

A note on derivations of Murray–von Neumann algebras

Richard V. Kadison^{a,1} and Zhe Liu^{b,1}

^aDepartment of Mathematics, University of Pennsylvania, Philadelphia, PA 19104; and ^bDepartment of Mathematics, University of Central Florida, Orlando, FL 32816

Edited by Mikael Rordam, University of Copenhagen, Copenhagen, Denmark, and accepted by the Editorial Board December 3, 2013 (received for review July 14, 2013)

A Murray–von Neumann algebra is the algebra of operators affiliated with a finite von Neumann algebra. In this article, we first present a brief introduction to the theory of derivations of operator algebras from both the physical and mathematical points of view. We then describe our recent work on derivations of Murray–von Neumann algebras. We show that the “extended derivations” of a Murray–von Neumann algebra, those that map the associated finite von Neumann algebra into itself, are inner. In particular, we prove that the only derivation that maps a Murray–von Neumann algebra associated with a factor of type II_1 into that factor is 0. Those results are extensions of Singer’s seminal result answering a question of Kaplansky, as applied to von Neumann algebras: The algebra may be noncommutative and may even contain unbounded elements.

quantum mechanics | finite von Neumann algebra | type II_1 factor | Murray–von Neumann algebra | derivation

Section 1: Derivations and Quantum Physics

A derivation of an algebra $\mathfrak{A}$ into an $\mathfrak{A}$ -bimodule $\mathfrak{M}$ is a linear mapping δ of $\mathfrak{A}$ into $\mathfrak{M}$ such that $\delta(AB) = A\delta(B) + \delta(A)B$ for each A and B in $\mathfrak{A}$. When we are talking about the natural $\mathfrak{A}$ -bimodule structure of $\mathfrak{A}$ (as an $\mathfrak{A}$ -bimodule) arising from the addition and multiplication operations on $\mathfrak{A}$, we say that δ is a derivation of $\mathfrak{A}$ into itself or, simply, a derivation of $\mathfrak{A}$. The study of derivations of operator algebras is a large, complicated topic that underwent a vast development in the 60s and 70s.

The term “derivation” is (visibly) closely related to “derivative.” Proceeding informally, if we consider the process of finding the derivative $\frac{df}{dt}$ of functions f , we see that it is linear and satisfies the Leibniz rule for products of functions, $\frac{d}{dt}(fg) = \frac{df}{dt}g + f\frac{dg}{dt}$. In effect, differentiation acts as a derivation on the ring of functions. Of course, something as basic as differentiation and its algebraic counterpart, derivations, must find their way into fundamental physics. And indeed they do; derivations appear as the generators of one-parameter groups that express the symmetries and dynamical evolution of quantum-mechanical systems. We can see this relation to derivations by examining Dirac’s Program (1) for a mathematical formulation of the fundamentals of quantum mechanics. Since this connection with quantum physics is a major motivation for the present study of derivations, we expand on it.

In the early chapters of ref. 1, Dirac is pointing out that Hilbert spaces and their orthonormal bases, if chosen carefully, can be used to simplify calculations and for determinations of probabilities, for example, finding the frequencies of the spectral lines in the visible range of the hydrogen atom (the Balmer series), that is, the spectrum of the operator corresponding to the energy “observable” of the system, the Hamiltonian. In mathematical terms, Dirac is noting that bases, carefully chosen, will simultaneously “diagonalize” self-adjoint operators in an abelian (or “commuting”) family. We shall be doing precisely that in the proof of *Theorem 12*, one of our main results.

The early experimental work that led to quantum mechanics made it clear that, when dealing with systems at the atomic scale, where the measurement process interferes with what is being measured, we are forced to model the physics of such systems at a single instant of time, as an algebraic mathematical structure that is not commutative. Dirac thinks of his small, physical system as an algebraically structured family of observables—elements of the system to be observed when studying the system, for example, the position of a particle in the system would be an observable Q (a “canonical coordinate”) and the (conjugate) momentum of that particle as another observable P —and they are independent of time. As the particle moves under the “dynamics” of the system, the position Q and momentum P become time dependent. By analogy with classical mechanics, Dirac refers to them, in this case, as “dynamical variables.” He recalls the Hamilton equation of motion for a general dynamical variable that is a function of the canonical coordinates $\{q_r\}$ and their conjugate momenta $\{p_r\}$:

$$\frac{dq_r}{dt} = \frac{\partial H}{\partial p_r}, \quad \frac{dp_r}{dt} = -\frac{\partial H}{\partial q_r},$$

where H is the energy expressed as a function of the q_r and p_r and, possibly, of t . This H is the Hamiltonian of the system. Hence, with v a dynamical variable that is a function of the q_r and p_r , but not explicitly of t ,

Significance

In this article, derivations of algebras of unbounded operators acting on a Hilbert space are discussed. Derivations appear as the generators of one-parameter groups that express the symmetries and dynamical evolution of quantum-mechanical systems. One can see this relation to derivations by examining Dirac’s Program for a mathematical formulation of the fundamentals of quantum mechanics.

Author contributions: R.V.K. and Z.L. designed research, performed research, and wrote the paper.

The authors declare no conflict of interest.

This article is a PNAS Direct Submission. M.R. is a guest editor invited by the Editorial Board.

¹To whom correspondence may be addressed. E-mail: kadison@math.upenn.edu or zhe.liu@ucf.edu.

$$\frac{dv}{dt} = \sum_r \left(\frac{\partial v}{\partial q_r} \frac{dq_r}{dt} + \frac{\partial v}{\partial p_r} \frac{dp_r}{dt} \right) = \sum_r \left(\frac{\partial v}{\partial q_r} \frac{\partial H}{\partial p_r} - \frac{\partial v}{\partial p_r} \frac{\partial H}{\partial q_r} \right) = [v, H],$$

where $[v, H]$ is the classical Poisson bracket of v and H . Dirac is using Lagrange's idea of introducing canonical coordinates and their conjugate momenta, in terms of which the dynamical variables of interest for a given system may be expressed, even though those q_r and p_r may not be associated with actual particles in the system. Noting the fundamental nature of the Poisson bracket in classical mechanics, and establishing its Lie bracket properties, Dirac defines a quantum Poisson bracket $[u, v]$ by analogy with the classical bracket. So, it must be "real." Dirac then argues "quasi" mathematically, to show that $uv - vu$ must be $i\hbar[u, v]$, where the real constant $\hbar$ has to be set by the basic quantum mechanical experiments (giving $\hbar = \frac{h}{2\pi}$, with h Planck's constant). Again using classical analogy, the classical coordinates and their conjugate momenta have Poisson brackets

$$[q_r, q_s] = [p_r, p_s] = 0, \quad [q_r, p_s] = \delta_{r,s},$$

where $\delta_{r,s}$ is the Kronecker delta, 1 when $r = s$ and 0 otherwise. So, Dirac assumes that the quantum Poisson brackets of the position Q s and the momentum P s satisfy these same relations. In the case of one degree of freedom, that is, one Q (and its conjugate momentum P), $QP - PQ = i\hbar I$, the basic Heisenberg relation. This relation encodes the noncommutativity needed to produce the so-called "ad hoc quantum assumptions" made by the early workers in quantum physics. At the same time, this relation gives us a numerical grip on "uncertainty" and "indeterminacy" in quantum mechanics. In addition, the Heisenberg relation makes it clear (regrettably) that quantum mechanics cannot be modeled using finite matrices alone. The trace of $QP - PQ$ is 0 when Q and P are such matrices, whereas the trace of $i\hbar I$ is not 0 (no matter how we normalize the trace). It can be shown that the Heisenberg relation cannot be satisfied even with bounded operators on an infinite-dimensional Hilbert space. Unbounded operators are needed, even unavoidable for "representing" (that is, "modeling") the Heisenberg relation mathematically. An extended and thorough study of this modeling appears in ref. 2, where, among other things, the result that the Heisenberg relation cannot be satisfied with self-adjoint operators, unbounded and affiliated with factors of type II_1 , appears. (Such operators and operator algebras will be described presently.) Of course, we do not abandon finite matrices and finite factors on this account. They can still play a crucial role in describing key aspects of quantum physics.

