

THE CONNES EMBEDDING PROBLEM: A GUIDED TOUR

ISAAC GOLDBRING

ABSTRACT. The Connes embedding problem (CEP) is a problem in the theory of tracial von Neumann algebras and asks whether or not every tracial von Neumann algebra embeds into an ultrapower of the hyperfinite II_1 factor. The CEP has had interactions with a wide variety of areas of mathematics, including C^* -algebra theory, geometric group theory, free probability, and non-commutative real algebraic geometry, to name a few. After remaining open for over 40 years, a negative solution was recently obtained as a corollary of a landmark result in quantum complexity theory known as $\text{MIP}^* = \text{RE}$. In these notes, we introduce all of the background material necessary to understand the proof of the negative solution of the CEP from $\text{MIP}^* = \text{RE}$. In fact, we outline two such proofs, one following the “traditional” route that goes via Kirchberg’s QWEP problem in C^* -algebra theory and Tsirelson’s problem in quantum information theory and a second that uses basic ideas from logic.

CONTENTS

1. Introduction	503
2. Equivalent reformulations of CEP	507
3. A crash course in operator algebras	510
4. A crash course in complexity theory	527
5. A quantum detour	534
6. From $\text{MIP}^* = \text{RE}$ to the failure of CEP: the traditional route	541
7. From $\text{MIP}^* = \text{RE}$ to the failure of CEP: a model-theoretic shortcut	545
8. The enforceable II_1 factor (should it exist)	553
Acknowledgments	558
About the author	558
References	558

1. INTRODUCTION

1.1. What is this all about? The story told in this tour is (in this author’s humble opinion) absolutely fascinating! It can also be completely confusing and terrifying to an outsider. It contains a seemingly infinite number of acronyms (CEP, WEP, QWEP, LLP, MIP^* , RE, . . .), all sorts of tensor products ($\bar{\otimes}$, $\otimes_{\max}$, $\otimes_{\min}$), entangled particles, and even good friends Einstein and Gödel both make an appearance (the latter twice).

Received by the editors October 4, 2021.

2020 *Mathematics Subject Classification.* Primary 46L10, 46L06, 81P45, 03C66.

The author was partially supported by NSF grant DMS-2054477.

At one end of the story is the *Connes embedding problem* (CEP), a problem in the field of *von Neumann algebras* first posed by Alain Connes in his famous 1976 paper “Classification of Injective Factors” [15] (the paper mainly responsible for his being awarded the Fields Medal in 1982). Roughly speaking, a von Neumann algebra is a collection of bounded operators on a Hilbert space containing the identity operator, closed under addition, composition, scalar multiplication, and adjoint, and which is closed in a certain topology known as the weak operator topology. The von Neumann algebras Connes was considering came equipped with a trace functional that shares many of the nice properties enjoyed by the (normalized) trace functional on matrices.

Here is the passage from [15] which led to the establishment of the CEP:

We now construct an approximate imbedding of N in $\mathcal{R}$. Apparently such an imbedding ought to exist for all II_1 factors because it does for the regular representation of free groups. However, the construction below relies on condition 6.

What is this quote trying to convey? $\mathcal{R}$ is the *hyperfinite II_1 factor*, arguably the most important tracial von Neumann algebra. For now, one should just think of $\mathcal{R}$ as an appropriate limit of matrix algebras $M_n(\mathbb{C})$ of increasing sizes. We will have much to say about this algebra throughout this paper. A II_1 factor is just a particular kind of tracial von Neumann algebra, and the N appearing in the passage is a particular II_1 factor satisfying a certain list of properties. By an approximate imbedding of N in $\mathcal{R}$, Connes means that any finite amount of *information* about elements of N (that is, the trace of finitely many $*$ -polynomials with elements from N plugged in) can be *simulated* by appropriate elements of $\mathcal{R}$. Connes later shows that such approximate imbeddings correspond to actual embeddings of N into a so-called *ultrapower* of $\mathcal{R}$, denoted $\mathcal{R}^u$. He comments that such an embedding “ought” to always exist since it does for a particular von Neumann algebra, namely the group von Neumann algebra associated to the free group, denoted $L(\mathbb{F}_2)$ (see Subsection 3.7). Why that embedding “ought to exist” is not quite clear. Nevertheless, Connes is only able to show that the N under consideration can be embedded in $\mathcal{R}^u$ using one of the conditions (namely the sixth one) he has assumed about this particular algebra.

Thus, the Connes embedding problem (CEP) states: every tracial von Neumann algebra embeds (in a trace-preserving way) in an ultrapower $\mathcal{R}^u$ of $\mathcal{R}$. We will say this slightly more precisely in Subsection 3.6. Many prefer to call this a “problem” rather than a “conjecture” since “ought to” is not a very strong sentiment.

The robustness of the CEP lies in its many reformulations and from the many areas of mathematics it has touched upon; see Section 2 for some examples.

At the other end of this story (and seemingly a world far, far away) is a landmark theorem in *quantum complexity theory* known as $\text{MIP}^* = \text{RE}$ [39]. Like most theorems in complexity theory, it compares two *complexity classes*. Roughly speaking, a complexity class consists of a collection of *problems* that all share some common level of *difficulty* with which one can solve or verify these problems. The class RE denotes those problems for which there is a computer program so that, if you left the program running long enough, would list all instances for which the problem has a positive answer (but you would never know about instances with a negative answer). Usually, complexity theorists are more interested in levels of efficiency, and the class RE is hardly ever discussed. The other complexity class in the above

equation is MIP^* , which denotes those problems for which a *verifier* interacting with multiple cooperating (but noncommunicating) *provers* who share a source of quantum entanglement can reliably verify a positive instance of a problem. The result $\text{MIP}^* = \text{RE}$ states that these two classes coincide! This is a monumental result, for it shows the power of quantum ideas in computational complexity. One particular instance of this result is that the (in)famous *halting problem*, which asks if a particular computer program will halt on, say, the empty input (which is known to be an undecidable problem), can actually be efficiently and reliably verified by two provers sharing some quantum entanglement. Here, “efficiently” means in polynomial time and “reliably” means that if the machine halts, then the verifier will accept the provers’ proof of that fact with probability 1, while if it does not halt, then only half the time will they accept a (necessarily incorrect) proof purporting to show that the machine halts. (An execution of the protocol has a probabilistic outcome, whence here the condition is that there is acceptance with probability at most $\frac{1}{2}$ over the verifiers’ and provers’ random choices in the case of a Turing machine that does not halt. If making a mistake half of the time seems unacceptable, the reader will find comfort in knowing that the chance of mistake can be made as close to 0 as one wishes upon repetition of the protocol.) This is an astounding result!

Even more amazing than the sheer statement of the result is that the equality $\text{MIP}^* = \text{RE}$ actually yields a negative solution to CEP!

1.2. Connecting the dots. But how could these seemingly unrelated topics be so tightly connected? The answer lies within a series of previously known connections. First, in a fundamental paper of Kirchberg [44], it was shown that CEP is equivalent to an important problem (Kirchberg even used the word *conjecture*) in the theory of C^* -algebras stemming from the complexity of C^* -tensor products known now as Kirchberg’s QWEP problem (see Subsection 3.8). Later, Fritz [24] and, independently, Junge et al. [41] demonstrated that a positive answer to the QWEP problem would yield a positive answer to a problem in quantum information theory known as *Tsirelson’s problem* which, roughly speaking, asks whether the usual quantum mechanical framework and that coming from quantum field theory yield the same set of quantum correlations corresponding to Bell experiments. While the jump from the QWEP problem to Tsirelson’s problem might seem like quite a leap, once one unravels the definitions, this is actually a fairly straightforward argument, which will be given in Subsection 6.1. Both sets of authors almost proved that the Kirchberg and Tsirelson problems were actually equivalent; Ozawa succeeded in connecting the last dots in [47].

Now we are at least in the same arena: quantum information theory and quantum complexity theory (both areas at least have “quantum” in their names). The last step in the puzzle is to use a result of Fritz, Netzer, and Thom [25] about the computability of the operator norm for universal group C^* -algebras and the analysis leading to the equivalence of QWEP and Tsirelson to show that if Tsirelson’s problem has a positive answer, then every language in MIP^* would actually be decidable, contradicting $\text{MIP}^* = \text{RE}$. (See Subsection 6.1 for the complete argument.)

Okay, so that was a mouthful!

1.3. Why another treatment of CEP? Numerous accounts of the CEP and its many equivalents can be found in the literature. In fact, Pisier [52] recently wrote a fascinating account (coming in just shy of 500 pages) on the CEP and its equivalences with QWEP and Tsirelson (and so, so much more). Much trimmer accounts were given by Ozawa [47, 48] and Capraro and Lupini [13].

So if there are so many accounts of the CEP, why write another? We have several good reasons:

First, all of the above accounts were written pre- $\text{MIP}^* = \text{RE}$, so none of them actually explain how the story resolves itself.

Second, all of the above accounts go into an extreme amount of detail and assume a fair amount of background knowledge in operator algebras. We envision the reader in, say, quantum physics or complexity theory wanting to understand the main thread of the story and being overburdened by the overhead needed to enter the fray. In this survey, we try very hard to at least state all of the necessary definitions. On the other hand, we offer very few details or proofs in the interest of space and refer the reader to the above references if they are interested in the gritty details. Also, since we are focusing on the one-way implications (as opposed to the equivalences the other accounts present), we save ourselves some complications.

Third, the operator algebra community may know very little quantum theory or complexity theory, so we offer brief introductions to these areas to at least paint the picture for them.

Finally, and most certainly gratuitously, we offer an *alternative* and, in this author's biased opinion, *simpler* path from $\text{MIP}^* = \text{RE}$ to the failure of CEP than that outlined above using basic methods from mathematical logic. This path also offers some extra bells and whistles to the failure of CEP, including a Gödelian refutation of the CEP and a proof of the existence of *many* counterexamples to the CEP. While we have our logician hats on, we take advantage of the fact that we have the readers' attention to describe a model-theoretic weakening of the CEP that is still open and quite fascinating (at least to us!).

1.4. A quick guide to this guide. In Section 2, we briefly describe some of the known equivalents of the CEP. The reader may benefit from coming back to this section after having read some of the definitions, but this is supposed to whet the reader's appetite and convince them that the rest of the paper is worth reading.

Section 3 is a crash course in operator algebras, assuming some basic functional analysis that someone in quantum physics should probably be familiar with. We cover both the C^* and von Neumann algebra background needed as well as topics such as states and traces, the ultrapower construction, operator algebras arising from groups, and finally, what is so darn complicated about C^* -algebra tensor products, culminating in a discussion of why a positive solution to the CEP implies a positive solution to the QWEP problem.

Section 4 is a similar crash course but this time in complexity theory. We start from the definition of Turing machines, defining some of the basic complexity classes, and then work our way up to the class MIP of languages verifiable by a verifier interacting with multiple cooperating provers.

In Section 5, we make a quantum detour for those unfamiliar with the basic tenets of quantum mechanics and even take a digression on superdense coding just for fun (and to indicate the power of entanglement). This section culminates with the definition of the complexity class MIP^* , the analogue of MIP where the provers

are allowed to share quantum entanglement as a resource, and the precise statement of the result $\text{MIP}^* = \text{RE}$.

Section 6 contains the details of the proof of the failure of the QWEP problem from $\text{MIP}^* = \text{RE}$ by first showing how the latter yields a negative solution to Tsirelson's problem and then by establishing how a negative solution to Tsirelson's problem yields a negative solution to the QWEP problem. Combined with our derivation of a positive solution of the QWEP problem from a positive solution to the CEP, this completes the proof of the negative solution to the CEP from $\text{MIP}^* = \text{RE}$.

Section 7 offers the alternative proof alluded to above using basic ideas from logic. We present the appropriate logic for studying tracial von Neumann algebras and discuss the main contribution from logic, namely Gödel's completeness theorem. We also describe the extra information about the CEP gleaned from the logical perspective mentioned above, including a completely operator-algebraic reformulation of our main model-theoretic contribution in terms of the undecidability of a certain *moment approximation problem*. We also offer an alternative proof of the failure of Tsirelson from $\text{MIP}^* = \text{RE}$ using the completeness theorem. Most of the material in this section represents joint work with Bradd Hart [31, 32].

Finally, in Section 8, we discuss the open problem around the existence of the so-called *enforceable factor*, which is the model-theoretic weakening of the CEP referred to above.

2. EQUIVALENT REFORMULATIONS OF CEP

One of the aspects of the CEP that makes it such an interesting problem is its numerous equivalences spanning many seemingly different areas of mathematics. In the main text, the equivalences with Kirchberg's QWEP conjecture in C^* -algebra theory and Tsirelson's problem in quantum information theory will be expounded on in more detail due to their relevance to the current story. In this section, we briefly mention some of the other well-known equivalences.

2.1. Free probability theory. In free probability theory, one considers *noncommutative* probability spaces, such as tracial von Neumann algebras (M, τ) , where the elements of M act as noncommutative random variables and the trace τ is the analogue of the integral. Voiculescu demonstrated the robustness of this theory, establishing free analogues of many familiar facts from ordinary probability theory and giving applications to operator algebras and random matrices (to name a few). A nice introduction to free probability is Speicher's lecture notes [57].

In classical probability theory, the entropy of a random variable is an important numerical value measuring the amount of information obtained when measuring the random variable. One method of calculating the entropy of a discrete random variable with probability distribution $\{p_1, \dots, p_n\}$ is to approximate the distribution using *microstates*, which are functions $f : \{1, \dots, N\} \rightarrow \{1, \dots, n\}$ for which the fraction of $j \in \{1, \dots, N\}$, where $f(j) = k$ is within ϵ of p_k for all $k = 1, \dots, n$. By taking the logarithm of the number of such functions divided by N for a given pair (N, ϵ) of parameters and then letting $N \rightarrow \infty$ and $\epsilon \rightarrow 0$, we obtain the entropy $H(p_1, \dots, p_n)$ of the distribution. A more general version of this works for a wider class of random variables.

When faced with the task of defining the free entropy of a tuple $(a_1, \dots, a_n)$ of self-adjoint elements in a tracial von Neumann algebra (M, τ) , Voiculescu proceeds analogously by considering those tuples $(A_1, \dots, A_n)$ of self-adjoint matrices in some matrix algebra $M_k(\mathbb{C})$ for which a certain finite number of *moments* approximate the corresponding moments in the tracial von Neumann algebra; that is, $\tau(p(a_1, \dots, a_n))$ and $\text{tr}(p(A_1, \dots, A_n))$ differ by at most ϵ for finitely many noncommutative $*$ -polynomials $p(X_1, \dots, X_n)$ in n -variables. Now one has to calculate the volume of the set of those matrices and let the various parameters involved tend to infinity or 0. With this definition of free entropy, one can prove a number of results which are the *free* analogues of the corresponding result in the classical theory. For example, it is known that a tuple of classical random variables has maximal entropy if and only if they are independent and have Gaussian distribution. In the free theory, the free entropy of a tuple is maximal if and only if the elements of the tuple are freely independent and have *semicircular distributions* (which are known to be the free analogue of the Gaussian distribution). The paper [63] is a survey of free entropy by Voiculescu himself.

This definition of free entropy leads to an interesting feature: if there are no such tuples $(A_1, \dots, A_n)$ that *simulate* $(a_1, \dots, a_n)$, then the free entropy of $(a_1, \dots, a_n)$ equals $-\infty$. It is well-known (see Subsection 3.6) that, for a given tracial von Neumann algebra (M, τ) , the set of such moments is nonempty for all such tuples $(a_1, \dots, a_n)$ from M if and only if M embeds into $\mathcal{R}^u$ (in a trace-preserving way). Thus, CEP is equivalent to all tuples of self-adjoint elements in tracial von Neumann algebras having nonnegative free entropy.

2.2. Hyperlinear groups. Okay, so this one really is not an equivalence but rather an equivalence with a *special case* of the CEP. An important notion in group theory is that of a *sofic group*. Roughly speaking, a countable discrete group G is sofic if, for every finite subset F of G , there is a symmetric group S_n and a function $\phi : F \rightarrow S_n$ that is an *approximately injective approximate homomorphism*. For example, if $g, h, gh \in F$, then one would like to say that $\phi(gh)$ is close to $\phi(g)\phi(h)$, where closeness is measured with respect to the normalized Hamming distance between permutations (which calculates what fraction of elements the permutations disagree on). The importance of this class of groups is that many important conjectures in group theory are known to hold when restricted to the class of sofic groups. Surprisingly, there is no known example of a nonsofic group! One can make a similar definition, replacing symmetric groups S_n with unitary groups U_n , equipped with their normalized Hilbert–Schmidt metric; the resulting class of groups is called the class of *hyperlinear groups*. Every sofic group is hyperlinear, and since we do not know if every group is sofic, we do not know if this inclusion is proper. Moreover, there is no known example of a nonhyperlinear group. We refer the reader to [13] for more information on sofic and hyperlinear groups.

The connection with CEP comes via an observation of Radulescu [53], who showed that G is hyperlinear if and only if the group von Neumann algebra $L(G)$ of G (see Subsection 3.7) embeds into $\mathcal{R}^u$. In other words, if CEP is true just for group von Neumann algebras, then every group is hyperlinear!

Interestingly enough, even though we now know that CEP is false, we still do not know if its special case for group von Neumann algebras holds; that is, we still do not know if every group is hyperlinear.

2.3. Embeddability of general von Neumann algebras. The CEP is about tracial von Neumann algebras. But there is a much wider class of von Neumann algebras out there. Is there a reformulation of the CEP that addresses them? The answer is yes, and it was established by Ando, Haagerup, and Winslow in [1]. There is a so-called type III (in the sense of Subsection 3.5) version of $\mathcal{R}$, called the *Araki–Woods factor* $\mathcal{R}_\infty$, which is the unique hyperfinite type III₁ factor. Moreover, there is a generalization of the tracial ultraproduct construction, known as the *Ocneanu ultraproduct*, that covers the much larger class of σ -finite von Neumann algebras, of which $\mathcal{R}_\infty$ is one of them. The main result of [1] states that CEP is equivalent to the assertion that *every* separably acting von Neumann algebra embeds *with expectation* into the Ocneanu ultrapower $\mathcal{R}_\infty^u$. The notion of an embedding with expectation is defined in Subsection 3.9. In the case of tracial von Neumann algebras, the embedding is automatically with expectation, but in the general case, it is a necessary nontriviality condition.

2.4. Existentially closed factors. The model-theoretic notion of an *existentially closed* (e.c.) structure is the generalization of the notion of algebraically closed field to an arbitrary structure (see Subsection 8.3 for a precise definition). In particular, it makes sense to study e.c. groups, e.c. graphs, and, yes, even e.c. tracial von Neumann algebras. One can prove many general facts about the class of e.c. tracial von Neumann algebras, such as they must be II₁ factors with McDuff’s property and with only approximate inner automorphisms. There are a plethora of e.c. tracial von Neumann algebras; in particular, every tracial von Neumann algebra embeds in an e.c. one. However, can one actually name a concrete e.c. tracial von Neumann algebra? It turns out that a positive solution to CEP is equivalent to the statement that $\mathcal{R}$ is an e.c. tracial von Neumann algebra. A proof of this fact will be given in Subsection 8.3.

2.5. Noncommutative real algebraic geometry. A *Positivstellensatz* is a theorem that declares that certain elements that are positive in some way are so for some *good reason*. Perhaps the best-known such result is the positive solution to Hilbert’s 17th Problem, due to Artin [2] (although this author is unabashedly fond of Abraham Robinson’s model-theoretic solution [54]): if $f(X_1, \dots, X_n) \in \mathbb{R}\langle X_1, \dots, X_n \rangle$ is a positive semidefinite rational function (that is, a rational function such that $f(x_1, \dots, x_n) \geq 0$ for all $x_1, \dots, x_n \in \mathbb{R}$), then f is a sum of squares of rational functions, providing a good reason that f is positive semidefinite.

One can ponder noncommutative versions of Artin’s theorem. First, we set $\mathbb{R}\langle X_1, \dots, X_n \rangle$ to be the set of polynomials in n noncommuting variables. Consider the *positivity* statement that $f(A_1, \dots, A_n) \geq 0$ for all self-adjoint matrices $A_1, \dots, A_n \in M_m(\mathbb{R})$ of operator norm at most 1, for all $m \in \mathbb{N}$. Then a theorem of Helton and McCullough [36] tells us that there is a good reason for this kind of positivity, namely that, for all $\epsilon \in \mathbb{R}^{>0}$, $f + \epsilon$ belongs to the *quadratic module* generated by $1 - X_i^2$, $i = 1, \dots, n$. Here, a quadratic module is a subset M of $\mathbb{R}\langle X_1, \dots, X_n \rangle$ containing 1, closed under addition and closed under the function $a \mapsto g^*ag$, where $a \in M$ and $g \in \mathbb{R}\langle X_1, \dots, X_n \rangle$ (and where g^* is the result of reversing the orders of the variables in each monomial of g). Note indeed that all functions in the quadratic module generated by the $1 - X_i^2$ ’s must be positive in the above sense, and the Helton–McCullough result says that this is (approximately) the good reason that any such noncommutative polynomial might be positive.

Now suppose instead that we assume that f is merely *trace positive*, that is, $\text{tr}(f(A_1, \dots, A_n)) \geq 0$ for all such $A_1, \dots, A_n$ as in the previous paragraph. Clearly, the operators in the Helton and McCullough result are trace positive. But now you can also add finite sums of commutators $[A, B] := AB - BA$ since the trace of a commutator vanishes. One can ask if this new class of noncommutative polynomials gives a necessary and sufficient condition to be trace positive; that is, if f is tracially positive, must it be the case that, for every $\epsilon > 0$, $f + \epsilon$ differs from an element of the quadratic module generated by the $1 - X_i^2$'s by a sum of commutators? It turns out that this tracial version of the *Positivstellensatz* from the previous paragraph is actually equivalent to the CEP, a result proven by Klep and Schweighofer in [45].

3. A CRASH COURSE IN OPERATOR ALGEBRAS

In this long section, we explain all of the background material in operator algebras one needs to know to understand the statements of both the CEP and the QWEP problem as well as to understand how a positive solution to the former implies a positive solution to the latter. Nearly everything discussed here can be found in Pisier's book [52]. Brown and Ozawa's book [12] is another nice reference.

3.1. Introducing C^* -algebras. A $*$ -algebra is an algebra $\mathcal{A}$ over $\mathbb{C}$ satisfying, for all $x, y \in \mathcal{A}$ and $\lambda \in \mathbb{C}$,

- $(x + y)^* = x^* + y^*$,
- $(xy)^* = y^*x^*$,
- $(x^*)^* = x$,
- $(\lambda x)^* = \bar{\lambda}x^*$.

If $\mathcal{A}$ is actually a unital algebra over $\mathbb{C}$ with unit 1 for which $1^* = 1$, we say that $\mathcal{A}$ is a *unital $*$ -algebra*. There are obvious notions of $*$ -subalgebra of a $*$ -algebra and unital $*$ -subalgebra of a unital $*$ -algebra.

A $*$ -homomorphism between $*$ -algebras is an algebra homomorphism that also preserves the $*$ -operation. If $\phi : \mathcal{A} \rightarrow \mathcal{B}$ is a $*$ -homomorphism between unital $*$ -algebras, then we implicitly assume that ϕ maps the unit of $\mathcal{A}$ to the unit of $\mathcal{B}$.

In this paper, the most relevant (unital) $*$ -algebras are $\mathcal{B}(\mathcal{H})$, the $*$ -algebra of bounded operators on the Hilbert space $\mathcal{H}$, and its (unital) $*$ -subalgebras. Recall that for a Hilbert space $\mathcal{H}$, a linear operator $T : \mathcal{H} \rightarrow \mathcal{H}$ is *bounded* if its *operator norm* $\|T\| := \sup\{\|T\xi\| : \|\xi\| \leq 1\}$ is finite. $\mathcal{B}(\mathcal{H})$ is a $*$ -algebra with the algebra operations being addition, composition, and scalar multiplication and with the $*$ -operation being given by the *adjoint*, where, for $T \in \mathcal{B}(\mathcal{H})$, we have that $T^* \in \mathcal{B}(\mathcal{H})$ is the unique operator for which $\langle T\xi, \eta \rangle = \langle \xi, T^*\eta \rangle$ for all $\xi, \eta \in \mathcal{H}$. (In connection with this formula, we follow the convention that inner products are linear in the first argument and conjugate-linear in the second argument; this is the opposite of the convention used in the physics literature.) $\mathcal{B}(\mathcal{H})$ is a unital $*$ -algebra with identity operator $I_{\mathcal{H}}$ acting as the unit.

We now define our first class of operator algebras, namely the class of C^* -algebras. For both classes of operator algebras, there are two approaches to their definition, namely the concrete and the abstract. A *concrete C^* -algebra* is a $*$ -subalgebra $\mathcal{A}$ of $\mathcal{B}(\mathcal{H})$ that is closed in the operator norm topology. If, moreover, $\mathcal{A}$ contains the identity $I_{\mathcal{H}}$, then we say that $\mathcal{A}$ is a *unital concrete C^* -algebra*.

We now present the abstract approach to C^* -algebras. Suppose that $\mathcal{A}$ is a $*$ -algebra. A C^* -norm on $\mathcal{A}$ is a norm on $\mathcal{A}$ satisfying the following identities for all $x, y \in \mathcal{A}$:

- $\|xy\| \leq \|x\|\|y\|$,
- $\|x^*\| = \|x\|$,
- $\|x^*x\| = \|x\|^2$.

The first two identities are the usual axioms for defining a *normed* $*$ -algebra; the last axiom, called the C^* -identity, is what makes a C^* -norm a C^* -norm. An *abstract C^* -algebra* is a $*$ -algebra equipped with a complete C^* -norm. If $\mathcal{A}$ is a $*$ -algebra equipped with a C^* -norm, then the $*$ -algebra operations extend naturally to the completion of the $*$ -algebra, which is then an abstract C^* -algebra. An *abstract unital C^* -algebra* is an abstract C^* -algebra that is a unital $*$ -algebra; in this case, we have $\|1\| = 1$.

It is an important fact that a $*$ -homomorphism between abstract C^* -algebras is necessarily contractive; it is an isometric embedding if and only if it is injective. In particular, given any $*$ -algebra $\mathcal{A}$, there is at most one norm on $\mathcal{A}$ which makes $\mathcal{A}$ into a C^* -algebra.

It is an easy exercise to see that the operator norm on $\mathcal{B}(\mathcal{H})$ is a C^* -norm, whence every concrete C^* -algebra is an abstract C^* -algebra. On the other hand, the *Gelfand–Naimark theorem* states that every abstract C^* -algebra is isomorphic (as abstract C^* -algebras) to a concrete C^* -algebra. This result can be reformulated in terms of representations of C^* -algebras. Given a C^* -algebra $\mathcal{A}$, a *representation* of $\mathcal{A}$ is a $*$ -homomorphism $\pi : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ for some Hilbert space $\mathcal{H}$. Usually a nondegeneracy condition is assumed on a representation, namely that $\{\pi(a)(\xi) : a \in \mathcal{A}, \xi \in \mathcal{H}\}$ is dense in $\mathcal{H}$; if $\mathcal{A}$ is unital, this is equivalent to assuming that $\pi(1) = I_{\mathcal{H}}$. The representation π is *faithful* if it is moreover injective (equivalently isometric). Thus, the Gelfand–Naimark theorem states that every abstract C^* -algebra admits a faithful representation.

From here on out, we no longer make a distinction between concrete and abstract C^* -algebras and simply take either perspective whenever it is convenient.

Unless stated otherwise, in the rest of this paper, we restrict our attention to unital C^* -algebras; we might often repeat this convention for emphasis.

A C^* -algebra is *commutative* (or *abelian*) if its multiplication is commutative. Given a compact Hausdorff space X , the set $C(X)$ of continuous, complex-valued functions is a unital commutative C^* -algebra under the pointwise operations of addition, multiplication, and scalar multiplication, with the $*$ -operation being given by $f^* := \bar{f}$ (complex conjugate), and with norm given by $\|f\| := \sup_{x \in X} |f(x)|$. In fact, all unital commutative C^* -algebras are of this form, and there is a dual equivalence of categories (known as *Gelfand duality*) between compact Hausdorff spaces with continuous maps and unital commutative C^* -algebras with $*$ -homomorphisms. For this reason, C^* -algebra theory is often dubbed “noncommutative topology.”

There are special kinds of elements in C^* -algebras that will be important throughout this paper. If $\mathcal{A}$ is a C^* -algebra, $x \in \mathcal{A}$ is called

- *self-adjoint* if $x^* = x$,
- *positive* if $x = y^*y$ for some $y \in \mathcal{A}$,
- a *projection* if x is self-adjoint and $x^2 = x$,
- *unitary* if $x^*x = xx^* = 1$.

In the case that $\mathcal{A} = \mathcal{B}(\mathcal{H})$, the self-adjoint (resp., positive) elements are those with spectrum contained in the reals (resp., the positive reals) while the projections correspond to orthogonal projections onto closed subspaces. In the case that $\mathcal{A} = C(X)$ for X a compact space, the self-adjoint (resp., positive) elements correspond to the real-valued (resp., positive real-valued) functions while the projections correspond to those functions which take only the values 0 or 1.

3.2. Let's be positive. An important role in this story is played by maps between C^* -algebras which are not necessarily $*$ -homomorphisms but which still preserve some remnants of the C^* -algebra structure. It is hard to truly appreciate the importance of these maps without getting into the details of the results to follow, but we introduce the terminology in order to be able to follow the definitions and theorems.

First, we say that a linear map $\phi : \mathcal{A} \rightarrow \mathcal{B}$ between C^* -algebras is *positive* if it maps positive elements to positive elements. Note that a $*$ -homomorphism is positive: $\phi(a^*a) = \phi(a)^*\phi(a)$.

