

Allgemeine Eigenwerttheorie Hermitescher Funktionaloperatoren

General Eigenvalue Theory of Hermitian Functional Operators

J. von Neumann

Mathematische Annalen 102, 49–131

Introduction

I.

If we consider all functions on a certain metric space Ω (for example, the real line, the complex plane, the surface of the unit sphere, the interval $0, 1$, and so forth) that satisfy certain regularity conditions (for example, that are continuous and, apart from finitely many corners, continuously differentiable; that are twice continuously differentiable; that have a finite integral of the square of the absolute value over Ω ¹; and so forth), and that may in addition be subject to certain boundary conditions, then, as is well known, a *linear operator* means the following: a correspondence R that assigns to every function f in our class a function Rf (which need no longer belong to this class), but does so in such a way that always

$$R(a_1 f_1 + \cdots + a_k f_k) = a_1 Rf_1 + \cdots + a_k Rf_k \quad (a_1, \dots, a_k \text{ complex constants}).$$

¹The functions may take complex values.

If a general notion of measure (for example, in the sense of Lebesgue) exists in Ω , if dv is the volume element in Ω (on the line: dx ; in the plane: $dx dy$; on the surface of the unit sphere: $\sin \vartheta d\vartheta d\varphi$; and so forth), and if integration over this entire space is denoted by $\int_{\Omega}$, then a linear operator R is called *self-adjoint* or *Hermitian* if, for all admissible functions f, g ,

$$\int_{\Omega} f \overline{Rg} dv = \int_{\Omega} Rf \overline{g} dv$$

holds.²

Examples of Hermitian linear operators (which we shall abbreviate as H.O.) are the following. We now also display the arguments of the functions, and let P denote a general point of Ω :

$$Rf(P) = f(P),$$

$$Rf(P) = x f(P),$$

(x any coordinate of P in Ω),

$$Rf(P) = \int_{\Omega} \varphi(P', P) f(P') dv',$$

($\varphi(P', P) = \overline{\varphi(P, P')}$, i.e. the kernel is Hermitian),

$$Rf(P) = Lf(P),$$

(L any self-adjoint differential or partial differential expression; cf. note 2).

As is well known, the last two types of H.O. determine an eigenvalue problem which, when the kernel $\varphi(P, P')$, respectively the coefficients of L , possess sufficient

² $\overline{f}$ is the function whose value at every point is the complex conjugate of the value of f . Our terminology differs somewhat from the customary one. A differential operator R , for example, is usually called self-adjoint when

$$f \overline{Rg} - Rf \overline{g}$$

is the complete differential of an expression in f, g and their derivatives, so that

$$\int_{\Omega} f \overline{Rg} dv - \int_{\Omega} Rf \overline{g} dv$$

is an expression depending only on the boundary values of f, g and their derivatives. Our self-adjointness (Hermitian character) follows from this only when that expression vanishes as a consequence of suitable boundary conditions. Thus a differential operator that is self-adjoint in the usual sense may, for us, become self-adjoint only on a sufficiently restricted domain of definition (cf. [Introduction VII](#)).

regularity (namely, when only a “point spectrum” is present), can be formulated as follows.³

An eigenvalue is a number λ for which there exists a function $f \neq 0$ such that

$$Rf = \lambda f;$$

f is then an eigenfunction. In general there are infinitely many eigenvalues $\lambda_1, \lambda_2, \dots$, to which eigenfunctions $\varphi_1, \varphi_2, \dots$ can be assigned so that they form a complete system of orthonormal functions. In the following display, the mathematically intended order of the two cases is used.^{T1} That is,

$$\int_{\Omega} \varphi_m \overline{\varphi_n} dv = \begin{cases} 1, & m = n, \\ 0, & m \neq n \end{cases} \quad (\text{orthonormality}),$$

and, for every pair of admissible functions f, g ,

$$\int_{\Omega} f \overline{g} dv = \sum_{n=1}^{\infty} \left(\int_{\Omega} f \overline{\varphi_n} dv \right) \overline{\left(\int_{\Omega} g \overline{\varphi_n} dv \right)} \quad (\text{completeness}).$$

As for every complete orthonormal system, the expansion theorem with convergence in the mean holds:

$$\int_{\Omega} \left| f - \sum_{n=1}^N a_n \varphi_n \right|^2 dv \longrightarrow 0, \quad a_n = \int_{\Omega} f \overline{\varphi_n} dv,$$

as $N \rightarrow \infty$.