When we move to physical systems with infinitely many degrees of freedom, fields, or statistical mechanical systems, infinite systems of finite matrices, now of arbitrarily large orders, give us the Glimm algebras (3), and from the Glimm algebras, the Powers factors (4), and the complete theoretical description of the representations of the infinite, canonical, anticommutation relation (5)(IV; p. 663–669). Included in this is one of the most useful factors of type II_1 , the "hyperfinite II_1 factor." The key component of the structural description of all factors is a factor of type II_1 . The main Murray–von Neumann algebras we shall study consist of operators affiliated with a factor of type II_1 (see refs. 2, 6). Their basic algebraic properties follow from the pioneering results of Murray and von Neumann in ref. 7.

We continue our description of Dirac's program, and incorporate some of the techniques and advances from the theory of operator algebras, especially those from the sources just cited (toward which Dirac was working in the latter part of his life). Associated with the physical system is a family of observables having some algebraic structure and representable by self-adjoint operators on an infinite-dimensional complex Hilbert space $\mathcal{H}$. Along with this family of observables is a family of states (of the system). Loosely, each state is an "attitude" of the system in which a set of measurements can be performed during an experiment. (Much more austere, a state is an assignment of a probability measure to the "spectrum" of each observable.) Dirac, in a more tentative manner, associates a unit vector in $\mathcal{H}$ (up to a complex multiple of modulus 1, a "phase factor") with each state. If A is an observable and x corresponds to a state of interest, $\langle Ax, x \rangle$, the inner product of the two vectors Ax and x , is the real number we get by taking the average of many measurements of A with the system in the state corresponding to x . Each such measurement yields a real number in the spectrum of A . The probability that that measurement will lie in a given subset of the spectrum is the measure of that set, using the probability measure that the state assigns to A . The "expectation" of the observable A in the state corresponding to x is $\langle Ax, x \rangle$.

With this part of the model in place, Dirac assigns a self-adjoint operator H as the energy observable and, by analogy with classical mechanics, assumes that it will "generate" the dynamics, the time-evolution of the system. This time-evolution can be described in two ways, either as the states evolving in time, the "Schrödinger picture" of quantum mechanics, or the observables evolving in time, the "Heisenberg picture" of quantum mechanics. The prescription for each of these pictures is given in terms of the one-parameter unitary group $t \rightarrow U_t$, where $t \in \mathbb{R}$, the additive group of real numbers, and U_t is the unitary operator $\exp(itH)$, formed by applying the spectral-theoretic, function-calculus to the self-adjoint operator H , the Hamiltonian of our system. If the initial state of our system corresponds to the unit vector x , then at time t , the system will have evolved to the state corresponding to the unit vector $U_t x$. If the observable corresponds to the self-adjoint operator A at time 0, at time t , it will have evolved to $U_t^* A U_t (= \alpha_t(A))$, where, as can be seen easily, $t \rightarrow \alpha_t$ is a one-parameter group of automorphisms of the "algebra" (perhaps, "Jordan algebra") $\mathcal{R}$ of observables. In any event, the numbers we hope to measure are $\langle AU_t x, U_t x \rangle$, the expectation of the observable A in the state (corresponding to) $U_t x$, as t varies, and/or $\langle (U_t^* A U_t) x, x \rangle$, the expectation of the observable $\alpha_t(A)$ in the state x , as t varies. Of course, the two varying expectations are the same, which explains why Heisenberg's "matrix mechanics" and Schrödinger's "wave mechanics" gave the same results. (In Schrödinger's picture, x is a vector in $\mathcal{H}$, viewed as $L_2(\mathbb{R}^3)$, so that x is a function, the "wave function" of the state, evolving in time as $U_t x$, whereas in Heisenberg's picture, the "matrix" coordinates of the operator A evolve in time as $\alpha_t(A)$.) Loosely speaking, the symmetries of the system (and the associated conservation laws) are modeled by the corresponding symmetry groups as groups of automorphisms of $\mathcal{R}$. The time evolution of the system, with a given dynamics, corresponds to a one-parameter group of automorphisms, $t \rightarrow \alpha_t$ of $\mathcal{R}$. Again, very loosely, α_t will be $\exp(it\delta)$ for some linear mapping δ (of the algebra of observables). Thus,

$$\left. \frac{d(\alpha_t(A))}{dt} \right|_{t=0} = \left. \frac{d}{dt} e^{-itH} A e^{itH} \right|_{t=0} = -iH e^{-itH} A e^{itH} + e^{-itH} A e^{itH} (iH) \Big|_{t=0} = -iHA + iAH = i[A, H],$$

while

$$\left. \frac{d(\alpha_t(A))}{dt} \right|_{t=0} = \left. \frac{d}{dt} e^{it\delta(A)} \right|_{t=0} = i\delta(A) e^{it\delta(A)} \Big|_{t=0} = i\delta(A).$$

Thus, $\delta(A) = [A, H]$. Compare this with what we discussed in the case of Hamilton mechanics, time differentiation of the dynamical variable is Poisson bracketing with the Hamiltonian (the total energy). In quantum mechanics, differentiation of the "evolving

observable” is Lie bracketing with the (quantum) Hamiltonian. Of course, this bracketing, δ , is a derivation of the system as the other generators of the one-parameter automorphism groups of the “operator algebras” that describe our physical system and its symmetries are likely to be—hence, our interest in studying those derivations.

Section 2: Derivations and Hochschild’s Cohomology

At a conference held in 1953, Kaplansky asked Singer if he had an idea of what the derivations of $C(X)$, the algebra of continuous functions on a compact Hausdorff space X , might be. A day later, Singer gave Kaplansky a short, clever argument that such derivations are the 0-mapping (that is, must map all of $C(X)$ to 0) (see ref. 8). Kaplansky’s paper (9) and the strong interest in derivations of operator algebras grew out of Singer’s result. Kaplansky showed that each derivation of a type I von Neumann algebra (for example, $B(\mathcal{H})$, the algebra of all bounded operators on $\mathcal{H}$) into itself is “inner” (that is, has the form $\text{Ad}(B)$, where $\text{Ad}(B)(A) = AB - BA$). In the course of his argument, Kaplansky proves that each such derivation is (norm-)continuous and conjectures that that “automatic” continuity is true for all C^* -algebras. This conjecture was proved a few years later by Sakai (10) and extended by Ringrose to derivations of a C^* -algebra into a Banach bimodule (11). These were among the earliest automatic-continuity results. In refs. 12 and 13 (see also refs. 14 and 15) it was proved that each derivation of a C^* -algebra acting on a Hilbert space $\mathcal{H}$ extends to a derivation of the strong-operator closure of that algebra, a von Neumann algebra, and that each derivation of a von Neumann algebra is inner. The proof of this last result is not simple. Surprisingly enough, this theorem is an extension of Singer’s result. Of course, the von Neumann algebra is a C^* -algebra. If it is abelian, it is isomorphic to a $C(X)$, and each inner derivation, $\text{Ad}(B)$ is the 0-mapping. One may object that not all abelian C^* -algebras are von Neumann algebras, but this can be easily remedied by adding the possibility of extending a derivation of a C^* -algebra to its strong-operator closure. It is not, however, in this primitive sense that we see the von Neumann algebra derivation theorem as an extension of Singer’s derivation theorem, but, rather, in the sense that it tells us that each such derivation is 0 as an element of the 1-cohomology group of the von Neumann algebra (16).