For many purposes, a stronger version of positivity is needed. First, for any C^* -algebra $\mathcal{A}$ and any $n \geq 1$, we let $M_n(\mathcal{A})$ denote the set of $n \times n$ matrices with entries from $\mathcal{A}$. We can view $M_n(\mathcal{A})$ as a $*$ -subalgebra of $\mathcal{B}(\bigoplus_{i=1}^n \mathcal{H})$ and it is readily verified that, under this identification, $M_n(\mathcal{A})$ is closed in the operator norm topology, that is, $M_n(\mathcal{A})$ is a C^* -algebra once again. Note also that a linear map $\phi : \mathcal{A} \rightarrow \mathcal{B}$ induces a linear map $\phi_n : M_n(\mathcal{A}) \rightarrow M_n(\mathcal{B})$ given by $\phi_n(a_{ij}) = (\phi(a_{ij}))$. We say that ϕ is *completely positive* (cp) if each ϕ_n is a positive map. If, in addition, $\phi(1) = 1$, we say that ϕ is *unital, completely positive* (ucp). A $*$ -homomorphism $\phi : \mathcal{A} \rightarrow \mathcal{B}$ is easily seen to induce $*$ -homomorphisms $\phi_n : M_n(\mathcal{A}) \rightarrow M_n(\mathcal{B})$, whence $*$ -homomorphisms are ucp. It can be shown that if $\mathcal{A}$ is commutative, then any positive map $\phi : \mathcal{A} \rightarrow \mathcal{B}$ is automatically completely positive.

In a similar vein, one says that ϕ as above is *completely bounded* (resp., *completely contractive*) if each ϕ_n is bounded (resp., contractive).

A fundamental theorem of Stinespring says that ucp maps are not too far away from being $*$ -homomorphisms. More precisely, consider the following situation: suppose that $\mathcal{H}$ and $\mathcal{K}$ are Hilbert spaces and $V : \mathcal{H} \rightarrow \mathcal{K}$ is an isometry, that is, a linear map for which $V^*V = I_{\mathcal{H}}$ (or, in other words, $\langle V\xi, V\xi \rangle = \langle \xi, \xi \rangle$ for all $\xi \in \mathcal{H}$). Then for any representation $\pi : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{K})$ of $\mathcal{A}$, we have a map $\phi : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ given by $\phi(a)(\xi) := V^*(\pi(a)(V\xi))$ which is readily verified to be ucp. The Stinespring dilation theorem says that all ucp maps $\phi : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ arise in this way. A particular consequence of this theorem is that ucp maps are completely contractive. The relevance of Stinespring's theorem to our story is that certain results that hold somewhat immediately for $*$ -homomorphisms will also hold for ucp maps (see Subsection 3.8).

We mention in passing that cp maps play an important role in quantum information theory. Indeed, one perspective on a quantum state (say on a finite-dimensional state space) is that of a positive matrix of trace 1, corresponding to the density matrix of some mixed state. A *quantum channel* is a linear map that is to represent some allowable physical transformation on quantum states. In particular, if it is to map density matrices to density matrices, then it should be trace-preserving and positive. However, often one needs to add *ancilla bits* to a given state and then apply the corresponding quantum channel. The desire to have the resulting matrix

be a density matrix again is equivalent to the requirement that the quantum channel be completely positive instead of merely positive. The reader can find more details in, for example, Vern Paulsen's lecture notes [50].

3.3. Introducing von Neumann algebras. We now turn our attention to the other kind of operator algebra, the von Neumann algebra. Once again, we have a choice between a concrete definition and an abstract definition. A *concrete von Neumann algebra* is a unital $*$ -subalgebra of $\mathcal{B}(\mathcal{H})$ closed in the *weak operator topology* (WOT), where a subbasic open neighborhood of $T \in \mathcal{B}(\mathcal{H})$ in the WOT has the form $\{S \in \mathcal{B}(\mathcal{H}) : |\langle T\xi, \eta \rangle - \langle S\xi, \eta \rangle| < \epsilon\}$ for some $\xi, \eta \in \mathcal{H}$ and some $\epsilon > 0$. Said differently, the WOT is the smallest topology on $\mathcal{B}(\mathcal{H})$ making the maps $T \mapsto \langle T\xi, \eta \rangle : \mathcal{B}(\mathcal{H}) \rightarrow \mathbb{C}$ continuous for all $\xi, \eta \in \mathcal{H}$. It is easy to check that the WOT is a finer topology on $\mathcal{B}(\mathcal{H})$ than the operator norm topology, whence every concrete von Neumann algebra is a (unital) concrete C^* -algebra. In general, von Neumann algebras are much *bigger* than general C^* -algebras due to the, well, weakness of the WOT.

A theorem of Sakai allows for an abstract reformulation: a (unital) abstract C^* -algebra $\mathcal{M}$ is isomorphic (as an abstract C^* -algebra) to a concrete von Neumann algebra if and only if $\mathcal{M}$ is isometrically isomorphic to a dual Banach space, that is, if and only if there is a closed subspace $X \subseteq \mathcal{M}^{**}$ such that $\mathcal{M} = X^*$ isometrically. In this case, X is unique and is called the *predual* of $\mathcal{M}$, denoted $\mathcal{M}_*$.

Von Neumann's *bicommutant theorem* allows for a purely algebraic reformulation of being a von Neumann algebra that is incredibly important to the theory. First, given a subset $\mathcal{S} \subseteq \mathcal{B}(\mathcal{H})$, set $\mathcal{S}' := \{T \in \mathcal{B}(\mathcal{H}) : TS = ST \text{ for all } S \in \mathcal{S}\}$, the so-called *commutant* of $\mathcal{S}$. Note that for any set $\mathcal{S}$, we have that $\mathcal{S}'$ is a von Neumann subalgebra of $\mathcal{B}(\mathcal{H})$ and that $\mathcal{S} \subseteq \mathcal{S}'' := (\mathcal{S}')'$. Von Neumann's bicommutant theorem states that for any unital $*$ -subalgebra $\mathcal{A}$ of $\mathcal{B}(\mathcal{H})$, $\mathcal{A}''$ coincides with the WOT-closure of $\mathcal{A}$ in $\mathcal{B}(\mathcal{H})$; this common algebra is the von Neumann algebra generated by $\mathcal{A}$. In fact, the bicommutant theorem shows that both of these coincide with the closure of $\mathcal{A}$ in the *strong operator topology* (SOT), where now a subbasic open neighborhood of $T \in \mathcal{B}(\mathcal{H})$ in the SOT has the form $\{S \in \mathcal{B}(\mathcal{H}) : \|(T - S)\xi\| < \epsilon\}$ for some $\xi \in \mathcal{H}$ and some $\epsilon > 0$. A consequence of the bicommutant theorem is that a unital $*$ -subalgebra $\mathcal{M}$ of $\mathcal{B}(\mathcal{H})$ is a von Neumann algebra if and only if $\mathcal{M} = \mathcal{M}''$.

A von Neumann algebra $\mathcal{M}$ is called *separable* if its bidual is a separable Banach space. This is equivalent to having a concrete representation of $\mathcal{M}$ on a separable Hilbert space, whence one sometimes calls such a von Neumann algebra *separably acting*.

Just as in the case of C^* -algebras, we can completely characterize the commutative von Neumann algebras. Given a σ -finite measure space (X, μ) , we can view $L^\infty(X, \mu) \subseteq \mathcal{B}(L^2(X, \mu))$ by identifying $f \in L^\infty(X, \mu)$ with $M_f \in \mathcal{B}(L^2(X, \mu))$ given by $M_f(g) := fg$. It is an exercise to check that $L^\infty(X, \mu) = L^\infty(X, \mu)'$, whence $L^\infty(X, \mu)$ is a commutative von Neumann algebra. Moreover, all commutative von Neumann algebras have this form, whence von Neumann algebra theory is often dubbed “noncommutative measure theory.”

While $*$ -homomorphisms between von Neumann algebras are automatically contractive with respect to the operator norm (as they are C^* -algebras), since the relevant topology for defining von Neumann algebras is the WOT, the appropriate continuity condition relates to this latter topology. More precisely, a positive linear map $\Phi : \mathcal{M} \rightarrow \mathcal{N}$ between von Neumann algebras is *normal* if the restriction of Φ

to the operator norm unit ball of $\mathcal{M}$ is continuous when $\mathcal{M}$ and $\mathcal{N}$ are equipped with their WOT topologies. (The restriction to the operator norm unit ball may seem slightly unsightly; this is equivalent to saying that Φ is continuous when both $\mathcal{M}$ and $\mathcal{N}$ are equipped with their weak*-topologies when viewed as dual Banach spaces.) One can reformulate normality in an intrinsic way that does not refer to the particular realization of $\mathcal{M}$ and $\mathcal{N}$: $\Phi : \mathcal{M} \rightarrow \mathcal{N}$ is normal if and only if it is positive, linear, and $\Phi(\sup_{i \in I} x_i) = \sup_{i \in I} \Phi(x_i)$ for every bounded increasing net $(x_i)_{i \in I}$ of positive elements in $\mathcal{M}$.

3.4. States and traces. Fix a compact Hausdorff space X . Given a complex Borel measure μ on X , we can consider the associated integration functional $I_\mu \in C(X)^*$ given by $I_\mu(f) := \int_X f d\mu$ which satisfies $\|I_\mu\| = \|\mu\|$, where $\|\mu\|$ denotes the total variation norm of μ . Setting $M(X)$ to be the Banach space of complex Borel measures on X , the Riesz representation theorem implies that this association yields an isomorphism $C(X)^* \cong M(X)$. Moreover, the probability measures μ on X correspond to those $I \in C(X)^*$ for which I is a positive map satisfying $I(1) = 1$.

More generally, given a unital C*-algebra $\mathcal{A}$, a *state* on $\mathcal{A}$ is a positive linear functional ϕ on $\mathcal{A}$ with $\phi(1) = 1$. Thus, the states on a unital abelian C*-algebra $C(X)$ correspond to the integration functionals associated to probability measures on X , and one thinks of states on arbitrary C*-algebras as the abstract analogue of such an integral.

The states on $\mathcal{A}$ form a convex, closed subset $\mathfrak{S}(\mathcal{A})$ of $\mathcal{A}^*$. The extreme points of $\mathfrak{S}(\mathcal{A})$ are referred to as *pure states*. By the Krein–Milman theorem, finite convex combinations of pure states are dense in the space of all states. The pure states on $\mathcal{B}(\mathcal{H})$ are the *vector states*, that is, the states of the form $T \mapsto \langle T\xi, \xi \rangle$ for some $\xi \in \mathcal{H}$.

A consequence of the Hahn–Banach theorem is the fact that, for any C*-algebra $\mathcal{A}$ and any self-adjoint element $a \in \mathcal{A}$, we have $\|a\| = \sup_{\phi \in \mathfrak{S}(\mathcal{A})} |\phi(a)|$. Another consequence of the Hahn–Banach theorem is that whenever $\mathcal{A}$ is a subalgebra of $\mathcal{B}$, any state on $\mathcal{A}$ can be extended to a state on $\mathcal{B}$.

When $\mathcal{H}$ is finite dimensional, every state on $\mathcal{B}(\mathcal{H})$ is of the form $T \mapsto \text{Tr}(T\rho)$ for a unique positive operator ρ of trace 1, where Tr denotes the trace of an operator. The operator ρ is often called the *density matrix* for the state. When $\mathcal{H}$ is not necessarily finite dimensional, the same result holds true for states on $\mathcal{B}(\mathcal{H})$ that are continuous with respect to the weak*-topology on $\mathcal{B}(\mathcal{H})$, except that ρ is now stipulated to be a trace-class operator (see, for example, [35, Theorem 19.9]).

The term “state” comes from quantum mechanics. A first introduction to quantum mechanics will introduce a state of a physical system as simply a unit vector ξ in the Hilbert space $\mathcal{H}$ associated to the physical system. This usage of the word state corresponds to the vector states described above. Later, one encounters the notion of *mixed state* (to accommodate the fact that results of quantum measurements are merely probabilistic *ensembles* of pure states), which is often defined in terms of the density matrix as defined above. The role of a state in quantum mechanics is simply to assign expected values of observables (see [35, Section 19]).

An important construction associated to a state is the *Gelfand–Naimark–Segal (GNS) construction*, which associates a representation of the C*-algebra to the state. Suppose that ϕ is a state on $\mathcal{A}$. We define a sesquilinear form $\langle \cdot, \cdot \rangle_\phi$ on $\mathcal{A}$ by $\langle x, y \rangle_\phi := \phi(y^*x)$. It is straightforward to check that this is a so-called pre-inner product on $\mathcal{A}$ in that it satisfies all of the properties of being an inner product except

that $\langle x, x \rangle_\phi = 0$ need not necessarily imply that $x = 0$; when $\langle \cdot, \cdot \rangle_\phi$ is actually an inner product, we say that ϕ is *faithful*. Associated to $\langle \cdot, \cdot \rangle_\phi$ is the seminorm $\| \cdot \|_\phi$ on $\mathcal{A}$ given by $\|x\|_\phi := \sqrt{\langle x, x \rangle_\phi}$. We obtain a Hilbert space, denoted $L^2(\mathcal{A}, \phi)$, by quotienting out by the closed subspace of vectors with $\| \cdot \|_\phi = 0$ and then taking the completion. Given $a \in \mathcal{A}$, we let $\hat{a}$ denote its equivalence class in $L^2(\mathcal{A}, \phi)$. It follows that there is a representation $\pi_\phi : \mathcal{A} \rightarrow L^2(\mathcal{A}, \phi)$ uniquely determined by the condition $\pi_\phi(a)(\hat{b}) := \hat{ab}$ for all $a, b \in \mathcal{A}$. The representation π_ϕ is *cyclic*, meaning that there is a vector $\xi \in L^2(\mathcal{A}, \phi)$ for which $\overline{\{\pi_\phi(a)\xi : a \in \mathcal{A}\}} = L^2(\mathcal{A}, \phi)$, namely $\xi = \hat{1}$. Moreover, the vector state $\langle \cdot, \hat{1} \rangle_\phi$ on $L^2(\mathcal{A}, \phi)$ restricts to ϕ on the image of $\mathcal{A}$. Note that π_ϕ is a faithful representation precisely when ϕ is faithful.

There is a converse to the above construction: if $\pi : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ is a cyclic representation of $\mathcal{A}$ with cyclic vector $\xi \in \mathcal{H}$, then one obtains a state ϕ_π on $\mathcal{A}$ by $\phi_\pi(a) := \langle \pi(a)\xi, \xi \rangle$ and the GNS representation associated to ϕ_π is unitarily equivalent to π .

Define $\mathcal{H}_u := \bigoplus_{\phi \in \mathfrak{S}(\mathcal{A})} L^2(\mathcal{A}, \phi)$ and set $\pi_u = \bigoplus_{\phi \in \mathfrak{S}(\mathcal{A})} \pi_\phi : \mathcal{A} \rightarrow \mathcal{H}_u$, which we call the *universal representation* of $\mathcal{A}$. Since $\|a\| = \sup_{\phi \in \mathfrak{S}(\mathcal{A})} |\phi(a)|$ for any self-adjoint $a \in \mathcal{A}$, it follows that π_u is a faithful representation of $\mathcal{A}$. Since any representation of $\mathcal{A}$ is a direct sum of cyclic representations and since every cyclic representation of $\mathcal{A}$ is, up to unitary equivalence, of the form π_ϕ for some $\phi \in \mathfrak{S}(\mathcal{A})$, we see that every representation of $\mathcal{A}$ is unitarily equivalent to a subrepresentation of π_u , whence the name!

An important ingredient in this story is the von Neumann algebra $\pi_u(\mathcal{A})''$ generated by the image of $\mathcal{A}$ in $\mathcal{B}(\mathcal{H}_u)$. It can be shown that this von Neumann algebra is isometrically isomorphic to the Banach space $\mathcal{A}^{**}$, whence it is this notation that is usually used.

A state ϕ on $\mathcal{A}$ is called a *tracial state* if $\phi(ab) = \phi(ba)$ for all $a, b \in \mathcal{A}$. For example, the normalized trace tr on $M_n(\mathbb{C})$ given by $\text{tr}(a) := \frac{1}{n} \text{Tr}(a)$ is a tracial state on $M_n(\mathbb{C})$.

Since $\mathbb{C}$ is a von Neumann algebra, it makes sense to speak of *normal* states on von Neumann algebras. A faithful normal tracial state on a von Neumann algebra is simply referred to as a *trace*. A von Neumann algebra is called *finite* if it admits a trace. (This terminology makes much more sense if you introduce Murray–von Neumann equivalence of projections.) A *tracial von Neumann algebra* is a pair $(\mathcal{M}, \tau)$, where $\mathcal{M}$ is a von Neumann algebra and τ is a trace on $\mathcal{M}$. An embedding of tracial von Neumann algebras is a normal, injective $*$ -homomorphism that preserves the trace. The normalized trace tr on $M_n(\mathbb{C})$ is a trace (in the von Neumann algebra sense) on $M_n(\mathbb{C})$. However, when $\mathcal{H}$ is infinite dimensional, there is no trace on $\mathcal{B}(\mathcal{H})$.

Suppose that τ is a trace on $\mathcal{M}$. The corresponding representation $\pi_\tau : \mathcal{M} \rightarrow \mathcal{B}(L^2(\mathcal{M}, \tau))$ is normal and the image of $\pi_\tau(\mathcal{M})$ is WOT-closed in $\mathcal{B}(L^2(\mathcal{M}, \tau))$. Also, $\mathcal{M}$ is separable if and only if it is separable with respect to the metric stemming from the norm $\| \cdot \|_\tau$. When restricted to the unit ball of $\mathcal{M}$, the topology induced by $\| \cdot \|_\tau$ coincides with the SOT on $\mathcal{M}$ it inherits from $L^2(\mathcal{M}, \tau)$.

We end this section by showing how traces can be used to define the hyperfinite II_1 factor, the star of this paper!

Given $n \geq 1$, there is a natural embedding of tracial von Neumann algebras $M_{2^n}(\mathbb{C}) \hookrightarrow M_{2^{n+1}}(\mathbb{C})$ given by $A \mapsto \begin{pmatrix} A & 0 \\ 0 & A \end{pmatrix}$. In this way, we obtain a directed system of tracial von Neumann algebras whose union is a $*$ -algebra we denote by $\mathcal{M} := \bigcup_n M_{2^n}(\mathbb{C})$. The fact that the embeddings preserve the normalized traces on the individual matrix algebras implies that $\mathcal{M}$ has a tracial state τ on it. We apply the GNS construction to τ (which still works even though the original algebra is not necessarily complete) and take the von Neumann algebra generated by $\pi_\tau(M)$ inside of $\mathcal{B}(L^2(\mathcal{M}, \tau))$. This von Neumann algebra is called the *hyperfinite II_1 factor* $\mathcal{R}$. We will see the reason for the “ II_1 factor” in the name in the next section but the terminology “hyperfinite” can be explained now. A separable von Neumann algebra is called *hyperfinite* if it contains an increasing union of finite-dimensional subalgebras whose union is WOT-dense. Murray and von Neumann showed that there is a unique separable hyperfinite II_1 factor. Consequently, if we started the above construction with any $M_n(\mathbb{C})$ instead of $M_2(\mathbb{C})$, we would have arrived at the same II_1 factor, namely $\mathcal{R}$.

3.5. More on von Neumann algebras. The *center of a von Neumann algebra* $\mathcal{M}$ is $Z(\mathcal{M}) := \mathcal{M} \cap \mathcal{M}' = \{x \in \mathcal{M} : xy = yx \text{ for all } y \in \mathcal{M}\}$. A von Neumann algebra M is called a *factor* when its center is trivial, that is when $Z(\mathcal{M}) = \mathbb{C} \cdot 1$. It is quite easy to see that $\mathcal{B}(\mathcal{H})$ is a factor; in particular, each $M_n(\mathbb{C})$ is a factor. This makes it plausible that $\mathcal{R}$ is also a factor, given that it is the completion of an increasing sequence of factors—one just needs to check that no elements snuck into the center at the completion stage.

The interest in factors comes from the fact that they are the *building blocks* of all von Neumann algebras in the sense that every von Neumann algebra can be written as a direct integral (a generalization of direct sum) of factors, and thus the study of arbitrary von Neumann algebras can usually be reduced to studying factors.

Murray and von Neumann divided the collection of factors into three types, (creatively) called types I, II, and III. They further split the first two types into subtypes as follows. First, for each $n \in \mathbb{N}$, there is a unique factor of type I_n , namely $M_n(\mathbb{C})$. The unique factor of type I_∞ is $\mathcal{B}(\mathcal{H})$ for $\mathcal{H}$ infinite dimensional. Next, a II_1 factor is an infinite-dimensional finite factor, that is, an infinite-dimensional factor that admits a trace. Thus, the hyperfinite II_1 factor is indeed a II_1 factor. A II_∞ factor is one that can be written as a proper increasing union of type II_1 factors. Equivalently, a II_∞ factor can be written in the form $\mathcal{M} \otimes \mathcal{B}(\mathcal{H})$ for some II_1 factor $\mathcal{M}$ (see Subsection 3.8 for the definition of tensor products of von Neumann algebras). There is also division of type III factors into subtypes III_λ for $\lambda \in [0, 1]$, but we will not need to get into that here.

It is important to note that, while an arbitrary finite von Neumann algebra may have many traces, the trace on a finite factor is unique. We also note the crucial fact that $\mathcal{R}$ embeds into any II_1 factor.

As alluded to above when defining finite von Neumann algebras, the above type classification makes more sense in the context of Murray–von Neumann equivalence. However, we can still see this division using traces. Given a von Neumann algebra $\mathcal{M}$, let $\mathcal{P}(\mathcal{M})$ denote the set of projections in $\mathcal{M}$. If τ is a trace on $\mathcal{M}$, set $\tau(\mathcal{P}(\mathcal{M})) := \{\tau(p) : p \in \mathcal{P}(\mathcal{M})\}$. Note that, for the unique type I_n factor $M_n(\mathbb{C})$, we have $\tau(\mathcal{P}(M_n(\mathbb{C}))) = \{0, \frac{1}{n}, \dots, \frac{n-1}{n}, 1\}$, one value for every dimension being projected onto. On the other hand, for a II_1 factor $\mathcal{M}$, one can show that $\tau(\mathcal{P}(\mathcal{M})) = [0, 1]$.

Thus, we think of II_1 factors being like matrix factors in that they admit traces, but now we have a *continuous* dimension for projections.

One can explain the cases I_∞ and II_∞ using traces if one considers the unnormalized trace Tr on $\mathcal{B}(\mathcal{H})$ given by $\text{Tr}(T) := \sum_{i \in I} \langle T\xi_i, \xi_i \rangle$, where $(\xi_i)_{i \in I}$ is any orthonormal basis for $\mathcal{H}$. Note that Tr can take the value ∞ . We then have that the possible traces of projections for the I_∞ factor are $\{0, 1, 2, \dots\} \cup \{\infty\}$ while in the type II_∞ case they are $[0, \infty]$.

3.6. The tracial ultrapower construction and the official statement of the CEP. In general, an ultraproduct of a collection of *similar* structures is a structure of the same type that represents some sort of *limit* of these structures. There are ways of making this precise using model theory, and we will discuss this in Subsection 7.2. In this section, we show how to carry this construction out in the case of tracial von Neumann algebras and see how this allows us to precisely state the CEP.

First, one needs to introduce the notion of an *ultrafilter*. (For an entire book devoted to ultrafilters and their applications, see the author's [27]). Given an index set I , an ultrafilter $\mathcal{U}$ on I is simply a $\{0, 1\}$ -valued finitely additive probability measure on I . One often identifies $\mathcal{U}$ with its set of measure 1 sets and writes $X \in \mathcal{U}$ instead of $\mathcal{U}(X) = 1$. From this perspective, one can alternatively view an ultrafilter $\mathcal{U}$ on I as a division of the subsets of I into two categories— $\mathcal{U}$ -large and $\mathcal{U}$ -small—satisfying the following requirements: I is $\mathcal{U}$ -large while $\emptyset$ is $\mathcal{U}$ -small; the intersection of two $\mathcal{U}$ -large sets is $\mathcal{U}$ -large; a superset of a $\mathcal{U}$ -large set is once again $\mathcal{U}$ -large; for any subset J of I , exactly one of J or $I \setminus J$ is $\mathcal{U}$ -large.

Following typical measure-theoretic terminology, given a property P that may or may not hold of elements of I , we may write “for $\mathcal{U}$ -almost all $i \in I$, $P(i)$ holds” when $\{i \in I : P(i) \text{ holds}\}$ belongs to $\mathcal{U}$.

Given a bounded sequence $(z_i)_{i \in I}$ of complex numbers, it is straightforward to show that there is a unique complex number z such that, for every $\epsilon > 0$, we have $|z - z_i| < \epsilon$ for $\mathcal{U}$ -almost all $i \in I$. This unique complex number z is called the $\mathcal{U}$ -ultralimit of the sequence $(z_i)_{i \in I}$, denoted $\lim_{i \in I, \mathcal{U}} z_i$ or simply $\lim_{\mathcal{U}} z_i$. The fact that *every* bounded sequence has a limit (in this generalized sense) is but one indication of the utility of ultrafilters.

Given $j \in I$, the unique ultrafilter $\mathcal{U}$ on I for which $\mathcal{U}(\{j\}) = 1$ is called the *principal ultrafilter generated by j* , and is denoted $\mathcal{U}_j$. An ultrafilter $\mathcal{U}$ on I is called *nonprincipal* if it is not principal. Equivalently, $\mathcal{U}$ is nonprincipal if $\mathcal{U}(X) = 0$ for all finite $X \subseteq I$. It is easy to check that $\lim_{\mathcal{U}_j} z_i = z_j$, whence ultralimits along principal ultrafilters do not really capture a genuine notion of limit. It is a basic fact that, for any infinite set I , there is a nonprincipal ultrafilter $\mathcal{U}$ on I (in fact there are many!). However, any proof of the existence of such ultrafilters is necessarily nonconstructive, and thus it is impossible to explicitly describe a nonprincipal ultrafilter (see [27, Chapter 5] for more on these foundational matters concerning ultrafilters). In what follows, the specific choice of nonprincipal ultrafilter will be irrelevant for our purposes.

We now come to the tracial ultraproduct construction. Fix a family $(\mathcal{M}_i, \tau_i)_{i \in I}$ of tracial von Neumann algebras and an ultrafilter $\mathcal{U}$ on I . We first set

$$\ell^\infty(\mathcal{M}_i) := \left\{ x \in \prod_{i \in I} \mathcal{M}_i : \sup_{i \in I} \|x(i)\| < \infty \right\},$$

that is, $\ell^\infty(\mathcal{M}_i)$ collects all those sequences from the Cartesian product $\prod_{i \in I} \mathcal{M}_i$ for which the operator norms of the coordinates are uniformly bounded. It is readily verified that $\ell^\infty(\mathcal{M}_i)$ is a C*-algebra under the supremum norm. It is tempting to try to define a tracial state τ on $\ell^\infty(\mathcal{M}_i)$ by declaring $\tau(x) := \lim_{\mathcal{U}} \tau_i(x(i))$, which makes sense given that the sequence $(\tau_i(x(i)))_{i \in I}$ is a uniformly bounded sequence of complex numbers (a consequence of the uniform bound on the operator norms of the coordinates of x). Unfortunately, if each $x(i)$ is a positive element of $\mathcal{M}_i$, whence x is a positive element of $\ell^\infty(\mathcal{M}_i)$, with the property that $\lim_{\mathcal{U}} \tau_i(x(i)) = 0$, then we have that $\tau(x) = 0$ even though x may not be zero. In other words, this definition leads to a tracial state on $\ell^\infty(\mathcal{M}_i)$ that is not faithful.

We fix the above problem by defining $c_{\mathcal{U}} := \{x \in \ell^\infty(\mathcal{M}_i) : \lim_{\mathcal{U}} \tau_i(x(i)) = 0\}$. While there are a lot of things to check, we have that:

- $c_{\mathcal{U}}$ is a two-sided ideal in $\ell^\infty(\mathcal{M}_i)$,
- $\ell^\infty(\mathcal{M}_i)/c_{\mathcal{U}}$ is a von Neumann algebra,
- the induced tracial state τ on $\ell^\infty(\mathcal{M}_i)/c_{\mathcal{U}}$ given by $\tau([x]_{\mathcal{U}}) := \lim_{\mathcal{U}} \tau_i(x(i))$ is a trace (that is, a faithful, normal tracial state) on $\ell^\infty(\mathcal{M}_i)/c_{\mathcal{U}}$, where $[x]_{\mathcal{U}}$ denotes the coset of x modulo $c_{\mathcal{U}}$.

The resulting tracial von Neumann algebra is denoted $(\prod_{\mathcal{U}} \mathcal{M}_i, \lim_{\mathcal{U}} \tau_i)$ and is called the *tracial ultraproduct* of the family $(\mathcal{M}_i, \tau_i)$ with respect to $\mathcal{U}$. In a certain sense, one should think of this ultraproduct as some sort of *limit* of the constituent tracial von Neumann algebras, an analogy that will become clearer as we proceed. When each $\mathcal{M}_i$ is a finite factor, then the trace on each factor is unique and we simplify the notation to $\prod_{\mathcal{U}} \mathcal{M}_i$. When each $(\mathcal{M}_i, \tau_i) = (\mathcal{M}, \tau)$, we simply write $(\mathcal{M}, \tau)^{\mathcal{U}}$ and speak of the *ultrapower* of $(\mathcal{M}, \tau)$ with respect to $\mathcal{U}$. Similarly, the ultrapower of a finite factor is denoted $\mathcal{M}^{\mathcal{U}}$. We view any tracial von Neumann algebra $(\mathcal{M}, \tau)$ as a subalgebra of $(\mathcal{M}, \tau)^{\mathcal{U}}$ via the *diagonal embedding* which maps an element $a \in \mathcal{M}$ to the coset of the diagonal sequence $(i \mapsto a)$ modulo $c_{\mathcal{U}}$.

