics, according to which an observer inside the sealed box is entangled with the cat and would collapse the wave function of the cat, which is taken as an objective physical entity. The “bubbles” referred to by Brukner [32] in the above quote are, according to the NAFL interpretation, the minds of observers, ideally human minds. It is certainly the case that the observations / axiomatic declarations in a human mind (bubble) are inaccessible to other human minds (bubbles), as required by Brukner.

We conclude that the NAFL interpretation of quantum superposition correctly answers all the questions raised by other interpretations, such as, the Copenhagen interpretation (see Sec. 5.3). Indeed, the NAFL interpretation can be thought of as upholding the implied non-realist stance of the Copenhagen interpretation. Also note that the notions of past, present and future are indispensable for the NAFL interpretation. The axiomatic nature of NAFL truth provides the correct logical framework for handling the time-dependent truths

of quantum mechanics.

5.4 Quantum entanglement and the EPR paradox

Quantum entanglement occurs, for example, in a pair of particles, when the quantum state of either of the particles cannot be described independently of the state of the other particle. According to standard interpretations of quantum mechanics, such as, the Copenhagen interpretation, a measurement of a given property (such as, position, momentum or spin) on one of the particles collapses the wave function of the entangled pair of particles, so that both particles instantaneously acquire a definite value of that property and the outcome of a measurement of the same property on the other particle is predicted with certainty. The EPR paradox [33] occurs when the entangled pair of particles is separated by a large distance, so that information about a measurement made on one of the particles is seemingly communicated to the other particle at a velocity exceeding the speed of light, which is prohibited by special relativity theory.

5.4.1 The NAFL interpretation of quantum entanglement

Let X and Y be a pair of spatially separated entangled particles with properties A and B respectively, such that the proof syntax (p-syntax) of the NAFL theory **NQM** proves the equivalence $A \leftrightarrow B$. For example, if a spin-zero particle decays into a pair of spin-1/2 particles X and Y , then conservation of angular momentum would require that the total spin before and after decay should be zero. In this case, A (B) could denote the property “The spin of particle X (Y) is spin up (spin down) on some axis”. Since A and B are undecidable in **NQM**, $A \leftrightarrow B$ is not a legitimate proposition, and hence, not a theorem, in the theory syntax (t-syntax) of **NQM** (see the rules for constructing the t-syntax in Sec. 3.1.2). Hence prior to any measurement, both A and B are in the superposed state of “neither true nor false” in a nonclassical model of **NQM**. If at time t a measurement confirming the property A of the particle X is made, then the observer adds the axiom A to **NQM** to obtain the interpretation **NQM*** (see Sec. 2). It follows that **NQM*** proves both A and B , so that the observer instantly knows, at time t , that both A and B are true (with respect to **NQM**). Thus in the NAFL interpretation, the measurement made on the particle X is also simultaneously a measurement on the particle Y , and nothing gets communicated between the distant particles because the collapse of the wave function is not a physical occurrence. Further, the fact that the observer instantly knows a property of a distant particle without any local measurement on that particle implies that the concept of absolute simultaneity is inherent in **NQM**, which is not a problem because NAFL rejects special relativity theory as infinitary (see Sec. 5.2).

5.5 Wave-particle duality and the Afshar experiment

In the double-slit interferometer experiment, one observes an interference pattern at the detection screen whenever path information is not available for the photons, which is taken as confirmation of the wave nature of light. When path information is available, one observes the expected diffraction pattern at the detection screen, which is a confirmation of particle-like behavior of the photons. According to Bohr’s complementarity principle (BCP), one can observe (or infer) either the wave nature or the particle nature of light, but not both at the same time in any given experiment, such as, the double slit interferometer. Afshar [34] performed a variation of the double slit experiment in which a wire grid is placed at the minima of the interference pattern. From the fact that the wire grid does not alter the beams, Afshar infers the existence of an interference pattern, while suitably placed detectors confirm the path information. Afshar claimed that the detection of an interference pattern and path information in his experiment is a violation of BCP and the Englebert-Greenberger duality relation. There have been many mutually contradictory criticisms of Afshar’s argument, which has not been conclusively refuted by these authors.

5.5.1 The NAFL interpretation of the Afshar experiment

Srinivasan [35] has argued that Afshar’s claim that both an interference pattern and path information are present can be granted to be correct, but this does not entail a violation of BCP and the Englebert-Greenberger duality relation, essentially because of two reasons. Firstly, Afshar does not take into account the time dependence of formal truth, which is present in the NAFL interpretation. Secondly, Afshar, like many others and unlike the NAFL interpretation, takes a realist view of the superposed state of the photons and the supposedly consequent wave nature of light,

Let **NQM** be the theory that formalizes quantum mechanics in NAFL and let the proposition P ($\neg P$) denote “The photon took path A (path B) in the Afshar experiment”. Note that P is undecidable in **NQM** and let **NQM*** = **NQM** be the interpretation of **NQM** (see Sec. 2). In the NAFL interpretation, the superposed state of the photon is represented by $P \wedge \neg P$ in a nonclassical model of **NQM** generated by **NQM***, where P is interpreted as “There is no proof of $\neg P$ in **NQM***” and $\neg P$ is interpreted as “There is no proof of P in **NQM***”. Here both P and $\neg P$ have a nonclassical truth value of “neither true nor false” in the said nonclassical model of **NQM**.

Thus the NAFL interpretation does not ascribe any physical reality to the wave nature of the photon, as represented by the superposed state $P \wedge \neg P$, which has a purely logical meaning given above (amounting to an assertion of the absence of path information for the particle-like photon). Note that consistency of **NQM** requires that it does not *prove* $P \wedge \neg P$. Hence the NAFL interpretation does not preclude the possibility of path information becoming available at a later time, say, via observations in the detectors of the Afshar experiment, in which case the observer can retroactively assert P ($\neg P$) by taking

$\mathbf{NQM}^* = \mathbf{NQM} + P(\neg P)$. Note that consistency requires that the observer cannot assert $P \wedge \neg P$ as an axiom and this corresponds to the requirement that the wave nature of the photon can never be directly observed; any measurement can only yield path information for particle-like photons.