In Hochschild’s cohomology of associative algebras (17, 18), an n -linear mapping φ (an “ n -cochain”) of an associative algebra $\mathfrak{A}$ into an $\mathfrak{A}$ -bimodule $\mathfrak{M}$ is transformed by a precisely defined process, the (n -coboundary) operator Δ_n , into an $n + 1$ cochain $\Delta_n(\varphi)$:

$$(\Delta_n \varphi)(A_0, A_1, \dots, A_n) = A_0 \varphi(A_1, \dots, A_n) + \sum_{j=1}^n (-1)^j \varphi(A_0, \dots, A_{j-2}, A_{j-1} A_j, A_{j+1}, \dots, A_n) + (-1)^{n+1} \varphi(A_0, \dots, A_{n-1}) A_n.$$

If $\Delta_n(\varphi) = 0$, φ is said to be an “ n -cocycle.” In any event, $\Delta_{n-1}(\varphi)$ is said to be an “ n -coboundary” and is an n -cocycle (as $\Delta_n \Delta_{n-1} = 0$, the main property of coboundary operations). The coboundary operators are linear, from which the n -cocycles form a linear subspace of the linear space of n -cochains (“on $\mathfrak{A}$ with coefficients in $\mathfrak{M}$ ”) and the n -coboundaries form a linear subspace of the n cocycles whose quotient (as additive groups) is the “ n th cohomology group” of $\mathfrak{A}$ with coefficients in $\mathfrak{M}$.

$\text{Ker} \Delta_n$: n -cocycles; $\text{Im} \Delta_{n-1}$: n -coboundaries; $\Delta_n \Delta_{n-1} = 0$; $\text{Im} \Delta_{n-1} \subseteq \text{Ker} \Delta_n$;

$$\text{Ker} \Delta_n / \text{Im} \Delta_{n-1} = H^n(\mathfrak{A}, \mathfrak{M}).$$

As it relates to our derivations, the Leibniz rule for derivations “embodies” the coboundary operator

$$((\Delta_1(\varphi))(A, B) = A\varphi(B) - \varphi(AB) + \varphi(A)B,$$

which is 0 for all A and B in $\mathfrak{A}$ precisely when φ is a derivation. At the same time, by convention, $C^0(\mathfrak{A}, \mathfrak{M})$ is $\mathfrak{M}$, and $\Delta_0 : C^0(\mathfrak{A}, \mathfrak{M}) \rightarrow C^1(\mathfrak{A}, \mathfrak{M})$ is defined by

$$(\Delta_0 m)(A) = Am - mA,$$

for $A \in \mathfrak{A}$ and $m \in \mathfrak{M}$. (Note that $(\Delta_0 m)$ is an inner derivation of $\mathfrak{A}$.) Therefore,

$$H^1(\mathfrak{A}, \mathfrak{M}) = \text{Ker} \Delta_1 / \text{Im} \Delta_0 = \text{“derivations”} / \text{“inner derivation.”}$$

The theorem of refs. 12, 13 is the statement that the first cohomology group of a von Neumann algebra (with coefficients in itself) is 0 (that is, that each cocycle is a coboundary—that each derivation is an inner derivation). Singer’s theorem tells us that insisting that a derivation apply to all functions in $C(X)$ (that is, in a commutative C^* -algebra) to yield functions, once more, forces the derivation to be the 0-mapping (“numerically”) on $C(X)$. This same insistence for a derivation of a noncommutative C^* -algebra (or its extension to a von Neumann closure of that algebra), again, forces the derivation to be “0” (“cohomologically”). The view of the basic derivation theory of operator algebras from the vantage point of Singer’s seminal answer to Kaplansky’s question and the corresponding result for noncommutative von Neumann algebras raises a number of highly provocative, related questions. For the present article, we concentrate on the questions referring to derivations of the algebras of unbounded operators. The central questions in this connection are as follows: Are there cohomological and numerical 0-nullification results for those algebras? There are, and these are the two main results of this paper.

Section 3: Murray–von Neumann Algebras and Derivations

Returning to the physics discussed in Sec. 1, note that the (physical) Hamiltonian will, in general, correspond to an unbounded operator on our Hilbert space $\mathcal{H}$ as will likely be the case for the other operators K such that $\text{Ad}(K)$ generates a group of symmetries of the quantum system. Of course, these unbounded operators will not lie in a von Neumann algebra, but they may be “affiliated” with the von Neumann algebra corresponding to our quantum system. This makes it very desirable to study derivations of algebras that include such unbounded operators. Regrettably, the tendency of unbounded operators not to combine effectively under the operations of addition and multiplication severely limits the possibility of forming algebras that include these affiliated operators, and along with that, we cannot speak of “their derivations.” There is, however, one intriguing exception discovered by Murray and von Neumann, the “finite” von Neumann algebras and their families of affiliated operators. We say that a closed densely defined operator T on a Hilbert space $\mathcal{H}$ is affiliated with a von Neumann algebra $\mathcal{R}$ when $U'T = TU'$ for each unitary operator U' in $\mathcal{R}'$, the commutant of $\mathcal{R}$. Murray and von Neumann show, at the end of ref. 7, that the family of operators affiliated with a factor of type II_1 (or,

more generally, affiliated with a finite von Neumann algebra, those in which the identity operator is finite) admits surprising operations of addition and multiplication that suit the formal algebraic manipulations used by the founders of quantum mechanics in their mathematical model. (Unbounded operators, even those that are closed and densely defined, can often neither be added nor multiplied usefully. They may not have common dense domains.) In ref. 2, it is proved that the family of operators affiliated with a finite von Neumann algebra is a $*$ algebra (with unit I , the identity operator) under the operations of addition $+$ and multiplication $\cdot$. (If operators S and T are affiliated with $\mathcal{R}$, then $S + T$ and ST are densely defined, preclosed and their closures, denoted by " $S + T$ " and " ST ", respectively, are affiliated with $\mathcal{R}$.) We refer to such algebras as Murray–von Neumann algebras.