If $\varphi(P, P')$, respectively the coefficients of L , cease to behave quite regularly, complications arise that are characterized by the occurrence of a continuous spectrum. Nevertheless, the concepts of eigenvalue and eigenfunction can be retained through suitable generalizations (in particular, by weakening the regularity requirements on the eigenfunction, or even by considering so-called differential

³The eigenvalue problem for regular integral operators was solved by Hilbert (*Göttinger Nachrichten*, 1904, pp. 49–91); his method was later considerably simplified by Courant, Hellinger, F. Riesz, E. Schmidt, Toeplitz, and others. In solving the eigenvalue problem, differential operators (at least self-adjoint operators of second order) can be replaced by integral operators (Hilbert, *Göttinger Nachrichten*, 1904, pp. 213–259) having the same eigenfunctions but reciprocal eigenvalues.

^{T1}*Translator’s note:* In the source, the two values in this displayed case distinction are interchanged.

eigenfunctions⁴). The completeness theorems likewise remain valid under suitable generalization (replacement of the sum by an integral, and so forth); but the entire construction becomes substantially more complicated and loses much of its original intuitive character.⁵

II.

If we replace the continuous spaces Ω considered in I by the “discrete space” $1, 2, 3, \dots$ of positive integers, replacing functions $f(P)$ (P ranging over Ω) by sequences a_n ($n = 1, 2, \dots$) and integrals $\int_{\Omega} f dv$ by sums $\sum_{n=1}^{\infty} a_n$ in the corresponding fashion, then we arrive at notions entirely analogous to those in I.

The simplest examples of H.O. (indeed, as can be shown after the concept is made more precise, the only ones) are linear transformations in infinitely many variables with a Hermitian matrix:

$$R(x_1, x_2, \dots) = (y_1, y_2, \dots), \quad y_n = \sum_{m=1}^{\infty} \alpha_{mn} x_m \quad (n = 1, 2, \dots), \quad \alpha_{mn} = \overline{\alpha_{nm}}.$$

(Here the role played by the “regularity conditions” for functions in Ω is played by conditions such as the requirement that all $\sum_{m=1}^{\infty} \alpha_{mn} x_m$ converge, and the like. We shall consider these matrices in greater detail in [Appendix III](#).) These linear transformations are plainly analogues of the integral-kernel transformations in continuous spaces Ω :

$$Rf(P) = \int_{\Omega} \varphi(P', P) f(P') dv'.$$

For these transformations, or for the corresponding Hermitian forms

$$\sum_{m,n=1}^{\infty} \alpha_{mn} x_m \overline{y_n},$$

⁴Hellinger, Göttingen inaugural dissertation, 1907, p. 7; Carleman, *Sur les équations intégrales à noyau réel et symétrique*, Uppsala, 1923, p. 82.

⁵The eigenvalue problem for integral equations with a “bounded” kernel was solved by Weyl (Göttingen inaugural dissertation, 1908) with the aid of Hilbert’s theory of bounded bilinear forms (*Göttinger Nachrichten*, 1906, pp. 157–209). Weyl also settled the eigenvalue problem for an extensive class of self-adjoint differential operators of second order whose coefficients may be singular (*Mathematische Annalen* 68 (1910), pp. 220–269). A still more general class of integral operators was investigated by Carleman; cf. note 4.

the eigenvalue problem was likewise posed and solved by Hilbert, provided the matrix $\{\alpha_{mn}\}$ (respectively the Hermitian form $\sum_{m,n=1}^{\infty} \alpha_{mn} x_m \overline{y_n}$) satisfies certain conditions corresponding to the “regularity” of the kernel $\varphi(P, P')$ in integral-kernel transformations.

The same phenomena occur. Under high regularity (complete continuity), there is only a point spectrum, with the same properties as in I, except that $\int_{\Omega}$ must be replaced by $\sum_{n=1}^{\infty}$. Thus λ is an eigenvalue if there is a sequence $(x_1, x_2, \dots)$ such that

$$R(x_1, x_2, \dots) = \lambda(x_1, x_2, \dots), \quad \text{i.e.} \quad \sum_{m=1}^{\infty} \alpha_{mn} x_m = \lambda x_n \quad (n = 1, 2, \dots),$$

without $(x_1, x_2, \dots) = (0, 0, \dots)$, and with finiteness of $\sum_{n=1}^{\infty} |x_n|^2$ also required as a “regularity condition”; $(x_1, x_2, \dots)$ is then an eigensequence. If $\lambda_1, \lambda_2, \dots$ are all the eigenvalues, eigensequences

$$(x_{11}, x_{12}, \dots), \quad (x_{21}, x_{22}, \dots), \quad \dots$$

can be assigned to them so that they form a complete orthonormal system:

$$\sum_{m=1}^{\infty} x_{pm} \overline{x_{qm}} = \begin{cases} 1, & p = q, \\ 0, & p \neq q \end{cases} \quad (\text{orthonormality}),$$

and⁶

$$\sum_{m=1}^{\infty} x_{mp} \overline{x_{mq}} = \begin{cases} 1, & p = q, \\ 0, & p \neq q \end{cases} \quad (\text{completeness}).$$