In the NAFL interpretation, while the presence of an interference pattern at any given time is a logical consequence of absence of path information *at that time*, the reverse implication does not hold, *i.e.*, the absence of path information (represented by $P \wedge \neg P$) cannot be inferred from the presence of an interference pattern. Indeed, this is a requirement of the consistency of $\mathbf{NQM}$. It is this logical fact, along with the time dependence of NAFL truth, that enables the retroactive assertion of path information in the Afshar experiment. Indeed, an interference pattern is present even in a single-photon version of the Afshar experiment [36, 37], and the NAFL interpretation asserts that a single photon is indeed particle-like and it is absurd to claim any “reality” for the “self-interference” of a single photon. At least in the single photon case, one must grant that the interference pattern can only be interpreted as a probability distribution for particle-like photons.

The time evolution of the state of the photons in the Afshar experiment is encoded by the following multivalued function in $\mathbf{NQM}^*$:

$$f(\tau, t) = 2, \quad \tau \in [0, 1], \quad t \in [0, 1], \quad (42)$$

$$= \{1, 2\}, \quad \tau = t = 1, \quad \text{dom}(f) \text{ not specified}, \quad (43)$$

$$= 1, \quad \tau \geq 1, \quad t \geq 0. \quad (44)$$

In analogy with (39) - (41), $t = \tau$ represents the present and $t < \tau$ ($t > \tau$) is the past (future). Here $f = 2$ encodes the superposed state of the photon ($P \wedge \neg P$) and $f = 1$ indicates the availability of path information ($P \vee \neg P$). The photons pass through the dual pinholes at $t = 0$ and strike the detectors at $t = 1$. When $\tau \in [0, 1]$, (42) indicates that the photons are (predicted to be) in the superposed state, for all $t \in [0, 1]$. At $t = \tau = 1$, the state of the photons is represented in (43) by the superclass $\{1, 2\}$, which essentially indicates that it is impossible to determine if $f = 1$ or $f = 2$ when the direction of approach of t to τ (from the past or future) is not specified. When $\tau \geq 1$, (44) asserts path information for the photons retroactively for $t \geq 0$. This retroactive assertion, which is valid only for $\tau \geq 1$ (*after* the photons have impinged on the detectors) does not entail a violation of the Bohr complementarity principle and the Englebert-Greenberger duality relation because at any given time $t = \tau$, the photon is in only one state, and no reality can be ascribed to the superposed state. Note again that the notions of past, present and future are indispensable for the NAFL interpretation, which is in violation of relativity theory.

5.6 The quantum Zeno effect and Zeno’s arrow paradox

When frequent measurements are made on a quantum system, its evolution will be slowed down and transition to states different from the initial state will be hindered. This phenomenon, known as the quantum Zeno effect (QZE), has

attracted considerable attention in the literature. [38, 39]. The analogy with Zeno’s arrow paradox was first noted by Misra and Sudarshan [40].

5.6.1 The NAFL interpretation of the quantum Zeno effect

A detailed analysis of the NAFL interpretation of the QZE would have to await a formalization of quantum mechanics in NAFL. However, at this stage, one can already conclude that the NAFL interpretation supports the theoretical and experimental findings on the QZE in two important ways.

- When the number of measurements is finite, we quote the following conclusions of Pascazio and Namiki, [41] which are corroborated by other authors [42, 43]:

“We have shown the QZE is liable to a purely dynamical explanation, which does not involve any projection operator. We claim therefore, contrary to widespread belief, that a quantum Zeno-type dynamics is *not* an argument in support of the collapse of the wave function, provided we observe the same state as the initial one at the final detector D_0 . The Schrödinger equation alone can yield a satisfactory explanation of the phenomenon. [...] We believe that a projection does not correspond to any *physical* operation and therefore should be regarded only as a convenient expedient (a ‘working rule’) in order to account for the loss of quantum mechanical coherence (the ‘collapse’ of the wave function). In this sense, von Neumann’s projection postulate is to be considered as purely mathematical and no physical meaning should be ascribed to it.”

This conclusion is fully supported by the NAFL interpretation of quantum superposition, in which no reality can be ascribed to the collapse of the wave function (see Sec. 5.3).

- The limit of infinitely many measurements (in which the quantum system is “frozen” in its initial state) is not physically attainable in finite time, *even in principle*, as seen from the following conclusion of Nakazato *et al.*: [44]

“We stress that in any conceivable experiment, only the QZE, with N finite (and rather small), can be observed. Our aim is to show that the $N \rightarrow \infty$ limit is physically unattainable, *even as a matter of principle*, and is rather to be regarded as a mathematical limit (although a very interesting one). In this sense, we shall say that the quantum Zeno *effect*, with N finite, becomes the quantum Zeno *paradox* when $N \rightarrow \infty$.”

Again the NAFL interpretation fully supports this conclusion. There is no quantum Zeno paradox in the NAFL interpretation because the notion

of infinitely many measurements cannot be formalized in any consistent NAFL theory of quantum mechanics. In particular, a superclass of infinitely many measurements in the QZE does not exist in any consistent NAFL theory because superclasses must include their limit points (see Remark 26). But in this case the $N \rightarrow \infty$ limit is physically unattainable and therefore cannot be considered as a valid limiting measurement.

Note also that, according to the main postulate of NAFL semantics (see Sec. 2) each measurement or observation is added as an axiomatic declaration to the NAFL theory formalizing quantum mechanics. Again, infinitely many such axiomatic declarations (as a function of time) cannot be conceivably be made in finite time because no such mapping exists, although the observer can add an infinite (but finitely expressed) axiom scheme at any given time.