If $\mathcal{R}$ is a finite von Neumann algebra, we denote by " $\mathcal{A}_f(\mathcal{R})$ " its associated Murray–von Neumann algebra. The complete cohomological 0-nullification result would say that each derivation of $\mathcal{A}_f(\mathcal{R})$ is inner (that is, is $\text{Ad}(T)$ for some T in $\mathcal{A}_f(\mathcal{R})$). The authors feel that this is true, but it is still open. (It is a work in progress for us.) In this article, we prove that the extended derivations of $\mathcal{A}_f(\mathcal{R})$ (those that map $\mathcal{R}$ into $\mathcal{R}$) are inner (Theorem 5). In Theorem 12, we prove that each derivation of $\mathcal{A}_f(\mathcal{M})$ with $\mathcal{M}$ a factor of type II_1 that maps $\mathcal{A}_f(\mathcal{M})$ into $\mathcal{M}$ is 0. In other words, the restriction that the range of the derivation is in $\mathcal{M}$, the "bounded" part of $\mathcal{A}_f(\mathcal{M})$, allows us to recapture Singer's numerical 0-nullification in the (noncommutative, unbounded) case of $\mathcal{A}_f(\mathcal{M})$. For the general result when $\mathcal{M}$ is a von Neumann algebra of type II_1 , a proof appears elsewhere (19) (the "transcription" from a II_1 factor to a II_1 von Neumann algebra is not an easy one in this case).

Section 4: Matrix Representation of Murray–von Neumann Algebras

Let $\mathcal{R}$ be a ring with unit I , and involution $A \rightarrow A^*$ ($A \in \mathcal{R}$).

Definition 1: We call a set $\{E_{ab}\}_{a,b \in \mathbb{A}}$ a matrix-unit system in $\mathcal{R}$ when each $E_{ab} \neq 0$, $E_{ab}E_{cd}$ is 0 if $b \neq c$ and $E_{ab}E_{bd} = E_{ad}$, for all a, b, c , and d in $\mathbb{A}$. If, in addition, $E_{ab}^* = E_{ba}$, we say that $\{E_{ab}\}$ is a self-adjoint matrix-unit system. If $\{F_{cd}\}_{c,d \in \mathbb{B}}$ is a matrix-unit system in $\mathcal{R}$ such that $\mathbb{A} \subseteq \mathbb{B}$ and $\{E_{ab}\}_{a,b \in \mathbb{A}} \subseteq \{F_{cd}\}_{c,d \in \mathbb{B}}$, we say that $\{F_{cd}\}$ is a larger matrix-unit system than $\{E_{ab}\}$. If $\{E_{ab}\}$ is maximal relative to this partial ordering of matrix-unit systems in $\mathcal{R}$, we call $\{E_{ab}\}_{a,b \in \mathbb{A}}$ a complete matrix-unit system for $\mathcal{R}$. Each E_{ab} in a matrix-unit system is said to be a matrix unit (in the system). The matrix units E_{aa} , $a \in \mathbb{A}$, are said to be principal (or diagonal) matrix units in the system $\{E_{ab}\}_{a,b \in \mathbb{A}}$.

The classic example of a system of matrix units is that of the set of $n \times n$ matrices each of which has a single nonzero entry 1. If that entry is in the j th row and k th column, the resulting matrix is E_{jk} of our matrix-unit system for $\mathcal{M}_n(\mathbb{C})$, the algebra of $n \times n$ matrices with complex entries (in which it is complete). The examples that are most relevant for our present purposes are the finite, complete, self-adjoint matrix-unit systems for factors of type II_1 . If $\mathcal{M}$ is such a factor, the principal matrix units $E_{11}, \dots, E_{nn}$ are equivalent projections (self-adjoint idempotents) and each E_{jk} is a partial isometry with initial projection E_{kk} and final projection E_{jj} . The key result that allows us to begin the process of constructing matrix-unit systems is in ref. 5 (II; sec. 6.5). Lemma 6.5.6 asserts that each projection in a von Neumann algebra $\mathcal{R}$ with no central portion of type I (equivalently, with no nonzero abelian projections), in particular, in a factor of type II_1 , is the sum of n equivalent (orthogonal) projections in $\mathcal{R}$, where n is any preassigned positive integer. In ref. 20, corollary 3.15, it is proved, among other such results, that each maximal abelian, self-adjoint subalgebra of a von Neumann algebra of type II_1 has n orthogonal equivalent projections with sum I . This possibility for choosing the principal matrix units for special purposes directed by spectral analysis is a technique that will be vital to our proof of Theorem 12.

With the ring $\mathcal{R}$ and a finite, self-adjoint matrix-unit system $\{E_{jk}\}_{j,k \in \{1, \dots, n\}}$, such that $\sum_{j=1}^n E_{jj} = I$, there is a procedure for associating a ring of matrices whose entries lie in the subring $\mathcal{T}$ of $\mathcal{R}$ consisting of the elements of $\mathcal{R}$ that commute with all of the matrix units of our system. This procedure is described in ref. 5 (II, lemma 6.6.3). That lemma directs us to assign to T in $\mathcal{R}$ the $n \times n$ matrix whose (j, k) entry T_{jk} is $\sum_{r=1}^n E_{rj}TE_{kr}$. That this element lies in $\mathcal{T}$ follows from $E_{st}T_{jk} = E_{st}(\sum_{r=1}^n E_{rj}TE_{kr}) = E_{st}E_{rj}TE_{kr} = E_{sj}TE_{kt} = E_{sj}TE_{kt} = E_{sj}TE_{ks}E_{st} = (\sum_{r=1}^n E_{rj}TE_{kr})E_{st} = T_{jk}E_{st}$ for $j, k \in \{1, \dots, n\}$.

If we denote by φ the mapping that assigns to T the matrix $[T_{jk}]$ in the $n \times n$ matrix ring $n \otimes \mathcal{T}$ over $\mathcal{T}$, then $\varphi(E_{jk})$ is the matrix with 1 at the (j, k) entry and 0 at all other entries, as the following calculation shows. The (s, t) entry for $\varphi(E_{jk})$ is $\sum_{r=1}^n E_{rs}E_{jk}E_{tr} = 0$ unless $s = j$ and $k = t$, in which case that entry is $\sum_{r=1}^n E_{rj}E_{jk}E_{kr} = \sum_{r=1}^n E_{rr} = I$, by assumption. With the present notation:

Theorem 2. The mapping φ is a $*$ isomorphism of $\mathcal{R}$ onto $n \otimes \mathcal{T}$.

The proof of Theorem 2 appears in ref. 19. Note that $\mathcal{R}$ in the theorem is a general $*$ ring (with unit I).

Section 5: Derivations—Main Results

Let $\mathcal{R}$ be a finite von Neumann algebra acting on a Hilbert space $\mathcal{H}$.

Definition 3: We say that δ , a derivation of $\mathcal{A}_f(\mathcal{R})$, is an extended derivation of $\mathcal{A}_f(\mathcal{R})$ if δ maps $\mathcal{R}$ into $\mathcal{R}$.

Lemma 4. Let T be an operator affiliated with $\mathcal{R}$. Suppose that there is a sequence $\{F_n\}$ of operators in $\mathcal{R}$ with strong-operator limit I , the identity operator, such that $TF_n x = 0$ for all x in $\mathcal{D}(TF_n)$, the domain of TF_n , and for each n . Then $Tx = 0$ for all x in $\mathcal{H}$.