When $\mathcal{U} = \mathcal{U}_j$ is principal, one can verify that $\prod_{\mathcal{U}} (\mathcal{M}_i, \tau_i) \cong (\mathcal{M}_j, \tau_j)$, whence this is not a terribly interesting procedure. The true power of the ultraproduct construction comes when one uses a nonprincipal ultrafilter, for then the ultraproduct is sort of an average or limit of the constituent tracial von Neumann algebras.

If $\lim_{\mathcal{U}} \dim(\mathcal{M}_j) < \infty$, then $\prod_{\mathcal{U}} (\mathcal{M}_i, \tau_i)$ is also finite dimensional. Otherwise, $\prod_{\mathcal{U}} (\mathcal{M}_i, \tau_i)$ is quite large, in fact, nonseparable, even if each $\mathcal{M}_i$ is separable. It is quite common to hear expressions such as “every II_1 factor in this paper (or talk) is separable unless it isn’t.” This tautology refers to the fact that often researchers are only interested in separable tracial von Neumann algebras, and the only nonseparable II_1 factors that one might encounter are those obtained from a nonprincipal ultraproduct of a family of separable II_1 factors. (We are being a bit sloppy—nonprincipality only guarantees nonseparability when the index set is countable; otherwise, one needs the mild assumption of *countable incompleteness*.)

We can now officially state the CEP.

Connes embedding problem. Given any nonprincipal ultrafilter $\mathcal{U}$ on $\mathbb{N}$, every separable tracial von Neumann algebra embeds into $\mathcal{R}^{\mathcal{U}}$; that is, it admits a trace-preserving injective *-homomorphism into $\mathcal{R}^{\mathcal{U}}$.

Let us make several remarks on variations of the statement of the CEP:

- It can be shown using some basic model theory that the CEP is equivalent to the statement that every separable tracial von Neumann algebra embeds into $\mathcal{R}^{\mathcal{U}}$ for *some* nonprincipal ultrafilter $\mathcal{U}$ on $\mathbb{N}$. This fact can also be witnessed by using a simple ultrafilter-free equivalent reformulation of the CEP known as the *microstate conjecture*, discussed below.
- The restriction to separable tracial von Neumann algebras is not necessary. It can be shown that ultrapowers of $\mathcal{R}$ with respect to certain kinds of ultrafilters on larger index sets, known as *good ultrafilters*, lead to larger ultrapowers that can embed tracial von Neumann algebras of larger density character. In other words, we can reformulate CEP by saying every tracial von Neumann algebra embeds into some ultrapower of $\mathcal{R}$.
- The validity of the CEP does not change if we restrict ourselves to embedding II_1 factors into an ultrapower of $\mathcal{R}$. The reason for this is due to the fact that every tracial von Neumann algebra $(\mathcal{M}, \tau)$ embeds into a II_1 factor, say $(\mathcal{M} * L(\mathbb{Z}), \tau * \tau_{L(\mathbb{Z})})$. Here $L(\mathbb{Z})$ is the group von Neumann algebra of the group of integers (see Subsection 3.7) and $*$ denotes the *free product* of tracial von Neumann algebras.
- One can replace $\mathcal{R}^{\mathcal{U}}$ with a nonprincipal ultraproduct of matrix algebras without changing the validity of the CEP. More precisely, CEP is equivalent to the statement that, for any nonprincipal ultrafilter $\mathcal{U}$ on $\mathbb{N}$, every separable tracial von Neumann algebra embeds into $\prod_{\mathcal{U}} M_n(\mathbb{C})$. This follows from the fact that each $M_n(\mathbb{C})$ embeds in $\mathcal{R}$, whence $\prod_{\mathcal{U}} M_n(\mathbb{C})$ embeds in $\mathcal{R}^{\mathcal{U}}$, while there are conditional expectations $\Phi_n : \mathcal{R} \rightarrow M_n(\mathbb{C})$, and the ultralimit of these expectations yields an embedding $\lim_{\mathcal{U}} \Phi_n : \mathcal{R}^{\mathcal{U}} \hookrightarrow \prod_{\mathcal{U}} M_n(\mathbb{C})$. (See Subsection 3.9 for the definition of conditional expectation.)

The last alternate reformulation makes the equivalence with the so-called microstate conjecture more apparent. The microstate conjecture states that for any tracial von Neumann algebra $(\mathcal{M}, \tau)$, any finite collection $p_1(x), \dots, p_m(x)$ of $*$ -polynomials in the noncommuting variables $x = (x_1, \dots, x_n)$, any $a_1, \dots, a_n \in \mathcal{M}$ in the operator norm unit ball of $\mathcal{M}$, and any $\epsilon > 0$, there is $k \in \mathbb{N}$ and $b_1, \dots, b_n \in M_k(\mathbb{C})$ in the operator norm unit ball such that

$$\max_{1 \leq i \leq n} |\tau(p_i(a)) - \text{tr}(p_i(b))| < \epsilon.$$

In other words, any *finite configuration* that can be obtained in some tracial von Neumann algebra can be approximately obtained in some matrix algebra. It is this formulation of CEP that appeared in connection with free entropy as discussed in Subsection 2.1.

3.7. Operator algebras coming from groups. A large source of operator algebras arise from groups and these algebras play an important role in our story.

First, a *unitary representation* of a (discrete) group G is a group homomorphism $\pi : G \rightarrow \mathcal{U}(\mathcal{A})$, where $\mathcal{A}$ is a C^* -algebra and $\mathcal{U}(\mathcal{A})$ denotes the group of unitary elements of $\mathcal{A}$.

Suppose that G is a group. Let $\ell^2(G)$ be the Hilbert space formally generated by an orthonormal basis ζ_h for all $h \in G$. For any $g \in G$, define u_g to be the linear operator on $\ell^2(G)$ determined by $u_g(\zeta_h) = \zeta_{gh}$ for all $h \in G$. Notice that u_g is

unitary for all $g \in G$ (since $u_g^* = u_g^{-1} = u_{g^{-1}}$) and so $\lambda : G \rightarrow \mathcal{U}(\ell^2(G))$ given by $\lambda(g) := u_g$ is a unitary representation of G , called the *left regular representation* of G .

Recall that group algebra $\mathbb{C}[G]$ consists of formal linear combinations $\sum_{g \in G} c_g g$ with only finitely many nonzero coefficients. There is a natural $*$ -algebra structure on $\mathbb{C}[G]$, the addition and multiplication being the obvious ones and the $*$ -operation being given by $(\sum_{g \in G} c_g g)^* = \sum_{g \in G} \overline{c_g} g^{-1}$. $\mathbb{C}[G]$ is in fact a unital $*$ -algebra with unit e , where e denotes the identity of the group.

The left regular representation λ of G extends by linearity to a unital $*$ -algebra homomorphism $\pi : \mathbb{C}[G] \rightarrow \mathcal{B}(\ell^2(G))$.

The *reduced group C^* -algebra* of G , denoted $C_r^*(G)$, is the closure of $\pi(\mathbb{C}[G])$ in the operator norm topology on $\mathcal{B}(\ell^2(G))$. The *group von Neumann algebra* of G , denoted $L(G)$, is the closure of $\pi(\mathbb{C}[G])$ in the WOT on $\mathcal{B}(\ell^2(G))$. Moreover, the vector state on $\mathcal{B}(\ell^2(G))$ corresponding to ξ_e yields a tracial state on $C_r^*(G)$ and a trace on $L(G)$.

When G is finite, $C_r^*(G) = L(G) = \mathbb{C}[G]$ and is generally considered uninteresting (to operator algebraists). When G is infinite, $L(G)$ is a II_1 factor precisely when G is an *infinite conjugacy class (ICC) group*, that is, when all nontrivial conjugacy classes of G are infinite.

The procedure of taking the group von Neumann algebra of a group can “forget” a lot of the algebraic structure of the group. For example, it follows from Connes’ fundamental work [15] that all ICC amenable groups have group von Neumann algebra isomorphic to $\mathcal{R}$.

In what follows, the reduced group C^* -algebra of a group is not quite as important as a second C^* -algebra associated to a group, the so-called *universal (or maximal) group C^* -algebra*. To define this, we define a norm on $\mathbb{C}[G]$ by defining

$$\left\| \sum_{g \in G} c_g u_g \right\| = \sup \left\{ \left\| \sum_{g \in G} c_g \pi(g) \right\| : \pi : G \rightarrow U(\mathcal{A}) \text{ a unitary representation of } G \right\}.$$

It is readily verified that this is a well-defined (that is, finite) C^* -norm on $\mathbb{C}[G]$. The completion of $\mathbb{C}[G]$ with respect to this norm is thus a C^* -algebra, called the *universal C^* -algebra* associated to G , denoted $C^*(G)$. Since the above norm is easily seen to be the maximal C^* -norm on $\mathbb{C}[G]$, it is sometimes called the *maximal norm* on $\mathbb{C}[G]$ and the completion the *maximal group C^* -algebra* of G . It follows immediately from the definition that any unitary representation $\pi : G \rightarrow \mathcal{U}(\mathcal{A})$ of G extends uniquely to a $*$ -homomorphism $\pi : C^*(G) \rightarrow \mathcal{A}$.

A particular corollary of this universal property of the universal group C^* -algebra is that $C^*(\mathbb{F}_\infty)$ is *surjectively universal*, where $\mathbb{F}_\infty$ is the free group on a countably infinite set of generators. More precisely, given any separable C^* -algebra $\mathcal{A}$, there is a surjective $*$ -homomorphism $C^*(\mathbb{F}_\infty) \rightarrow \mathcal{A}$. To see that this is the case, just note that there is a countable set $\{u_i : i \in \mathbb{N}\}$ of unitaries that generates $\mathcal{A}$ (as a C^* -algebra); now apply the universal property to the surjective unitary representation $\mathbb{F}_\infty \rightarrow \mathcal{U}(\mathcal{A})$ obtained by mapping the i th basis element of $\mathbb{F}_\infty$ onto u_i .

Another consequence of the universal property is that if $f : G \rightarrow H$ is a group homomorphism, then we get an induced $*$ -algebra homomorphism $C^*(f) : C^*(G) \rightarrow C^*(H)$. A less obvious fact is that if H is a subgroup of G , then $C^*(H)$ is naturally a C^* -subalgebra of $C^*(G)$. This follows immediately from the definitions

once one knows that any unitary representation of H can be extended to a unitary representation of G (via a technique known as *induction*; see [22, Chapter 6]).

By considering the left-regular representation of G , we immediately see that there is a canonical surjective $*$ -homomorphism $C^*(G) \rightarrow C_r^*(G)$. In general, this map is not an isomorphism; that is, it often has nontrivial kernel. In fact, the canonical map $C^*(G) \rightarrow C_r^*(G)$ is an isomorphism precisely when G is amenable.

One final fact will prove useful later: for any two groups G and H , we have $C^*(G * H) \cong C^*(G) * C^*(H)$, where $G * H$ denotes the free product of groups and $C^*(G) * C^*(H)$ denotes the *unital free product of C^* -algebras*, which is slightly annoying to define but whose properties can be guessed from the terminology.

3.8. The problem with C^* -algebra tensor products. Before discussing the issues associated with defining tensor products of C^* -algebras, we first recall the tensor product construction for vector spaces. Let V and W be vector spaces over the same field $\mathbb{K}$. We let $\mathbb{F}(V \times W)$ be the free $\mathbb{K}$ -vector space on the set $V \times W$, that is, all formal linear combinations $\sum_{(v,w) \in V \times W} c_{(v,w)}(v,w)$ with only finitely many nonzero coefficients. $\mathbb{F}(V \times W)$ carries an obvious $\mathbb{K}$ -vector space structure. The tensor product of V and W , denoted $V \odot W$, is the quotient of $\mathbb{F}(V \times W)$ by the subspace generated by elements of the following form, for $v, v' \in V$, $w, w' \in W$, and $\alpha \in \mathbb{K}$:

- $(v + v', w) - (v, w) - (v', w)$,
- $(v, w + w') - (v, w) - (v, w')$,
- $(\alpha v, w) - \alpha(v, w)$,
- $(v, \alpha w) - \alpha(v, w)$.

While it is more common to write $V \otimes W$ instead of $V \odot W$, we will reserve $\otimes$ for *analytic* tensor products (to be defined shortly) and will use $\odot$ for the above *algebraic* tensor product.

The equivalence class of (v, w) in $V \odot W$ is denoted $v \otimes w$. Thus, an arbitrary element of $V \odot W$ may be written as a formal linear combination $\sum_{i=1}^n \alpha_i v_i \otimes w_i$, but not necessarily uniquely.

If V and W are both finite dimensional, then so is $V \odot W$ with $\dim(V \odot W) = \dim(V) \cdot \dim(W)$; if $\{v_1, \dots, v_m\}$ is a basis for V and $\{w_1, \dots, w_n\}$ is a basis for W , then $\{v_i \otimes w_j : 1 \leq i \leq m, 1 \leq j \leq n\}$ is a basis for $V \odot W$.

It is clear from the construction that if $S : V_1 \rightarrow V_2$ and $T : W_1 \rightarrow W_2$ are $\mathbb{K}$ -linear maps, then there is a $\mathbb{K}$ -linear map $S \odot T : V_1 \odot W_1 \rightarrow V_2 \odot W_2$ uniquely determined by $(S \odot T)(v \otimes w) = S(v) \otimes T(w)$.

If $\mathcal{H}$ and $\mathcal{K}$ are Hilbert spaces, then the algebraic tensor product $\mathcal{H} \odot \mathcal{K}$ comes naturally equipped with an inner product uniquely determined by

$$\langle \xi_1 \otimes \eta_1, \xi_2 \otimes \eta_2 \rangle = \langle \xi_1, \xi_2 \rangle \cdot \langle \eta_1, \eta_2 \rangle.$$

The completion of $\mathcal{H} \odot \mathcal{K}$ with respect to this inner product is then a Hilbert space, denoted $\mathcal{H} \otimes \mathcal{K}$ and called the *Hilbert space tensor product* of $\mathcal{H}$ and $\mathcal{K}$. If $\{e_i : i \in I\}$ and $\{f_j : j \in J\}$ are orthonormal bases for $\mathcal{H}$ and $\mathcal{K}$, respectively, then $\{e_i \otimes f_j : i \in I, j \in J\}$ is an orthonormal basis for $\mathcal{H} \otimes \mathcal{K}$. Moreover, if $S : \mathcal{H}_1 \rightarrow \mathcal{H}_2$ and $T : \mathcal{K}_1 \rightarrow \mathcal{K}_2$ are bounded linear maps, then the algebraic tensor product map $S \odot T$ extends uniquely to a bounded linear map $S \otimes T : \mathcal{H}_1 \otimes \mathcal{K}_1 \rightarrow \mathcal{H}_2 \otimes \mathcal{K}_2$.

We now come to the task of defining tensor products of operator algebras. We first note that if $\mathcal{A}$ and $\mathcal{B}$ are two $*$ -algebras, there is a natural $*$ -algebra operation

on their algebraic tensor product $\mathcal{A} \odot \mathcal{B}$ determined by

- $(x_1 \otimes y_1) \cdot (x_2 \otimes y_2) = (x_1 x_2) \otimes (y_1 y_2),$
- $(x \otimes y)^* = x^* \otimes y^*.$

If $\mathcal{A}$ and $\mathcal{B}$ are both unital, then so is $\mathcal{A} \odot \mathcal{B}$ with unit $1 \otimes 1$.

The tensor product of von Neumann algebras is fairly uncontroversial. Consider concretely represented von Neumann algebras $\mathcal{M} \subseteq \mathcal{B}(\mathcal{H})$ and $\mathcal{N} \subseteq \mathcal{B}(\mathcal{K})$. It is straightforward to check that the algebraic tensor product $\mathcal{M} \odot \mathcal{N}$ is naturally a subset of $\mathcal{B}(\mathcal{H} \otimes \mathcal{K})$ (using the tensor product of linear transformation construction above) and that the $*$ -algebra structure induced by this identification agrees with the one placed on it in the previous paragraph. The *von Neumann algebra tensor product* $\mathcal{M} \bar{\otimes} \mathcal{N}$ of $\mathcal{M}$ and $\mathcal{N}$ is then the WOT closure of $\mathcal{M} \odot \mathcal{N}$ in $\mathcal{B}(\mathcal{H} \otimes \mathcal{K})$. One can verify that this construction is indeed independent of the choice of representations of $\mathcal{M}$ and $\mathcal{N}$.

The story for C^* -algebras, on the other hand, is far more complicated in general. Fix C^* -algebras $\mathcal{A}$ and $\mathcal{B}$. We seek C^* -norms on $\mathcal{A} \odot \mathcal{B}$, for then the completion of $\mathcal{A} \odot \mathcal{B}$ with respect to such a C^* -norm will be a C^* -algebra tensor product of $\mathcal{A}$ and $\mathcal{B}$.

One natural choice is to proceed as in the case of von Neumann algebras, that is, fix concrete representations $\mathcal{A} \subseteq \mathcal{B}(\mathcal{H})$ and $\mathcal{B} \subseteq \mathcal{B}(\mathcal{K})$, and to consider the operator norm on $\mathcal{A} \odot \mathcal{B} \subseteq \mathcal{B}(\mathcal{H} \otimes \mathcal{K})$. One can verify that this norm on $\mathcal{A} \odot \mathcal{B}$ is a C^* -norm and is independent of the choice of representations. This norm is called the *minimal tensor norm*, denoted $\|\cdot\|_{\min}$. The justification for the name comes from a theorem of Takesaki showing that $\|\cdot\|_{\min}$ is indeed the minimal C^* -norm on $\mathcal{A} \odot \mathcal{B}$. The completion of $\mathcal{A} \odot \mathcal{B}$ with respect to $\|\cdot\|_{\min}$ is denoted $\mathcal{A} \otimes_{\min} \mathcal{B}$ and is called the *minimal tensor product* of $\mathcal{A}$ and $\mathcal{B}$. One should be aware of the fact that some authors simply write $\mathcal{A} \otimes \mathcal{B}$ instead of $\mathcal{A} \otimes_{\min} \mathcal{B}$.

A useful property of the minimal tensor product is the following result, which follows from the independence of the choice of representation. If $\pi_{\mathcal{A}} : \mathcal{A}_1 \rightarrow \mathcal{A}_2$ and $\pi_{\mathcal{B}} : \mathcal{B}_1 \rightarrow \mathcal{B}_2$ are $*$ -homomorphisms, then the linear map $\pi_{\mathcal{A}} \odot \pi_{\mathcal{B}} : \mathcal{A}_1 \odot \mathcal{B}_1 \rightarrow \mathcal{A}_2 \odot \mathcal{B}_2$ extends uniquely to a $*$ -homomorphism $\pi_{\mathcal{A}} \otimes \pi_{\mathcal{B}} : \mathcal{A}_1 \otimes_{\min} \mathcal{B}_1 \rightarrow \mathcal{A}_2 \otimes_{\min} \mathcal{B}_2$. Using the Stinespring dilation theorem, one can generalize the conclusion of the previous sentence to ucp maps as follows. If $\Phi_{\mathcal{A}} : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H}_{\mathcal{A}})$ and $\Phi_{\mathcal{B}} : \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H}_{\mathcal{B}})$ are ucp maps, then there is a unique ucp map $\Phi_{\mathcal{A}} \otimes \Phi_{\mathcal{B}} : \mathcal{A} \otimes_{\min} \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H}_{\mathcal{A}} \otimes \mathcal{H}_{\mathcal{B}})$ determined by $(\Phi_{\mathcal{A}} \otimes \Phi_{\mathcal{B}})(a \otimes b) = \Phi_{\mathcal{A}}(a) \otimes \Phi_{\mathcal{B}}(b)$.

Another natural C^* -norm to consider on $\mathcal{A} \odot \mathcal{B}$ is the so-called *maximal norm* defined by

$$\|x\|_{\max} := \sup\{\|\pi(x)\| : \pi : \mathcal{A} \odot \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H}) \text{ a } * \text{-homomorphism}\}.$$

In connection with this formula, it is useful to observe that a $*$ -homomorphism $\pi : \mathcal{A} \odot \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ restricts to $*$ -homomorphisms $\pi_{\mathcal{A}} : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ and $\pi_{\mathcal{B}} : \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ with commuting ranges and, conversely, any two $*$ -homomorphisms $\pi_{\mathcal{A}} : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ and $\pi_{\mathcal{B}} : \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ with commuting ranges yield a $*$ -homomorphism $\pi_{\mathcal{A}} \odot \pi_{\mathcal{B}} : \mathcal{A} \odot \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ uniquely determined by $(\pi_{\mathcal{A}} \odot \pi_{\mathcal{B}})(x \otimes y) := \pi_{\mathcal{A}}(x) \pi_{\mathcal{B}}(y)$. (The commutativity of the ranges of $\pi_{\mathcal{A}}$ and $\pi_{\mathcal{B}}$ ensure that this map is in fact a $*$ -homomorphism.) It is clear that $\|\cdot\|_{\max}$ is a C^* -norm on $\mathcal{A} \odot \mathcal{B}$; the completion of $\mathcal{A} \odot \mathcal{B}$ with respect to $\|\cdot\|_{\max}$ is called the *maximal tensor product* of $\mathcal{A}$ and $\mathcal{B}$, denoted $\mathcal{A} \otimes_{\max} \mathcal{B}$. Any pair of $*$ -homomorphisms $\pi_{\mathcal{A}} : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ and $\pi_{\mathcal{B}} : \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ with commuting ranges yields a $*$ -homomorphism $\pi_{\mathcal{A}} \otimes \pi_{\mathcal{B}} : \mathcal{A} \otimes_{\max} \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ that

extends $\pi_{\mathcal{A}} \odot \pi_{\mathcal{B}}$. Consequently, $\|\cdot\|_{\max}$ really is the largest C^* -norm on $\mathcal{A} \odot \mathcal{B}$. Using a more complicated Stinespring argument than the one mentioned above, one can show that any pair of ucp maps $\Phi_{\mathcal{A}} : \mathcal{A} \rightarrow \mathcal{B}(\mathcal{H})$ and $\Phi_{\mathcal{B}} : \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ with commuting ranges yields a ucp map $\Phi_{\mathcal{A}} \otimes \Phi_{\mathcal{B}} : \mathcal{A} \otimes_{\max} \mathcal{B} \rightarrow \mathcal{B}(\mathcal{H})$ uniquely determined by $(\Phi_{\mathcal{A}} \otimes \Phi_{\mathcal{B}})(a \otimes b) = \Phi_{\mathcal{A}}(a)\Phi_{\mathcal{B}}(b)$.

Before moving forward, we notice the following two facts, which are readily verified from the definitions. For any pair of groups G and H , we have

- $C_r^*(G \times H) \cong C_r^*(G) \otimes_{\min} C_r^*(H)$,
- $C^*(G \times H) \cong C^*(G) \otimes_{\max} C^*(H)$.

Returning to the general discussion, we have defined two *extreme* C^* -norms on $\mathcal{A} \odot \mathcal{B}$. In general, they can be different. For example, it can be shown that the maximal and minimal norms on $C_r^*(\mathbb{F}_2) \odot C_r^*(\mathbb{F}_2)$ are distinct. The corresponding question for $C^*(\mathbb{F}_2) \odot C^*(\mathbb{F}_2)$ turns out to be equivalent to CEP, as we will soon see! Another somewhat surprising result is that the maximal and minimal norms on $\mathcal{B}(\mathcal{H}) \odot \mathcal{B}(\mathcal{H})$ (for $\mathcal{H}$ infinite dimensional) are also distinct, a result due to Junge and Pisier [42]. In fact, Ozawa and Pisier [49] showed that there exist at least continuum many different C^* -norms on $\mathcal{B}(\mathcal{H}) \odot \mathcal{B}(\mathcal{H})$ when $\mathcal{H}$ is infinite dimensional.

We say that $(\mathcal{A}, \mathcal{B})$ form a *nuclear pair* if there is a unique C^* -norm on $\mathcal{A} \odot \mathcal{B}$, that is, if the minimal and maximal norms on $\mathcal{A} \odot \mathcal{B}$ coincide. We also say that $\mathcal{A}$ is *nuclear* if $(\mathcal{A}, \mathcal{B})$ is a nuclear pair for every C^* -algebra $\mathcal{B}$. There are many interesting examples of nuclear C^* -algebras. For example, $M_n(\mathbb{C})$ is nuclear for all n , the reason being that $M_n \odot \mathcal{B} \cong M_n(\mathcal{B})$, which is already a C^* -algebra with a unique C^* -norm. A more interesting example coming from groups is that $C^*(G)$ is nuclear if and only if $C_r^*(G)$ is nuclear if and only if G is amenable (in which case $C^*(G) = C_r^*(G)$).

The following theorem of Kirchberg [44] will be central moving forward.

Theorem 3.1. $(C^*(\mathbb{F}_{\infty}), \mathcal{B}(\mathcal{H}))$ is a nuclear pair.

Note that neither of these algebras are nuclear. The importance of Kirchberg's theorem stems from the fact that $C^*(\mathbb{F}_{\infty})$ is surjectively universal while $\mathcal{B}(\mathcal{H})$ is injectively universal.

We also note that if H is a subgroup of G and $\mathcal{A}$ is any C^* -algebra for which $(C^*(G), \mathcal{A})$ is a nuclear pair, then so is $(C^*(H), \mathcal{A})$. In particular, whether or not $(C^*(\mathbb{F}_k), C^*(\mathbb{F}_k))$ is a nuclear pair is independent of the choice of $k \in \{2, 3, \dots\} \cup \{\infty\}$, a fact that will come up in our discussion of Kirchberg's QWEP problem. We will also need the fact that if $p : G \rightarrow H$ is a surjective group morphism for which the canonical surjection $C^*(p) : C^*(G) \rightarrow C^*(H)$ has a ucp lift (meaning a ucp map $\Phi : C^*(H) \rightarrow C^*(G)$ for which $C^*(p)\Phi$ is the identity on $C^*(H)$), then $(C^*(G), C^*(G))$ being a nuclear pair implies $(C^*(H), C^*(H))$ is a nuclear pair.

3.9. Kirchberg's QWEP problem. It follows from the definition of the minimal tensor product that for any inclusion $\mathcal{A} \subseteq \mathcal{B}$ of C^* -algebras, one has that $\mathcal{A} \otimes_{\min} \mathcal{C} \subseteq \mathcal{B} \otimes_{\min} \mathcal{C}$ (isometrically) for any other C^* -algebra $\mathcal{C}$. On the other hand, with the same setup, while there will always be a $*$ -homomorphism $\mathcal{A} \otimes_{\max} \mathcal{C} \rightarrow \mathcal{B} \otimes_{\max} \mathcal{C}$, this homomorphism need not be injective, that is, isometric. One case, however, where this does hold is the inclusion $\mathcal{A} \subseteq \mathcal{A}^{**}$. That is, it follows from the definitions that $\mathcal{A} \otimes_{\max} \mathcal{C} \subseteq \mathcal{A}^{**} \otimes_{\max} \mathcal{C}$ isometrically for all C^* -algebras $\mathcal{A}$ and $\mathcal{C}$.

Suppose again that $\mathcal{A} \subseteq \mathcal{B}$ and further suppose, for the sake of argument, that there is a ucp map $\Phi : \mathcal{B} \rightarrow \mathcal{A}^{**}$ with $\Phi(a) = a$ for all $a \in \mathcal{A}$. By a fact pointed out in the previous subsection, we obtain a ucp (and thus contractive) map $\Phi \otimes I_{\mathcal{C}} : \mathcal{B} \otimes_{\max} \mathcal{C} \rightarrow \mathcal{A}^{**} \otimes_{\max} \mathcal{C}$. By the observation made in the previous paragraph, it follows that the canonical map $\mathcal{A} \otimes_{\max} \mathcal{C} \rightarrow \mathcal{B} \otimes_{\max} \mathcal{C}$ is an isometric inclusion.

If $\mathcal{A}$ is a C^* -subalgebra of $\mathcal{B}$ and $\Phi : \mathcal{B} \rightarrow \mathcal{A}$ is a linear map for which $\Phi(a) = a$ for all $a \in \mathcal{A}$, then a theorem of Tomiyama says that the following are equivalent:

- Φ is cp,
- Φ is contractive,
- Φ is a *conditional expectation*, that is, $\Phi(axb) = a\Phi(x)b$ for all $a, b \in \mathcal{A}$ and $x \in \mathcal{B}$.

When such a map exists, we say that $\mathcal{A}$ is *cp-complemented* in $\mathcal{B}$. The nomenclature comes from Banach space theory, for a Banach space X is complemented in a superspace Y if and only if there is a contractive linear map $\Phi : Y \rightarrow X$ that is the identity on X . In the previous paragraph, we merely had to posit the existence of a ucp map $\Phi : \mathcal{B} \rightarrow \mathcal{A}^{**}$, whence we call Φ a *weak conditional expectation* and say that $\mathcal{A}$ is *weakly cp-complemented* in $\mathcal{B}$. Consequently, we proved that if $\mathcal{A}$ is weakly cp-complemented in $\mathcal{B}$, then $\mathcal{A} \otimes_{\max} \mathcal{C} \subseteq \mathcal{B} \otimes_{\max} \mathcal{C}$ isometrically for any other C^* -algebra $\mathcal{C}$. With more work, one can actually show that the converse of this observation holds as well. In fact, by the surjective universality of $C^*(\mathbb{F}_{\infty})$, we see that $\mathcal{A}$ is weakly cp-complemented in $\mathcal{B}$ if and only if $\mathcal{A} \otimes_{\max} C^*(\mathbb{F}_{\infty}) \subseteq \mathcal{B} \otimes_{\max} C^*(\mathbb{F}_{\infty})$ isometrically.

There are two notable examples of cp-complemented inclusions worth pointing out now.