5.6.2 The NAFL interpretation of Zeno's arrow paradox

Motion, say, of an arrow in flight, must necessarily occur in between instants of time. Zeno asserted that if time consisted *only* of instants, motion is impossible. To see precisely what Zeno had in mind, we quote Aristotle ([45], VI: 9, 239b5), where the emphasis is ours:

If everything when it occupies an equal space is at rest at that instant of time, and if that which is in locomotion is always occupying such a space at any moment, the flying arrow is therefore motionless at that instant of time and *at the next instant of time* but if both instants of time are taken as the same instant or continuous instant of time then it is in motion.

Clearly, Zeno required that if time consisted only of instants, then given any instant, there must exist a *next* instant and between those two instants motion cannot occur. This process, when iterated infinitely many times, leads to Zeno's conclusion that motion is impossible. The analogy with the quantum Zeno effect comes from the fact that if one could observe the evolution of time from one instant to the next, motion cannot conceivably occur and the arrow would remain frozen at its initial location.

Zeno's argument is refuted classically by the modern concept of the continuum: between any two instants, there must exist an uncountable infinity of instants and there is no such thing as a next instant. In other words, time does not evolve by passage from one instant to the next, as Zeno required; intervals of time must necessarily pass. The classical continuum is an infinitary object and is by no means uncontroversial. The intuitionists, for example, reject the classical notion of the continuum. Therefore, one can sympathize with Zeno's point of view.

In the logic NAFL, there must necessarily exist a superclass of instants between any two instants of time, which corresponds geometrically to a line segment, *i.e.*, a closed time interval (see. Sec. 4.4). The NAFL interpretation

rejects both the classical continuum and Zeno's arrow paradox by requiring that a time interval cannot be broken up into an uncountably infinite set of discrete points. The infinitely many observations or measurements of the arrow's position at every conceivable instant of time that Zeno was looking for cannot be made, *even in principle*, in analogy with the refutation of the quantum Zeno paradox in the NAFL interpretation.

5.7 Zeno's dichotomy paradox and supertasks

Achilles starts at $x = 0$ and runs at uniform velocity $dx/dt = 1$ along the x -axis. Before reaching his target at $(x, t) = (1, 1)$, Achilles would first have to cover half of the distance to reach $x = 1/2$ at $t = 1/2$. Assuming he does so, he would then have to cover half of the remaining distance to reach $x = 3/4$ at $t = 3/4$, and so on, *ad infinitum*. Zeno's ingenious and insightful observation was that in order to reach the target, Achilles would have to complete the infinite task (supertask) [46] of reaching infinitely many points $x \in \{0, 1/2, 3/4, \dots\}$ at infinitely many times $t \in \{0, 1/2, 3/4, \dots\}$. Zeno concluded that the supertask can never be completed in finite time, and this is the dichotomy paradox. Note that there is a purely temporal version of the dichotomy paradox, namely, that Achilles can never experience the instant $t = 1$ because to do so, he would have to complete the supertask of experiencing infinitely many instants $t \in \{0, 1/2, 3/4, \dots\}$.

Srinivasan [24] demonstrated that, contrary to widespread belief, modern mathematics does not really refute Zeno's paradox and its variants, which lead to meta-inconsistencies in classical infinitary reasoning. Metatheorem 1, Metatheorem 2 and Corollary 1 of Ref. [24], reproduced below as Metatheorem 13, Metatheorem 14 and Corollary 3 respectively, are the main results that establish these meta-inconsistencies (see Sec. 2, in particular, Remarks 10 and 11, of Ref. [24]).

Metatheorem 13. *Consider a physical process in which the time t passes through a strictly increasing, convergent sequence $(t_n)_{n \in \mathbb{N}}$ (also denoted as just (t_n)), within a finite time interval, and let $t_b = \lim_{n \rightarrow \infty} t_n$. Let S_1 be the supertask defined by t assuming, in sequence, each of the infinitely many values in (t_n) . Then the minimum value of the time at which S_1 gets completed is $t = t_b$.*

Proof. The proof consists of two simple steps which follow from the fact that t_b is the limit point of the strictly increasing sequence (t_n) .

- If $t < t_b$, then S_1 is incomplete.

This is seen from:

$$t < t_b \leftrightarrow \exists n (t < t_n < t_b) \leftrightarrow (S_1 \text{ is incomplete}).$$

- If $t = t_b$, then S_1 is complete.

This is obvious because t_b exceeds every value in (t_n) .

□

Remark 34. It follows from Metatheorem 13 that if the supertask S_1 is complete, that is, if t sequentially passes through all the infinitely many values in (t_n) , then t must necessarily attain the limiting value t_b . This establishes that a time of completion, namely, $t = t_b$, exists for the supertask S_1 .

Metatheorem 13 can be generalized to a wider range of supertasks, as will be seen from Metatheorem 14 below.

Metatheorem 14. Suppose that, in a given physical process, an object has a time-dependent state $S(t)$, where $t \in [t_a, t_c]$ and $t_c > t_a$. Here t is defined on the standard real number system $\mathbb{R}$, while, for our purposes, $S(t)$ is defined on the extended real number system $\bar{\mathbb{R}}$. [16] At a given time t_b , where $t_a < t_b \leq t_c$, suppose that $\lim_{t \rightarrow t_b^-} S(t)$ exists. Then $S(t_b) = \lim_{t \rightarrow t_b^-} S(t)$.