Under lower regularity (boundedness), a continuous spectrum can again occur, causing the same complications as those sketched in I for continuous spaces. Since we shall later have to concern ourselves with these in considerable detail, we shall not discuss them further here.⁷ Some sort of “regularity” assumption has,

⁶This is, as is well known, equivalent to the proper analogue of the completeness relation (cf. I),

$$\sum_{n=1}^{\infty} a_n \overline{b_n} = \sum_{p=1}^{\infty} \left(\sum_{n=1}^{\infty} a_n \overline{x_{pn}} \right) \overline{\left(\sum_{n=1}^{\infty} b_n \overline{x_{pn}} \right)}.$$

⁷The completely continuous and bounded forms (respectively operators) were discovered by Hilbert, who solved their eigenvalue problem (*Göttinger Nachrichten*, 1906, pp. 157–209). The bounded bilinear forms were investigated further by Hellinger; cf. note 4, and *Journal für die reine und angewandte Mathematik* **136** (1909), pp. 210–273.

moreover, always been made in the literature to date, both for matrices and for integral operators; the eigenvalue problem has never been discussed on the basis of Hermitian character alone.

III.

The analogous behavior of the continuous spaces Ω among themselves, and also with the discrete space $1, 2, \dots$ considered here, is known to rest on the following fact.

If $\varphi_1, \varphi_2, \dots$ is a complete system of orthonormal functions in the space Ω under consideration, then the correspondence

$$f \longleftrightarrow (x_1, x_2, \dots), \quad x_n = \int_{\Omega} f \overline{\varphi_n} dv \quad (n = 1, 2, \dots)$$

assigns to every f with finite integral $\int_{\Omega} |f|^2 dv$ a sequence $(x_1, x_2, \dots)$ with finite sum $\sum_{n=1}^{\infty} |x_n|^2$. Distinct f correspond to distinct $(x_1, x_2, \dots)$; and, by the theorem of Fischer and F. Riesz,⁸ every $(x_1, x_2, \dots)$ with finite $\sum_{n=1}^{\infty} |x_n|^2$ does in fact correspond to an f with finite $\int_{\Omega} |f|^2 dv$.

The correspondence is linear; that is,

$$\begin{aligned} f \longleftrightarrow (x_1, x_2, \dots) &\implies af \longleftrightarrow (ax_1, ax_2, \dots), \\ \left. \begin{aligned} f &\longleftrightarrow (x_1, x_2, \dots), \\ g &\longleftrightarrow (y_1, y_2, \dots) \end{aligned} \right\} &\implies f + g \longleftrightarrow (x_1 + y_1, x_2 + y_2, \dots), \end{aligned}$$

and,⁹

$$\int_{\Omega} f \overline{g} dv = \sum_{n=1}^{\infty} x_n \overline{y_n}.$$

Thus all the notions used in describing eigenvalues and eigenfunctions are invariant under this correspondence; they must therefore exhibit the same behavior in the continuous spaces considered in I and in the discrete space of II.

It is essentially by these considerations, as is well known, that Hilbert reduced the eigenvalue theory of integral operators in continuous spaces to that of H.O. for sequences $(x_1, x_2, \dots)$ (i.e. functions on the discrete space $1, 2, \dots$).¹⁰

⁸For example, *Göttinger Nachrichten*, 1907, pp. 116–122.

⁹We shall prove these generally known facts, together with the Fischer–F. Riesz theorem, in [Chapter I](#) and [Appendix I](#).

¹⁰This is historically inaccurate insofar as the Fischer–F. Riesz theorem was discovered only

IV.

The subject of this paper is the construction of a general eigenvalue theory for *all* H.O. We shall proceed abstractly, arranging matters from the outset so that our results apply uniformly to all the continuous spaces named in [I](#), as well as to the discrete space in [II](#) (that is, to the operators on the functions in these spaces).

To this end, in [Chapter I](#) we shall give a general characterization (by five properties A through E; see there) that applies to each of the following function spaces:

All complex functions f , measurable in the Lebesgue sense, on a space Ω (the number line, the complex plane, the surface of the unit sphere, the interval $0, 1$, and so forth; cf. [Appendix I](#)), with volume element dv , for which

$$\int_{\Omega} |f|^2 dv$$

is finite. All sequences $(x_1, x_2, \dots)$ of complex numbers for which

$$\sum_{n=1}^{\infty} |x_n|^2$$

is finite.