- If $\mathcal{M}$ is a finite von Neumann algebra, then any von Neumann subalgebra $\mathcal{N}$ of $\mathcal{M}$ is cp-complemented. To see this, fix a trace τ on $\mathcal{M}$ and note that $L^2(\mathcal{N}, \tau)$ is a closed subspace of $L^2(\mathcal{M}, \tau)$. One shows that the orthogonal projection $L^2(\mathcal{M}, \tau) \rightarrow L^2(\mathcal{N}, \tau)$ actually restricts to a conditional expectation $\mathcal{M} \rightarrow \mathcal{N}$.
- If H is a subgroup of G , then $C^*(H)$ is cp-complemented in $C^*(G)$. To see this, one shows that the map $\Phi : G \rightarrow \mathbb{C}[H]$, defined by setting $\Phi(g) = g$ for all $g \in H$ while $\Phi(g) = 0$ for all $g \in G \setminus H$, extends to a conditional expectation $C^*(G) \rightarrow C^*(H)$.

Returning to the general situation, if $\mathcal{A}$ is weakly cp-complemented in every superalgebra (or equivalently in $\mathcal{B}(\mathcal{H})$), then we say that $\mathcal{A}$ has the *weak expectation property* (WEP). An insight of Kirchberg [44] was to use his theorem proving that $(C^*(\mathbb{F}_{\infty}), \mathcal{B}(\mathcal{H}))$ is a nuclear pair to provide an alternate *test* for having WEP:

Theorem 3.2. *For a C^* -algebra $\mathcal{A} \subseteq \mathcal{B}(\mathcal{H})$, the following are equivalent:*

- (1) $\mathcal{A}$ has the WEP.
- (2) For every C^* -algebra $\mathcal{C}$, $\mathcal{A} \otimes_{\max} \mathcal{C} \subseteq \mathcal{B}(\mathcal{H}) \otimes_{\max} \mathcal{C}$ isometrically.
- (3) $(\mathcal{A}, C^*(\mathbb{F}_{\infty}))$ is a nuclear pair.

Proof. We already observed the equivalence of (1) and (2). Now suppose that $\mathcal{A}$ has WEP. We then have

$$\mathcal{A} \otimes_{\max} C^*(\mathbb{F}_{\infty}) \subseteq \mathcal{B}(\mathcal{H}) \otimes_{\max} C^*(\mathbb{F}_{\infty}) = \mathcal{B}(\mathcal{H}) \otimes_{\min} C^*(\mathbb{F}_{\infty}) \supseteq \mathcal{A} \otimes_{\min} C^*(\mathbb{F}_{\infty}),$$

that is, $(\mathcal{A}, C^*(\mathbb{F}_\infty))$ is a nuclear pair. Conversely, suppose that $(\mathcal{A}, C^*(\mathbb{F}_\infty))$ is a nuclear pair. It suffices to show that $\mathcal{A} \otimes_{\max} C^*(\mathbb{F}_\infty) \subseteq \mathcal{B}(\mathcal{H}) \otimes_{\max} C^*(\mathbb{F}_\infty)$. However, this follows immediately from the assumption, Kirchberg's theorem, and the fact that $\otimes_{\min}$ preserves inclusions. $\square$

What are the ramifications of assuming that $C^*(\mathbb{F}_\infty)$ itself has the WEP, that is, $(C^*(\mathbb{F}_\infty), C^*(\mathbb{F}_\infty))$ is a nuclear pair? First, as observed above, this is equivalent to $(C^*(\mathbb{F}_k), C^*(\mathbb{F}_k))$ being a nuclear pair for any fixed $k \geq 2$. Next, if we define the *QWEP* to be the property that a C^* -algebra is a quotient of a C^* -algebra with WEP, then $C^*(\mathbb{F}_\infty)$ having the WEP implies that all separable C^* -algebras have the QWEP. We note that it is common to see both phrases, “ $\mathcal{A}$ has the QWEP” and “ $\mathcal{A}$ is QWEP” (although the latter is of course grammatically incorrect).

Conversely, suppose that all separable C^* -algebras have the QWEP. Then certainly $C^*(\mathbb{F}_\infty)$ has the QWEP. However, $C^*(\mathbb{F}_\infty)$ has another property, the so-called *lifting property* (*LP*) which, when combined with QWEP, actually implies the WEP. A C^* -algebra $\mathcal{A}$ has the lifting property if, for any ucp map $\Phi : \mathcal{A} \rightarrow \mathcal{B}/\mathcal{J}$, where $\mathcal{B}$ is a C^* -algebra and $\mathcal{J}$ is a closed, two-sided ideal in $\mathcal{B}$, there is a ucp map $\Psi : \mathcal{A} \rightarrow \mathcal{B}$ for which $\pi \circ \Psi = \Phi$, where $\pi : \mathcal{B} \rightarrow \mathcal{B}/\mathcal{J}$ is the canonical quotient map. Said more casually, $\mathcal{A}$ has the LP if every ucp map into a quotient C^* -algebra has a ucp lift. Now suppose that $\mathcal{A}$ has the LP and the QWEP as witnessed by the quotient map $q : \mathcal{B} \rightarrow \mathcal{A}$ with $\mathcal{B}$ having the WEP. Let $\Psi : \mathcal{A} \rightarrow \mathcal{B}$ be a ucp lift of the identity map $\mathcal{A} \rightarrow \mathcal{A}$, that is, $q \circ \Psi = I_{\mathcal{A}}$. Then by applying the ucp lift $\Psi \otimes I_{C^*(\mathbb{F}_\infty)}$ of $q \otimes I_{C^*(\mathbb{F}_\infty)}$, we see that $\mathcal{A}$ also has the WEP.

We have thus arrived at the following.

Theorem 3.3. *The following assertions are equivalent:*

- (1) $C^*(\mathbb{F}_\infty)$ has the WEP.
- (2) For some (equivalently any) $k \in \{2, 3, \dots\} \cup \{\infty\}$, $(C^*(\mathbb{F}_k), C^*(\mathbb{F}_k))$ is a nuclear pair.
- (3) Every separable C^* -algebra has the QWEP.

Any of the above equivalent statements are known as *Kirchberg's QWEP problem*. (We might be tempted to follow the CEP's lead and call this the QWEPP, but that looks a bit silly.) As mentioned above, QWEP combined with LP implies WEP. It turns out that it suffices to consider a *local* version of LP, aptly called the *local lifting property* (*LLP*) and the same argument works; that is, QWEP together with LLP implies WEP. Thus, another equivalent formulation of the QWEP problem is the statement that LLP implies WEP.

We mention one other equivalent formulation of the QWEP problem that is not relevant for our particular story but is fascinating nonetheless: the QWEP problem is equivalent to the statement that $C^*(\mathbb{F}_\infty \times \mathbb{F}_\infty)$ has a faithful tracial state. What makes this interesting is that this is true for the reduced group C^* -algebra $C_r^*(\mathbb{F}_\infty \times \mathbb{F}_\infty)$ (simply because it is true for any reduced group C^* -algebra) and it is true for $C^*(\mathbb{F}_\infty)$ (a result due to Choi).

The classes of WEP and QWEP algebras enjoy a number of closure properties relevant to the proofs that follow. Rather than enumerate them all now, we will simply quote them when we need them later in the paper.

3.10. From CEP to QWEP. In this section, we show how a positive solution to the CEP implies a positive solution to the QWEP problem. While these statements

are indeed equivalent, we focus on the direction that we need in order to give a negative solution to CEP.

So how does CEP get involved in a story about C^* -algebras? The first clue is that, for von Neumann algebras, the WEP had already been well studied and is referred to as *injectivity*. A not so trivial result is that any hyperfinite von Neumann algebra is injective, whence $\mathcal{R}$ is injective. The extremely deep work of Connes in [15], where the CEP originally comes from, proved the converse, namely any separable injective II_1 factor must be hyperfinite, and thus isomorphic to $\mathcal{R}$. Thankfully we do not need this result in our story, although the proof that $\mathcal{R}$ is injective (and thus has WEP) is difficult enough.

Now that we know that $\mathcal{R}$ is injective, so is $\ell^\infty(\mathcal{R})$ as WEP is closed under the formation of direct sums. Since $\mathcal{R}^u$ is a C^* -algebra quotient of $\ell^\infty(\mathcal{R})$, we see that $\mathcal{R}^u$ is QWEP! Okay, it smells like we are getting closer.

Now suppose that $\mathcal{M}$ is a tracial von Neumann algebra that embeds in $\mathcal{R}^u$ in a trace-preserving manner. Without loss of generality, let us assume that $\mathcal{M}$ is simply a subalgebra of $\mathcal{R}^u$. By a fact pointed out above, this means that $\mathcal{M}$ is cp-complemented in $\mathcal{R}^u$. Since QWEP is preserved by (weakly) cp-complemented inclusions, we conclude that $\mathcal{M}$ is also QWEP.

We have thus arrived at the statement: a positive solution to CEP implies all finite von Neumann algebras are QWEP!

But we are still talking about von Neumann algebras. How do we bridge the gap into talking about C^* -algebras? Well, recall that every C^* -algebra $\mathcal{A}$ has a canonically associated von Neumann algebra $\mathcal{A}^{**}$. Since $\mathcal{A}$ is tautologically weakly cp-complemented in $\mathcal{A}^{**}$, in order to show that $\mathcal{A}$ has QWEP, it suffices to show that $\mathcal{A}^{**}$ has QWEP (again using the closure of QWEP under weakly cp-complemented subalgebras).

While $\mathcal{A}^{**}$ is a separable von Neumann algebra, it may not be finite. How do we get CEP to help us with nonfinite von Neumann algebras?

Given any von Neumann algebra $\mathcal{M}$, there is an important one-parameter group (σ_t^φ) of automorphisms of $\mathcal{M}$, known as the *modular group*. When $\mathcal{M}$ is finite, the modular automorphism group is trivial and thus plays no role. But in the general theory, it is an indispensable tool. (For all of the fancy type III material discussed in this paragraph, Takesaki's book [58] is the canonical reference.) Akin to the semidirect product construction in group theory, there is a *crossed product construction* that associates to any group acting on a von Neumann algebra a larger von Neumann algebra where this action is implemented by unitaries. Thus, we are entitled to consider the crossed product algebra $\mathcal{M} \rtimes_{\sigma_t^\varphi} \mathbb{R}$ corresponding to the action of $\mathbb{R}$ on $\mathcal{M}$ via the modular automorphism group. A serious theorem of Takesaki states that $\mathcal{M} \rtimes_{\sigma_t^\varphi} \mathbb{R}$ is semifinite. We came across semifinite factors above. For a general von Neumann algebra, we can take semifinite to mean that the algebra contains an increasing union of finite subalgebras whose union generates the von Neumann algebra. Since QWEP is preserved under unions and a von Neumann algebra is QWEP if it contains a WOT-dense $*$ -subalgebra with QWEP, we see that $\mathcal{M} \rtimes_{\sigma_t^\varphi} \mathbb{R}$ has QWEP. An alternative approach is to use the fact that a von Neumann algebra is QWEP if and only if all of the factors involved in its direct integral decomposition are QWEP. Thus, to show that a semifinite von Neumann algebra is QWEP, it suffices to consider the case of factors. But then a semifinite factor is of the form $\mathcal{M} \bar{\otimes} \mathcal{B}(\mathcal{H})$ for a finite factor $\mathcal{M}$, and one can use the fact that

the von Neumann algebra tensor product of QWEP von Neumann algebras is again QWEP. Either way, we now know that $\mathcal{M} \rtimes_{\sigma_t^\varphi} \mathbb{R}$ is QWEP.

Finally, it is a general fact that any von Neumann algebra $\mathcal{M}$ is always cp-complemented in any crossed product $\mathcal{M} \rtimes_\alpha G$; since $\mathcal{M} \rtimes_{\sigma_t^\varphi} \mathbb{R}$ is QWEP for any von Neumann algebra $\mathcal{M}$, it follows that $\mathcal{M}$ itself is also QWEP. Applying this fact to $\mathcal{M} = \mathcal{A}^{**}$, we see that $\mathcal{A}^{**}$, and thus $\mathcal{A}$, are also QWEP for any C*-algebra $\mathcal{A}$. This finishes the proof that a positive solution to CEP implies a positive solution to the QWEP problem.

As mentioned before, a positive solution to the QWEP problem implies a positive solution to the CEP. The proof involves the theory of *amenable traces*, which we will not go into now, but which will be important in our alternate derivation of the failure of CEP from $\text{MIP}^* = \text{RE}$ given in Subsection 7.5.

4. A CRASH COURSE IN COMPLEXITY THEORY

In this section, we introduce the basic notions from (classical) complexity theory needed to understand the statement of the result $\text{MIP}^* = \text{RE}$. Essentially all of this material (apart from the business about nonlocal games) was taken from the book [3].

4.1. Turing machines. A *Turing machine* is one of the more popular mathematical formulations of an idealized computing device. Formally, a Turing machine is a pair $\mathbf{M} = (Q, \delta)$, where Q is a finite set of *states* of the machine and $\delta : Q \times \{0, 1, \square, \triangle\}^3 \rightarrow Q \times \{0, 1, \square, \triangle\}^2 \times \{L, S, R\}^3$ is the *transition function*; here $\square$ and $\triangle$ are two special symbols whose significance will be seen shortly. We always assume that Q contains two special states, namely the *start state* q_{start} and the *halting state* q_{halt} .

Throughout, for any $n \in \mathbb{N}$, $\{0, 1\}^n$ denotes the set of binary strings of length n while $\{0, 1\}^* := \bigcup_{n \in \mathbb{N}} \{0, 1\}^n$ denotes the set of all finite binary strings. Given $z \in \{0, 1\}^*$, $|z|$ denotes the length of the string z .

Here is how one should envision the computation performed by the Turing machine $\mathbf{M}$ upon some input $z \in \{0, 1\}^*$. The machine contains three *tapes*, which are one-way infinite strips containing boxes which, at any given moment in the computation, contain exactly one symbol from $\{0, 1, \square, \triangle\}$. The first tape is the *input tape*, the second tape is the *work tape*, and the last tape is the *output tape*. At the beginning of the computation, the input tape has the *start symbol* $\triangle$ in the first box, then the input string z in the next $|z|$ boxes, and then the remainder of the boxes contain the *blank symbol* $\square$. Both the work tape and the output tape contain the start symbol $\triangle$ in the first box and then blank symbols $\square$ in the remaining boxes. One envisions each tape having a “tape head” which is placed over exactly one box in the tape at any given moment during the computation; the tape head for the input tape can read the symbol in that box while the tape head for the other two tapes can both read the symbol in that box and potentially change it to a new symbol.

The Turing machine begins the computation in the start state q_{start} with the tape head above the leftmost box (which contains the start symbol $\triangle$) for each tape. In general, at any given moment during the computation, the Turing machine is in some state $q \in Q$ with tape heads reading boxes $k_1, k_2, k_3 \in \mathbb{N}$ (representing how far they are from the beginning of their respective tape) and with symbols $s_1, s_2, s_3 \in \{0, 1, \square, \triangle\}$ inside of each of the boxes being read. The Turing machine then computes $\delta(q, s_1, s_2, s_3)$, obtaining the tuple $(q', s'_2, s'_3, I_1, I_2, I_3)$, which should

be interpreted as follows:

- The box in the work tape (resp., output tape) that the tape head is reading should have its contents replaced by s'_2 (resp., s'_3).
- The tape head for the input tape should move to the left if $I_1 = L$, to the right if $I_1 = R$, and should stay in the same place if $I_1 = S$. Similar actions should be taken corresponding to I_2 for the work tape and I_3 for the output tape. If any tape head is at the leftmost box and the instruction is L , then the tape head should also stay in the same place.
- After executing these acts, the Turing machine should now enter state q' .

The machine continues *running* in this fashion. If the machine ever enters the state q_{halt} , then the machine *stops running*, that is, no further modification of the three tapes will take place. In this case, the *output* of the computation upon input z is the longest initial string on the output tape not containing any blank symbols. (If all the symbols are blank, then the output is considered the empty string).

Every Turing machine $\mathbf{M}$ *computes* a partial function $f^{\mathbf{M}} : \{0, 1\}^* \rightarrow \{0, 1\}^*$ whose domain consists of those strings $z \in \{0, 1\}^*$ for which $\mathbf{M}$ halts upon input z ; in this case, we define $f^{\mathbf{M}}(z)$ to be the corresponding output. We sometimes abuse notation and identify $f^{\mathbf{M}}$ with $\mathbf{M}$ itself, that is, we may write $\mathbf{M}(z)$ instead of $f^{\mathbf{M}}(z)$. We say a partial function $f : \{0, 1\}^* \rightarrow \{0, 1\}^*$ is *computable* if $f = f^{\mathbf{M}}$ for some Turing machine $\mathbf{M}$.

Given a function $T : \mathbb{N} \rightarrow \mathbb{N}$, we say that the Turing machine $\mathbf{M}$ *runs in* $T(n)$ -*time* if, for any input $z \in \{0, 1\}^*$, upon input z , $\mathbf{M}$ halts in at most $T(|z|)$ steps. Note that if $\mathbf{M}$ runs in $T(n)$ -time for some function T , then $f^{\mathbf{M}}$ is a total function. We say that $\mathbf{M}$ is a *polynomial time* (resp., *exponential time*) Turing machine if $\mathbf{M}$ runs in Cn^c -time (resp., $C2^{n^c}$ -time) for some constants $C > 0$ and $c \geq 1$.

A *language* is simply a subset $\mathbf{L} \subseteq \{0, 1\}^*$. We identify a language $\mathbf{L}$ with its characteristic function $\chi_{\mathbf{L}} : \{0, 1\}^* \rightarrow \{0, 1\}$. Consequently, it makes sense to speak of $\mathbf{L}$ being computable by a Turing machine. Usually, a language is described in terms of some mathematical problem under consideration, e.g., the set of finite graphs that can be 3-colored. The implicit assumption is that there is some natural (and effective) way of coding the set of such graphs as a set of finite binary strings. In the following, for all languages introduced in this manner, we assume that the reader can figure out how such a coding might be performed.

Turing machines are one of several mathematically precise models for computation; other alternatives include register machines and the class of recursive functions. However, all known *reasonable* models of computation lead to precisely the same class of computable functions. This is evidence for the *Church–Turing thesis*, which states that this common class of functions coincides with our heuristic notion of what a computable function should be. (See [18, Chapter 3] for more on this.) One can even formulate a stronger version of the thesis, which states that even when taking into account the efficiency of computations, that is, how *fast* one can compute a function, the choice of model is still irrelevant. (It is plausible that quantum computers could pose a serious threat to the strong Church–Turing thesis.) The import of the strong Church–Turing thesis for us in these notes is that, in what follows, when claiming that a certain problem can be solved in a certain efficient manner, we never need to actually write down the Turing machine that witnesses this fact. Instead, one can write down an argument using *pseudo-code*

and the reader can (if they choose to) convert the pseudo-code into an actual Turing machine program.

4.2. Some basic complexity classes. A *complexity class* is simply a collection of languages. The most interesting complexity classes are those defined by some sort of condition saying that the languages in the class represent efficiently computable (or verifiable, as we shall shortly see) problems.

The complexity class P is defined to be the class of languages L such that membership in L can be decided by a Turing machine in polynomial time, that is, χ_L can be computed by a polynomial time Turing machine. For example, the set of connected graphs is a language that belongs to P (as witnessed by, say, the breadth first search algorithm).

The complexity class EXP is defined in the same manner as P , replacing polynomial time by exponential time. The *time hierarchy theorem* implies that $P \subsetneq EXP$.

Sometimes it is too difficult to come up with an algorithm that efficiently decides membership in a particular language while it is the case that if someone were to hand you a proof that a certain string belonged to the language, then you could efficiently verify that the proof was indeed correct. The complexity class NP captures this idea. More precisely, the complexity class NP consists of those languages L for which there is a polynomial time Turing machine M and a polynomial $p(n)$ such that:

- for all $z \in L$, there is $w \in \{0, 1\}^{p(|z|)}$ for which $M(z, w) = 1$.
- for all $z \notin L$ and for all $w \in \{0, 1\}^{p(|z|)}$, $M(z, w) = 0$.

In the above definition, one thinks of w as the *proof* that $z \in L$; other commonly used terms for w are “witness” and “certificate.” One often envisions this situation using two fictitious players, a verifier and a prover. If $z \in L$, the prover hands the verifier a proof w that z indeed belongs to L ; the prover has unlimited computation power in this regard. In order for the verifier to be able to efficiently check that the proof indeed works, the proof cannot be too long (or else the verifier will not even be able to read the entire proof), hence the polynomial length requirement. Moreover, if $z \notin L$, then there should be no proof that z belongs to L , whence the second condition.

It is clear that $P \subseteq NP$. While intuitively it seems clear that this inclusion should be proper (there *ought* to be problems that are impossible to efficiently decide, yet there are always proofs that are efficiently verifiable), this fact has yet to be established and remains one of the more famous open problems in mathematics.

We also note that $NP \subseteq EXP$, as one can check all of the exponentially many possible certificates for a given string in exponential time.

An example of a language in NP is the set of codes for pairs (G, k) , where G is a finite graph that contains an independent set of size k ; a certificate for a given pair is simply an independent set of size k . This language is unlikely to be in P . Indeed, this is an example of a so-called *NP-complete problem*, meaning that it is as difficult as any other problem in NP (in a precise sense), whence if it belongs to P , then so do all languages in NP and $P = NP$. Another example of a language in NP is the set of codes for pairs (G_1, G_2) of finite graphs that are isomorphic; the certificate here is the isomorphism between the graphs. This problem, however, is unlikely to be NP -complete (see [3, Section 8.4]).

An alternative way of defining the class NP is to use *nondeterministic Turing machines*, which is actually the original definition and explains the terminology (NP stands for nondeterministic polynomial time). A nondeterministic Turing machine is defined exactly like a deterministic one except that it has two transition functions rather than one. Consequently, rather than there being a single (deterministic) sequence of steps during a computation upon a given input, there is an entire binary tree of such computations, for at every step during a computation, one can apply either of the two transition functions. One additional difference is that instead of a single halting state q_{halt} , we now have two halting states called q_{accept} and q_{reject} . We say that the nondeterministic Turing machine $\mathbf{M}$ outputs 1 on input z if there is *some* sequence of steps which causes the machine to reach q_{accept} . If every sequence of steps causes the machine to reach q_{reject} , then we say that $\mathbf{M}$ outputs 0. If every sequence of computations results in either q_{accept} or q_{reject} in time $T(|z|)$, then we say that $\mathbf{M}$ runs in time $T(n)$. It thus makes sense to speak of polynomial time (resp., exponential time) nondeterministic Turing machines.

It can easily be verified that a language $\mathbf{L}$ belongs to NP if and only if there is a polynomial time nondeterministic Turing machine $\mathbf{M}$ such that $f^{\mathbf{M}} = \chi_{\mathbf{L}}$. Moreover, using exponential time nondeterministic Turing machines, we can also define the complexity class NEXP. Of course, using nondeterministic Turing machines that run in doubly exponential time, one can also define NEEXP (this will come up later). A nondeterministic version of the time hierarchy theorem guarantees $\text{NP} \subsetneq \text{NEXP} \subsetneq \text{NEEXP}$.

At this point, we have $P \subseteq \text{NP} \subseteq \text{EXP} \subseteq \text{NEXP}$ with $P \subsetneq \text{EXP}$ and $\text{NP} \subsetneq \text{NEXP}$. One can also use a *padding* argument to show that if $P = \text{NP}$, then $\text{EXP} = \text{NEXP}$.

So far we have only been concerned with time efficiency. One can instead consider *space efficiency*. We will only consider the class PSPACE, which consists of all languages $\mathbf{L}$ for which there is a Turing machine such that, upon input z , it decides whether or not $z \in \mathbf{L}$ using only a polynomial amount of work space. It is clear that $P \subseteq \text{PSPACE}$. It is also fairly easy to see that $\text{NP} \subseteq \text{PSPACE}$, for one can simply check all possible certificates, erasing one's work after each individual check, thus using only a polynomial amount of space. A slightly less obvious inclusion is $\text{PSPACE} \subseteq \text{EXP}$; the proof uses the notion of a configuration graph for a computation. So, to update our state of knowledge, we have $P \subseteq \text{NP} \subseteq \text{PSPACE} \subseteq \text{EXP} \subseteq \text{NEXP}$. It is not known if either inclusion $\text{NP} \subseteq \text{PSPACE}$ or $\text{PSPACE} \subseteq \text{NEXP}$ is proper (although one of them must be as $\text{NP} \subsetneq \text{NEXP}$); as with $P \subseteq \text{NP}$, it is believed that both of these inclusions are proper.

We end this section with the definition of the class BPP. Although it will not play a direct role in the story to follow, it will make a later pill easier to swallow. We return to the setting of nondeterministic Turing machines, but this time we count the proportion of computations that output $\mathbf{M}(z) = 1$. We say that the language $\mathbf{L}$ belongs to the class BPP if there is a nondeterministic Turing machine $\mathbf{M}$ such that, upon $z \in \{0, 1\}^*$, the probability that a random nondeterministic computation agrees with $\chi_{\mathbf{L}}(z)$ is at least $\frac{2}{3}$. Just as in the case of NP, there is a formulation using deterministic Turing machines: $\mathbf{L}$ belongs to BPP if and only if there is a Turing machine $\mathbf{M}$ and a polynomial $p(n)$ such that, for every string $z \in \{0, 1\}^*$,

the probability that a random $r \in \{0, 1\}^{p(|z|)}$ is such that $\mathbf{M}(z, r) = \chi_{\mathbf{L}}(z)$ is at least $\frac{2}{3}$. We remark that the choice of $\frac{2}{3}$ is fairly arbitrary; by repeating the computation several (but still a reasonable number of) times and taking the majority result of the computations, we can replace $\frac{2}{3}$ with a probability as close to 1 as one desires. The class BPP is contained in EXP as one can check all random bits and compute the probability that a random choice yields 1 or 0.

If $\mathbf{L}$ is a language in BPP and one repeats the computation a sufficient number of times to achieve a probability of, say, 0.99, then one can be fairly certain that the result of the probabilistic computation is the truth, and thus BPP seems like a fairly good substitute for P. In fact, there are complexity-theoretic reasons for believing that BPP might coincide with P (see [3, Chapter 16]).

4.3. Interactive proofs. We now imagine the situation where rather than the prover just handing the verifier a proof, the prover and the verifier are allowed to interact. Given an input, the verifier can ask the prover a question, the prover can answer that question, then (based on that answer) the verifier can ask the prover another question to which the prover can reply, and so on, for a certain number of rounds. Each time, the verifier's question and the prover's answer depend on the entire sequence of questions and answers obtained up to that point (as well as the input). After this discussion, the verifier can decide whether or not to accept. Once again, the verifier uses a polynomial-time Turing machine to choose which questions to ask and whether or not to accept at the end of the conversation while the prover has no computational limitations.

It is not too difficult to verify that, with this description of interactive proof, the corresponding complexity class would simply be NP in disguise. Indeed, one can just use the conversation, or *transcript* as it is usually called, as the certificate. However, combining this idea with a randomized process as discussed at the end of the last subsection does lead to a class with more computational power (although, interestingly enough, it is a class we have already seen before).

In order to define this class, we fix $k \in \mathbb{N}$ (although one could actually work with a polynomial-time computable $k : \mathbb{N} \rightarrow \mathbb{N}$ instead) and a polynomial $p(n)$. Assume that we also have a Turing machine $\mathbf{V}$ (now that we are really viewing the machine as a verifier, we have replaced $\mathbf{M}$ with $\mathbf{V}$) such that, for all $z \in \{0, 1\}^*$, all $r \in \{0, 1\}^{p(|z|)}$, and all strings $a_1, \dots, a_{2k} \in \{0, 1\}^*$, $\mathbf{V}$ halts upon input $(z, r, a_1, \dots, a_{2i})$ in time polynomial in $|z|$ for all $i = 0, \dots, k$. We then imagine a prover $P : \{0, 1\}^* \rightarrow \{0, 1\}^*$ interacting with $\mathbf{V}$ as follows. First, the verifier randomly selects $r \in \{0, 1\}^{p(|z|)}$ and computes $a_1 := \mathbf{V}(z, r)$; this is $\mathbf{V}$'s first *question* to P . P then responds with the *answer* $a_2 := P(z, a_1)$. (Note that P does not have access to the random string r ; one says that $\mathbf{V}$ is using *private coins*. It turns out that for what we are going to define below, one can also use *public coins* that the prover is aware of.) This constitutes the first *round* of their interaction. The verifier then asks P their second question $a_3 := \mathbf{V}(z, r, a_1, a_2)$ and P responds with $a_4 := P(z, a_1, a_2, a_3)$, completing the second round of interaction. This is repeated for a total of k rounds. $\mathbf{V}$ then returns their decision $\mathbf{V}(z, r, a_1, \dots, a_{2k}) \in \{0, 1\}$, indicating whether or not they accept the prover's answers as constituting evidence that z indeed belongs to $\mathbf{L}$.

The complexity class IP is defined to be the collection of those languages $\mathbf{L}$ for which there is a Turing machine $\mathbf{V}$ as above such that:

- If $z \in \mathbf{L}$, then there is a prover P such that the probability that a random string r causes $\mathbf{V}$ to accept is at least $\frac{2}{3}$.
- If $z \notin \mathbf{L}$, then no prover can cause $\mathbf{V}$ to accept more than $\frac{1}{3}$ of the time.

The probabilities $\frac{2}{3}$ and $\frac{1}{3}$ above are called the *completeness* and *soundness* parameters, respectively. As in the case of BPP, they are somewhat arbitrary; any choice of completeness parameter strictly larger than one's choice of soundness parameter will define the same class. It turns out that one can even replace the completeness parameter by 1 without changing the class; this property of IP is called *perfect completeness*.