Proof. Consider any strictly increasing Cauchy sequence of times (t_n) , defined within the time interval $[t_a, t_b)$ such that $\lim_{n \rightarrow \infty} t_n = t_b$. For example, taking $(t_a, t_b, t_c) = (0, 1, 2)$, consider the sequence (t_n) defined as:

$$t_n = 1 - (1/2)^n, \quad n = 0, 1, 2, \dots, \quad (45)$$

Given the arrow of time in a physical process, the object in question completes the supertask S_2 of attaining, in sequence, an actual infinity of states $(S(t_n))_{n \in \mathbb{N}}$. Clearly, S_2 occurs simultaneously with the supertask S_1 defined in Metatheorem 13, in the sense that there is a one-to-one correspondence in time between the steps of S_1 and S_2 . It follows that a time of completion, namely, $t = t_b$, exists for the supertask S_2 (see Remark 34). We claim that as t attains its limiting value t_b upon completion of S_1 , the completion of S_2 must also happen with $S(t)$ attaining its limiting value $\lim_{n \rightarrow \infty} S(t_n)$, which exists, because $\lim_{n \rightarrow \infty} S(t_n) = \lim_{t \rightarrow t_b^-} S(t)$. Clearly, as t passes through an actual infinity of values t_n to attain the limiting value t_b , the following must hold:

$$\forall \epsilon > 0 \exists n (|S(t_n) - \lim_{n \rightarrow \infty} S(t_n)| < \epsilon). \quad (46)$$

Here $|S(t_n) - \lim_{n \rightarrow \infty} S(t_n)|$ can be viewed as a distance on the real line, which, upon the completion of the supertask S_2 , must shrink to less than every positive real number ϵ . It is a direct mathematical consequence of (46) that the completion of S_2 must happen with $S(t)$ jumping to the value $\lim_{n \rightarrow \infty} S(t_n)$. What is being asserted here is that both the values $t = t_b$ and $S(t) = \lim_{n \rightarrow \infty} S(t_n)$ must be physically realized upon completion of the supertask S_2 . It follows that $S(t_b) = \lim_{n \rightarrow \infty} S(t_n) = \lim_{t \rightarrow t_b^-} S(t)$. $\square$

Corollary 3. In the physical process defined in Metatheorem 14, suppose

$$\lim_{t \rightarrow t_b^-} S(t) \neq \lim_{t \rightarrow t_b^+} S(t).$$

Then such a discontinuous physical process is not time reversible.

Proof. Without loss of generality, assume $(t_a, t_b, t_c) = (0, 1, 2)$. Metatheorem 14, as applied in the physical time variable t , implies that $S(1) = \lim_{t \rightarrow 1^-} S(t)$. Make the change of variable $t' = 2 - t$ and let $S'(t') = S(t)$. To arrive at a contradiction, apply Metatheorem 14 in the variable t' , with $(t'_a, t'_b, t'_c) = (0, 1, 2)$. We find that $S(1) = S'(1) = \lim_{t' \rightarrow 1^-} S'(t') = \lim_{t \rightarrow 1^+} S(t)$, which contradicts the previous application of Metatheorem 14 in the variable t . $\square$

From Metatheorem 14, we may infer that in general, supertasks cannot be completed:

Metatheorem 15. *Let $S(t)$ denote the state of an object in an arbitrary physical process defined in Metatheorem 14, where we take, without loss of generality, $t \in [0, 2]$ and $(t_a, t_b, t_c) = (0, 1, 2)$. Let the sequence $(S(t_n))_{n \in \mathbb{N}}$ represent the supertask generated by t_n assuming all the values in the sequence $(t_n)_{n \in \mathbb{N}}$, defined as in (45). Then $(S(t_n))_{n \in \mathbb{N}}$ cannot be completed.*

Proof. Consider a physical process in which the state function $\bar{S}(t)$ is defined on $t \in [0, 2]$ by

$$\bar{S}(t) = 0 \text{ if } t \in [0, 1) \text{ and } \bar{S}(t) = 1 \text{ if } t \in (1, 2]. \quad (47)$$

We may consider $\bar{S}(t)$ as defining a test physical process, which is used to *measure* the completion status of an arbitrary physical process whose state function is $S(t)$, as follows. We take “ $\bar{S}(t) = 0$ ” to encode (or measure) “ $\exists n \in \mathbb{N}(t_n > t)$ ”, which is equivalent to “ $(S(t_n))_{n \in \mathbb{N}}$ is incomplete at time t ”. Similarly, “ $\bar{S}(t) = 1$ ” encodes (or measures) “ $\forall n \in \mathbb{N}(t_n < t)$ ”, which is equivalent to “ $(S(t_n))_{n \in \mathbb{N}}$ is complete at time t ”. Metatheorem 14, with $(t_a, t_b, t_c) = (0, 1, 2)$, requires that $\bar{S}(1) = 0$, which encodes (or measures) “ $(S(t_n))_{n \in \mathbb{N}}$ is incomplete at $t = 1$ ”. This contradiction proves that $(S(t_n))_{n \in \mathbb{N}}$ cannot be completed. $\square$

Remark 35. *Note that the proof of Metatheorem 15 does not make use of the failure of time reversal invariance [47] implied by Corollary 3, which can also be considered as a proof by contradiction that supertasks cannot be completed.*

Remark 36. *Metatheorem 15 implies a meta-inconsistency in classical infinitary reasoning, which requires that supertasks (such as, those arising in Zeno’s dichotomy paradox and its variants considered in Ref. [24]) can be completed. For example, in the case of Zeno’s dichotomy paradox, let $(S(t_n))_{n \in \mathbb{N}} = (x(t_n))_{n \in \mathbb{N}}$, where $x(t)$ is Achilles’s position at time t , as defined earlier. We know that $x(1) = 1$, i.e., Achilles’s constant velocity requires that he reach $x = 1$ at $t = 1$. But Metatheorem 15 implies that the supertask $(x(t_n))_{n \in \mathbb{N}}$ cannot be completed, i.e., Achilles cannot reach $x = 1$ at $t = 1$.*

5.7.1 Resolution of the dichotomy and related paradoxes in NAFL

In the NAFL theory of real numbers **NPAR** (see Sec. 4.5), the only collections of real numbers that are admissible are superclasses, defined in Sec. 4.4. In

particular, open / half-open intervals and a countable infinity of real numbers do not exist in **NPAR** (see Remark 26 and Sec. 4.5.1). This immediately implies that Zeno's dichotomy paradox, which requires the existence of the half-open interval of reals $[0, 1)$, cannot be formulated in NAFL. The classical definition of a supertask, which requires a countable infinity of real numbers, is also not admissible in NAFL. Hence paradoxes involving supertasks (such as, Thomson's lamp experiment, [48] the "beautiful supertask" of Laraudogoitia [49] and its variant wherein a baton is passed on with each elastic collision (see the infinite relay paradox in Ref. [24])) cannot be formulated in NAFL.