In the first case, it should be observed that functions f, g for which $f(P) = g(P)$ everywhere except on a Lebesgue null set are not to be regarded as distinct.

The elements of these function spaces admit all vector operations: they can be added and multiplied by complex numbers. Moreover, for two such elements the “inner product”

$$\int_{\Omega} f \bar{g} dv, \quad \text{respectively} \quad \sum_{n=1}^{\infty} x_n \bar{y}_n,$$

is defined. We denote the inner product by (f, g) (unless something special is said, we shall also regard sequences $(x_1, x_2, \dots), (y_1, y_2, \dots), \dots$ as functions, as in [II](#), and denote them by $f, g, \dots$). With its aid, the Hermitian character of the operators R can be defined in every case by

$$(f, Rg) = (Rf, g).$$

later. The same considerations have played an important role in the new quantum theory; cf. Schrödinger, *Annalen der Physik* **79**, no. 8 (1926), pp. 734–756.

(In I this was the definition, and it is equivalent to the one in II.) As with vectors, the inner product permits the definition of an absolute value

$$|f| = \sqrt{(f, f)} \quad \left(\sqrt{\int_{\Omega} |f|^2 dv} \text{ or } \sqrt{\sum_{n=1}^{\infty} |x_n|^2} \right),$$

and of a distance $|f - g|$ (thus the analogue of ordinary Euclidean distance). Our space is consequently “metrized” and “topologized” by it; that is, topological expressions such as closedness, continuity, and so forth become meaningful in the space.

It may not be superfluous to devote a word to the consistent restriction to functions f with finite $\int_{\Omega} |f|^2 dv$. In eigenvalue problems for sequences, it has been customary since the fundamental investigations of Hilbert (cf. note 7) and E. Schmidt¹¹ to require the finiteness of $\sum_{n=1}^{\infty} |x_n|^2$; for functions f on continuous spaces, $\int_{\Omega} |f|^2 dv$ is in any event finite for the eigenfunctions of the point spectrum. In the continuous spectrum, admittedly, neither $\sum_{n=1}^{\infty} |x_n|^2$ nor $\int_{\Omega} |f|^2 dv$ is finite (if eigensequences or eigenfunctions still exist at all and recourse need not be had to “differential solutions”). We shall nevertheless show, by generalizing Hilbert’s methods, that it is especially advantageous to treat the eigenfunctions in the continuous spectrum as improper objects; that is, to formulate the eigenvalue problem without mentioning them explicitly.

When the discrete space $1, 2, \dots$ is taken as its basis, our function space is the so-called complex Hilbert space (i.e. the infinite-dimensional Euclidean space): the set of all sequences $(x_1, x_2, \dots)$ with finite $\sum_{n=1}^{\infty} |x_n|^2$. Since, as was observed in III, all our manifolds of functions are isomorphic (that is, can be mapped one-to-one onto one another in such a way that the operations $f + g$, af , and (f, g) pass into their analogues), we shall call all of them Hilbert spaces. They are all to be regarded as special realizations of the general abstract Hilbert space, which is characterized by the “intrinsic” properties A through E alone (and, up to isomorphism, completely; cf. Chapter I).

We denote the abstract Hilbert space by $\mathfrak{H}$.

V.

In the abstract Hilbert space, by an H.O. we understand, as stated in I, a function R (defined for certain f in $\mathfrak{H}$ and taking values in $\mathfrak{H}$) that possesses the following

¹¹For example, *Rendiconti del Circolo Matematico di Palermo* **25** (1908), pp. 57–73.

properties:

$$R(a_1 f_1 + \cdots + a_k f_k) = a_1 R f_1 + \cdots + a_k R f_k \quad (a_1, \dots, a_k \text{ complex constants}), \quad (\alpha)$$

$$(f, Rg) = (Rf, g), \quad (\beta)$$

insofar as both sides have meaning. For the time being, moreover, we assume that the domain of R is a linear manifold; that is, whenever it contains $f_1, \dots, f_k$, it also contains $a_1 f_1 + \cdots + a_k f_k$. (The final definition of an H.O. will in fact be somewhat different, but this does not concern us at first.)