A nice example of a language in IP is the collection of pairs (G_1, G_2) of finite graphs that are *not* isomorphic. Note that this class is not obviously in NP for there are too many possible functions that could serve as an isomorphism. There is however a simple interactive proof for this class. Indeed, the verifier randomly selects $i \in \{1, 2\}$ and then randomly selects a permutation σ on the number of vertices of G_i , obtaining a graph H isomorphic to G_i . The verifier then sends the graph H to the prover as its question. The prover then responds with a bit $a \in \{1, 2\}$, which represents its guess as to which of the two graphs G_1 or G_2 the verifier selected randomly. The verifier accepts if and only if the prover guessed the chosen graph correctly. Note that if $G_1 \not\cong G_2$ (that is, if the pair (G_1, G_2) belongs to the class), then the prover can always respond correctly, for the prover can just figure out whether or not H is isomorphic to G_1 or to G_2 . (Do not forget that the prover is all-powerful!) However, if $G_1 \cong G_2$ (that is, (G_1, G_2) does not belong to the class), then H is isomorphic to both G_1 and G_2 , and thus the prover (regardless of its unlimited power) can do no better than simply guessing which graph was chosen by the verifier, thus only convincing the verifier at most half of the time.

Given a verifier $\mathbf{V}$ as in the definition of IP, one can compute in $\text{poly}(|z|)$ -space the optimal prover strategy. This shows that $\text{IP} \subseteq \text{PSPACE}$. A landmark theorem in the subject shows that in fact we have equality:

Theorem 4.1 (Lund, Fortnow, Karloff, Nisan [23]; Shamir [55]). $\text{IP} = \text{PSPACE}$.

Recall that randomization alone likely does not achieve anything new (earlier we remarked that $\text{P} = \text{BPP}$ is likely), and, similarly, interaction alone does not achieve anything too new (as we just recover NP). However, combining randomization with interaction bumps us up to PSPACE (which is likely bigger than NP).

But why stop at one prover? One can consider interactions as above but allowing for *multiple* provers to interact with the verifier. It should be emphasized that the provers are not allowed to interact with each other during the interaction, but only with the verifier. They can, however, have a meeting before the interaction starts and decide upon a strategy that they will use while interacting with the verifier. In other words, the provers are cooperating but noncommunicating. If the provers use deterministic strategies as above, we arrive at the complexity class MIP. By allowing different kinds of strategies (in particular, those that employ quantum methods), we arrive at variations of MIP, such as the famous MIP^* appearing in the equation $\text{MIP}^* = \text{RE}$.

It turns out that the complexity class MIP is unchanged if one restricts to just two provers and one round of interaction. We thus make that default assumption from now on. By ignoring one of the provers, we clearly have that $\text{IP} \subseteq \text{MIP}$. As with IP, one can also achieve perfect completeness.

With two provers, one can now utilize “police-style” interrogation tactics. This makes it possible for the verifier to read polynomially many random portions of an exponentially long proof and come to a conclusion that with high probability agrees with the truth. A formalization of this idea yields another major theorem in the subject:

Theorem 4.2 (Babai, Fortnow, Lund [5]). $\text{MIP} = \text{NEXP}$.

As mentioned before, it is believed that $\text{PSPACE} \subsetneq \text{NEXP}$. If this inclusion is indeed proper, then the jump from one to more than one prover does in fact lead to a larger collection of languages.

4.4. Nonlocal games. It will be useful to recast our description of the class MIP in terms of so-called *nonlocal games*, a certain collection of two-person games. (The terminology “nonlocal” comes from the connection with Bell’s theorem on quantum nonlocality, as we discuss later.)

Consider a language $\mathbf{L}$ in MIP as witnessed by the polynomial-time verifier $\mathbf{V}$. Given input z and a sequence of random bits r , by computing $\mathbf{V}(z, r)$, we are really computing the two questions x and y (sequences of bits of length polynomial in $|z|$) that are being sent to the two provers, whom we will call Alice and Bob, following typical quantum information nomenclature. Alice and Bob, employing their deterministic strategies A and B , then respond with their answers, say $a := A(x)$ and $b := B(y)$, and then the prover calculates $\mathbf{V}(z, r, x, y, a, b)$ to decide whether or not to accept their answers. Whether or not z belongs to $\mathbf{L}$ then corresponds to the expected value over a randomly chosen r that the verifier returns $\mathbf{V}(z, r, x, y, a, b) = 1$. Note that the polynomial time requirement on $\mathbf{V}$ allows us to assume that the set of possible answers only contains bits that are of size at most some fixed polynomial in $|z|$.

This reformulation leads us to the following notion: A *nonlocal game with k questions and n answers* is a pair $\mathfrak{G} = (\pi, D)$, where π is a probability distribution on $[k] \times [k]$ and $D : [k] \times [k] \times [n] \times [n] \rightarrow \{0, 1\}$ is the decision predicate for the game. Here, $[k] := \{1, \dots, k\}$ and similarly for $[n]$. A *strategy* for the players consists of a conditional probability $p(a, b|x, y)$ expressing the probability that Alice and Bob respond with answers a and b if they are asked questions x and y , respectively. We view such a strategy p as an element of $[0, 1]^{k^2 n^2}$. Above, we only considered *deterministic strategies*, namely those p for which there are functions $A, B : [k] \rightarrow [n]$ such that $p(A(x), B(y)|x, y) = 1$ for all $x, y \in [k]$. We let $C_{\text{det}}(k, n) \subseteq [0, 1]^{k^2 n^2}$ denote the set of such deterministic strategies. Later, we will consider several other sets of strategies.

Given a strategy p , the *value of the game $\mathfrak{G}$ with respect to the strategy p* is the expected value the players will win $\mathfrak{G}$ if they play according to p , that is,

$$\text{val}(\mathfrak{G}, p) := \sum_{(x, y) \in [k] \times [k]} \pi(x, y) \sum_{(a, b) \in [n] \times [n]} D(x, y, a, b) p(a, b|x, y).$$

We set $\text{val}(\mathfrak{G}) := \sup_{p \in C_{\text{det}}(k, n)} \text{val}(\mathfrak{G}, p)$ and refer to this as the *classical value* of the game $\mathfrak{G}$.

We can now rephrase the definition of MIP in terms of nonlocal games: a language $\mathbf{L}$ belongs to MIP if and only if there is an *efficient mapping* $z \mapsto \mathfrak{G}_z$ (in the precise sense described earlier in this subsection) so that:

- If $z \in \mathbf{L}$, then $\text{val}(\mathfrak{G}_z) \geq \frac{2}{3}$.
- If $z \notin \mathbf{L}$, then $\text{val}(\mathfrak{G}_z) \leq \frac{1}{3}$.

The class MIP^* appearing in the result $\text{MIP}^* = \text{RE}$ is defined in the analogous way except that we replace classical strategies by quantum strategies. But first, an interlude to explain all things quantum.

5. A QUANTUM DETOUR

In this section, we introduce the quantum prerequisites necessary to understand the definition of the complexity class MIP^* . Our presentation of quantum mechanics is fairly standard and can be found in any good textbook on quantum mechanics. As mentioned above, we also found Paulsen's lecture notes [50] very helpful as well.

5.1. Quantum measurements. In quantum mechanics, one associates to each physical system a corresponding Hilbert space $\mathcal{H}$. The *state* of the system at any given time is given by a unit vector $\xi \in \mathcal{H}$. The state of the system evolves linearly according to a certain partial differential equation (the *Schrödinger equation*) **until it is measured**. A measurement should be thought of as an experiment on the system which has a finite number, say n , possible outcomes. (There are also experiments that can have a countably infinite set of outcomes, say the infinite discrete set of energies of some particle, or even a continuum of outcomes, say when measuring the position or momentum of a particle. For the purposes of this paper, it suffices to focus on the case of finitely many outcomes.) Formally, a measurement with n outcomes consists of n bounded operators $M_1, \dots, M_n \in \mathcal{B}(\mathcal{H})$. The *Born rule* states that, if the state of the system is ξ upon measurement, then the probability that the i th outcome happens is given by $\|M_i \xi\|^2$. Furthermore, in case the i th outcome is measured, the *collapse dynamics* tells us that the state of the system instantaneously (and discontinuously) changes to $M_i(\xi)/\|M_i(\xi)\|$. Since the sum of the outcome probabilities must be 1, we see that

$$1 = \sum_{i=1}^n \|M_i \xi\|^2 = \sum_{i=1}^n \langle M_i^* M_i \xi, \xi \rangle.$$

Since this equality must hold true for all unit vectors $\xi \in \mathcal{H}$, it follows that $\sum_{i=1}^n M_i^* M_i = I_{\mathcal{H}}$. Consequently, any sequence $M_1, \dots, M_n \in \mathcal{B}(\mathcal{H})$ satisfying this latter property constitutes a measurement of the system.

If one is only interested in the probabilities of the outcomes rather than the outcomes themselves (as we will be when we return to our discussion of nonlocal games), then it simplifies matters by replacing a measurement as above by a sequence $P_1, \dots, P_n$ consisting of positive operators which sum up to $I_{\mathcal{H}}$ and interpret the probability that the i th outcome is obtained when the system is in state ξ to be given by $\langle P_i \xi, \xi \rangle$. Such a collection of positive operators is called a *positive operator-valued measure* (POVM) on $\mathcal{H}$ (this terminology comes from spectral theory). If one specializes even further to the case that each P_i is not only a positive operator but in fact a projection, then one speaks of *projection-valued measures* (PVMs) on $\mathcal{H}$. Note then that the projections are automatically pairwise

orthogonal, so a PVM on $\mathcal{H}$ with n outcomes corresponds to a decomposition of $\mathcal{H}$ into n orthogonal subspaces.

Many introductions to quantum mechanics discuss the measurements of *observables*. An observable for the physical system is a self-adjoint operator $\mathcal{O}$ on $\mathcal{H}$. Supposing for simplicity that $\mathcal{H}$ is finite dimensional, the spectral theorem implies that we can find a PVM $P_1, \dots, P_n$ on $\mathcal{H}$ such that the P_i 's correspond to the projections onto the various eigenspaces of $\mathcal{H}$ corresponding to $\mathcal{O}$. The self-adjointness assumption on $\mathcal{O}$ further implies that the corresponding eigenvalues are real numbers, whence we can interpret them as corresponding to actual possible physical measurements. Conversely, given any PVM $P_1, \dots, P_n$ on $\mathcal{H}$ and real numbers $\lambda_1, \dots, \lambda_n$, one has an observable $\mathcal{O} := \sum_{i=1}^n \lambda_i P_i$.

A simple example of the content of the previous paragraph is given by the *spin* of an electron. The spin of an electron along any choice of axis comes in one of two flavors: *up* or *down*. (By the way, this is what is “quantum” about quantum mechanics: many attributes of a physical system come in a discrete set of possibilities.) For the sake of completeness, let us say that we are measuring spin along the vertical axis. The state of the electron is given by a unit vector ψ in the Hilbert space $\mathbb{C}^2$. We view the usual orthonormal basis $\{e_1, e_2\}$ for $\mathbb{C}^2$ as representing the two possible spin values: so e_1 corresponds to up while e_2 corresponds to down. Now a general unit vector ψ in $\mathbb{C}^2$ can be written in the form $\psi = \alpha_1 e_1 + \alpha_2 e_2$ for unique complex numbers $\alpha_1, \alpha_2 \in \mathbb{C}$ for which $|\alpha_1|^2 + |\alpha_2|^2 = 1$. What is strange and new about quantum mechanics is that a given electron can be in a state that is neither up nor down. More specifically, when neither α_1 nor α_2 are 0, the electron is considered to be in a *superposition* of the two states and will only reveal one of these two states upon a measurement of the spin, that is, using the PVM P_1, P_2 on $\mathbb{C}^2$ consisting of the projections onto the coordinate axes. The state of the electron merely gives us probabilistic information as to which of the two outcomes will happen upon such a measurement. Moreover, once the measurement has been made, the new state of the electron instantaneously and discontinuously jumps to the unit vector e_1 or e_2 corresponding to the outcome of the measurement just made. This reflects the fact that if another measurement is made directly following the first measurement, the same outcome will occur. It is important to make the distinction between superpositions and definite measurement outcome with probabilities measuring ignorance of the actual value.

The above description of quantum mechanics is the standard or *Copenhagen* interpretation, and it is a mighty big pill to swallow upon a first reading. (Technically speaking, this is really the *von Neumann–Dirac* formulation of the theory; however, it has become common parlance to refer to this interpretation as the Copenhagen interpretation, even though Niels Bohr himself explicitly disagreed with this formulation.) Perhaps the biggest point of contention are the questions, “What constitutes a measurement?” together with the follow-up question “Why did the state of the electron *collapse* to one of the two basis states?” This is the so-called *measurement problem* and is a very popular topic of debate amongst philosophers and theoretical physicists. It has led to a plethora of alternate interpretations of quantum mechanics (often yielding mathematically equivalent predictions); a good introduction to these foundational issues is Barrett’s recent book [6].

To keep the strangeness coming, suppose that we want to measure spin in the horizontal direction instead of the vertical direction. It turns out that the appropriate basis to consider is now $\{v_1, v_2\}$, where $v_1 = \frac{1}{\sqrt{2}}e_1 + \frac{1}{\sqrt{2}}e_2$ and $v_2 = \frac{1}{\sqrt{2}}e_1 - \frac{1}{\sqrt{2}}e_2$. In other words, the PVM Q_1, Q_2 consisting of the orthogonal projections onto the lines spanned by v_1 and v_2 , respectively, constitutes a measurement of the spin of the electron in the horizontal direction. Suppose that an electron has a definite spin, say up, in the vertical direction, whence its state is e_1 . In the eigenbasis for the observable of spin in the horizontal direction, the state becomes $e_1 = \frac{1}{\sqrt{2}}v_1 + \frac{1}{\sqrt{2}}v_2$. Consequently, a measurement of an electron with an up spin in the vertical direction will yield a spin of either left or right in the horizontal direction with equal probability. Even more strangely, suppose that the electron that had a definite vertical spin that was up was then measured in the horizontal direction and the outcome was spin left; that is, the measurement led to an outcome state of v_1 . Suppose further that a subsequent measurement of the electron in the vertical direction was performed. Since $v_1 = \frac{1}{\sqrt{2}}e_1 + \frac{1}{\sqrt{2}}e_2$, we see that the outcome of the measurement now yields up or down with equal probability. Thus, the measurement in the horizontal direction destroyed the definite spin the electron had in the vertical direction!

5.2. The spookiness of entanglement. The postulates of quantum mechanics tell us that if $\mathcal{H}_A$ and $\mathcal{H}_B$ are the Hilbert spaces representing two physical systems, then the appropriate Hilbert space for studying the composite system is the tensor product space $\mathcal{H}_A \otimes \mathcal{H}_B$. The fact that elements of the tensor product need not be merely simple tensors leads to the fascinating concept of *entanglement*, which, in some sense, is the essence of this entire story!

In order to get an idea of the utility of entanglement as a resource in, say, quantum information theory, we present the example of *superdense coding*. We set $\psi_{EPR} := \frac{1}{\sqrt{2}}(e_1 \otimes e_1 + e_2 \otimes e_2) \in \mathbb{C}^2 \otimes \mathbb{C}^2 \cong \mathbb{C}^4$. This quantum state is known as the *EPR state*, named after Einstein, Podolsky, and Rosen. Shortly, we will have more to say about this state and why Einstein, Podolsky, and Rosen were considering it. Let us imagine that Alice and Bob each possess an electron and the joint state of the vertical spins of the two electrons is ψ_{EPR} , that is, the electrons are in an equal superposition of both spins being up or both spins being down. Furthermore, imagine that Alice and Bob are really (really) far away from each other. We show how Alice and Bob can utilize the fact that their electrons are in this entangled state in order for Alice to send two classical bits of information to Bob by just sending one *qubit* of information, that is, by Alice sending Bob her electron (after she has first done some work on it).

Depending on what two bits of information Alice wishes to send to Bob, she performs one of the following actions to her electron:

- $\psi_{11} := (I \otimes I)\psi_{EPR} = \frac{1}{\sqrt{2}}(e_1 \otimes e_1 + e_2 \otimes e_2)$,
- $\psi_{12} := (X \otimes I)\psi_{EPR} = \frac{1}{\sqrt{2}}(e_2 \otimes e_1 + e_1 \otimes e_2)$,
- $\psi_{21} := (Z \otimes I)\psi_{EPR} = \frac{1}{\sqrt{2}}(e_1 \otimes e_1 - e_2 \otimes e_2)$,
- $\psi_{22} := (ZX \otimes I)\psi_{EPR} = \frac{1}{\sqrt{2}}(e_1 \otimes e_2 - e_2 \otimes e_1)$.

Here, $X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix}$, the so-called *bit-flip operator*, and $Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$, the so-called *phase-flip operator*.

One can check that the four vectors $\psi_{11}, \psi_{12}, \psi_{21}, \psi_{22}$ form an orthonormal basis for $\mathbb{C}^2 \otimes \mathbb{C}^2 \cong \mathbb{C}^4$ known as the *Bell basis*. Consequently, any observable $\mathcal{O}$ on $\mathbb{C}^4$ with distinct eigenvalues and with the Bell basis vectors as eigenvectors can be used to distinguish these vectors, that is, when the state of the system is ψ_{ij} , a measurement of $\mathcal{O}$ will yield ψ_{ij} with probability 1, whence Bob knows which of the four actions above Alice took and thus knows which pair of bits she wished to send to Bob. (One can be explicit about the observable $\mathcal{O}$, namely $\mathcal{O} = (H \otimes I_{\mathbb{C}^2})C$, where $H := \frac{1}{\sqrt{2}}(X + Z)$ is the so-called *Hadamard operator* and C is the so-called *controlled not operator*.)

Notice something peculiar about the EPR state? If the state of two electrons is given by ψ_{EPR} , then they are in a superposition of either both electrons having spin up or both electrons having spin down (with equal probability). However, if Alice performs a measurement of the spin of her electron and sees a result of spin up, she knows, with absolute certainty, that a subsequent measurement of the spin of Bob's electron will also be spin up. Thus, while Bob's electron did not have a determinate spin before Alice's measurement, the result of Alice's measurement *instantaneously* gave a determinate value to the spin of Bob's electron.

Einstein was worried by this phenomenon, which he called “spooky action at a distance.” Together with Podolsky and Rosen [17], they used the EPR state to present an argument for the *incompleteness* of quantum mechanics. The gist of the argument is as follows: Suppose that Alice and Bob share a pair of electrons in the EPR state ψ_{EPR} , and that Alice and Bob are again really (really) far apart. Suppose that Alice measures her electron and sees the result spin up. Then Alice knows with 100% certainty that if Bob were to measure his electron, then it must also have a determinately up spin. The same holds for a measurement result of spin down. Since Alice can predict with certainty the outcome of Bob's measurement, and since her measurement could not possibly have altered the spin of Bob's electron, Bob's spin must have a definite value, independent of whether or not Alice were to measure it. This definite value must represent some element of physical reality and if quantum mechanics were to be complete, there must be some counterpart of this physical reality in the theory. Since there is nothing in the description of the EPR state which specifies a determinate value for Bob's spin, quantum mechanics must be incomplete.

It gets even worse, for if Alice were to decide to measure her spin along a different axis, say the horizontal axis, then once again the result of her measurement would allow her to definitively conclude the value of Bob's electron's spin in the horizontal axis. In this case, *both* the vertical and horizontal spins would have definite, predetermined values, which is a contradiction of the fact that knowing, say, the vertical spin of an electron forces us to be maximally uncertain about the horizontal spin of the electron. So in some sense, the EPR argument even posits that quantum mechanics is inconsistent!

The underlying philosophy that Einstein, Podolsky, and Rosen have in their argument is usually dubbed *local realism*: the term “local” refers to the assumption that Alice's measurement could not have affected Bob's electron since they are so far away and communication cannot travel faster than the speed of light, while the term “realism” refers to the statement that the fact that one can determine the spin of Bob's electron with certainty implies that there must be some real, predetermined value to the spin. Einstein, Podolsky, and Rosen believed that there

should be some *hidden variable* explaining this predetermined spin, allowing them to preserve their classical, locally real intuitions.

John Bell [7] set up a thought experiment to determine whether there could indeed be a formulation of quantum mechanics that was complete and adhered to the local realist philosophy. He showed that this is in fact impossible by showing that a small set of local realist assumptions leads to an inequality on the expected outcome of a certain experiment and that a particular quantum measurement could violate that inequality. Moreover, it is actually experimentally testable whether or not this inequality holds in nature. Spoiler alert: the inequality is violated by nature, whence quantum mechanics comes out victorious! Thus, while seemingly strange, quantum mechanics lies in contradistinction to the local realist assumptions.

Besides being an intellectually fascinating story, there turns out to be a direct link between these Bell inequalities and the phenomena of having quantum strategies for nonlocal games that exceed all possible classical values, which we now explain. (The idea of treating the violation of Bell-type inequalities as quantum strategies for nonlocal games that exceed the classical value of the game seems to have first been seriously studied by Cleve, Hoyer, Toner, and Watrous [14]).

We have already discussed deterministic strategies for nonlocal games. One may imagine incorporating a probabilistic component to these strategies by considering a probability space (Ω, μ) and deterministic strategies $A_\omega : [k] \rightarrow [n]$ and $B_\omega : [k] \rightarrow [n]$, one for each $\omega \in \Omega$. Consequently, the players can randomly (according to (Ω, μ)) select an ω and then play deterministically according to A_ω and B_ω . In terms of the EPR experiment, one may think of ω as the hidden variable for which we do not have perfect knowledge but that if we were to know it, then things would behave deterministically. The probability space (Ω, μ) represents our epistemic (lack of) knowledge of the hidden variable. Consequently, we now have probabilistic strategies

$$p(a, b|x, y) := \mu(\{\omega \in \Omega : A_\omega(x) = a \text{ and } B_\omega(y) = b\}),$$

which are called *local strategies*, the term “local” referring to the fact that each player’s output still only depends on their local environment. The set of such local strategies is denoted $C_{\text{loc}}(k, n)$. It is straightforward to see that $C_{\text{loc}}(k, n)$ is a compact, convex subset of $[0, 1]^{k^2 n^2}$ whose extreme points are the elements in $C_{\text{det}}(k, n)$. Moreover, it is clear that every element of $C_{\text{loc}}(k, n)$ is a convex combination of elements of $C_{\text{det}}(k, n)$, whence $\text{val}(\mathfrak{G}) = \sup_{p \in C_{\text{loc}}(k, n)} \text{val}(\mathfrak{G}, p)$ for any nonlocal game $\mathfrak{G}$ with k questions and n answers.

On the other hand, we can consider quantum strategies for nonlocal games as follows. We let $C_q(k, n)$ consist of those strategies p for which

$$p(a, b|x, y) = \langle (A_a^x \otimes B_b^y) \psi, \psi \rangle,$$

where, for each $x, y \in [k]$, $A^x = (A_a^x)_{a \in [n]}$ and $B^y = (B_b^y)_{b \in [n]}$ are POVMs with n outcomes on *finite-dimensional* Hilbert spaces $\mathcal{H}_A$ and $\mathcal{H}_B$, respectively. We call such a strategy p a *quantum strategy*. These strategies correspond to Alice and Bob sharing a (possibly entangled) state ψ of their composite system $\mathcal{H}_A \otimes \mathcal{H}_B$ and performing measurements A^x and B^y on their portion of the state upon receiving questions x and y , respectively. Using a technique known as *Naimark dilation* (a special case of the Stinespring dilation theorem from above), one can replace POVMs with the more convenient to use PVMs without altering the definition of

$C_q(k, n)$. It is a straightforward argument to show that $C_q(k, n)$ is a convex subset of $[0, 1]^{k^2 n^2}$.

We have that $C_{\text{loc}}(k, n) \subseteq C_q(k, n)$. Indeed, since every element of $C_{\text{loc}}(k, n)$ is a convex combination of deterministic strategies and $C_q(k, n)$ is convex, it suffices to show that every deterministic strategy p is contained in $C_q(k, n)$. However, this is quite easy: if $A : [k] \rightarrow [n]$ is the function determining Alice's strategy, let A^x be the POVM on $\mathbb{C}$ for which $A_{A(x)}^x = I$ and $A_a^x = 0$ for all $a \neq A(x)$. Bob's POVM B_y^b is defined in the analogous way. It follows that, for any state $\xi \in \mathbb{C} \otimes \mathbb{C}$, we have that $p(a, b|x, y) = \langle (A_a^x \otimes B_b^y) \xi, \xi \rangle$.

Given a nonlocal game $\mathfrak{G}$, we define its *entangled value* to be

$$\text{val}^*(\mathfrak{G}) := \sup_{p \in C_q(k, n)} \text{val}(\mathfrak{G}, p).$$

By the previous paragraph, we have that $\text{val}(\mathfrak{G}) \leq \text{val}^*(\mathfrak{G})$ for any nonlocal game $\mathfrak{G}$. The idea behind Bell's theorem, recast in the setting of nonlocal games, is that there are nonlocal games $\mathfrak{G}$ for which $\text{val}(\mathfrak{G}) < \text{val}^*(\mathfrak{G})$.

For example, we consider the following game, known as the *CHSH game*. (The acronym CHSH stands for Clauser, Horne, Shimony, and Holt, the researchers responsible for the CHSH inequality, a Bell-type inequality that was one of the first to be experimentally testable.) The CHSH game $\mathfrak{G}_{\text{CHSH}}$ is a game with $k = n = 2$. The question distribution is the uniform distribution on $[2] \times [2]$ and with decision predicate D given by the following conditions.

- If $x = 1$ or $y = 1$, then Alice and Bob win if and only if their answers agree.
- If $x = y = 2$, then Alice and Bob win if and only if their answers disagree.

By inspecting all deterministic strategies, one finds that $\text{val}(\mathfrak{G}_{\text{CHSH}}) = \frac{3}{4}$. However, the entangled value of the game satisfies $\text{val}^*(\mathfrak{G}_{\text{CHSH}}) = \cos^2(\frac{\pi}{8}) \approx 0.85 > \text{val}(\mathfrak{G}_{\text{CHSH}})$. The interested reader can find the details for this calculation in [14, Section 3.1]. We merely point out that a quantum strategy for achieving $\text{val}^*(\mathfrak{G})$ uses the EPR state ψ_{EPR} .

5.3. MIP*. Based on the nonlocal game definition of the complexity class MIP and our recent discussion of quantum strategies for nonlocal games, it should be clear how to define the complexity class MIP*: the language $\mathbf{L}$ belongs to MIP* if there is an efficient mapping (in the precise sense from Subsection 4.4) $z \mapsto \mathfrak{G}_z$ from strings to nonlocal games such that:

- If $z \in \mathbf{L}$, then $\text{val}^*(\mathfrak{G}_z) \geq \frac{2}{3}$.
- If $z \notin \mathbf{L}$, then $\text{val}^*(\mathfrak{G}_z) \leq \frac{1}{3}$.

We remark that the definition of MIP* first appeared in the aforementioned paper [14].

To be fair, we are really defining the complexity class $\text{MIP}^*(2, 1)$, which only has two provers and one round of interactions. There are ways to define similar classes that allow more verifiers and rounds, but the eventual result $\text{MIP}^* = \text{RE}$ will show they yield the same class anyways, so we will not bother.

So how do the classes MIP and MIP* relate? The lesson from the previous section was that provers that share entanglement can win some nonlocal games more often than they *rightfully should*. In other words, it seems that it might be the case that for a language $\mathbf{L}$ that belongs to MIP and for a string z that does

not belong to $\mathbf{L}$, the provers might have a strategy for the corresponding game $\mathfrak{G}_z$ whose value exceeds $\frac{1}{3}$.

Nevertheless (and perhaps somewhat surprisingly), Ito and Vidick [38] showed that $\text{MIP} \subseteq \text{MIP}^*$. The rough idea behind this inclusion is that it suffices to show that $\text{NEXP} \subseteq \text{MIP}^*$, and the games involved in the proof that $\text{MIP} = \text{NEXP}$ are such that their classical and quantum values are approximately the same.

Later, Natarajan and Wright [46] showed that $\text{NEEXP} \subseteq \text{MIP}^*$. Recalling that $\text{MIP} = \text{NEXP} \subsetneq \text{NEEXP}$, this shows that $\text{MIP} \subsetneq \text{MIP}^*$, whence adding entanglement does indeed strictly increase the computational power of the verifier.

So exactly how much extra power does entanglement give us? Besides the result mentioned in the last paragraph, there was only an a priori seemingly silly upper bound on MIP^* , namely $\text{MIP}^* \subseteq \text{RE}$, where RE (which is short for *recursively enumerable*) is the complexity class which consists of those languages $\mathbf{L}$ for which there is a Turing machine $\mathbf{M}$ (with absolutely no efficiency requirements whatsoever) whose domain is $\mathbf{L}$; that is, $\mathbf{L}$ consists of the set of inputs for which $\mathbf{M}$ halts. An alternative formulation of RE is helpful: $\mathbf{L}$ belongs to RE if there is a total computable function whose range is $\mathbf{L}$ (this is why modern computability theorists refer to this as being *computably enumerable* or *CE*). To see the inclusion $\text{MIP}^* \subseteq \text{RE}$, note first that, given any dimension d , one can effectively enumerate a countable set of quantum strategies of dimension d that is dense in the set of such strategies and for which one can effectively compute $\text{val}(\mathfrak{G}, p)$ for any such quantum strategy p . By letting d tend to ∞ , if one ever finds such a strategy p for which $\text{val}(\mathfrak{G}_z, p) > \frac{1}{3}$, one knows that $z \in \mathbf{L}$ (and one is guaranteed that this will happen for some such p if $z \in \mathbf{L}$). Note that if $z \notin \mathbf{L}$, this procedure will never convince us that $z \in \mathbf{L}$ because maybe we did not wait long enough, and a higher-dimensional strategy would indeed have convinced us if we were just a bit more patient.

The amazing fact proven in [39] is that this upper bound is actually tight! That is, $\text{MIP}^* = \text{RE}$ holds! More specifically, the authors prove that there is an effective mapping $\mathbf{M} \mapsto \mathfrak{G}_{\mathbf{M}}$ from (codes for) Turing machines to nonlocal games such that:

- If $\mathbf{M}$ halts on the empty tape, then $\text{val}^*(\mathfrak{G}_{\mathbf{M}}) = 1$.
- If $\mathbf{M}$ does not halt on the empty tape, then $\text{val}^*(\mathfrak{G}_{\mathbf{M}}) \leq \frac{1}{2}$.