Metatheorem 14, Corollary 3 and Metatheorem 15 cannot be formalized in NAFL because they make use of the classical notion of supertasks. This raises the issue of the definition of supertasks in the NAFL theory **NPAR**, and whether they can be completed. As noted in Sec. 4.5.1, when the natural numbers are embedded in the extended real line of **NPAR**, one obtains the superclass $\mathbb{N}_{\mathbb{R}} = \{0, 1, \dots, \infty\}$. Consider a physical process with state $S(t)$ defined as in Metatheorem 14. Consider any strictly increasing Cauchy sequence of times $(t_n)_{n \in \mathbb{N}}$, defined within the time interval $[t_a, t_b)$ such that $\lim_{n \rightarrow \infty} t_n = t_b$. The NAFL supertask (denoted as n-supertask) corresponding to the classical supertask $(S(t_n))_{n \in \mathbb{N}}$ is defined as

$$(S(t_n))_{n \in \mathbb{N}_{\mathbb{R}}} = (S(t_0), S(t_1), \dots, \lim_{n \rightarrow \infty} S(t_n)). \quad (48)$$

Note that $\lim_{n \rightarrow \infty} S(t_n) = \lim_{t \rightarrow t_b^-} S(t)$, which is assumed to exist.

Metatheorem 16. *In the physical process with state $S(t)$ defined as in Metatheorem 14, the n-supertask $(S(t_n))_{n \in \mathbb{N}_{\mathbb{R}}}$ given by (48) must necessarily get completed, i.e., the limiting value $\lim_{n \rightarrow \infty} S(t_n) = \lim_{t \rightarrow t_b^-} S(t)$ gets realized physically.*

Proof. By assumption, every finite subsequence of the n-supertask $(S(t_n))_{n \in \mathbb{N}_{\mathbb{R}}}$ can be completed. The metatheorem follows from this assumption and the following two requirements (see Sec. 4.5.1 and Remark 26):

- The n-supertask $(S(t_n))_{n \in \mathbb{N}_{\mathbb{R}}}$, which is a sequence of real numbers, is encoded in the theory **NPAR** as a superclass, which must necessarily include its limit points.
- The domain and range of the function S must be superclasses in the regions where S is continuous.

□

Remark 37. *Note that Metatheorem 16 is a finitistically valid result (unlike the completion of the classical supertask $(S(t_n))_{n \in \mathbb{N}}$) because the n-supertask $(S(t_n))_{n \in \mathbb{N}_{\mathbb{R}}}$ is finitarily encoded as a superclass, as noted in the above proof. Metatheorem 16 implies that $S(t_b) = \lim_{t \rightarrow t_b^-} S(t)$ whenever the domain of the function S is restricted to $t \in [t_a, t_b]$. Because the NAFL theory **NPAR** permits multivalued functions (see Sec. 4.5.2), it does not follow that this value of $S(t_b)$*

holds when the domain of S is specified as $t \in [t_b, t_c]$, in which case $S(t_b) = \lim_{t \rightarrow t_b^+} S(t)$, and it is possible that $\lim_{t \rightarrow t_b^-} S(t) \neq \lim_{t \rightarrow t_b^+} S(t)$. It follows that Corollary 3, which cannot be formalized in **NPAR**, does not hold and time reversal invariance is upheld in the NAFL version of the physical process defined by $S(t)$, even at points of discontinuity of S .

Consider Zeno’s dichotomy paradox as stated in Sec. 5.7, where $t \in [0, 2]$, and add the requirement that Achilles dies at $t = 1$ upon reaching his target (this is the “spatial dichotomy paradox” discussed in Sec. 1.5.1 of Ref. [24]). From Metatheorem 14, we may deduce the contradiction that Achilles must be alive at $t = 1$. Assuming Achilles is dead for times $t > 1$, Corollary 3 implies that in the time reversed version of Achilles’s run with the change of variable $t' = 2 - t$, Achilles must be dead at $t' = t = 1$, with the consequent failure of time reversal invariance, which is another contradiction. However, in the NAFL theory **NPAR**, these contradictions are avoided and time reversal invariance holds. Achilles’s state of “alive” or “dead” is specified in **NPAR** by the following multivalued function:

$$\begin{aligned} f(t) &= 0, & t \in [0, 1], \\ &= \{0, 1\}, & t = 1, \text{ dom}(f) \text{ unspecified}, \\ &= 1, & t \in [1, 2]. \end{aligned}$$

Here $f = 0$ ($f = 1$) means that Achilles is alive (dead), $\{0, 1\}$ is a superclass and $\text{dom}(f)$ is the domain of the relevant branch of the function f .

5.8 Benacerraf’s shrinking genie

Benacerraf [50] proposed a version of Zeno’s dichotomy paradox in which Achilles is replaced by a genie whose height $h(t)$ shrinks continuously in proportion to the distance covered on $x \in [0, 1]$, as follows.

$$h(t) = h_0(1 - t), \quad x(t) = t, \quad t \in [0, 1], \quad (49)$$

where we assume that the genie is a one-dimensional creature and $h_0 > 0$ is the initial height of the genie. Clearly, the genie, which runs at unit velocity, reaches every point $x \in [0, 1]$, but does not reach $x = 1$, where it vanishes. Here it is assumed that the genie no longer exists when its height has reduced to zero.