Theorem 5. Suppose that δ is an extended derivation of $\mathcal{A}_f(\mathcal{R})$. Then there is an operator B in $\mathcal{R}$ such that, for each operator A in $\mathcal{A}_f(\mathcal{R})$, $\delta(A) = \text{Ad}(B)(A) = A \cdot B \cdot B^* \cdot A$.

Proof: By definition of extended derivations of $\mathcal{A}_f(\mathcal{R})$, the restriction of δ on $\mathcal{R}$ is a derivation of $\mathcal{R}$. Since every derivation of a von Neumann algebra is inner (12, 13), there is an operator B in $\mathcal{R}$ such that $\delta(A) = AB - BA$ for all $A \in \mathcal{R}$.

Define $\text{Ad}(B): \mathcal{A}_f(\mathcal{R}) \rightarrow \mathcal{A}_f(\mathcal{R})$ by $\text{Ad}(B)(A) = A \cdot B \cdot B^* \cdot A$, ($A \in \mathcal{A}_f(\mathcal{R})$). Note that for every A in $\mathcal{R}$, $\text{Ad}(B)(A) = A \cdot B \cdot B^* \cdot A = AB - BA = \delta(A)$. Let $\delta_0 = \delta - \text{Ad}(B)$. Then δ_0 is a derivation of $\mathcal{A}_f(\mathcal{R})$ and $\delta_0(\mathcal{R}) = 0$. We shall show that $\delta_0(\mathcal{A}_f(\mathcal{R})) = 0$, which will complete the proof.

For any operator A in $\mathcal{A}_f(\mathcal{R})$, let VH be the polar decomposition of A and let E_n be the spectral projection for H corresponding to the interval $[-n, n]$ for each positive integer n . Then, the sequence $\{E_n\}$ is strong-operator convergent to I , and for each n , AE_n is a bounded everywhere-defined operator in $\mathcal{R}$. Moreover,

$$0 = \delta_0(AE_n) = A\delta_0(E_n) + \delta_0(A)E_n = \delta_0(A)E_n.$$

From the preceding lemma, $\delta_0(A) = 0$ ($A \in \mathcal{A}_f(\mathcal{R})$).

We shall prove (Theorem 12, Corollary 13) that the only derivation of $\mathcal{A}_f(\mathcal{M})$, with $\mathcal{M}$ a factor of type II_1 , that maps $\mathcal{A}_f(\mathcal{M})$ into $\mathcal{M}$ is 0. Recall that factors are von Neumann algebras whose centers consist of scalar multiples of the identity operator I . A von Neumann algebra is said to be finite when the identity operator I is finite. Factors without minimal projections in which I is finite are said to be of type II_1 . The following results (whose proofs appear in ref. 19) are used in the proof of Theorem 12, where the harder argumentation occurs.

Definition 6: We say that a von Neumann algebra $\mathcal{R}$ is diffuse if it has no projection minimal in $\mathcal{R}$.

Lemma 7. Each von Neumann algebra $\mathcal{R}$ with no central portion of type I, in particular, a von Neumann algebra of type II_1 , is diffuse.

Proposition 8. Every maximal abelian self-adjoint subalgebra (masa) $\mathcal{A}$ in a diffuse von Neumann algebra $\mathcal{R}$ is diffuse.

Lemma 9. Suppose that B is an operator in $\mathcal{R}$, a finite von Neumann algebra, and that B is not in the center of $\mathcal{R}$. Then, if there is an operator T in $\mathcal{A}_f(\mathcal{R})$ such that $\text{Ad}(B)(T) \notin \mathcal{R}$, there is a self-adjoint operator S in $\mathcal{A}_f(\mathcal{R})$ such that $\text{Ad}(B)(S) \notin \mathcal{R}$.

Lemma 10. Suppose that B is an operator in $\mathcal{R}$, a finite von Neumann algebra, and that B is not in the center of $\mathcal{R}$. If $\text{Ad}(B)(T)$ is in $\mathcal{R}$ for every self-adjoint operator T in $\mathcal{A}_f(\mathcal{R})$, then there is a self-adjoint operator S in $\mathcal{R}$, not in the center of $\mathcal{R}$, such that $\text{Ad}(S)(T)$ is in $\mathcal{R}$ for every self-adjoint operator T in $\mathcal{A}_f(\mathcal{R})$.

Proposition 11. Let $\mathcal{A}$ be an abelian von Neumann algebra acting on a Hilbert space $\mathcal{H}$. Suppose $\{F_a\}_{a \in \mathbb{A}}$ is a family of mutually orthogonal, nonzero projections in $\mathcal{A}$ with sum F , and $\{H_a\}_{a \in \mathbb{A}}$ is a family of self-adjoint operators affiliated with $\mathcal{A}$ such that $H_a F_a = H_a$ for each a in $\mathbb{A}$. Let $\mathcal{D}_a$ be $\mathcal{D}(H_a) \cap F_a(\mathcal{H})$ and $\mathcal{D}$ be the linear span of $\{\{\mathcal{D}_a\}_{a \in \mathbb{A}}, (I - F)(\mathcal{H})\}$. If H_0 is the linear operator with domain $\mathcal{D}$ that maps x_a in $\mathcal{D}_a$ to $H_a x_a$ and x' in $(I - F)(\mathcal{H})$ to 0, then H_0 is closable with closure a self-adjoint operator affiliated with $\mathcal{A}$.

The theorem that follows is formulated in terms of a II_1 factor rather than a general II_1 von Neumann algebra (that appears in ref. 19) to simplify a complicated argument to a certain extent. In the case of a general II_1 von Neumann algebra, quite a bit of difficulty resides in the nature of the center of the von Neumann algebra. This should not be surprising; we are dealing with derivations and (Lie) bracketing and the crucial hypothesis in our main result (following this discussion) is that the operator B about which the assertion is made does not lie in the center. Before we can succeed in constructing what we need in the case where the von Neumann algebra has a robust center, we must transform the condition of “noncentrality” into detailed spectral information about B . Manipulation of central carriers to find a nonzero central projection over which B has distinct spectrum (bounded apart) is necessary. This was quite a difficult task, accomplished by making use of Stone’s characterization of norm-closed, self-adjoint subalgebras of $C(X)$ (21).

Theorem 12. If $\mathcal{M}$ is a factor of type II_1 and B is an operator in $\mathcal{M}$ and B is not a scalar multiple of the identity operator (that is, B is not in the center of $\mathcal{M}$), then there is an operator H in $\mathcal{A}_f(\mathcal{M})$ such that $\text{Ad}(B)(H) \notin \mathcal{M}$.

Proof: Of course, if $\text{Ad}(B)(H) \notin \mathcal{M}$ with B in $\mathcal{M}$ and H in $\mathcal{A}_f(\mathcal{M})$, then $H \notin \mathcal{M}$. From Lemma 9 and Lemma 10, it suffices to consider the case in which B is a self-adjoint element in $\mathcal{M}$, and even a stronger result should be true, that is, a self-adjoint H can be found for each self-adjoint B in $\mathcal{M}$ (not in the center of $\mathcal{M}$) such that $\text{Ad}(B)(H) \notin \mathcal{M}$. We reduce our problem further. Since $\text{Ad}(B)$ and $\text{Ad}(B + aI)$ are the same mapping of $\mathcal{A}_f(\mathcal{M})$, for each a in $\mathbb{C}$, by appropriate choice of a , we may assume that both $\|B\|$ and $-\|B\|$ are in $\text{sp}(B)$, the spectrum of B . Again, since $\text{Ad}(aB)(H) = \text{Ad}(B)(aH) = a\text{Ad}(B)(H)$, for each positive real a , by appropriate choice of a , we may assume that the maximum $\|B\|$ of the spectrum of B is 1, and, with the present reduction, the minimum $-\|B\|$ is -1 .