The language consisting of codes for Turing machines that halt on the empty tape is known as the *halting problem* **HALT**. Since the halting problem is complete for the class RE, the inclusion $\text{RE} \subseteq \text{MIP}^*$ holds.

Regardless of your interest in CEP, the equality $\text{MIP}^* = \text{RE}$ is an amazing fact. The halting problem is an undecidable problem (this follows from a simple diagonalization argument together with the fact that there is a so-called *universal Turing machine*). Nevertheless, if two cooperating but noncommunicating provers share some quantum entanglement, they can reliably convince a verifier whether or not a given Turing machine halts! This is a landmark intellectual achievement.

The proof of $\text{MIP}^* = \text{RE}$ is very complicated, and we will not discuss it here. The introduction to [39] does a great job outlining the essence of the proof.

The story of MIP^* is about allowing quantum resources but keeping the computational model classical. Further, it is interesting to ask what happens if we also replace the computational model we are using (i.e., the Turing machine) with a quantum computational model (e.g., quantum circuits). It turns out that there is nothing to be gained here: by prefixing the corresponding classical complexity class with a “Q” to denote its counterpart defined using a quantum computational

model, we have $\text{QIP} = \text{IP}$, $\text{QMIP} = \text{MIP}$, and $\text{QMIP}^* = \text{MIP}^* = \text{RE}$; see [62] for the details.

6. FROM $\text{MIP}^* = \text{RE}$ TO THE FAILURE OF CEP: THE TRADITIONAL ROUTE

The derivation of the negative solution to the CEP from $\text{MIP}^* = \text{RE}$ now proceeds in two steps: We first show how $\text{MIP}^* = \text{RE}$ leads to a negative solution to *Tsirelson's problem* in quantum information theory; we show this in the first subsection. In the second subsection, we then show how a negative solution to Tsirelson's problem naturally leads to a negative solution to Kirchberg's QWEP problem. As we already observed in Subsection 3.10, this leads to a negative solution to the CEP.

6.1. A negative solution to Tsirelson's problem. In order to explain Tsirelson's problem, we need to introduce some more collections of strategies. First, we define $C_{qs}(k, n)$ exactly as in the definition of $C_q(k, n)$, except that we remove the finite-dimensionality assumptions on Alice's and Bob's state spaces $\mathcal{H}_A$ and $\mathcal{H}_B$; a strategy in this larger class is called a *quantum spatial strategy*. Quantum spatial strategies still correspond to the idea that Alice and Bob each have their own physical system and the state of their composite system is given by the tensor product. It can be checked that there is no loss of generality in restricting our attention to separable Hilbert spaces in the definition of $C_{qs}(k, n)$. Moreover, by considering projections onto larger and larger finite-dimensional subspaces, we see that $C_{qs}(k, n) \subseteq \overline{C_q(k, n)}$, the closure of $C_q(k, n)$ in the usual topology it inherits from being a subset of $[0, 1]^{k^2 n^2}$. Like $C_q(k, n)$, one can check that $C_{qs}(k, n)$ is convex.

There is another model that is natural to consider which arises in quantum field theory. In quantum field theory, one usually considers a large quantum system (maybe the system describing the whole universe!) and then it may be difficult to separate Alice and Bob's systems as isolated subsystems of the larger system. The state of the large system is now given by some unit vector ξ in a single Hilbert space $\mathcal{H}$, and Alice's and Bob's measurements are now given by families of POVMs $(A^x)_{x \in [k]}$ and $(B^y)_{y \in [k]}$ acting on this single Hilbert space $\mathcal{H}$. Since we are still assuming that Alice and Bob are far away and so they cannot interact with each other, it is natural to assume that either of them can measure first without affecting the value of the other's measurements (or even that they can perform their measurements simultaneously). According to von Neumann, the mathematical way of modeling this situation is to assume that Alice's and Bob's measurements commute with one another, that is, $A_a^x B_b^y = B_b^y A_a^x$ for all $x, y \in [k]$ and all $a, b \in [n]$. The corresponding strategy is given by $p(a, b | x, y) = \langle A_a^x B_b^y \xi, \xi \rangle$ and is called a *quantum commuting strategy*. (Commutativity ensures that this a priori complex value lies in $[0, 1]$.) The set of quantum commuting strategies is denoted $C_{qc}(k, n)$. Note that there is no requirement that $\mathcal{H}$ be finite dimensional (although one can take it to be separable) and, in fact, requiring $\mathcal{H}$ to be finite dimensional yields another description of the set $C_q(k, n)$ (see [16]). Later, we will see that $C_{qc}(k, n)$ is a closed convex subset of $[0, 1]^{k^2 n^2}$ and that, like $C_q(k, n)$, one can use PVMs instead of POVMs without changing the definition.

It is clear that $C_{qs}(k, n) \subseteq C_{qc}(k, n)$. In [59], Boris Tsirelson claimed (without proof) that equality holds for all (k, n) . After he was questioned about this, he realized that he could not prove this claim. In fact, upon further reflection, he could not even establish whether or not $C_{qs}(k, n)$ was closed nor whether or not $C_{qa}(k, n) := \overline{C_{qs}(k, n)} = \overline{C_q(k, n)}$ coincided with $C_{qc}(k, n)$ (see his note [60]). The question of whether or not $C_{qa}(k, n) = C_{qc}(k, n)$ for all (k, n) is known as *Tsirelson's problem*. (Incidentally, in [56] Slofstra showed that Tsirelson's original claim was false, that is, there is a pair (k, n) such that $C_{qs}(k, n) \neq C_{qc}(k, n)$, and he even strengthened this result to show that $C_{qs}(k, n)$ need not be closed, that is, there is (k, n) for which $C_{qs}(k, n) \subsetneq C_{qa}(k, n)$.)

Fix a nonlocal game $\mathfrak{G}$ with k questions and n answers. It is clear that

$$\sup_{p \in C_{qa}(k, n)} \text{val}(\mathfrak{G}, p) = \sup_{p \in C_{qs}(k, n)} \text{val}(\mathfrak{G}, p) = \text{val}^*(\mathfrak{G}).$$

However, we can also use elements of $C_{qc}(k, n)$ to define values of games; namely, we define the *commuting value* of $\mathfrak{G}$ to be $\text{val}^{co}(\mathfrak{G}) := \sup_{p \in C_{qc}(k, n)} \text{val}(\mathfrak{G}, p)$. It is clear that $\text{val}^*(\mathfrak{G}) \leq \text{val}^{co}(\mathfrak{G})$ and that equality holds for all nonlocal games if Tsirelson's problem has an affirmative answer. In fact, it can be shown that an affirmative answer to Tsirelson's problem is equivalent to the statement $\text{val}^*(\mathfrak{G}) = \text{val}^{co}(\mathfrak{G})$ for all nonlocal games $\mathfrak{G}$.

Recall that in our discussion of the inclusion $\text{MIP}^* \subseteq \text{RE}$, we discussed how $\text{val}^*(\mathfrak{G})$ can be effectively approximated from below. On the other hand, it turns out that $\text{val}^{co}(\mathfrak{G})$ can be effectively approximated from above. This result follows from two facts:

- There is a finitely presented group $G_{\mathfrak{G}}$ (which in fact only depends on the number of questions and answers in $\mathfrak{G}$) and an element $\eta_{\mathfrak{G}} \in C^*(G_{\mathfrak{G}})$ such that $\text{val}^{co}(\mathfrak{G}) = \|\eta_{\mathfrak{G}}\|$ (see Corollary 6.3).
- For any finitely presented group G , one can always find effective upper bounds on the operator norm of $C^*(G)$ (a result due to Fritz, Netzer, and Thom [25, Corollary 2.2]).

In Subsection 7.8, we offer a simple model-theoretic proof of the fact that $\text{val}^{co}(\mathfrak{G})$ can be approximated from above, although, to be fair, we really establish a slightly different version of this fact sufficient to derive the failure of Tsirelson's problem from $\text{MIP}^* = \text{RE}$. In any event, if $\text{val}^*(\mathfrak{G}) = \text{val}^{co}(\mathfrak{G})$, that is, if Tsirelson's problem has an affirmative answer, then we can effectively approximate $\text{val}^*(\mathfrak{G}) = \text{val}^{co}(\mathfrak{G})$ both from below and above, which would then imply that all languages in MIP^* are decidable, contradicting $\text{MIP}^* = \text{RE}$!

By the way, the argument in the preceding paragraph shows that $\text{MIP}^{co} \subseteq \text{coRE}$, where MIP^{co} is defined exactly like MIP^* but using the commuting value val^{co} of games instead of the entangled value val^* , and coRE denotes the class of languages whose complement lies in RE . It is currently unknown if this upper bound is sharp.

6.2. A negative solution to Kirchberg's QWEP problem. In this subsection, we show how a negative solution to Tsirelson's problem yields a negative solution to Kirchberg's QWEP problem. We follow Fritz's presentation [24] closely.

We begin by considering the abelian C^* -algebra $\mathbb{C}^n$. For each $a = 1, \dots, n$, we let e_a denote the a th standard basis element of $\mathbb{C}^n$. (We are using a as the index since we are using the notation from nonlocal games.) For any $k \geq 1$, we also

consider the k -fold free product $\ast_{x=1}^k \mathbb{C}^n$ and denote by e_a^x the version of e_a in the x th copy of $\mathbb{C}^n$.

Proposition 6.1.

- (1) *There is a 1-1 correspondence between n -outcome POVMs $\{A_1, \dots, A_n\}$ in $\mathcal{B}(\mathcal{H})$ and ucp maps $\Phi : \mathbb{C}^n \rightarrow \mathcal{B}(\mathcal{H})$ given by $\Phi(e_a) := A_a$.*
- (2) *There is a 1-1 correspondence between k -tuples $\{A_1^x, \dots, A_n^x\}_{x=1}^k$ of n -outcome POVMs in $\mathcal{B}(\mathcal{H})$ and ucp maps $\Phi : \ast_{x=1}^k \mathbb{C}^n \rightarrow \mathcal{B}(\mathcal{H})$ given by $\Phi(e_a^x) := A_a^x$.*

Proof. The proof of (1) is easy to check, using that a positive map with commutative domain is automatically completely positive. Part (2) follows from (1) and a theorem of Florin Boca [11], which implies that the individual ucp maps $\Phi^x : \mathbb{C}^n \rightarrow \mathcal{B}(\mathcal{H})$ given by $\Phi^x(e_a^x) := A_a^x$ can be jointly extended to a single ucp map $\Phi : \ast_{x=1}^k \mathbb{C}^n \rightarrow \mathcal{B}(\mathcal{H})$. $\square$

We now bring group C^* -algebras into the picture, getting us closer to the QWEP problem. We first note that $\mathbb{C}^n \cong C^*(\mathbb{Z}_n)$, where $\mathbb{Z}_n$ denotes the additive group of integers modulo n . Indeed, let u be a generator of $\mathbb{Z}_n$ and consider the element $z := \sum_{a=1}^n \exp(\frac{2\pi i a}{n}) e_a \in \mathbb{C}^n$. It is readily verified that z is an element of $\mathcal{U}(\mathbb{C}^n)$ of order n , whence the assignment $u \mapsto z$ yields a unitary representation $\mathbb{Z}_n \rightarrow \mathcal{U}(\mathbb{C}^n)$, extending to a $*$ -homomorphism $C^*(\mathbb{Z}_n) \rightarrow \mathbb{C}^n$ that can be checked to be an isomorphism. (This identification usually goes under the name *discrete Fourier transform*.)

Let $\mathbb{F}(k, n) := \ast_{x=1}^k \mathbb{Z}_n$ denote the group freely generated by k elements of order n . We then have

$$C^*(\mathbb{F}(k, n)) = C^*(\ast_{x=1}^k \mathbb{Z}_n) \cong \ast_{x=1}^k C^*(\mathbb{Z}_n) \cong \ast_{x=1}^k \mathbb{C}^n.$$

We abuse notation slightly and let e_a^x denote the element of $C^*(\mathbb{F}(k, n))$ corresponding to $e_a \in \ast_{x=1}^k \mathbb{C}^n$. (Another viewpoint is that $(e_a^x)_{a=1}^n$ denote the spectral projections corresponding to the x th unitary element of $C^*(\mathbb{F}(k, n))$.)

Here is the main result connecting the QWEP problem and Tsirelson's problem:

Theorem 6.2. *Fix $k, n \geq 2$ and a strategy $p \in [0, 1]^{k^2 n^2}$. We then have:*

- (1) *$p \in C_{qa}(k, n)$ if and only if there is a state ϕ on $C^*(\mathbb{F}(n, k)) \otimes_{\min} C^*(\mathbb{F}(n, k))$ for which $p(a, b|x, y) = \phi(e_a^x \otimes e_b^y)$.*
- (2) *$p \in C_{qc}(k, n)$ if and only if there is a state ϕ on $C^*(\mathbb{F}(n, k)) \otimes_{\max} C^*(\mathbb{F}(n, k))$ for which $p(a, b|x, y) = \phi(e_a^x \otimes e_b^y)$.*

Proof. For the forward direction of (1), we may assume, without loss of generality, that $p \in C_{qs}(k, n)$, say $p(a, b|x, y) = \langle (A_a^x \otimes B_b^y) \xi, \xi \rangle$, where the POVMs A^x and B^y act on the Hilbert spaces $\mathcal{H}_A$ and $\mathcal{H}_B$, respectively. By Proposition 6.1 and the above identification $C^*(\mathbb{F}(k, n)) \cong \ast_{x=1}^k \mathbb{C}^n$, we have ucp maps $\Phi_A : C^*(\mathbb{F}(k, n)) \rightarrow \mathcal{B}(\mathcal{H}_A)$ and $\Phi_B : C^*(\mathbb{F}(k, n)) \rightarrow \mathcal{B}(\mathcal{H}_B)$ corresponding to these POVMs. These two ucp maps combine to yield a ucp map $\Phi = \Phi_A \otimes \Phi_B : C^*(\mathbb{F}(k, n)) \otimes_{\min} C^*(\mathbb{F}(k, n)) \rightarrow \mathcal{B}(\mathcal{H}_A \otimes \mathcal{H}_B)$. Consequently, we can define a state ϕ on $C^*(\mathbb{F}(k, n)) \otimes_{\min} C^*(\mathbb{F}(k, n))$ by setting $\phi(w \otimes z) := \langle (\Phi(w \otimes z) \xi, \xi) \rangle$. It is clear that this state ϕ implements p as in the statement of (1).

Conversely, suppose that ϕ is a state on $C^*(\mathbb{F}(n, k)) \otimes_{\min} C^*(\mathbb{F}(n, k))$ for which $p(a, b|x, y) = \phi(e_a^x \otimes e_b^y)$. Concretely represent $C^*(\mathbb{F}(k, n)) \subseteq \mathcal{B}(\mathcal{H})$ so that $C^*(\mathbb{F}(k, n)) \otimes_{\min} C^*(\mathbb{F}(k, n)) \subseteq \mathcal{B}(\mathcal{H} \otimes \mathcal{H})$. Extend ϕ to a state on $\mathcal{B}(\mathcal{H} \otimes \mathcal{H})$, which

we will continue to denote ϕ . Since convex combination of vector states are dense in the state space of $\mathcal{B}(\mathcal{H} \otimes \mathcal{H})$, given $\epsilon > 0$ there are vectors $\xi_1, \dots, \xi_n \in \mathcal{H} \otimes \mathcal{H}$ for which $|\phi(e_a^x \otimes e_b^y) - \sum_{i=1}^n \langle (e_a^x \otimes e_b^y) \xi_i, \xi_i \rangle| < \epsilon$ for all $x, y \in [k]$ and $a, b \in [n]$. This shows that p can be approximated by convex combinations of elements of $C_{qs}(k, n)$. Since $C_{qa}(k, n)$ is closed and convex, we have that $p \in C_{qa}(k, n)$.

The proof of the forward direction of (2) is identical to the proof of the forward direction of (1), using the fact that one can combine ucp maps with commuting ranges into a ucp map on the maximal tensor product. For the converse, suppose that ϕ is a state on $C^*(\mathbb{F}(n, k)) \otimes_{\max} C^*(\mathbb{F}(n, k))$ for which $p(a, b|x, y) = \phi(e_a^x \otimes e_b^y)$. Let $\pi_\phi : C^*(\mathbb{F}(n, k)) \otimes_{\max} C^*(\mathbb{F}(n, k)) \rightarrow \mathcal{B}(\mathcal{H})$ be the GNS representation corresponding to the state ϕ with cyclic vector ξ . Set $A_a^x := \pi_\phi(e_a^x \otimes I)$ and $B_b^y := \pi_\phi(I \otimes e_b^y)$. It is clear that A_a^x and B_b^y commute for all x, y, a, b and that $p(a, b|x, y) = \langle A_a^x B_b^y \xi, \xi \rangle$, whence $p \in C_{qc}(k, n)$. $\square$

We note that the proof above fulfills a few promises made earlier, namely that elements of $C_{qc}(k, n)$ can always be taken to arise from PVMs (instead of just POVMs) and that $C_{qc}(k, n)$ is closed and convex (being the continuous image of the compact convex set of states on $C^*(\mathbb{F}(k, n)) \otimes_{\max} C^*(\mathbb{F}(k, n))$).

Given a nonlocal game $\mathfrak{G} = (\pi, D)$ with k questions and n answers, set

$$\eta_{\mathfrak{G}} := \sum_{x, y \in [k]} \pi(x, y) \sum_{a, b \in [n]} D(x, y, a, b) (e_a^x \otimes e_b^y) \in C^*(\mathbb{F}(k, n)) \odot C^*(\mathbb{F}(k, n)).$$

Corollary 6.3. *For any nonlocal game $\mathfrak{G}$, we have $\text{val}^*(\mathfrak{G}) = \|\eta_{\mathfrak{G}}\|_{\min}$ and $\text{val}^{\text{co}}(\mathfrak{G}) = \|\eta_{\mathfrak{G}}\|_{\max}$.*

We remind the reader that Corollary 6.3 is responsible for the negative solution to Tsirelson's problem from $\text{MIP}^* = \text{RE}$. Indeed, $\|\eta_{\mathfrak{G}}\|_{\max}$ corresponds to the operator norm of $\eta_{\mathfrak{G}}$ when viewed as an element of $C^*(\mathbb{F}(k, n) \times \mathbb{F}(k, n))$; by [25, Corollary 2.2], one can effectively compute upper bounds of the operator norm of elements of $C^*(\mathbb{F}(k, n) \times \mathbb{F}(k, n))$, whence one can effectively compute upper bounds for $\text{val}^{\text{co}}(\mathfrak{G})$.

Corollary 6.4. *For any $k, n \geq 2$, if $(C^*(\mathbb{F}(k, n)), C^*(\mathbb{F}(k, n)))$ is a nuclear pair, then Tsirelson's problem has a positive solution for scenarios of dimension (k, n) ; that is, $\text{val}^*(\mathfrak{G}) = \text{val}^{\text{co}}(\mathfrak{G})$ for all nonlocal games $\mathfrak{G}$ with k questions and n answers.*

The quotient map $\mathbb{Z} \rightarrow \mathbb{Z}_n$ yields a quotient map $\mathbb{F}_k \rightarrow \mathbb{F}(k, n)$, leading to a surjective $*$ -homomorphism $C^*(\mathbb{F}_k) \rightarrow C^*(\mathbb{F}(k, n))$. One can show that this map has a ucp lift $C^*(\mathbb{F}(k, n)) \rightarrow C^*(\mathbb{F}_k)$ (see, for example, [24, Lemma D.3]), whence $(C^*(\mathbb{F}(k, n)), C^*(\mathbb{F}(k, n)))$ is a nuclear pair if $(C^*(\mathbb{F}_k), C^*(\mathbb{F}_k))$ is a nuclear pair. (It can be shown that $\mathbb{F}(k, n)$ contains a copy of $\mathbb{F}_2$ if $(k, n) \neq (2, 2)$, whence the converse to the previous sentence also holds in this case.) Consequently, we have the following.

Corollary 6.5. *If Kirchberg's QWEP problem has a positive answer, then Tsirelson's problem has a positive answer.*

In the last subsection, we saw that $\text{MIP}^* = \text{RE}$ implies that Tsirelson's problem has a negative solution. Combined with Corollary 6.5, we now have that the QWEP problem has a negative solution, and thus, coupled with the discussion in Subsection 3.10, we finally have the desired negative solution to the CEP!

7. FROM $\text{MIP}^* = \text{RE}$ TO THE FAILURE OF CEP: A MODEL-THEORETIC SHORTCUT

In this section, we show how ideas from logic yield (in this author's opinion) a more elementary derivation of a negative solution of CEP from $\text{MIP}^* = \text{RE}$. Much of the material presented in this section represents joint work of the author and Bradd Hart [31] and [32].

7.1. A continuous logic for studying tracial von Neumann algebras. The negative solution to CEP from $\text{MIP}^* = \text{RE}$ presented in this section uses techniques from logic. Consequently, we need to describe an appropriate first-order language in a certain *continuous logic* for studying tracial von Neumann algebras. (We apologize for the double use of the word “language” in this paper. The complexity-theoretic languages have been denoted using bold letters $\mathbf{L}$; we will use Roman letters L for languages in the sense of logic.)

For a von Neumann algebra $\mathcal{M}$, we let $\mathcal{M}_1$ denote the operator norm unit ball. Recall that by a $*$ -polynomial $p(x_1, \dots, x_n)$ in the indeterminates $x_1, \dots, x_n$ we mean an expression built from the indeterminates using the $*$ -algebra operations. Let $\mathcal{F}$ denote the set of all $*$ -polynomials $p(x_1, \dots, x_n)$ ($n \geq 0$) such that, for *any* von Neumann algebra $\mathcal{M}$, we have $p(\mathcal{M}_1^n) \subseteq \mathcal{M}_1$. For example, the following functions belong to $\mathcal{F}$:

- the *constant symbols* 0 and 1 (thought of as 0-ary functions),
- $x \mapsto x^*$,
- $x \mapsto \lambda x$ ($|\lambda| \leq 1$),
- $(x, y) \mapsto xy$,
- $(x, y) \mapsto \frac{x+y}{2}$.

We then work in the formal *language* $L_{vNa} := \mathcal{F} \cup \{\text{tr}_{\mathbb{R}}, \text{tr}_{\mathbb{S}}, d\}$, where $\text{tr}_{\mathbb{R}}$ (resp., $\text{tr}_{\mathbb{S}}$) denotes the real (resp., imaginary) parts of the trace and d denotes the metric on the operator norm unit ball given by $d(x, y) := \|x - y\|_{\tau}$. We can then formulate certain properties of tracial von Neumann algebras using the language L_{vNa} as follows.

Basic L_{vNa} -formulae will be formulae of the form $\text{tr}_{\mathbb{R}}(p(\vec{x}))$ or $\text{tr}_{\mathbb{S}}(p(\vec{x}))$ for $p \in \mathcal{F}$. *Quantifier-free L_{vNa} -formulae* are formulae of the form $f(\varphi_1, \dots, \varphi_m)$, where $f : \mathbb{R}^m \rightarrow \mathbb{R}$ is a continuous function and $\varphi_1, \dots, \varphi_m$ are basic L_{vNa} -formulae. Finally, an arbitrary L_{vNa} -formula is of the form

$$Q_{x_{i_1}}^1 \cdots Q_{x_{i_k}}^k \varphi(x_1, \dots, x_n),$$

where each $i_j \in \{1, \dots, n\}$, $\varphi(x_1, \dots, x_n)$ is a quantifier-free L_{vNa} -formula, and each Q^i is either sup or inf. We think of these Q_i 's as *quantifiers over the unit ball of the algebra*.

For those keeping score at home, our setup here is a bit more specialized than the general treatment of continuous logic in [8] (or even the version [20] presented for operator algebraists), but a dense set of the formulae in [8] are logically equivalent to formulae in the above form, so there is no loss of generality in our treatment here.

Also, in order to keep the set of formulae *separable* and *computable*, when forming the set of quantifier-free formulae, we really should restrict ourselves to a computable dense subset of the set of all continuous functions $\mathbb{R}^m \rightarrow \mathbb{R}$ as m ranges over $\mathbb{N}$. (See [31, Section 2].)

Suppose that $\varphi(\vec{x})$ is a formula, $\mathcal{M}$ is a tracial von Neumann algebra, and $\vec{a} \in \mathcal{M}_1^n$, where n is the length of the tuple $\vec{x}$. We let $\varphi(\vec{a})^{\mathcal{M}}$ denote the real number obtained by replacing the variables $\vec{x}$ with the tuple $\vec{a}$; we may think of $\varphi(\vec{a})^{\mathcal{M}}$ as the truth value of $\varphi(\vec{x})$ in $\mathcal{M}$ when $\vec{x}$ is replaced by $\vec{a}$. For example, if $\varphi(x_1)$ is the formula $\sup_{x_2} d(x_1 x_2, x_2 x_1)$, then $\varphi(a)^{\mathcal{M}} = 0$ if and only if a is in the center of $\mathcal{M}$.

If φ has no *free variables* (that is, all variables occurring in φ are bounded by some quantifier), then we say that φ is a *sentence*, and we observe that $\varphi^{\mathcal{M}}$ is a real number. Given a tracial von Neumann algebra, the *theory* of $\mathcal{M}$ is the function $\text{Th}(\mathcal{M})$ which maps the sentence φ to the real number $\varphi^{\mathcal{M}}$. Sometimes authors define $\text{Th}(\mathcal{M})$ to consist of the set of sentences φ for which $\varphi^{\mathcal{M}} = 0$; since $\text{Th}(\mathcal{M})$, as we have defined it, is determined by its zeroset, these two formulations are equivalent.

If $\varphi(\vec{x})$ is a formula, then there is a bounded interval $[m_\varphi, M_\varphi] \subseteq \mathbb{R}$ called the range of φ such that, for any tracial von Neumann algebra $\mathcal{M}$ and any $\vec{a} \in \mathcal{M}_1^n$, we have $\varphi(\vec{a})^{\mathcal{M}} \in [m_\varphi, M_\varphi]$.

At this point we need to mention an important if not seemingly pedantic point (to a nonlogician). We have been focusing our attention on those structures in the language L_{vNa} that actually correspond to (unit balls of) tracial von Neumann algebras. This is a perfectly legitimate thing to do because the class of tracial von Neumann algebras form an *elementary class*. Perhaps a simpler example from classical logic will help illustrate the point. Let $L_{grp} = \{\cdot, e\}$ consist of a single binary function symbol $\cdot$ and constant symbol e . Of course, the intended L_{grp} -structures are the ones that interpret these symbols as the multiplication and identity of a group. However, there are perfectly reasonable, if not silly, L_{grp} -structures, such as, for example, one that interprets $\cdot$ as a constant function. The key point is that we can write down a collection of axioms, that is, a set of L_{grp} -sentences T_{grp} , that single out the class of groups in the sense that an L_{grp} -structure G is a group if and only if every sentence in T_{grp} is true in G . This is the definition of what it means for the class of groups to be an elementary class in the language L_{grp} .

A similar situation is true in our context, namely, there is a collection T_{vNa} of L_{vNa} -sentences such that an L_{vNa} -structure is the unit ball of a von Neumann algebra if and only if each sentence in T_{vNa} evaluates to 0 in the structure. In fact, one can add to these axioms a couple of extra sentences in order to obtain the theory T_{II_1} whose models are all (unit balls of) II_1 factors.

7.2. The model-theoretic reformulation of CEP. An L_{vNa} -sentence σ of the form

$$\sup_{x_1} \cdots \sup_{x_n} \varphi(x_1, \dots, x_n)$$

is called *universal* if φ is quantifier-free and the range of φ is nonnegative and similarly *existential* if all the quantifiers are $\inf$. This terminology is justified if one thinks of the value 0 as *true* for then $\sigma^{\mathcal{M}} = 0$ if and only if $\varphi(a_1, \dots, a_n) = 0$ for all $a_1, \dots, a_n \in \mathcal{M}_1$. If we restrict the function $\text{Th}(\mathcal{M})$ to the set of all universal (resp., existential) sentences, the resulting function is defined to be the *universal* (resp., *existential*) theory of $\mathcal{M}$, denoted $\text{Th}_\forall(\mathcal{M})$ (resp., $\text{Th}_\exists(\mathcal{M})$).

It is fairly easy to see that $\text{Th}_\forall(\mathcal{M}) = \text{Th}_\forall(\mathcal{M}^{\mathcal{U}})$ for any tracial von Neumann algebra $\mathcal{M}$ and any ultrafilter $\mathcal{U}$. (The *Los theorem* [20, Proposition 4.3] shows that $\text{Th}(\mathcal{M}) = \text{Th}(\mathcal{M}^{\mathcal{U}})$, but we will not need this more general fact.) Consequently, if $\mathcal{N}$ embeds into $\mathcal{M}^{\mathcal{U}}$, then we have that $\text{Th}_\forall(\mathcal{N}) \leq \text{Th}_\forall(\mathcal{M})$ (as a function). It turns

out that the converse is also true. To see this, assume that $\text{Th}_V(\mathcal{N}) \leq \text{Th}_V(\mathcal{M})$. To show that $\mathcal{N}$ embeds into an ultrapower of $\mathcal{M}$, it suffices (by standard ultrapower arguments) to show, given any finitely many $a_1, \dots, a_n \in \mathcal{N}_1$, any atomic formula $\varphi(x_1, \dots, x_n)$, and any $\epsilon > 0$, that there are $b_1, \dots, b_n \in \mathcal{M}_1$ such that $|\varphi(\vec{a})^{\mathcal{N}} - \varphi(\vec{b})^{\mathcal{M}}| < \epsilon$. Set $r := \varphi(\vec{a})^{\mathcal{N}}$ and $\sigma := \inf_{\vec{x}} |r - \varphi(\vec{x})|$. It is clear that $\sigma^{\mathcal{N}} = 0$. The assumption that $\text{Th}_V(\mathcal{N}) \leq \text{Th}_V(\mathcal{M})$ implies that $\text{Th}_{\exists}(\mathcal{M}) \leq \text{Th}_{\exists}(\mathcal{N})$, whence $\sigma^{\mathcal{M}} = 0$, which easily implies the existence of the desired tuple $\vec{b} \in \mathcal{M}_1$. (None of this is particular to the case of tracial von Neumann algebras and holds for any pair of structures in the same language.)