We call an H.O. R *bounded* if $|f| \leq 1$ implies $|Rf| \leq C$ for some fixed constant C independent of f . Since generally $|af| = |a||f|$, this is equivalent to

$$|Rf| \leq C|f|, \quad \text{or also} \quad |Rf - Rg| \leq C|f - g|;$$

thus it expresses a kind of continuity of R .¹²

In Hilbert's theory of bounded Hermitian forms (i.e. H.O.), one may always assume that R is defined for every f in $\mathfrak{H}$: if this is not the case, it can be extended to all of $\mathfrak{H}$ by virtue of its continuity. But if we wish to establish a general theory of H.O. that also goes beyond the bounded ones, we may not require R to be defined everywhere in $\mathfrak{H}$. (Indeed, by a theorem of Toeplitz, every H.O. meaningful everywhere must be bounded; cf. [Theorem 48](#), α), γ), and [note 61](#).) For example, let Ω be the interval $0, 1$, and let R be the operator

$$Rf = i \frac{df}{dx}$$

for all differentiable functions vanishing at 0 and 1. Then R is an H.O., but plainly it cannot be defined for nondifferentiable functions. Likewise, for a highly singular integral kernel $\varphi(P, P')$, the expression

$$\int_{\Omega} \varphi(P', P) f(P') dv'$$

will not converge for every f with finite $\int_{\Omega} |f|^2 dv$. Another example, of a different but equally characteristic kind, is this: let Ω be the number line and $Rf(x) = xf(x)$.

¹²If we regard $\mathfrak{H}$ as a sequence space (Hilbert space proper), then the wording of our definition of boundedness differs from Hilbert's: in Hilbert's terminology we require the boundedness of R folded with itself. Nevertheless, the two formulations are equivalent; we discuss this more closely in [Chapter II, Theorem 12](#). It should also be noted that the entire topology of $\mathfrak{H}$ is always to be understood in the sense of so-called "strong convergence"; cf. Weyl, dissertation, p. 9.

This R too is an H.O.; but since

$$\int_{-\infty}^{\infty} |f(x)|^2 dx$$

may be finite while

$$\int_{-\infty}^{\infty} x^2 |f(x)|^2 dx$$

is infinite, this R also does not always have meaning.

The restriction that (α) and (β) are to hold only insofar as both sides have meaning is therefore quite essential. As a minimum, however, we shall require that

the domain of R be everywhere dense in $\mathfrak{H}$. (γ)

What this means can be seen without difficulty from, for example, $i d/dx$ and multiplication by x : the former is always meaningful for all polynomials vanishing at 0 and 1, and the latter for all functions that vanish outside an arbitrary but finite interval; both sets of functions are everywhere dense.

With this arrangement, however, we run the danger of requiring too little: our H.O. could already be meaningless at points at which they could still be defined naturally, merely because the domain was narrowed arbitrarily. (For example, for $i d/dx$ one could, despite the everywhere-denseness of the domain, have required analyticity rather than mere differentiability.) To avoid this difficulty, we introduce the concept of a maximal H.O.

An H.O. S is called an *extension* of the H.O. R if, wherever Rf has meaning, Sf also has meaning and agrees with Rf . If $S \neq R$ (that is, if S has meaning at more points than R), then S is a *proper extension*. An H.O. R having no further proper extensions is called *maximal*. (Abbreviations: ext., max.)

Since, as we shall see in [Theorem 35](#), every H.O. can be extended to a max. H.O., we must investigate max. H.O. first and foremost. To do this, however, it is necessary to give a general formulation of the eigenvalue problem, in the abstract Hilbert space, that takes account of point and continuous spectra on an equal footing. This will be done in the next paragraph, following Hilbert's formulation of the problem.

VI.

Consider first the k -dimensional Euclidean space. If an H.O. is given there, that is, a correspondence

$$R(x_1, \dots, x_k) = (y_1, \dots, y_k),$$

$$y_n = \sum_{m=1}^k \alpha_{mn} x_m \quad (n = 1, \dots, k),$$

$$\alpha_{mn} = \overline{\alpha_{nm}}.$$

(there are of course no convergence questions here), then, as is well known, R can always be put into diagonal form by introducing new Cartesian coordinates, that is, by a unitary transformation. Thus there is a unitary transformation matrix $\{u_{mn}\}$ such that

$$\sum_{p=1}^k u_{mp} \overline{u_{np}} = \begin{cases} 1, & m = n, \\ 0, & m \neq n, \end{cases} \quad \sum_{p=1}^k u_{pm} \overline{u_{pn}} = \begin{cases} 1, & m = n, \\ 0, & m \neq n, \end{cases}$$

and, after the transformation

$$\xi_n = \sum_{m=1}^k u_{mn} x_m, \quad \eta_n = \sum_{m=1}^k u_{mn} y_m \quad (n = 1, \dots, k),$$