Let $\mathcal{A}$ be a maximal abelian, self-adjoint subalgebra (masa) of $\mathcal{M}$ containing B . From ref. 5 (I; Theorem 5.2.1), $\mathcal{A} \cong C(X)$, with X an extremely disconnected compact Hausdorff space. Suppose that the operator B corresponds to $\hat{B}$ in $C(X)$. Since 1 and -1 are the maximum and minimum of $\text{sp}(B)$, there are x and x' in X such that $\hat{B}(x) = 1$ and $\hat{B}(x') = -1$. Let S_0 be the closure of the open set on which $\hat{B}$ takes value greater than $\frac{7}{8}$, and let S'_0 be the closure of the open set on which $\hat{B}$ takes value less than $-\frac{7}{8}$. These sets, S_0 and S'_0 , are nonnull, since $x \in S_0$ and $x' \in S'_0$. Let E_0 and E'_0 be the projections in $\mathcal{A}$ corresponding to the characteristic functions of S_0 and S'_0 , respectively. Then from the function representation in $C(X)$, $BE_0 \geq \frac{7}{8}E_0$ and $BE'_0 \leq -\frac{7}{8}E'_0$. If $E \in \mathcal{A}$ is a subprojection of E_0 , then $BE = BE_0 E \geq \frac{7}{8}E_0 E = \frac{7}{8}E$. Similarly, if $E' \in \mathcal{A}$ is a subprojection of E'_0 , then $BE' \leq -\frac{7}{8}E'$.

Without loss of generality, let us assume that $\tau(E_0) \leq \tau(E'_0)$, where τ is the trace on $\mathcal{M}$. Applying corollary 3.14 of ref. 20, there is, for a suitably large positive integer n with $\frac{1}{n} < \tau(E_0)$, a subprojection E in $\mathcal{A}$ of E_0 with $\tau(E) = \frac{1}{n}$. Similarly, there is a subprojection E' in $\mathcal{A}$ of E'_0 with $\tau(E') = \frac{1}{n}$.

Let E be E_1 and let E' be E_n with n the positive integer in the preceding paragraph. From corollary 3.15 of ref. 20, there are $n - 2$ orthogonal equivalent projections each with trace $\frac{1}{n}$ in $\mathcal{A}$, $E_2, E_3, \dots, E_{n-1}$, with sum $I - E_1 - E_n$. (Let $F = I - E_1 - E_n$. According to the corollary, there are $n - 2$ orthogonal equivalent projections in $\mathcal{A}$ with sum F , the identity of $\mathcal{A}$.)

Let V_j be the partial isometry with initial projection E_1 and final projection E_j . Then $V_j^* V_j = E_1$ and $V_j V_j^* = E_j$. Let $E_{jk} = V_j V_k^*$. Then E_{jk} is a partial isometry with initial projection E_k and final projection E_j , and $E_{jj} = V_j V_j^* = E_j$ ($j = 1, 2, \dots, n$) and $\sum_{j=1}^n E_{jj} = \sum_{j=1}^n E_j = I$. Moreover, $E_{jk} E_{kl} = V_j V_k^* V_k V_l^* = V_j E_1 V_l^* = V_j V_l^* = E_{jl}$, $E_{jk} E_{lm} = V_j V_k^* V_l V_m^* = 0$ (if $k \neq l$), and $E_{jk}^* = E_{kj}$. Hence, $\{E_{jk}\}_{j,k=1,\dots,n}$ is a self-adjoint system of $n \times n$ matrix units for $\mathcal{M}$ (and for $\mathcal{A}_f(\mathcal{M})$ as well).

Employing the discussion, results, and notation of Sec. 4, when we compute the matrix in $n \otimes \mathcal{T}$ of the matrix unit E_{1n} in $\mathcal{M}$, the result is the $n \times n$ matrix with I at the $(1, n)$ position and 0 at all other positions. The mapping from $\mathcal{M}$ to $n \otimes \mathcal{T}$ described in ref. 5 (II; sec. 6.6) and in Sec. 4 is a $*$ isomorphism of $\mathcal{M}$ onto $n \otimes \mathcal{T}$.

Returning to the operator B , a self-adjoint element in the masa $\mathcal{A}$, with 1 and -1 as maximum and minimum of its spectrum, respectively; from our construction, $\mathcal{A}$ contains the principal matrix units $E_{11}, \dots, E_{nn}$ of our matrix unit system $\{E_{jk}\}_{j,k=1,\dots,n}$, and $BE_{11} \geq \frac{7}{8}E_{11}$, $BE_{nn} \leq -\frac{7}{8}E_{nn}$. Suppose, also, that we have chosen H , a self-adjoint operator in $\mathcal{A}_f(\mathcal{M})$ as well as in the algebra of operators affiliated with $\mathcal{A}$. Without specifying H precisely, at this point, we assume that $HE_{11} \geq E_{11}$ and $H \hat{\cdot} B (= HB) \notin \mathcal{M}$. Our goal, now, is to show that $HE_{1n} (= H \hat{\cdot} E_{1n})$ and B form a commutator $(\text{Ad}(B)(HE_{1n}))$ that is not in $\mathcal{M}$ (hence, is in $\mathcal{A}_f(\mathcal{M}) \setminus \mathcal{M}$).

The final step is a precise construction of the operator H . For this step, we make use of the fact that each masa in a factor of type II_1 is diffuse (see Proposition 8). Using this, we construct a sequence of nonzero mutually orthogonal subprojections $F_1, F_2, \dots$ of E_{11} in $\mathcal{A}$. We note, from Proposition 11, that $2F_1 + 3F_2 + 4F_3 + \dots$ is an operator with closure H affiliated with $\mathcal{A}$ (here, $\mathbb{A} = \{1, 2, \dots\}$, $H_j = (j+1)F_j$, $\mathcal{D}(H_j) = \mathcal{H}$, $\mathcal{D}_j = F_j(\mathcal{H})$), and that $HE_{11} = H$. Moreover, $E_{11}F_j = F_j$, and $F_j \hat{\cdot} H = HF_j = (j+1)F_j$, because $F_j F_k = 0$ when $j \neq k$. (Recall that, if T is a closed operator and B is a bounded operator on the Hilbert space $\mathcal{H}$, then the operator TB is closed. So, we write HF_j instead of $H \hat{\cdot} F_j$.) Now, F_j and B are in $\mathcal{A}$. Thus,

$$F_j \hat{\cdot} HB = (j+1)F_j B = (j+1)BF_j = (j+1)BE_{11}F_j \geq (j+1)\frac{7}{8}E_{11}F_j = \frac{7}{8}(j+1)F_j$$

for each j . As F_j is a nonzero projection, $\|F_j \hat{\cdot} HB\| \geq \frac{7}{8}(j+1)\|F_j\| = \frac{7}{8}(j+1)$ for each j . Thus, HB is unbounded and affiliated with $\mathcal{A}$. At the same time,

$$\|F_j \hat{\cdot} HBE_{1n}\| \geq \left\| \frac{7}{8}(j+1)F_jE_{1n} \right\| = \frac{7}{8}(j+1),$$

because E_{1n} is a partial isometry with final space $E_{11}(\mathcal{H})$, containing $F_j(\mathcal{H})$. It follows that HBE_{1n} is an unbounded operator in $\mathcal{A}_f(\mathcal{M})$. We shall use this construction to provide us with the desired commutator $\text{Ad}(B)(HE_{1n})$ in $\mathcal{A}_f(\mathcal{M}) \setminus \mathcal{M}$.