Consider a hypothetical universe in which the only objects are identical shrinking genies, so that the genies have no length scale available other than their own height. In addition to the above traveling genie, let there be a stationary genie at $x = 0$ and a target genie at $x = 1$ in the postulated hypothetical universe. The stationary genie would consider its height to be fixed at h_0 (assumed to be the only available length scale) and would therefore observe that both the traveling genie and the target genie are also at fixed height h_0 , at a fixed unit distance from each other, and are both receding at an increasing velocity proportional to $(1 - t)^{-2}$, as can be seen from the following. Recall that the traveling genie starts at $x = 0$ and travels at unit velocity (as stated

in (49)). Let $(x_1(t), h_1(t))$ and $(x_2(t), h_2(t))$ be the coordinates of the traveling genie and the target genie respectively, from the point of view of the stationary genie, which fixes its height at h_0 as the standard unit of length. It is easy to see that

$$\begin{aligned} x_1(t) &= \frac{t}{1-t}, & x_2(t) &= \frac{1}{1-t}, & t \in [0, 1) \rightarrow (h_1(t) = h_2(t) = h_0), \\ \frac{dx_1}{dt} &= \frac{dx_2}{dt} = \frac{1}{(1-t)^2}, & t \in [0, 1) \rightarrow (x_2(t) - x_1(t) = 1), \\ x_1(1) &= x_2(1) = \frac{dx_1}{dt}(1) = \frac{dx_2}{dt}(1) = +\infty, \end{aligned} \quad (50)$$

where $x_1(t) \in \bar{\mathbb{R}}$, $x_2(t) \in \bar{\mathbb{R}}$, $t \in [0, 1]$ and $\bar{\mathbb{R}}$ denotes the extended real number system, [16] which has been used to remove the singularity at $t = 1$. From the point of view of an observer external to the postulated hypothetical universe, the infinities occur here because the length scale of the stationary genie has shrunk to zero at $t = 1$. Observe that $(x_2(1) - x_1(1))$, being of the form $(\infty - \infty)$, is undefined in the extended real number system. We also need to define $h_1(1)$ and $h_2(1)$. There are three possibilities here, all of which lead to contradictions, as follows. A natural definition, which follows from an application of Metatheorem 14, would be

$$x_2(1) - x_1(1) = 1, \quad h_1(1) = h_2(1) = h_0, \quad (51)$$

which implies that the traveling genie does not complete the two supertasks of catching up with the target genie and also vanishing (along with the target genie) at $t = 1$. Alternatively, we could define, in partial violation of Metatheorem 14,

$$x_2(1) - x_1(1) = 0, \quad h_1(1) = h_2(1) = h_0. \quad (52)$$

Here the viewpoint of the stationary genie would confirm that the traveling genie catches up with the target genie at nonzero height h_0 and collides with it at $t = 1$, in contradiction to (49). Thirdly, we could set, in violation of Metatheorem 14,

$$h_1(1) = h_2(1) = 0. \quad (53)$$

In this case, by assumption, neither the genies nor their coordinates $x_1(1)$ and $x_2(1)$ exist at $t = 1$ and therefore the traveling genie does not complete the supertask of catching up with the target genie before both vanish.

Applying time reversal to (50) will lead to further absurdities at $t = 1$ and Corollary 3 shows that time reversal invariance is not upheld. Similar contradictions can be deduced from an application of Metatheorem 14 and Corollary 3 to (49). These contradictions must be attributed to the underlying classical assumption that supertasks can be completed.

5.8.1 Newtonian kinematics versus relativistic kinematics

It should be emphasized that from the point of view of Newtonian kinematics, (49) and (50) are entirely equivalent. An observer external to the postulated

hypothetical universe would hold the view that the genies are shrinking as in (49), while the genies would maintain that their universe is expanding with respect to their fixed height h_0 , as in (50). Purely from the point of view of kinematics, it *ought* to be impossible to say which of these two scenarios is the underlying “reality”, if one rejects Platonism as a philosophy of physics. However, this equivalence between (49) and (50) does not hold in the kinematics of special relativity theory (**SR**), wherein the Lorentz transformations apply and according to which (50) is illegal, essentially due to an illegitimate choice of the standard of length by the genies. Indeed, **SR** requires that the length standard be chosen such that the velocity of light in vacuum is a *defined* constant c , in order to conform with the light postulate of **SR**. Assuming that the standard of length in the coordinate system of (49) is chosen to uphold the defined value c of the velocity of light, the time τ_g taken by light to traverse the height of the genies would be measured, by an observer external to the postulated hypothetical universe, as

$$\tau_g = \frac{h_0(1-t)}{c}, \quad t \in [0, 1]. \quad (54)$$

Hence the external observer would conclude that the genies are shrinking, and that their height is $h_0(1-t)$. The genies in the postulated hypothetical universe would also measure exactly the same value of τ_g as in (54) for light to traverse their height, which they fix at the value h_0 as their standard of length. It follows that the velocity of light in vacuum (c_g) in the coordinate system of (50) would be *measured* by the genies as

$$c_g \approx \frac{c}{1-t}, \quad t \in [0, 1), \quad (55)$$

where the approximation is valid to leading order as $c \rightarrow \infty$ or as $h_0 \rightarrow 0$ or, in general, as $h_0/c \rightarrow 0$. Clearly, (55) is in violation of the light postulate of **SR** and it follows that (50) is illegitimate in **SR**. Note also that the velocities of the genies in (50) exceed the velocity of light c_g in (55) by an order of magnitude as $t \rightarrow 1^-$, which is also not allowed in **SR**. Thus relativistic kinematics picks out (49) as a Platonic reality and one may conclude that Platonism, rejected in the proposed finitistic logic NAFL, is inherent in **SR**. The inability of relativistic kinematics to support (50) in the postulated hypothetical universe suggests a possible meta-inconsistency in **SR**, because our contention is that (50) ought only to be rejected, if at all, by invoking alternative theories, such as, dynamics.