R is described simply by

$$\eta_n = \lambda_n \xi_n \quad (n = 1, \dots, k),$$

where $\lambda_1, \dots, \lambda_k$ are real constants, the eigenvalues of R .¹³ That is,

$$\alpha_{mn} = \sum_{p=1}^k \lambda_p u_{mp} \overline{u_{np}} \quad (m, n = 1, \dots, k).$$

Here $\lambda_1, \dots, \lambda_k$, apart from their order, are uniquely determined by R , but the u_{mn} are not. As is immediately seen, they may be replaced by $\vartheta_n u_{mn}$, where

¹³As one sees, $\lambda_1, \dots, \lambda_k$ are eigenvalues, and

$$(\overline{u_{11}}, \dots, \overline{u_{k1}}), \dots, (\overline{u_{1k}}, \dots, \overline{u_{kk}})$$

are corresponding eigenvectors; cf. II.

$\vartheta_1, \dots, \vartheta_k$ are arbitrary numbers of absolute value 1; and when several of the λ_n are equal, even the corresponding rows $(u_{1n}, \dots, u_{kn})$ may be transformed unitarily among themselves in an arbitrary manner.

An ambiguous formulation of the eigenvalue problem such as this one (transformability to diagonal form) is ill suited for passage from k -dimensional Euclidean space to infinite-dimensional, that is, Hilbert space. The correct formulation of the finite-dimensional problem, which is easy to generalize, is obtained by stating it uniquely after Hilbert. This is done as follows.

Set

$$e_{mn}(\lambda) = \sum_{\lambda_p \leq \lambda} u_{mp} \overline{u_{np}} \quad (\lambda \text{ an arbitrary real number});$$

then the matrices $\{e_{mn}(\lambda)\}$ are H.O. $E(\lambda)$ with the readily verified property

$$E(\lambda)E(\mu) = E(\mu)E(\lambda) = E(\lambda) \quad (\lambda \leq \mu). \quad (\alpha)$$

(In particular, $E(\lambda)^2 = E(\lambda)$.)¹⁴ Let $l_1, \dots, l_h$ ($h \leq k$) be the mutually distinct values among $\lambda_1, \dots, \lambda_k$, arranged in increasing order. Then $E(\lambda)$, as a function of λ , is constant on the intervals

$$\lambda < l_1, \quad l_1 \leq \lambda < l_2, \quad \dots, \quad l_{h-1} \leq \lambda < l_h, \quad l_h \leq \lambda,$$

and has jumps from the left at $l_1, \dots, l_h$. In the first of these intervals it is plainly 0; in the last, by the unitarity of $\{u_{mn}\}$, it is the identity matrix, hence the identity operator 1. The last formula for α_{mn} immediately gives, with l_0 arbitrary but $l_0 < l_1$,¹⁵

$$\alpha_{mn} = \sum_{p=1}^h l_p (e_{mn}(l_p) - e_{mn}(l_{p-1})) = \int_{-\infty}^{\infty} \lambda de_{mn}(\lambda).$$

It follows at once, if, as in IV, $(f, g) = \sum_{n=1}^k x_n \overline{y_n}$ when $f = (x_1, \dots, x_k)$ and $g = (y_1, \dots, y_k)$, that

$$(f, Rg) = \int_{-\infty}^{\infty} \lambda d(f, E(\lambda)g). \quad (\beta)$$

It is easy to see that (α) and (β) , together with $E(\lambda) = 0$, respectively 1, for small, respectively large, λ , and with right-continuity of $E(\lambda)$ as a function of λ , determine the operators $E(\lambda)$ completely. Since the diagonal form of R can easily

¹⁴The meaning of AB for operators A, B that are meaningful everywhere is clear.

¹⁵The final expression is a Stieltjes integral.

be recovered from the $E(\lambda)$ (respectively the $e_{mn}(\lambda)$), this is the unique formulation of the eigenvalue problem.

Let us now return to the abstract Hilbert space $\mathfrak{H}$. Conditions (α) and (β) may be carried over verbatim, provided only that the $E(\lambda)$ are defined precisely: they are to be H.O. meaningful everywhere.¹⁶ We formulate the additional requirements concerning the dependence of $E(\lambda)$ on λ as follows:

$$\begin{aligned} E(\lambda) &\longrightarrow 0 \quad (\lambda \rightarrow -\infty), & E(\lambda) &\longrightarrow 1 \quad (\lambda \rightarrow +\infty), \\ E(\lambda) &\longrightarrow E(\lambda_0) \quad \text{as } \lambda \rightarrow \lambda_0 \text{ with } \lambda \geq \lambda_0. \end{aligned} \tag{\gamma}$$

¹⁷ (That precisely this is the correct generalization will be shown by its success.)