The operator $B \hat{\cdot} HE_{1n}$ corresponds to the $n \times n$ matrices over $\mathcal{T}$ with $\sum_{r=1}^n E_{rj} B \hat{\cdot} HE_{1n} E_{kr}$ at the (j, k) entry. Since B is in $\mathcal{A}$ and H is affiliated with $\mathcal{A}$, they commute with all of the principal matrix units E_{kk} ($k=1, \dots, n$), this (j, k) entry is $\sum_{r=1}^n E_{rj} B \hat{\cdot} HE_{1n} E_{kr} E_{kk}$, which is 0 unless $j=1$ and $k=n$. It follows that the (j, k) entry for the $n \times n$ matrix corresponding to $B \hat{\cdot} HE_{1n}$ is 0 at all entries except, possibly, the $(1, n)$ entry, which is $\sum_{r=1}^n E_{r1} B \hat{\cdot} HE_{1r}$. At the same time, each of B , H , and $B \hat{\cdot} H$, has diagonal matrix in $n \otimes \mathcal{T}$ corresponding to it. To see this, note that the (j, k) entry of the matrix corresponding to B is $(B_{jk} =) \sum_{r=1}^n E_{rj} B E_{kr}$, which is $\sum_{r=1}^n E_{rj} E_{jj} B E_{kk} E_{kr} = \sum_{r=1}^n E_{rj} B E_{jj} E_{kk} E_{kr}$. Since $E_{jj} E_{kk}$ is 0 unless $j=k$, in which case $E_{jj} E_{kk} = E_{jj}$, the (j, k) entry of the matrix corresponding to B is 0 unless $j=k$, in which case the (j, j) entry is $(B_{jj} =) \sum_{r=1}^n E_{rj} B E_{jr}$ for each j . Thus, B corresponds to the diagonal matrix with B_{jj} at the diagonal positions (j, j) ($j=1, \dots, n$), and 0 at every off-diagonal position. If we compute $\text{Ad}(B)(HE_{1n})$ ($= (HE_{1n}) \hat{\cdot} B \hat{\cdot} B \hat{\cdot} (HE_{1n})$) in terms of the $n \times n$ matrices corresponding to it, we have that $\text{Ad}(B)(HE_{1n})$ corresponds to the $n \times n$ matrix with (j, k) entry,

$$\sum_{r=1}^n E_{rj} \hat{\cdot} HE_{1n} B E_{kr} \hat{\cdot} \sum_{r=1}^n E_{rj} B \hat{\cdot} HE_{1n} E_{kr} = \sum_{r=1}^n E_{rj} \hat{\cdot} H E_{jj} E_{1n} E_{kk} B E_{kr} \hat{\cdot} \sum_{r=1}^n E_{rj} B \hat{\cdot} H E_{jj} E_{1n} E_{kk} E_{kr},$$

which is 0 unless $j=1$ and $k=n$, in which case it is the $(1, n)$ entry,

$$\sum_{r=1}^n E_{r1} \hat{\cdot} HE_{1n} B E_{nr} \hat{\cdot} \sum_{r=1}^n E_{r1} B \hat{\cdot} HE_{1r} = \left(\sum_{r=1}^n E_{r1} \hat{\cdot} HE_{1r} \right) \left(\sum_{s=1}^n E_{sn} B E_{ns} \right) \hat{\cdot} \left(\sum_{r=1}^n E_{r1} B E_{1r} \right) \hat{\cdot} \left(\sum_{s=1}^n E_{s1} \hat{\cdot} HE_{1s} \right) = H_{11} B_{nn} \hat{\cdot} B_{11} \hat{\cdot} H_{11}.$$

We want to show that this entry is not in $\mathcal{M}$ (and is, hence, unbounded). If this $(1, n)$ entry is in $\mathcal{M}$, then multiplying it on the left by $-E_{11}$ and on the right by E_{11} results in

$$-E_{11} \left(H_{11} B_{nn} \hat{\cdot} B_{11} \hat{\cdot} H_{11} \right) \hat{\cdot} E_{11} = -E_{11} \left(\sum_{r=1}^n E_{r1} \hat{\cdot} HE_{1n} B E_{nr} \hat{\cdot} \sum_{r=1}^n E_{r1} B \hat{\cdot} HE_{1r} \right) E_{11} = B \hat{\cdot} HE_{11} \hat{\cdot} HE_{1n} B E_{n1},$$

which is also in $\mathcal{M}$. We argue, by contradiction, to show that this is not the case.

In the construction of H , we defined projections F_j in $\mathcal{A}$ such that $F_j \hat{\cdot} H = (j+1)F_j$. Thus,

$$\|B \hat{\cdot} HE_{11} \hat{\cdot} HE_{1n} B E_{n1}\| = \|F_j\| \|B \hat{\cdot} HE_{11} \hat{\cdot} HE_{1n} B E_{n1}\| \|F_j\| \geq (j+1) \|BF_j E_{11} F_j - F_j E_{1n} B E_{n1} F_j\| \geq (j+1) \|BF_j - F_j E_{1n} B E_{n1} F_j\|.$$

Now, by choice of E_{11} ,

$$BF_j = BE_{11} F_j \geq \left(\frac{7}{8} E_{11} \right) F_j = \frac{7}{8} F_j,$$

while

$$-F_j E_{1n} B E_{n1} F_j = -F_j E_{1n} B E_{nn} E_{n1} F_j \geq F_j E_{1n} \left(\frac{7}{8} E_{nn} \right) E_{n1} F_j = \frac{7}{8} F_j E_{11} F_j = \frac{7}{8} F_j.$$

Hence,

$$BF_j - F_j E_{1n} B E_{n1} F_j \geq \frac{14}{8} F_j, \quad \|BF_j - F_j E_{1n} B E_{n1} F_j\| \geq \frac{14}{8},$$

and

$$\|B \hat{\cdot} HE_{11} \hat{\cdot} HE_{1n} B E_{n1}\| \geq \frac{14}{8} (j+1) > j,$$

for each positive integer j . It follows that $B \hat{\cdot} HE_{11} \hat{\cdot} HE_{1n} B E_{n1}$ is not bounded, not in $\mathcal{M}$, and that $\text{Ad}(B)(HE_{1n}) \in \mathcal{A}_f(\mathcal{M}) \setminus \mathcal{M}$.

Corollary 13. Suppose that δ is a derivation of $\mathcal{A}_f(\mathcal{M})$ that maps $\mathcal{A}_f(\mathcal{M})$ into $\mathcal{M}$, where $\mathcal{M}$ is a factor of type II_1 . Then $\delta(A) = 0$ for every A in $\mathcal{A}_f(\mathcal{M})$.