We can thus reformulate the CEP as follows. For every tracial von Neumann algebra $\mathcal{M}$, we have that $\text{Th}_V(\mathcal{M}) \leq \text{Th}_V(\mathcal{R})$. Recalling that $\mathcal{R}$ embeds into every II_1 factor, we can further reformulate the CEP—there is a unique universal theory of II_1 factors, namely $\text{Th}_V(\mathcal{R})$.

7.3. The completeness theorem for (continuous) first-order logic. Before discussing the completeness theorem for II_1 factors in the context of continuous logic, we consider the simpler example of groups in classical logic.

Consider the following theorem in group theory: every group has a unique identity element. This theorem can be written as an L_{grp} -sentence σ_e defined by $\forall x(\forall y(x \cdot y = y \cdot x = x) \rightarrow x = e)$. What exactly does it mean for σ_e to be a theorem of group theory?

Syntactically, what this means is that in some formal proof system for first-order logic, there is a formal proof of σ_e from T_{grp} , denoted $T_{grp} \vdash \sigma_e$. This formal proof is simply a finite list of sentences, each of which is either an element of T_{grp} or can be obtained from earlier elements of the list using the rules of the proof system, with σ_e being the last element of the list.

Semantically, we might say that in every model of T_{grp} , that is, in every group, the sentence σ_e is true. We denote this relationship by $T_{grp} \models \sigma_e$.

For any reasonable proof system, it is fairly easy to prove that $\vdash$ implies $\models$; this is called the *soundness theorem for first order logic*. A much less obvious result is that the converse also holds, namely that any time $T_{grp} \models \sigma$, then in fact $T_{grp} \vdash \sigma$. (There is obviously nothing special here about T_{grp} and this works for any classical first-order theory.) This is called the *completeness theorem for first-order logic* and is due to Einstein's pal Kurt Gödel. (See [18, Section 2.5] for a nice treatment.)

The relevance of the completeness theorem is that if we use some effective coding of the symbols of L_{grp} and the logical symbols, then we can start a computer program running all proofs from T_{grp} and outputting all theorems of T_{grp} . In other words, the language (in the sense of complexity theory) consisting of codes for theorems of group theory belongs to RE. Note that all that was used about T_{grp} is that the set of codes for axioms in T_{grp} itself belongs to RE.

There are corresponding soundness and completeness theorems for continuous logic due to Ben Yaacov and Pedersen [10]. Due to the approximate nature of continuous logic, the completeness theorem takes a slightly different form. Restricted to our case of interest, namely T_{II_1} , it reads: for every L_{vNa} -sentence σ , we have

$$\sup\{\sigma^{\mathcal{M}} : \mathcal{M} \text{ a } \text{II}_1 \text{ factor}\} = \inf\{r \in \mathbb{Q}^{>0} : T_{\text{II}_1} \vdash \sigma \dot{-} r\}.$$

Here, $\dot{-}$ is the function given by $r \dot{-} s := \max(r - s, 0)$. Consequently, the completeness theorem tells us that the largest truth value that σ could take in a II_1 factor is the smallest upper bound for σ that you could prove from the axioms T_{II_1} .

The proof system for continuous logic is still of the form that you can effectively enumerate theorems from an effectively enumerated set of axioms. In particular, since the set of axioms for T_{II_1} (given in [20]) is easily checked to be effectively enumerated, we see that the set of theorems of T_{II_1} belongs to RE.

7.4. CEP and the computability of the universal theory of $\mathcal{R}$. Given a tracial von Neumann algebra $\mathcal{M}$, we say that the universal theory of $\mathcal{M}$ is *computable* if there is an algorithm such that, upon input an L_{vNa} -sentence σ and a rational $\epsilon > 0$, returns $a, b \in \mathbb{Q}^{>0}$ with $a < b$ and $b - a < \epsilon$ and for which $\sigma^{\mathcal{M}} \in (a, b)$.

The ideas in Subsection 7.3 allow us to prove the following.

Theorem 7.1 (Goldbring and Hart [31]). *If CEP has a positive answer, then the universal theory of $\mathcal{R}$ is computable.*

Proof. If CEP holds, then, recalling that $\mathcal{R}$ embeds into any II_1 factor, we have, for any universal sentence σ , that $\sup\{\sigma^{\mathcal{M}} : \mathcal{M} \text{ a } II_1 \text{ factor}\} = \sigma^{\mathcal{R}}$. Consequently, if we start enumerating all proofs from T_{II_1} and record all instances of theorems of the form $\sigma \div r$, then we know that $\sigma^{\mathcal{R}} \leq r$ and this allows us to effectively enumerate better and better upper bounds for $\sigma^{\mathcal{R}}$.

On the other hand, we can also effectively enumerate better lower bounds for $\sigma^{\mathcal{R}}$. There are two ways that one can go about this. One way is to write $\sigma = \sup_x \varphi(x)$, where $\varphi(x) = f(\tau(p_1(x)), \dots, \tau(p_n(x)))$, and each p_i is a $*$ -polynomial and f is a *computable* continuous function, that is, generated from a computable set of connectives. One can then approximately calculate φ on matrices of larger dimensions with rational coordinates. (Technically, x is restricted to range over matrices of operator norm at most one, but one can efficiently verify this too.) Another option is to consider the existential sentence $\sigma_0 := M_\sigma \div \sigma$, where M_σ is an upper bound for σ (uniform over all II_1 factors) that is effectively computable from σ itself. Since $\mathcal{R}$ embeds in every II_1 factor, we once again have $\sup\{\sigma_0^{\mathcal{M}} : \mathcal{M} \text{ a } II_1 \text{ factor}\} = \sigma_0^{\mathcal{R}}$ (this does not use CEP), and thus we can enumerate all proofs from T_{II_1} and every time we see that $\sigma_0 \div r$, we know that $\sigma^{\mathcal{R}} \geq M_\sigma - r$.

We run both the upper and lower bound algorithms simultaneously and wait until they output numbers within ϵ of each other. $\square$

7.5. Synchronous strategies, definable sets, and finishing the proof. Based on Theorem 7.1, in order to refute CEP, it suffices to prove Theorem 7.2:

Theorem 7.2 (Goldbring and Hart [32]). *The universal theory of $\mathcal{R}$ is not computable.*

We will use $MIP^* = RE$ to prove this result. But how? Given a nonlocal game $\mathfrak{G}$, the definition of $\text{val}^*(\mathfrak{G})$ resembles a universal sentence in the language L_{vNa} except that it is not a priori clear how to view the set of correlations C_{qa} as something that we can quantify over a II_1 factor.

Thankfully, a specific subset of C_{qa} can be characterized by a formula that our logic can handle. A strategy p is called *synchronous* if $p(a, b|x, x) = 0$ for all $x \in [k]$ and distinct $a, b \in [n]$. That is, p is synchronous if, whenever both players are asked the same question, they always answer with the same answer. We let $C_{qa}^s(k, n)$ (resp., $C_{qc}^s(k, n)$) denote the synchronous elements of $C_{qa}(k, n)$ (resp., $C_{qc}(k, n)$). Given a nonlocal game $\mathfrak{G}$ with k questions and n answers, we set its

synchronous entangled value to be

$$\text{sval}^*(\mathfrak{G}) := \sup_{p \in C_{qa}^s(k, n)} \text{val}(\mathfrak{G}, p).$$

One defines the *synchronous commuting value* $\text{sval}^{co}(\mathfrak{G})$ in the obvious way. In general, we have that $\text{sval}^*(\mathfrak{G}) \leq \text{val}^*(\mathfrak{G})$ and $\text{sval}^{co}(\mathfrak{G}) \leq \text{val}^{co}(\mathfrak{G})$.

Paulsen et al. [51, Corollary 5.6] showed that $p \in C_{qc}^s(k, n)$ if and only if there is a tracial state τ on $C^*(\mathbb{F}(k, n))$ such that $p(a, b|x, y) = \tau(e_a^x e_b^y)$. Contrast this result with Theorem 6.2: gone is the maximal tensor product of $C^*(\mathbb{F}(k, n))$ with itself, but instead the state is required to be a tracial state. Later, Kim, Paulsen, and Schaufhauser [43, Theorem 3.6] showed that $p \in C_{qa}^s(k, n)$ if and only if there is an *amenable* tracial state τ on $C^*(\mathbb{F}(k, n))$ such that $p(a, b|x, y) = \tau(e_a^x e_b^y)$. One definition of an amenable tracial state τ on a C^* -algebra $\mathcal{A}$ is that there is a $*$ -homomorphism $\theta : \mathcal{A} \rightarrow \mathcal{R}^u$ with a ucp lift $\mathcal{A} \rightarrow \ell^\infty(R)$ for which $\tau = \tau_{\mathcal{R}}^u \circ \theta$, where $\tau_{\mathcal{R}}$ is the unique trace on $\mathcal{R}$. There are many alternate characterizations of being an amenable trace showing that this is indeed a robust condition. In fact, Kirchberg used the notion of amenable trace in his proof that, for finite von Neumann algebras, being QWEP is equivalent to being isomorphic to a $*$ -subalgebra of $\mathcal{R}^u$.

In any event, the above characterization of $C_{qa}^s(k, n)$ can be used to show that $p \in C_{qa}^s(k, n)$ if and only if there are n -outcome PVMs $(f^x)_{x \in [k]}$ in $\mathcal{R}^u$ such that $p(a, b|x, y) = \tau_{\mathcal{R}}^u(f_a^x f_b^y)$. In what follows, we let X_n denote the set of PVMs in $\mathcal{R}^u$ of length n . By the previous paragraph, we have

$$\text{sval}^*(\mathfrak{G}) = \sup_{f^1, \dots, f^k \in X_n} \left(\sum_{x, y \in [k]} \pi(x, y) \sum_{a, b \in [n]} D(x, y, a, b) \tau(f_a^x f_b^y) \right)^{\mathcal{R}^u}.$$

We now note two very important facts:

- (1) The value contained in the parentheses is a legitimate first-order formula evaluated in the ultrapower $\mathcal{R}^u$ of $\mathcal{R}$.
- (2) The *proof* of $\text{MIP}^* = \text{RE}$ actually shows that the reduction $\mathbf{M} \mapsto \mathfrak{G}_{\mathbf{M}}$ from Turing machines to nonlocal games is such that if $\mathbf{M}$ halts, then $\text{sval}^*(\mathfrak{G}_{\mathbf{M}}) = 1$.

From these two facts, it seems like the proof of Theorem 7.2 is complete, for if we could approximately compute the value of any universal sentence in $\mathcal{R}$, then we could approximate $\text{sval}^*(\mathfrak{G}_{\mathbf{M}})$ for any Turing machine $\mathbf{M}$ and thus be able to decide the halting problem!

There is one (not so minor) issue: the supremum in the above display is still not technically allowable in our logic! Indeed, we are taking the supremum over elements from a certain set X_n rather than just tuples from the unit ball. However, it turns out that the founders of continuous logic thought long and hard about such suprema and their efforts will pay off tremendously.

To explain this, we return to classical logic for one moment and the case of groups. Given any group G , its center $Z(G)$ can be defined by the formula $\gamma(x) := \forall y(xy = yx)$, that is, $Z(G) = \{g \in G : \gamma(g) \text{ is true in } G\}$. Consequently, given any formula $\theta(x)$, the formula $(\forall x \in Z(G))\theta(x)$ represents an actual sentence in classical logic, being shorthand for the more cumbersome $\forall x(\gamma(x) \rightarrow \theta(x))$.

We are faced with a similar situation in the above paragraph. The elements of X_n are those that make the formula

$$\max \left(\max_{i=1, \dots, n} d(p_i, p_i^*), \max_{i=1, \dots, n} d(p_i, p_i^2), d \left(\sum_{i=1}^n p_i, 1 \right) \right)$$

equal to 0. One would hope that we could thus take the supremum over this set of elements as a shorthand for a more complicated *legitimate* formula. Unfortunately, such a move is not always possible.

More generally, given L_{vNa} -formulae $\theta(x)$ and $\psi(x)$, the expression

$$\sup \{ \psi(x)^{\mathcal{R}} : \theta(x)^{\mathcal{R}} = 0 \}$$

is only equivalent to $\sigma^{\mathcal{R}}$ for an actual L_{vNa} -sentence σ if $\theta(x)$ satisfies a certain *almost-near* property; that is, for each $\epsilon > 0$, there is a $\delta > 0$ such that, for all $a \in \mathcal{R}_1$, if $\theta(a)^{\mathcal{R}} < \delta$, then there is $b \in \mathcal{R}_1$ with $\theta(b)^{\mathcal{R}} = 0$ and $d(a, b) < \epsilon$. In the operator algebraic literature, this is usually referred to as a *weak stability* phenomena. In the model theory literature, this is called being a *definable set*. (See the author's paper [28] for more on definability in continuous logic and its connection to operator-algebraic matters.)

Now here is the fantastic (and fortuitious) part: the formula defining X_n above does have this property! And here's the kicker: Kim, Paulsen, and Schaufhauser themselves proved it [43, Lemma 3.5] while establishing their above characterization of $C_{qa}^s(k, n)$. Thus, we are entitled to write the above *formula* for $\text{sval}^*(\mathfrak{G})$ as a shorthand for a legitimate sentence in the language of continuous logic. There is a little bit of fine print to check, namely that the resulting sentence is in fact universal and that this transformation can be done effectively, but the details can indeed be carried out. This completes the proof of Theorem 7.2 and thus the model-theoretic proof of the negative solution to CEP.

7.6. A Gödelian refutation of the CEP. Gödel's incompleteness theorem is one of the landmark intellectual achievements of the 20th century. It addresses a seemingly simple question: is there an algorithm such that, upon input a sentence in the language of number theory, returns the truth value of the sentence in the natural numbers? Surprisingly, Gödel proved that the answer is no [26]! (See also [18, Section 3.5] for a more modern treatment.)

Gödel actually proved something much stronger, namely he proved that any attempt to answer the previous question by giving an effective axiomatization of number theory is doomed to fail. More specifically, there is a natural (and effective) collection of axioms known as *Peano arithmetic* such that any effective extension T of Peano arithmetic is destined to be incomplete, meaning there will be sentence σ that is true in $\mathbb{N}$ but not provable from T . Since simply listing all true sentences of $\mathbb{N}$ as axioms is obviously complete, it follows that the set of true sentences is not effectively enumerable.

We can use the ideas in the previous subsection to give a Gödelian-style refutation of CEP. Indeed, one can view CEP as the question of asking whether or not the effective list of axioms for being a II_1 factor is enough to axiomatize the universal theory of $\mathcal{R}$. The failure of CEP shows that the answer to this is no. But perhaps there is a stronger, but still effective, list of axioms extending the axioms for being a II_1 factor so that any model of these axioms would then satisfy the conclusion of CEP. Our proof from the previous subsection shows that the answer is still no:

Theorem 7.3 (Goldbring and Hart [32]). *There does not exist an effectively enumerable list T of axioms extending T_{II_1} such that all models of T are embeddable in $\mathcal{R}^u$.*

In particular, this shows that the collection of axioms $\sigma \div r$ for which $\sigma^{\mathcal{R}} \leq r$ is not effectively enumerable.

Theorem 7.3 allows us to provide *many* counterexamples to CEP:

Corollary 7.4. *There is a sequence $\mathcal{M}_1, \mathcal{M}_2, \dots$, of separable II_1 factors, none of which embed into an ultrapower of $\mathcal{R}$, and such that, for all $i < j$, $\mathcal{M}_i$ does not embed into an ultrapower of $\mathcal{M}_j$.*

Proof. We construct the sequence inductively. Set $\mathcal{M}_1$ to be any separable II_1 factor that does not embed into an ultrapower of $\mathcal{R}$. Suppose now that $\mathcal{M}_1, \dots, \mathcal{M}_n$ have been constructed satisfying the conclusion of Corollary 7.4. For each $i = 1, \dots, n$, let σ_i be a nonnegative sentence such that $\sigma_i^{\mathcal{R}} = 0$ but $\sigma_i^{\mathcal{M}_i} > 0$. For each $i = 1, \dots, n$, fix a rational number $\delta_i \in (0, \sigma_i^{\mathcal{M}_i})$. Let T be the theory of II_1 factors together with the single condition $\max_{i=1, \dots, n} (\sigma_i \div \delta_i) = 0$. It is clear that T is an effectively enumerable subset of the theory of $\mathcal{R}$. Thus, by Theorem 7.3, there is a separable model $\mathcal{M}_{n+1}$ of T such that $\mathcal{M}_{n+1}$ does not embed into an ultrapower of $\mathcal{R}$. Given $i = 1, \dots, n$, since $\sigma_i^{\mathcal{M}_i} > \delta_i$ while $\sigma_i^{\mathcal{M}_{n+1}} \leq \delta_i$, it follows that $\mathcal{M}_i$ does not embed into an ultrapower of $\mathcal{M}_{n+1}$. This indicates how to continue the recursive construction, completing the proof. $\square$

7.7. The universal theory of $\mathcal{R}$ and the moment approximation problem.

In this subsection, we offer a purely operator-algebraic reformulation of the statement that $\text{Th}_{\forall}(\mathcal{R})$ is not computable, first proven in [32].

Given positive integers n and d , we fix variables $x_1, \dots, x_n$ and enumerate all $*$ -monomials in the variables $x_1, \dots, x_n$ of total degree at most d as $m_1, \dots, m_L$. (Of course, $L = L(n, d)$ depends on both n and d .) We consider the map $\mu_{n,d} : \mathcal{R}_1^n \rightarrow \mathbb{D}^L$ given by $\mu_{n,d}(\vec{a}) = (\tau(m_i(\vec{a}))) : i = 1, \dots, L$. (Here, $\mathbb{D}$ is the complex unit disk.)

Let $X(n, d)$ denote the range of $\mu_{n,d}$, and let $X(n, d, p)$ be the image of the unit ball of $M_p(\mathbb{C})$ under $\mu_{n,d}$. Notice that $\bigcup_{p \in \mathbb{N}} X(n, d, p)$ is dense in $X(n, d)$.

Theorem 7.5. *The following statements are equivalent:*

- (1) *The universal theory of $\mathcal{R}$ is computable.*
- (2) *There is a computable function $F : \mathbb{N}^3 \rightarrow \mathbb{N}$ such that, for every $n, d, k \in \mathbb{N}$, $X(n, d, F(n, d, k))$ is $\frac{1}{k}$ -dense in $X(n, d)$.*

Proof. First suppose that the universal theory of $\mathcal{R}$ is computable. We produce a computable function F as in (2). Fix n, d , and k , and set $\epsilon := \frac{1}{3k}$. Computably find $s_1, \dots, s_t$, an ϵ -net in $\mathbb{D}^L$. For each $i = 1, \dots, t$, ask the universal theory of $\mathcal{R}$ to compute intervals (a_i, b_i) with $b_i - a_i < \epsilon$ and with $(\inf_{\vec{x}} |\mu_{n,d}(\vec{x}) - s_i|)^{\mathcal{R}} \in (a_i, b_i)$. For each $i = 1, \dots, t$ such that $b_i < 2\epsilon$, let $p_i \in \mathbb{N}$ be the minimal p such that when you ask the universal theory of $M_p(\mathbb{C})$ to compute intervals of shrinking radius containing $(\inf_{\vec{x}} |\mu_{n,d}(\vec{x}) - s_i|)^{M_p(\mathbb{C})}$, there is a computation that returns an interval (c_i, d_i) with $d_i < 2\epsilon$. Let p be the maximum of these p_i 's. We claim that setting $F(n, d, k) := p$ is as desired. Indeed, suppose that $s \in X(n, d)$ and take $i = 1, \dots, t$ such that $|s - s_i| < \epsilon$. Then $(\inf_{\vec{x}} |\mu_{n,d}(\vec{x}) - s_i|)^{\mathcal{R}} < \epsilon$, whence $b_i < 2\epsilon$. It follows

that there is an interval (c_i, d_i) as above with $(\inf_{\vec{x}} |\mu_{n,d}(\vec{x}) - s_i|)^{M_p(\mathbb{C})} < d_i < 2\epsilon$. Let $a \in M_p(\mathbb{C})$ realize the infimum. Then $|\mu_{n,d}(\vec{a}) - s| < 3\epsilon = \frac{1}{k}$, as desired.

Now suppose that F is as in (2). We show that the universal theory of $\mathcal{R}$ is computable. Toward this end, fix a universal sentence

$$\sigma = \sup_{\vec{x}} f(\tau(m_1), \dots, \tau(m_\ell)),$$

where $\vec{x} = x_1, \dots, x_n, m_1, \dots, m_\ell$ are $*$ -monomials in $\vec{x}$ of total degree at most d , and f is a *computable* connective. Fix also rational $\epsilon > 0$. We show how to compute the value of $\sigma^{\mathcal{R}}$ to within ϵ . Since f is computable, it has a *computable modulus of continuity* δ , meaning that we can find $k \in \mathbb{N}$ computably so that $\frac{1}{k} \leq \delta(\epsilon)$. Set $p = F(n, d, 2k)$. Computably construct a sequence $\vec{a}_1, \dots, \vec{a}_t \in (M_p(\mathbb{C})_1)^n$ that is a $\frac{1}{2k}$ -cover of $(M_p(\mathbb{C})_1)^n$ (with respect to the ℓ^1 metric corresponding to the 2-norm). Consequently, $\mu_{n,d}(\vec{a}_1), \dots, \mu_{n,d}(\vec{a}_t)$ is a $\frac{1}{2k}$ -cover of $X(n, d, p)$. Set

$$r := \max_{i=1, \dots, t} f(\tau(m_1(\vec{a}_i)), \dots, \tau(m_\ell(\vec{a}_i))).$$

By assumption, $X(n, d, p)$ is $\frac{1}{2k}$ -dense in $X(n, d)$. It follows that $r \leq \sigma^{\mathcal{R}} \leq r + \epsilon$, as desired. $\square$

7.8. A negative solution to Tsirelson's problem from $\text{MIP}^* = \text{RE}$. In this subsection, we offer an alternative proof of the negative solution to Tsirelson's problem using $\text{MIP}^* = \text{RE}$ and the completeness theorem; we follow closely the treatment given in [32]. As in the definition of the synchronous entangled value $\text{sval}^*(\mathfrak{G})$ of a nonlocal game $\mathfrak{G}$ with k questions and n answers, we define its *synchronous commuting value* to be $\text{sval}^{co}(\mathfrak{G}) := \sup_{p \in C_{qc}^s(k, n)} \text{val}(\mathfrak{G}, p)$.

Definition 7.6. Fix $0 < r \leq 1$. We define $\text{MIP}_{0,r}^{co,s}$ to be the set of those languages $\mathbf{L}$ for which there is an efficient mapping $z \mapsto \mathfrak{G}_z$ from strings to nonlocal games such that $z \in \mathbf{L}$ if and only if $\text{sval}^{co}(\mathfrak{G}_z) \geq r$.

Theorem 7.7 (Goldbring and Hart [32]). *For any $0 < r \leq 1$, every language in $\text{MIP}_{0,r}^{co,s}$ belongs to the complexity class coRE .*

In other words, if $\mathbf{L} \in \text{MIP}_{0,r}^{co,s}$, then there is an algorithm which enumerates the complement of $\mathbf{L}$.

For the remainder of this subsection, we work in the first-order language $L_{\tau C^*}$ for tracial C^* -algebras, that is, C^* -algebras equipped with a distinguished tracial state, which is defined in a manner analogous to the language L_{vNa} used to study tracial von Neumann algebras. Fix a nonlocal game $\mathfrak{G}$ with k questions and n answers. Let $w = (w_{x,a})_{x \in [k], a \in [n]}$ denote a tuple of variables, and let $\psi_{\mathfrak{G}}(w)$ be the $L_{\tau C^*}$ -formula

$$\sum_{(x,y) \in [k] \times [k]} \pi(x, y) \sum_{(a,b) \in [n] \times [n]} D(x, y, a, b) \tau(w_{x,a} w_{y,b}),$$

which informally calculates the expected value of winning when playing according to a strategy from C_{qc}^s (recalling the Paulsen et al. characterization of elements of C_{qc}^s). We then let $\theta_{\mathfrak{G},r}$ be the $L_{\tau C^*}$ -sentence

$$\inf_w \max \left(\max_{x,a} (\|w_{x,a}^2 - w_{x,a}\|, \max_{x,a} \|w_{x,a}^* - w_{x,a}\|, \max_x \|\sum_i w_{x,a} - 1\|, r \div \psi_{\mathfrak{G}}(w)) \right),$$

which informally calculates $\text{sval}^{co}(\mathfrak{G})$.

Let $T_{\tau C^*}$ be the $L_{\tau C^*}$ -theory of tracial C^* -algebras. The following is immediate.

Proposition 7.8. *For any nonlocal game $\mathfrak{G}$, we have $\text{sval}^{co}(\mathfrak{G}) \geq r$ if and only if the theory $T_{\tau C^*} \cup \{\theta_{\mathfrak{G},r} = 0\}$ is satisfiable, that is, has a model.*

We will also need the following immediate consequence of the completeness theorem.

Lemma 7.9. *Let U be a continuous theory. Then U is satisfiable if and only if $U \not\vdash \perp$.¹*

We can now prove Theorem 7.7. Let $\mathbf{L}$ belong to $\text{MIP}_{0,r}^{co,s}$ and consider a string $z \notin \mathbf{L}$, with corresponding game $\mathfrak{G}_z$. By Proposition 7.8 and Lemma 7.9, we have that $T_{\tau C^*} \cup \{\theta_{\mathfrak{G}_z,r} = 0\} \vdash \perp$. Since this latter condition is recursively enumerable, the proof of Theorem 7.7 is complete.

One can now deduce the failure of Tsirelson's problem from $\text{MIP}^* = \text{RE}$ as follows. Suppose, towards a contradiction, that $C_{qa}^s(k, n) = C_{qc}^s(k, n)$ for every k and n . Let $\mathbf{M} \mapsto \mathfrak{G}_{\mathbf{M}}$ be the efficient mapping from Turing machines to nonlocal games provided by $\text{MIP}^* = \text{RE}$. Given a Turing machine $\mathbf{M}$, one simultaneously starts computing lower bounds on $\text{val}^*(\mathfrak{G}_{\mathbf{M}})$ while running proofs from $T_{\tau C^*} \cup \{\theta_{\mathfrak{G}_{\mathbf{M}},1} = 0\}$. Since we are assuming that $C_{qa}^s(k, n) = C_{qc}^s(k, n)$ (where k and n are the number of questions and answers of $\mathfrak{G}_z$), we have that either the first computation eventually yields the fact that $\text{val}^*(\mathfrak{G}_{\mathbf{M}}) > \frac{1}{2}$, in which case $\mathbf{M}$ halts, or else the second computation eventually yields the fact that $T_{\tau C^*} \cup \{\theta_{\mathfrak{G}_{\mathbf{M}},1} = 0\} \vdash \perp$, in which case $\text{sval}^*(\mathfrak{G}_{\mathbf{M}}) < 1$, and $\mathbf{M}$ does not halt. In this way, we can decide the halting problem, a contradiction. Note that we derived the a priori stronger statement that $C_{qa}^s(k, n) \neq C_{qc}^s(k, n)$ for some k and n .

8. THE ENFORCEABLE II_1 FACTOR (SHOULD IT EXIST)

In this section, we describe a model-theoretic weakening of CEP, namely the statement that the enforceable II_1 factor exists. In order to explain this statement and its connection to CEP, we first need to introduce a certain two-player game.

8.1. A different kind of game. We introduce a method for building tracial von Neumann algebras first introduced in [29] for an arbitrary structure in continuous logic (based on the discrete case presented in Hodges' book [37]). This method goes under many names, such as *Henkin constructions*, *model-theoretic forcing*, or *building models by games*.

We fix a countably infinite set C of distinct symbols that are to represent generators of a separable tracial von Neumann algebra that two players (traditionally named $\forall$ and $\exists$) are going to build together (albeit adversarially). The two players take turns playing finite sets Σ of expressions of the form $\|p(c)\|_{\tau} - r| < \epsilon$, where c is a tuple of variables from C , $p(x)$ is a $*$ -polynomial, and each player's move is required to extend (that is, contain) the previous player's move. These sets are called (open) *conditions*. The game begins with $\forall$'s move. Moreover, these conditions are required to be *satisfiable*, meaning that there should be some tracial von Neumann algebra $\mathcal{M}$ and some tuple a from $\mathcal{M}_1$ such that $\|p(a)\|_{\tau} - r| < \epsilon$ for each such expression in the condition. We play this game for countably many rounds. At the end of this game, we have enumerated some countable, satisfiable

¹ $\perp$ represents a contradiction, i.e., any continuous sentence which cannot evaluate to 0. For instance, the constant function 1.

set of expressions. Provided that the players address a *dense* set of moments infinitely often, they can ensure that the play is *definitive*, meaning that the final set of expressions yields complete information about all $*$ -polynomials over the variables C (that is, for each $*$ -polynomial $p(x)$ and each tuple c from C , there should be a unique r such that the play of the game implies that $\|p(c)\|_\tau = r$) and that this data describes a countable, dense $*$ -subalgebra of a unique tracial von Neumann algebra, which is called the *compiled structure*. In what follows, we assume all plays of the game are definitive.

8.2. Enforceable properties of tracial von Neumann algebras. Crucial to the connection between the above games and the CEP is the notion of an enforceable property:

Definition 8.1. Given a property P of tracial von Neumann algebras, we say that P is an *enforceable* property if there a strategy for $\exists$ so that, regardless of player $\forall$'s moves, if $\exists$ follows the strategy, then the compiled structure will have property P .