A family $E(\lambda)$ satisfying (α) and (γ) is called a *resolution of the identity* (abbreviated: r.i.); if R is related to it by (β) , R is the corresponding operator. Convergence questions do of course arise, but we shall see (cf. [Theorem 36](#)) that they are best resolved by strengthening (β) as follows: Rg is to have meaning if and only if

$$\int_{-\infty}^{\infty} \lambda^2 d|E(\lambda)g|^2 \tag{\beta'}$$

is finite; then, as we shall show, all the integrals in (β) are absolutely convergent, and (β) is to hold. (It follows from (α) that $|E(\lambda)f|^2$, as a function of λ , never decreases; cf. [Theorems 17](#) and [18](#). The integrand in (β') is therefore nonnegative; hence the integral in (β') is either absolutely convergent or diverges properly, namely to $+\infty$.) It is easy to see that an r.i. (that is, one satisfying (α) and (γ)) uniquely determines the corresponding operator, if one exists: the domain is given by (β') , and all (f, Rg) by (β) , hence also Rg itself.

The eigenvalue problem thus consists in deciding: Is there an associated H.O. for every r.i.? Is it maximal? Must different operators belong to different r.i.? Which H.O. belong to r.i.?¹⁸ We shall see that the first three questions are to be

¹⁶In the k -dimensional space $E(\lambda)^2 = E(\lambda)$; hence, as is easily verified, $|E(\lambda)f| \leq |f|$. One may therefore expect it to be bounded and thus take it to be meaningful everywhere; cf. also [Chapter III, Theorems 14, 15, and 16](#).

¹⁷For operators A_n, B meaningful everywhere, $A_n \rightarrow B$ means that $A_n f \rightarrow Bf$ for every f . The identity operator 1 is defined by $1f = f$.

¹⁸Hilbert proved that every bounded H.O. has one and only one r.i.; moreover, precisely those r.i. occur for which $E(\lambda)$ becomes constant, hence 0 or 1 , for all sufficiently small or sufficiently large λ . Condition (β') then drops out, since it is always satisfied. Citation: note [7](#). In our formulation the distinction between point and continuous spectrum is evidently not yet completed (in (γ) , $\lambda \geq \lambda_0$ is required); this is immaterial for our purposes.

answered affirmatively (cf. [Theorem 36](#)). For the fourth we must therefore ask further: the H.O. associated with r.i. form a subclass of the class of all max. H.O.; what is this subclass?

VII.

We shall see ([Theorems 36 and 38](#)) that this subclass does not comprise all max. H.O. There is one exceptional operator $\bar{R}$,¹⁹ and all other exceptions can be built from it and from the H.O. having an r.i. by certain trivial operations, not to be discussed here (cf. [Theorem 38](#)). For $\bar{R}$ we shall see that its properties in fact exclude every reasonable formulation of an eigenvalue problem, but that it is by no means “pathological”: it is easy to describe (cf. note 19).

Nevertheless, max. H.O. that are bounded, definite, or real²⁰ must possess an r.i.; thus certain simple properties are incompatible with the absence of an r.i. (cf. [Theorem 40](#)). We may also interpret the occurrence of this $\bar{R}$ as showing that the Hermitian character of infinite-dimensional operators does not prevent the exceptional cases of elementary-divisor theory from occurring, in contrast to the finite-dimensional case or, in infinite dimensions, to real, bounded, or definite operators; but only a single exceptional type can occur.

For the moment, this may suffice concerning max. H.O. (In particular, we shall not discuss further how, and to what extent, knowledge of the r.i. – that is, of the “Hilbert spectral form” – solves the eigenvalue problem as it was formulated in [I](#) and [II](#). In this connection we refer to the investigations cited in notes [4](#) and [5](#).) It remains, however, to discuss those H.O. whose maximality is not assured, or is absent. H.O. are usually given in precisely this way: for example, as self-adjoint differential operators defined for all n -times differentiable functions (n being their order) that satisfy certain boundary conditions. We mentioned in [V](#) that every H.O. can be extended to a max. H.O.; it was not said, however, that this can be done in only one way.

We shall also obtain simple results on this point, including the following: if an H.O. is not itself max., or is not max. after a certain trivial extension, it can be

¹⁹[Appendix IV](#) is devoted, among other things, to this operator.