Proof: Since δ maps $\mathcal{A}_f(\mathcal{M})$ into $\mathcal{M}$, δ maps $\mathcal{M}$ into $\mathcal{M}$. So, δ is an extended derivation of $\mathcal{A}_f(\mathcal{M})$. From Theorem 5, δ is inner, that is, there is an operator B in $\mathcal{M}$ such that, for each operator A in $\mathcal{A}_f(\mathcal{M})$, $\delta(A) = \text{Ad}(B)(A) = A \hat{\cdot} B \hat{\cdot} B \hat{\cdot} A$. If the operator B is in the center of $\mathcal{M}$, then B is in the center of $\mathcal{A}_f(\mathcal{M})$ (see proposition 30 of ref. 6) and hence for each operator A in $\mathcal{A}_f(\mathcal{M})$, $\text{Ad}(B)(A) = A \hat{\cdot} B \hat{\cdot} B \hat{\cdot} A = 0$. If B is not in the center of $\mathcal{M}$, from Theorem 12, there is an operator H in $\mathcal{A}_f(\mathcal{M}) \setminus \mathcal{M}$ such that $\text{Ad}(B)(H) \notin \mathcal{M}$, contrary to the assumption that δ maps $\mathcal{A}_f(\mathcal{M})$ into $\mathcal{M}$. Thus, the only derivation of $\mathcal{A}_f(\mathcal{M})$ into $\mathcal{M}$ is 0.

Section 6: Further Questions

The question of what is involved, mathematically, when $\text{Ad}(B)(T)$ is bounded in the circumstances where $B \in \mathcal{R}$ with $\mathcal{R}$ a finite von Neumann algebra and $T \in \mathcal{A}_f(\mathcal{R})$, especially when T is unbounded, is of vital interest to the program of describing the derivations of $\mathcal{A}_f(\mathcal{R})$. The simple observation that $\text{Ad}(B)(T)$ is bounded when T is bounded or when T is in the center of $\mathcal{A}_f(\mathcal{R})$ leads us, at once, to the guess that $\text{Ad}(B)(T)$ is bounded if and only if T is the sum of an operator in $\mathcal{A}_f(\mathcal{R})$ commuting with B (the set of such operators will be denoted by “ $(B)'$ ”) and an operator in $\mathcal{R}$. We prove this result, here, when B is a projection. For the general case, it is still open. It seems necessary to develop a calculus of which operators produce a bounded commutator with T , given that A does. So, for example, we have proved that polynomials in A do, as do A^t , where A is positive with spectrum in $(0, 1)$ and $t \in (0, 1)$. When this “calculus” has reached a certain state of development, one might begin to show that T is in $(B)' + \mathcal{R}$.

Theorem 14. *If E is a projection in $\mathcal{R}$, with $\mathcal{R}$ a finite von Neumann algebra, then $\text{Ad}(E)(T) \in \mathcal{R}$, for some T in $\mathcal{A}_f(\mathcal{R})$, if and only if $T \in (E)' \hat{+} \mathcal{R}$, where $(E)'$ is the commutant of E in $\mathcal{A}_f(\mathcal{R})$.*

Proof: Note that $T = E \hat{\circ} T \hat{\circ} E \hat{+} E \hat{\circ} T \hat{\circ} (I - E) \hat{+} (I - E) \hat{\circ} T \hat{\circ} E \hat{+} (I - E) \hat{\circ} T \hat{\circ} (I - E)$. It follows from this decomposition that

$$\text{Ad}(E)(T) = E \hat{\circ} T \hat{\circ} E \hat{+} (I - E) \hat{\circ} T \hat{\circ} E \hat{+} E \hat{\circ} T \hat{\circ} (I - E) \hat{+} (I - E) \hat{\circ} T \hat{\circ} (I - E) = (I - E) \hat{\circ} T \hat{\circ} E \hat{+} E \hat{\circ} T \hat{\circ} (I - E).$$

Assuming that $\text{Ad}(E)(T)$ is a bounded operator in $\mathcal{R}$, $\text{Ad}(E)(T)E = (I - E) \hat{\circ} T \hat{\circ} E$ and $-\text{Ad}(E)(T)(I - E) = E \hat{\circ} T \hat{\circ} (I - E)$ are bounded, as is $E \hat{\circ} T \hat{\circ} (I - E) \hat{+} (I - E) \hat{\circ} T \hat{\circ} E (= B)$. It follows that $T = E \hat{\circ} T \hat{\circ} E \hat{+} (I - E) \hat{\circ} T \hat{\circ} (I - E) \hat{+} B \in (E)' \hat{+} \mathcal{R}$.

1. Dirac PAM (1947) *The Principles of Quantum Mechanics* (Clarendon, Oxford), 3rd Ed.
2. Liu Z (2011) On some mathematical aspects of the Heisenberg relation. *Sci. China Math.* 54:2427–2452.
3. Glimm JG (1960) On a certain class of operator algebras. *Trans Am Math Soc* 95: 318–340.
4. Powers RT (1967) Representations of uniformly hyperfinite algebras and their associated von Neumann rings. *Ann Math* 86:138–171.
5. Kadison RV, Ringrose JR (1983, 1986, 1991, 1992) *Fundamentals of the Theory of Operator Algebras*, Vols. I and II (Grad Stud in Math, AMS, Providence, RI., reprints of the 1983 and 1986 originals), III and IV (Birkhäuser Boston, Inc., Boston).
6. Liu Z (2012) A double commutant theorem for Murray-von Neumann algebras. *Proc Natl Acad Sci USA* 109:7676–7681.
7. Murray FJ, von Neumann J (1936) On rings of operators. *Ann Math* 37:116–229.
8. Kadison RV (2000) Which Singer is that? *Surveys in Differential Geometry VII* ed Yau S-T (Int. Press, Somerville, MA), pp 347–373.
9. Kaplansky I (1953) Modules over operator algebras. *Am J Math* 75:839–858.
10. Sakai S (1960) On a conjecture of Kaplansky. *Tohoku Math J* 12:31–33.
11. Ringrose JR (1972) Automatic continuity of derivations of operator algebras. *J Lond Math Soc* 5:432–438.
12. Kadison RV (1966) Derivations of operator algebras. *Ann Math* 83:280–293.
13. Sakai S (1966) Derivations of W^* -algebras. *Ann Math* 83:273–279.
14. Kadison RV, Ringrose JR (1966) Derivations of operator group algebras. *Am J Math* 88: 562–576.
15. Johnson BE, Ringrose JR (1969) Derivations of operator algebras and discrete group algebras. *Bull Lond Math Soc* 1:70–74.
16. Kadison RV, Ringrose JR (1971) Cohomology of operator algebras I. Type I von Neumann algebras. *Acta Math* 126:227–243.
17. Hochschild G (1945) On the cohomology groups of an associative algebra. *Ann Math* 46:58–67.
18. Hochschild G (1946) On the cohomology theory for associative algebras. *Ann Math* 47: 568–579.
19. Kadison RV, Liu Z, Derivations of Murray-von Neumann algebras. *Math Scand*, in press.
20. Kadison RV (1984) Diagonalizing matrices. *Am J Math* 106:1451–1468.
21. Stone MH (1948) The generalized Weierstrass approximation theorem. *Math Mag* 21: 167–184.