5.8.2 The NAFL resolution of Benacerraf’s shrinking genie paradox

We have seen that in the NAFL theory of real numbers **NPAR**, Metatheorem 14, Corollary 3 and Metatheorem 15 cannot be formalized. Instead, Metatheorem 16 and Remark 37 apply. This entails that the traveling genie (or at least its center of mass) must continue to exist at $t = 1$ because existence of the genie cannot be formulated on the half-open interval $[0, 1)$, which is not

a superclass. In the coordinate system of (50), the distance $(x_2(t) - x_1(t))$ and the height $h(t)$ are specified in **NPAR** by the following multivalued functions:

$$\begin{aligned} f(t) = (x_2(t) - x_1(t)) &= 1, & t \in [0, 1], \\ &= \{0, 1\}, & t = 1, \text{ dom}(f) \text{ unspecified}, \\ &= 0, & t = 1 \text{ and dom}(f) = [1, \infty]. \end{aligned}$$

$$\begin{aligned} h(t) &= h_0, & t \in [0, 1], \\ &= \{0, h_0\}, & t = 1, \text{ dom}(h) \text{ unspecified}, \\ &= 0, & t = 1 \text{ and dom}(h) = [1, \infty]. \end{aligned}$$

Here $\text{dom}(\cdot)$ is the domain of the relevant branch of the indicated function, and $\{0, 1\}$ and $\{0, h_0\}$ are superclasses.

Note that $(x_2(t) - x_1(t))$ and $h(t)$ drop discontinuously to 0 at $t = 1$. These multivalued functions resolve the various contradictions reported in (51) - (53), uphold time reversal invariance and fully establish the kinematic equivalence of (49) and (50) in Newtonian mechanics; note that **SR** is inconsistent by the NAFL yardstick (see Sec. 5.2). Post $t = 1$, the genies are point masses and their evolution can be continued with a different length scale chosen as in (49), which will replace (50).

6 Concluding remarks

The exposition of non-Aristotelian finitary logic (NAFL) in this paper upholds virtually all the ideas expressed in the author's previous work, [6, 15, 24, 35] but presents these and several new results with much greater clarity and precision. We have demonstrated that the scope of finitism is expanded from just a theory (Primitive recursive arithmetic) [4] within classical logic to an entire logic (NAFL) that, in our opinion, rivals classical logic and intuitionism. The distinguishing feature of finitism in NAFL is that classical infinitary reasoning is not banned arbitrarily, but is *refuted* via logical principles deducible from finitary NAFL semantics, which correctly captures the notion of a potential infinity. NAFL, which provides a negative resolution of Hilbert's program, has profound consequences for the logical foundations of physics, both classical and quantum. The paraconsistent, time-dependent and axiomatic truths of NAFL provide the correct logical framework for resolution of the paradoxes of quantum mechanics, such as, superposition, Wigner's friend paradox, entanglement, wave-particle duality and the quantum Zeno effect, as well as classical paradoxes, such as, Zeno's dichotomy paradox and its variants, which have remained unresolved for almost 2500 years. Future work will focus on the logical foundations of quantum mechanics, which will require further development of real analysis in NAFL. For the sake of brevity, we had to omit discussion of philosophical paradoxes, such as, those arising from self-reference or the Platonic nature of classical truth. These will be elaborated upon in future work.

References

- [1] C. C. Chang and H. J. Keisler, *Model Theory*, 3rd ed. *Studies in Logic and the Foundations of Mathematics* **73** (North-Holland, Amsterdam, 1990).
- [2] W. Hodges, *A Shorter Model Theory* (Cambridge University Press, Cambridge, 1997).
- [3] E. Mendelson, *Introduction to Mathematical Logic*, 5th ed. (Taylor & Francis, 2009).
- [4] W. W. Tait, Finitism, *The Journal of Philosophy* **78** (1981) 524–546.
- [5] M. Davis, The incompleteness theorem, *Notices of the AMS* **53** (2006) 414–418.
- [6] R. Srinivasan, The quantum superposition principle justified in a new non-Aristotelian finitary logic, *Int. J. Quantum Inf.* **3** (2005) 263–267.
- [7] Ø. Linnebo and S. Shapiro, Actual and potential infinity, *Noûs* **53**, No. 1 (2019) 160–191.
- [8] S. G. Simpson, Foundations of Mathematics: an Optimistic Message. In: R. Kahle, M. Rathjen (eds) *The Legacy of Kurt Schütte* (Springer, Cham. (2020) https://doi.org/10.1007/978-3-030-49424-7_20).
- [9] M. Eberl, A model theory for the potential infinite, *Rep. Math. Log.* **57** (2022) 3–30.
- [10] A. Sloman, Rules of inference, or suppressed premisses?, *Mind* (New series) **73**, No. 289 (1964) 84–96.
- [11] P. M. Cohn, *Basic Algebra: Groups, Rings, and Fields* (Springer, 2003).
- [12] A. Wiles, Modular elliptic curves and Fermat’s Last Theorem, *Annals of Mathematics* **141**, No. 3 (1995) 443–551.
- [13] R. Taylor and A. Wiles, Ring theoretic properties of certain Hecke algebras, *Annals of Mathematics* **141**, No. 3 (1995) 553–572.
- [14] S. G. Simpson, *Subsystems of second order arithmetic* (Cambridge University Press, 2009).
- [15] R. Srinivasan and H. P. Raghunandan, Foundations of real analysis and computability theory in non-Aristotelian finitary logic, arXiv:math.LO/0506475.
- [16] C. D. Alprantis and O. Burkinshaw, *Principles of Real Analysis* (Academic Press, 1998).