Perhaps being an enforceable property seems so severe that there are in fact no enforceable properties. We will soon see that many interesting properties are in fact enforceable. First, we mention the *conjunction lemma* [29, Lemma 2.4]: if P_n is an enforceable property for each $n \in \mathbb{N}$, then so is the conjunction $\bigwedge_n P_n$.

As a first example of an enforceable property of tracial von Neumann algebras, we show that being a factor is enforceable. To see this, let $\theta(x)$ be the L_{vNa} -formula $\sqrt{\|x\|_\tau^2 - \tau(x)^2}$, and let $\eta(x)$ be the L_{vNa} -formula $\sup_y \|xy - yx\|_\tau$. Finally, let σ be the L_{vNa} -sentence $\sup_x (\theta(x) \div \eta(x))$. It was shown in [20] that a von Neumann algebra $\mathcal{M}$ is a factor if and only if $\sigma^\mathcal{M} = 0$. To see that being a factor is enforceable, by the conjunction lemma, it suffices to show that, given any $n \in \mathbb{N}$ and rational $\epsilon > 0$, the expression $(\theta(c_n) \div \eta(c_n)) < \epsilon$ is enforceable. To see that this is the case, suppose that player $\forall$ opened the game with the open condition Σ . Without loss of generality, we may suppose that c_n appears in Σ . Since Σ is satisfiable in some tracial von Neumann algebra, it is also satisfiable in some II_1 factor $\mathcal{M}$ (as every tracial von Neumann algebra embeds in a II_1 factor). Consequently, since $\sigma^\mathcal{M} = 0$, we see that $\mathcal{M}$ also witnesses that $\Sigma \cup \{\theta(c_n) \div \eta(c_n) < \epsilon\}$ is a condition, whence $\exists$ can respond with this condition, as desired.

We next show that being a II_1 factor is enforceable. Since being a factor is enforceable, it suffices to show that it is enforceable that, in the compiled structure, there is a projection of irrational trace (say $\frac{1}{\pi}$). By the conjunction lemma again, it suffices to show that, for any rational $\epsilon > 0$, there is some $n \in \mathbb{N}$ for which $\max(d(c_n, c_n^*), d(c_n, c_n^2), |\tau(c_n) - \frac{1}{\pi}|) < \epsilon$ is enforceable. Indeed, if this is enforceable, then since *almost* projections are near actual projections, there will be actual projections in the compiled structure whose trace approaches $\frac{1}{\pi}$, whence there will be an actual projection of trace $\frac{1}{\pi}$, as desired. However, this condition is clearly enforceable by the exact same argument used in the previous paragraph, this time, using a *fresh* constant, that is, some c_n which did not appear in player $\forall$'s opening play Σ .

One can go even further and show that being a *McDuff* II_1 factor is enforceable. A II_1 $\mathcal{M}$ factor is McDuff if $\mathcal{M} \bar{\otimes} \mathcal{R} \cong \mathcal{M}$. For example, $\mathcal{R}$ is McDuff. An alternate formulation for being McDuff will prove useful: $\mathcal{M}$ is McDuff if and only if there is a copy of $M_2(\mathbb{C})$ inside of $\mathcal{M}' \cap \mathcal{M}^u$. This amounts to showing that for any finite

$\mathcal{F} \subseteq \mathcal{M}$ and any rational $\epsilon > 0$, there are matrix units $(e_{ij})_{i,j=1,2}$ for $M_2(\mathbb{C})$ for which $\|[x, e_{ij}]\|_\tau < \epsilon$ for all $x \in \mathcal{F}$ and all $i, j = 1, 2$. Hopefully by now the strategy is apparent: given any open play Σ for player $\forall$, we realize Σ in some tracial von Neumann algebra $\mathcal{M}$. We then note that Σ is also realized in $\mathcal{M} \bar{\otimes} M_2(\mathbb{C})$ and then choose fresh constants $c_{n_1}, \dots, c_{n_4}$ and say that they are almost matrix units for $M_2(\mathbb{C})$ which almost commute with $c_1, \dots, c_n$. As we let n increase and ϵ decrease and using the fact that almost matrix units are near actual matrix units, the result follows using the conjunction lemma.

8.3. Existentially closed tracial von Neumann algebras (and yet another reformulation of CEP). One can push the line of reasoning in the previous subsection much further. First, it is helpful to introduce the notion of an *existentially closed (e.c.) tracial von Neumann algebra*. A tracial von Neumann algebra $\mathcal{M}$ is e.c. if: whenever $\mathcal{M} \subseteq \mathcal{N}$, there is an embedding $\mathcal{N} \hookrightarrow \mathcal{M}^{\mathcal{U}}$ into some ultrapower of $\mathcal{M}$ that restricts to the diagonal embedding of $\mathcal{M}$ into its ultrapower. If $\mathcal{N}$ is separable (whence so is $\mathcal{M}$), then this is equivalent to Definition 8.1 where we can use any nonprincipal ultrapower of $\mathcal{M}$. This version of the definition is the semantic version. Syntactically, $\mathcal{M}$ is e.c. if for any existential formula $\varphi(x)$ (where x is a finite tuple of variables), any $a \in \mathcal{M}_1$, and any tracial von Neumann algebra $\mathcal{N}$ containing $\mathcal{M}$, we have $\varphi(a)^{\mathcal{M}} = \varphi(a)^{\mathcal{N}}$. In other words, any phenomena that *could happen* in an extension of $\mathcal{M}$ approximately also happens in $\mathcal{M}$. It is for this reason that one should think of an e.c. tracial von Neumann algebra as being an analogue of an algebraically closed field.

Existentially closed tracial von Neumann algebras appear in abundance. Indeed, any tracial von Neumann algebra embeds into an e.c. one of the same density character. Moreover, we know many properties of an e.c. tracial von Neumann algebra: they must be McDuff II_1 factors, all of their automorphisms must be approximately inner, etc. . . . The reader interested in learning more about e.c. tracial von Neumann algebras can consult [19], [28], and [33].

But can we name a concrete e.c. tracial von Neumann algebra? Well:

Theorem 8.2 (Farah, Goldbring, Hart, and Sherman [19]). *$\mathcal{R}$ is an e.c. tracial von Neumann algebra if and only if CEP has a positive solution.*

Proof. We first note that if $\mathcal{R}$ is e.c., then CEP holds: given a II_1 factor $\mathcal{M}$, we have that $\mathcal{R} \subseteq \mathcal{M}$, whence, since $\mathcal{R}$ is e.c., we have that $\mathcal{M}$ embeds into $\mathcal{R}^{\mathcal{U}}$. Conversely, suppose that CEP holds; we show that $\mathcal{R}$ is e.c. To see this, suppose that $\mathcal{R} \subseteq \mathcal{M}$ with $\mathcal{M}$ separable. By CEP, $\mathcal{M}$ embeds into $\mathcal{R}^{\mathcal{U}}$. At the moment, this does not imply that $\mathcal{R}$ is e.c. as the composed embedding $\mathcal{R} \hookrightarrow \mathcal{R}^{\mathcal{U}}$ need not be the diagonal embedding. However, a nontrivial result of Kenley Jung [40] implies that every embedding $\pi : \mathcal{R} \hookrightarrow \mathcal{R}^{\mathcal{U}}$ is unitarily conjugate to the diagonal embedding, meaning that there is a unitary element $u \in \mathcal{R}^{\mathcal{U}}$ such that $\pi(a) = uau^*$ for all $a \in \mathcal{R}$ (viewing $\mathcal{R}$ as literally a subalgebra of $\mathcal{R}^{\mathcal{U}}$ via the diagonal embedding). It is straightforward to check that this finishes the job. $\square$

Following [21], we call a tracial von Neumann algebra $\mathcal{M}$ *locally universal* if every tracial von Neumann algebra embeds into an ultrapower of $\mathcal{M}$. In this terminology, CEP asks if $\mathcal{R}$ is locally universal. The proof of Theorem 8.2 shows the following.

Theorem 8.3. *Every e.c. tracial von Neumann algebra is locally universal. In particular, locally universal tracial von Neumann algebras exist.*

The latter conclusion was first reached (using a different argument) in [21, Example 6.4] and was referred to as a resolution to the *Poor Man's resolution to the Connes embedding problem*. Note also that CEP holds if and only if any locally universal tracial von Neumann algebra embeds in $\mathcal{R}^u$.

Another important fact for us is the following (see [29, Proposition 2.10] for a proof):

Theorem 8.4. *Being an e.c. tracial von Neumann algebra is an enforceable property.*

The proof of Theorem 8.4 is a more elaborate version of the arguments given in the last section.

One might ask: is there some first-order way of axiomatizing the e.c. tracial von Neumann algebras? The answer is no, a result first proven by Hart, Sinclair, and the author in [33] although we now know of some more elementary proofs (see [28, Corollary 5.19] for example).

8.4. CEP and enforceability. As we have seen in the previous subsections, while enforceability of a property seemed like it should rarely happen, we actually know of many interesting properties that are in fact enforceable. We now consider a real extreme version of this:

Definition 8.5. A tracial von Neumann algebra $\mathcal{M}$ is said to be *enforceable* if the property of being isomorphic to $\mathcal{M}$ is an enforceable property.

Clearly, if an enforceable tracial von Neumann algebra exists, then it is unique. Enforceable structures do exist in many other contexts. For example, the enforceable graph is the *random* or *Rado* graph and the enforceable field of a particular characteristic is the algebraic closure of the prime field. (Note, however, that the enforceable group does not exist; while somewhat implicit in [37], this is made explicit in [34].) On the analytic side, we have that the enforceable metric space is the *Urysohn space* [61], the unique Hilbert space of dimension $\aleph_0$ is the enforceable Hilbert space [8, Section 15], and the enforceable Banach space is the *Gurarij Banach space* [9].

So what about the enforceable tracial von Neumann algebra? Here is the connection to CEP:

Theorem 8.6 (Goldbring [29]). *The following statements are equivalent:*

- (1) *CEP has a positive solution.*
- (2) *The property of being hyperfinite is enforceable.*
- (3) *$\mathcal{R}$ is the enforceable tracial von Neumann algebra.*
- (4) *The property of being embeddable in $\mathcal{R}^u$ is enforceable.*

Proof. (1) implies (2): By CEP, every open condition is satisfied in $\mathcal{R}^u$ and hence in $\mathcal{R}$. Thus, given any n and rational $\epsilon > 0$, if player $\forall$ opens with Σ , then Σ is satisfied in $\mathcal{R}$ and thus $c_1, \dots, c_n$ are all within ϵ in $\|\cdot\|_\tau$ of some finite linear combination of approximate matrix units for some sufficiently large matrix algebra. Thus, player $\exists$ can respond with this extension of Σ . Now apply the conjunction lemma.

(2) implies (3): If being hyperfinite is enforceable, then since being a II_1 factor is also enforceable, we see by the conjunction lemma that being a hyperfinite II_1 factor is enforceable, whence $\mathcal{R}$ itself is enforceable.

(3) implies (4) is trivial: For (4) implies (1), if being embeddable in $\mathcal{R}^u$ is enforceable, then since being e.c. is also enforceable, we see that there is an e.c. tracial von Neumann algebra that embeds in $\mathcal{R}^u$. Since this e.c. tracial von Neumann algebra is necessarily locally universal, by the observation made in the previous subsection, we have that CEP has a positive solution. $\square$

Now that we know that CEP has a negative solution, we see that no e.c. tracial von Neumann algebra embeds in $\mathcal{R}^u$. Since being e.c. is enforceable, we see that the situation is pretty dire: it is enforceable that the compiled structure does not embed in $\mathcal{R}^u$, which should be seen as a *generic* negative solution to the CEP.

8.5. Properties of the enforceable II_1 factor (again, should it exist). Now that we know that CEP has a negative solution, we know that $\mathcal{R}$ is not enforceable. But there is still the possibility that the enforceable tracial von Neumann algebra $\mathcal{E}$ (which must necessarily be a II_1 factor) exists. It is this author's humble opinion that the existence of the enforceable II_1 factor is one of the most interesting open problems in the model theory of operator algebras. Indeed, if $\mathcal{E}$ exists, then it rivals $\mathcal{R}$ for being the most *canonical* II_1 factor. On the other hand, if $\mathcal{E}$ does not exist, then this can be seen as a strong negative solution to the CEP.

We first mention a theorem that might help us figure out whether or not it exists (see [29] for a proof):

Theorem 8.7 (Dichotomy theorem). *Exactly one of the following two conditions holds:*

- (1) *For every enforceable property P of tracial von Neumann algebras, there exist continuum many nonisomorphic separable tracial von Neumann algebras with property P .*
- (2) *The enforceable II_1 factor $\mathcal{E}$ exists.*

Consequently, one strategy for showing that $\mathcal{E}$ does exist is to find some enforceable property P such that fewer than continuum many tracial von Neumann algebras have property P .

On the other hand, in order to prove that $\mathcal{E}$ does not exist, it might prove useful to analyze some of its properties (should it exist). As mentioned above, being e.c. is an enforceable property, and thus $\mathcal{E}$, if it exists, has all of the properties common to e.c. factors, such as being McDuff and having only approximate inner automorphisms. Moreover, as shown in [29, Section 6], $\mathcal{E}$ would embed into every e.c. factor, which is reminiscent of the situation that $\mathcal{R}$ embeds into every II_1 factor.

Recalling that $\mathcal{R}$ has the McDuff property, we see that $\mathcal{R} \otimes \mathcal{R} \cong \mathcal{R}$. However, one can show that if $\mathcal{E}$ exists, then $\mathcal{E} \otimes \mathcal{E} \not\cong \mathcal{E}$. Indeed, it is possible to show that if the property of being isomorphic to $\mathcal{M} \otimes \mathcal{M}$ for some II_1 factor $\mathcal{M}$ is enforceable, then CEP holds (see [29, Remark 5.8]). Thus, $\mathcal{E} \not\cong \mathcal{M} \otimes \mathcal{M}$ for any tracial von Neumann algebra $\mathcal{M}$.

The theorem of Jung [40] (mentioned in the proof of Theorem 8.2) states that every embedding of $\mathcal{R}$ into $\mathcal{R}^u$ is unitarily conjugate to the diagonal embedding. We say that a II_1 factor $\mathcal{M}$ has the *Jung property* if every embedding of $\mathcal{M}$ into its ultrapower $\mathcal{M}^u$ is unitarily conjugate to the diagonal embedding. Atkinson and Kunnawalkam Elayavalli [4] showed that $\mathcal{R}$ is the only $\mathcal{R}^u$ -embeddable factor with the Jung property. However, in [30], we showed that $\mathcal{E}$, if it exists, also has the Jung property. One can use this fact to show that $\mathcal{E}$, should it exist, cannot even be

elementarily equivalent to $\mathcal{E} \bar{\otimes} \mathcal{E}$, meaning that there must be some L_{vNa} -sentence σ such that $\sigma^{\mathcal{E}} \neq \sigma^{\mathcal{E} \bar{\otimes} \mathcal{E}}$!

ACKNOWLEDGMENTS

We would like to thank the following people for giving us helpful comments and/or corrections regarding earlier versions of this manuscript: Alec Fox, Jeffrey Barrett, Michael Cranston, Vern Paulsen, Jennifer Pi, Thomas Vidick, and Henry Yuen.

ABOUT THE AUTHOR

Isaac Goldbring is associate professor of mathematics at the University of California, Irvine. His research interests are applications of logic, specifically model theory and nonstandard analysis, to a wide variety of areas of mathematics, including operator algebras, combinatorial number theory, and Lie theory.

REFERENCES

- [1] H. Ando, U. Haagerup, and C. Winsløw, *Ultraproducts, QWEP von Neumann algebras, and the Effros-Maréchal topology*, J. Reine Angew. Math. **715** (2016), 231–250, DOI 10.1515/crelle-2014-0005. MR3507925
- [2] E. Artin, *Über die Zerlegung definiter Funktionen in Quadrate* (German), Abh. Math. Sem. Univ. Hamburg **5** (1927), no. 1, 100–115, DOI 10.1007/BF02952513. MR3069468
- [3] S. Arora and B. Barak, *Computational complexity: A modern approach*, Cambridge University Press, Cambridge, 2009, DOI 10.1017/CBO9780511804090. MR2500087
- [4] S. Atkinson and S. Kunnawalkam Elayavalli, *On ultraproduct embeddings and amenability for tracial von Neumann algebras*, Int. Math. Res. Not. IMRN **4** (2021), 2882–2918, DOI 10.1093/imrn/rnaa257. MR4218341
- [5] L. Babai, L. Fortnow, and C. Lund, *Nondeterministic exponential time has two-prover interactive protocols*, Comput. Complexity **1** (1991), no. 1, 3–40, DOI 10.1007/BF01200056. MR1113533
- [6] J. A. Barrett, *The conceptual foundations of quantum mechanics*, Oxford University Press, Oxford, 2019. MR4206568
- [7] J. S. Bell, *On the Einstein Podolsky Rosen paradox*, Phys. Phys. Fiz. **1** (1964), no. 3, 195–200, DOI 10.1103/PhysicsPhysiqueFizika.1.195. MR3790629
- [8] I. Ben Yaacov, A. Berenstein, C. W. Henson, and A. Usvyatsov, *Model theory for metric structures*, Model theory with applications to algebra and analysis. Vol. 2, London Math. Soc. Lecture Note Ser., vol. 350, Cambridge Univ. Press, Cambridge, 2008, pp. 315–427, DOI 10.1017/CBO9780511735219.011. MR2436146
- [9] I. Ben Yaacov and C. W. Henson, *Generic orbits and type isolation in the Gurarij space*, Fund. Math. **237** (2017), no. 1, 47–82, DOI 10.4064/fm193-3-2016. MR3606398
- [10] I. Ben Yaacov and A. P. Pedersen, *A proof of completeness for continuous first-order logic*, J. Symbolic Logic **75** (2010), no. 1, 168–190, DOI 10.2178/jsl/1264433914. MR2605887
- [11] F. Boca, *Free products of completely positive maps and spectral sets*, J. Funct. Anal. **97** (1991), no. 2, 251–263, DOI 10.1016/0022-1236(91)90001-L. MR1111181
- [12] N. P. Brown and N. Ozawa, *C*-algebras and finite-dimensional approximations*, Graduate Studies in Mathematics, vol. 88, American Mathematical Society, Providence, RI, 2008, DOI 10.1090/gsm/088. MR2391387
- [13] V. Capraro and M. Lupini, *Introduction to sofic and hyperlinear groups and Connes’ embedding conjecture*, Lecture Notes in Mathematics, vol. 2136, Springer, Cham, 2015. With an appendix by Vladimir Pestov, DOI 10.1007/978-3-319-19333-5. MR3408561
- [14] R. Cleve, P. Hoyer, B. Toner, and J. Watrous, *Consequences and limits of nonlocal strategies*, CCC ’04 Proceedings of the 19th IEEE Annual Conference on Computational Complexity (2004), 236–249.
- [15] A. Connes, *Classification of injective factors. Cases II_1 , II_∞ , III_λ , $\lambda \neq 1$* , Ann. of Math. (2) **104** (1976), no. 1, 73–115, DOI 10.2307/1971057. MR454659

- [16] A. C. Doherty, Y.-C. Liang, B. Toner, and S. Wehner, *The quantum moment problem and bounds on entangled multi-prover games*, Twenty-Third Annual IEEE Conference on Computational Complexity, IEEE Computer Soc., Los Alamitos, CA, 2008, pp. 199–210, DOI 10.1109/CCC.2008.26. MR2513501
- [17] A. Einstein, B. Podolsky, and N. Rosen, *Can quantum-mechanical description of physical reality be considered complete?*, Physical Review **47** (1935), 777–780.
- [18] H. B. Enderton, *A mathematical introduction to logic*, 2nd ed., Harcourt/Academic Press, Burlington, MA, 2001. MR1801397
- [19] I. Farah, I. Goldbring, B. Hart, and D. Sherman, *Existentially closed II_1 factors*, Fund. Math. **233** (2016), no. 2, 173–196, DOI 10.4064/fm126-12-2015. MR3466004
- [20] I. Farah, B. Hart, and D. Sherman, *Model theory of operator algebras II: model theory*, Israel J. Math. **201** (2014), no. 1, 477–505, DOI 10.1007/s11856-014-1046-7. MR3265292
- [21] I. Farah, B. Hart, and D. Sherman, *Model theory of operator algebras III: elementary equivalence and II_1 factors*, Bull. Lond. Math. Soc. **46** (2014), no. 3, 609–628, DOI 10.1112/blms/bdu012. MR3210717
- [22] G. B. Folland, *A course in abstract harmonic analysis*, 2nd ed., Textbooks in Mathematics, CRC Press, Boca Raton, FL, 2016. MR3444405
- [23] C. Lund, L. Fortnow, H. Karloff, and N. Nisan, *Algebraic methods for interactive proof systems*, 31st Annual Symposium on Foundations of Computer Science, Vol. I, II (St. Louis, MO, 1990), IEEE Comput. Soc. Press, Los Alamitos, CA, 1990, pp. 2–10, DOI 10.1109/FSCS.1990.89518. MR1150685
- [24] T. Fritz, *Tsirelson’s problem and Kirchberg’s conjecture*, Rev. Math. Phys. **24** (2012), no. 5, 1250012, 67, DOI 10.1142/S0129055X12500122. MR2928100
- [25] T. Fritz, T. Netzer, and A. Thom, *Can you compute the operator norm?*, Proc. Amer. Math. Soc. **142** (2014), no. 12, 4265–4276, DOI 10.1090/S0002-9939-2014-12170-8. MR3266994
- [26] K. Gödel, *Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I* (German), Monatsh. Math. Phys. **38** (1931), no. 1, 173–198, DOI 10.1007/BF01700692. MR1549910
- [27] I. Goldbring, *Ultrafilters throughout mathematics*, Graduate Studies in Mathematics, vol. 220, American Mathematical Society, Providence, RI, 2022.
- [28] I. Goldbring, *Spectral gap and definability*, Beyond First Order Model Theory, Volume 2, (to appear).
- [29] I. Goldbring, *Enforceable operator algebras*, J. Inst. Math. Jussieu **20** (2021), no. 1, 31–63, DOI 10.1017/S1474748019000112. MR4205777
- [30] I. Goldbring, *Nonembeddable II_1 factors resembling the hyperfinite II_1 factor*, Journal of Noncommutative Geometry (to appear).
- [31] I. Goldbring and B. Hart, *Computability and the Connes embedding problem*, Bull. Symb. Log. **22** (2016), no. 2, 238–248, DOI 10.1017/bsl.2016.5. MR3532694
- [32] I. Goldbring and B. Hart, *The universal theory of the hyperfinite II_1 factor is not computable*, Preprint, [arXiv:2006.05629](https://arxiv.org/abs/2006.05629), 2020.
- [33] I. Goldbring, B. Hart, and T. Sinclair, *The theory of tracial von Neumann algebras does not have a model companion*, J. Symbolic Logic **78** (2013), no. 3, 1000–1004. MR3135510
- [34] I. Goldbring, S. Kunawalkam Elayavalli, and Y. Lodha, *Generic properties in spaces of enumerated groups* (manuscript in preparation).
- [35] B. C. Hall, *Quantum theory for mathematicians*, Graduate Texts in Mathematics, vol. 267, Springer, New York, 2013, DOI 10.1007/978-1-4614-7116-5. MR3112817
- [36] J. W. Helton and S. A. McCullough, *A Positivstellensatz for non-commutative polynomials*, Trans. Amer. Math. Soc. **356** (2004), no. 9, 3721–3737, DOI 10.1090/S0002-9947-04-03433-6. MR2055751
- [37] W. Hodges, *Building models by games*, London Mathematical Society Student Texts, vol. 2, Cambridge University Press, Cambridge, 1985. MR812274
- [38] T. Ito and T. Vidick, *A multi-prover interactive proof for NEXP sound against entangled provers*, 2012 IEEE 53rd Annual Symposium on Foundations of Computer Science—FOCS 2012, IEEE Computer Soc., Los Alamitos, CA, 2012, pp. 243–252. MR3186611
- [39] Z. Ji, A. Natarajan, T. Vidick, J. Wright, and H. Yuen, *$MIP^* = RE$* , Preprint, [arXiv:2001.04383](https://arxiv.org/abs/2001.04383), 2021.
- [40] K. Jung, *Amenability, tubularity, and embeddings into $\mathcal{R}^\omega$* , Math. Ann. **338** (2007), no. 1, 241–248, DOI 10.1007/s00208-006-0074-y. MR2295511

- [41] M. Junge, M. Navascues, C. Palazuelos, D. Perez-Garcia, V. B. Scholz, and R. F. Werner, *Connes embedding problem and Tsirelson's problem*, J. Math. Phys. **52** (2011), no. 1, 012102, 12, DOI 10.1063/1.3514538. MR2790067
- [42] M. Junge and G. Pisier, *Bilinear forms on exact operator spaces and $B(H) \otimes B(H)$* , Geom. Funct. Anal. **5** (1995), no. 2, 329–363, DOI 10.1007/BF01895670. MR1334870
- [43] S.-J. Kim, V. Paulsen, and C. Schafhauser, *A synchronous game for binary constraint systems*, J. Math. Phys. **59** (2018), no. 3, 032201, 17, DOI 10.1063/1.4996867. MR3776034
- [44] E. Kirchberg, *On nonsemsplit extensions, tensor products and exactness of group C^* -algebras*, Invent. Math. **112** (1993), no. 3, 449–489, DOI 10.1007/BF01232444. MR1218321
- [45] I. Klep and M. Schweighofer, *Connes' embedding conjecture and sums of Hermitian squares*, Adv. Math. **217** (2008), no. 4, 1816–1837, DOI 10.1016/j.aim.2007.09.016. MR2382741
- [46] A. Natarajan and J. Wright, $\text{NEEXP} \subseteq \text{MIP}^*$, Preprint, [arXiv:1904.05870v3](https://arxiv.org/abs/1904.05870), 2019.
- [47] N. Ozawa, *About the Connes embedding conjecture: algebraic approaches*, Jpn. J. Math. **8** (2013), no. 1, 147–183, DOI 10.1007/s11537-013-1280-5. MR3067294
- [48] N. Ozawa, *About the QWEP conjecture*, Internat. J. Math. **15** (2004), no. 5, 501–530, DOI 10.1142/S0129167X04002417. MR2072092
- [49] N. Ozawa and G. Pisier, *A continuum of C^* -norms on $\mathbb{B}(H) \otimes \mathbb{B}(H)$ and related tensor products*, Glasg. Math. J. **58** (2016), no. 2, 433–443, DOI 10.1017/S0017089515000257. MR3483590
- [50] V. Paulsen, *Entanglement and nonlocality*, Lecture notes available at https://www.math.uwaterloo.ca/~vpaulsen/EntanglementandNonlocalityLectureNotes_7.pdf
- [51] V. I. Paulsen, S. Severini, D. Stahlke, I. G. Todorov, and A. Winter, *Estimating quantum chromatic numbers*, J. Funct. Anal. **270** (2016), no. 6, 2188–2222, DOI 10.1016/j.jfa.2016.01.010. MR3460238
- [52] G. Pisier, *Tensor products of C^* -algebras and operator spaces—the Connes-Kirchberg problem*, London Mathematical Society Student Texts, vol. 96, Cambridge University Press, Cambridge, 2020, DOI 10.1017/9781108782081. MR4283471
- [53] F. Rădulescu, *The von Neumann algebra of the non-residually finite Baumslag group $\langle a, b | ab^3a^{-1} = b^2 \rangle$ embeds into R^ω* , Hot topics in operator theory, Theta Ser. Adv. Math., vol. 9, Theta, Bucharest, 2008, pp. 173–185. MR2436761
- [54] A. Robinson, *On ordered fields and definite functions*, Math. Ann. **130** (1955), 275–271, DOI 10.1007/BF01343896. MR75932
- [55] A. Shamir, $IP = PSPACE$, J. Assoc. Comput. Mach. **39** (1992), no. 4, 869–877, DOI 10.1145/146585.146609. MR1187216
- [56] W. Slofstra, *The set of quantum correlations is not closed*, Forum Math. Pi **7** (2019), e1, 41, DOI 10.1017/fmp.2018.3. MR3898717
- [57] R. Speicher, *Lectures notes on “Free probability”*, [arXiv:1908.08125](https://arxiv.org/abs/1908.08125), 2019.
- [58] M. Takesaki, *Theory of operator algebras. II*, Encyclopaedia of Mathematical Sciences, vol. 125, Springer-Verlag, Berlin, 2003. Operator Algebras and Non-commutative Geometry, 6, DOI 10.1007/978-3-662-10451-4. MR1943006
- [59] B. S. Tsirelson, *Some results and problems on quantum Bell-type inequalities*, Hadronic J. Suppl. **8** (1993), no. 4, 329–345. MR1254597
- [60] B. Tsirelson, *Bell inequalities and operator algebras*, available at <https://www.tau.ac.il/~tsirel/download/bellopalg.pdf>.
- [61] A. Usvyatsov, *Generic separable metric structures*, Topology Appl. **155** (2008), no. 14, 1607–1617, DOI 10.1016/j.topol.2007.08.023. MR2435152
- [62] T. Vidick and J. Watrous, *Quantum proofs*, Found. Trends Theor. Comput. Sci. **11** (2015), no. 1-2, 1–215, DOI 10.1561/04000000068. MR3499154
- [63] D. Voiculescu, *Free entropy*, Bull. London Math. Soc. **34** (2002), no. 3, 257–278, DOI 10.1112/S0024609301008992. MR1887698

DEPARTMENT OF MATHEMATICS, UNIVERSITY OF CALIFORNIA, IRVINE, 340 ROWLAND HALL (BLDG.# 400), IRVINE, CALIFORNIA 92697-3875

Email address: isaac@math.uci.edu

URL: <http://www.math.uci.edu/~isaac>