- [17] K. Knopp, Multiple-Valued Functions, Section II in *Theory of Functions Parts I and II, Two Volumes Bound as One* (New York: Dover, Part I p. 103 and Part II pp. 93-146, 1996).
- [18] P. M. Morse and H. Feshbach, Multivalued Functions, §4.4 in *Methods of Theoretical Physics, Part I* (New York: McGraw-Hill, pp. 398-408, 1953).
- [19] R. Zach, Hilbert's Program, *The Stanford Encyclopedia of Philosophy* (Fall 2019 Edition), Edward N. Zalta (ed.), <https://plato.stanford.edu/archives/fall2019/entries/hilbert-program/>.
- [20] D. Freeborn, M. Gilton and C. Mitsch, How Haag-tied is QFT, really? (Preprint (2022), PhilSci Archive, <http://philsci-archive.pitt.edu/21552/>).
- [21] J. Earman and D. Fraser, Haag's theorem and its implications for the foundations of quantum field theory, *Erkenntnis* **64** (2006) 305–344.
- [22] M. Mamone-Capria, On the incompatibility of special relativity and quantum mechanics, *Journal for Foundations and Applications of Physics* **5**, No. 2 (2018) 163–189.
- [23] J. X. Madarász, I. Németi and G. Székely, First-order logic foundation of relativity theories. In: Gabbay, D.M., Zakharyashev, M., Goncharov, S.S. (eds) *Mathematical Problems from Applied Logic II. International Mathematical Series* **5** (2007) 217–252 (Springer, New York, NY).
- [24] R. Srinivasan, Logical foundations of physics. Part I. Zeno's dichotomy paradox, supertasks and the failure of classical infinitary reasoning (Preprint (2019), PhilSci Archive, <http://philsci-archive.pitt.edu/15799/>).
- [25] E. Schrödinger, Die gegenwärtige Situation in der Quantenmechanik (The present situation in quantum mechanics), *Naturwissenschaften* **23**, No. 48 (1935) 807–812.
- [26] R. Omnès, The Copenhagen interpretation, in: *Understanding Quantum Mechanics* (Princeton University Press, 1999, pp. 41–54).
- [27] George Musser, Quantum paradox points to shaky foundations of reality, *Science*, (doi: 10.1126/science.abe3693), <https://www.science.org/content/article/quantum-paradox-points-shaky-foundations-reality/>.
- [28] K. W. Bong, A. Utreras-Alarcón, F. Ghafari *et al*, A strong no-go theorem on the Wigner's friend paradox, *Nat. Phys.* **16** (2020) 1199 –1205.
- [29] D. Frauchiger and R. Renner, Quantum theory cannot consistently describe the use of itself, *Nat. Commun.* **9** (2018) 3711.
- [30] Č. Brukner, A no-go theorem for observer-independent facts, *Entropy* **20** (2018),350.

- [31] M. Proietti *et al*, Experimental test of local observer independence, *Sci. Adv.* **5** (2019) eaaw9832.
- [32] Č. Brukner, Wigner’s friend and relational objectivity, *Nat Rev Phys* **4** (2022) 628–630.
- [33] A. Einstein, B. Podolsky and N. Rosen, Can Quantum-Mechanical Description of Physical Reality Be Considered Complete?, *Phys. Rev.* **47**, No. 10 (1935) 777–780.
- [34] S. S. Afshar, E. Flores, K. F. McDonald and E. Knoesel, Paradox in wave-particle duality, *Foundations of Physics* **37** (2007) 295–305.
- [35] R. Srinivasan, Logical analysis of the Bohr complementarity principle in Afshar’s experiment under the NAFL interpretation, *Int. J. Quantum Inf.* **8** (2010) 465–491.
- [36] V. Jacques *et al.*, Illustration of quantum complementarity using single photons interfering on a grating, *New Journal of Physics* **10**, No. 12 (2008) 123009.
- [37] D. D. Georgiev, Single photon experiments and quantum complementarity, *Progress in Physics* **2** (2007) 97–103.
- [38] P. Facchi and S. Pascazio, Quantum Zeno dynamics: mathematical and physical aspects, *J. Phys. A: Math. Theor.* **41** (2008) 493001.
- [39] W. M. Itano, The quantum Zeno paradox, 42 years on, *Current Science* **116**, No. 2 (2019) 201–204.
- [40] B. Misra and E. C. G. Sudarshan, The Zeno’s paradox in quantum theory, *J. Math. Phys.* **18** (1977) 756–763.
- [41] S. Pascazio and M. Namiki, Dynamical quantum Zeno effect, *Phys. Rev. A* **50** (1994) 4582–4592.
- [42] S. Pascazio, Dynamical origin of the quantum Zeno effect, *Found. Phys.* **27** (1997) 1655–1670.
- [43] T. Petrosky, S. Tasaki and I. Prigogine, Quantum Zeno effect, *Phys. Lett. A* **151** (1990) 109–113.
- [44] H. Nakazato, M. Namiki, S. Pascazio and H. Rauch, On the quantum Zeno effect, *Phys. Lett. A* **199** (1995) 27–32.
- [45] Aristotle (Translated by Robin Waterfield), *Physics* (Oxford University Press, 2008).
- [46] J. Earman and J. Norton, Infinite pains: The trouble with supertasks, in *Benacerraf and his Critics*, eds. A. Morton and S. P. Stich (Blackwell, 1996), pp. 231–261.

- [47] J. Earman, What time reversal invariance is and why it matters, *International Studies in the Philosophy of Science* **16** (2002) 245–264.
- [48] J. F. Thomson, Tasks and super-tasks, *Analysis* **15** (1954) 1–13.
- [49] J. P. Laraudogoitia, A beautiful supertask, *Mind* **105** (1996) 81–83.
- [50] P. Benacerraf, Tasks, super-tasks, and the modern eleatics, *The Journal of Philosophy* **59** (1962) 765–784.