²⁰Let $\varphi_1, \varphi_2, \dots$ be a complete orthonormal system in $\S$ (cf. [Definition 3](#)); then every f can be written uniquely in the form $\sum_{n=1}^{\infty} x_n \varphi_n$. It is called real if all the x_n are real. $\Re f$ is obtained from f by replacing every x_n by $\Re x_n$ (the real part). R is real if, whenever Rf has meaning, $R\Re f$ also has meaning and equals $\Re(Rf)$. (Thus an arbitrary choice of a complete orthonormal system enters into the concept of “reality.”)

made max. in a continuum of different ways (cf. the concluding remark of [Chapter VIII](#)). Bounded H.O. always fall in the first case.

A simple example may illustrate how this nonunique extendibility arises. Let Ω be the interval $0, 1$, and let R be the operator $i d/dx$, defined for all differentiable functions. Since

$$\begin{aligned}(f, Rg) - (Rf, g) &= -i \int_0^1 \{f(x)\overline{g'(x)} + f'(x)\overline{g(x)}\} dx \\ &= -i\{f(1)\overline{g(1)} - f(0)\overline{g(0)}\},\end{aligned}$$

R is not Hermitian; but the operators R_0 and R_{ϑ} , where ϑ is a number of absolute value 1, are Hermitian. They arise from R by the following restrictions of its domain:

$$f(0) = f(1) = 0, \quad \text{respectively} \quad f(1) = \vartheta f(0).$$

R_{ϑ} has maximal extensions, for example S_{ϑ} . All the S_{ϑ} are also maximal extensions of R_0 . They are all distinct: for if $S_{\vartheta_1} = S_{\vartheta_2} = S$, then S would have meaning for an f and a g satisfying

$$f(0) = g(0) = 1, \quad f(1) = \vartheta_1, \quad g(1) = \vartheta_2,$$

and consequently

$$(f, Sg) - (Sf, g) = -i(\vartheta_1\overline{\vartheta_2} - 1) \neq 0,$$

so that S would not be an H.O.^{T2} Thus the multiplicity of the maximal extensions here stems from the multiplicity of the possible boundary conditions. We shall see that, even in the most general case, the situation displays a certain analogy with this one (cf. [Theorems 26](#) and [27](#)).²¹

VIII.

The connection between the H.O. of $\S$ and matrices is, as is well known, the following. Let R be an H.O. and let $\varphi_1, \varphi_2, \dots$ be a complete orthonormal system

^{T2}*Translator's note:* The source prints $+i$ in the integration-by-parts formula and omits the factor i in this last display. With the inner-product convention of this paper and $R = i d/dx$, the boundary form has the factor $-i$; both signs and the omitted factor have been corrected.

²¹Carleman had already found signs of this nonuniqueness for singular integral operators, but in that case he could not solve the eigenvalue problem completely (the analogue of the decisive condition (α) from [VI](#) was lacking). Cf. note 4.

(cf. [Definition 3](#)); then the matrix $\{\alpha_{mn}\}$, where

$$\alpha_{mn} = (\varphi_m, R\varphi_n),$$

is assigned to R . For bounded H.O. this assignment is entirely clear; for unbounded ones it involves certain complications. It will occupy us in [Appendix III](#). The unitary transformation theory of unbounded matrices leads to a particularly peculiar pathology, but this will be investigated elsewhere.²²

Nor shall we discuss here the relation of these results to the new quantum theory (the so-called “transformation theory”). In any event, they show that the general theory of H.O. by no means everywhere exhibits the behavior generally presumed and anticipated in “transformation theory” on the basis of analogy with bounded, or even finite-dimensional, operators.

IX.

It remains to say something about the method and arrangement of this paper.

The essential device is to investigate a certain operator U , which could symbolically be denoted by

$$U = \frac{R + i1}{R - i1},$$

in place of R . For if R is an H.O., then U preserves length (that is, $|Uf| = |f|$), and is therefore continuous and hence easier to handle (cf. [Theorems 22, 23, and 24](#)). This “Cayley transformation”

$$U = \frac{R + i1}{R - i1}$$

is the core of our method.

All other methods fail: the elegant procedures of Hellinger and F. Riesz, because they require iteration of the operator R , which is straightforward only for H.O. meaningful everywhere (and hence bounded), but is initially questionable for arbitrary H.O.; all maximum–minimum methods, as well as E. Schmidt’s method, because they are restricted from the outset to cases without continuous spectrum. Only Hilbert’s original method, approximation by finite-dimensional section operators, permits certain results to be obtained, but only with very great difficulty

²²Cf. a forthcoming paper by the author in *Journal für die reine und angewandte Mathematik*, “Zur Theorie der unbeschränkten Matrizen.”

and only a fraction of what we can achieve.²³

The arrangement of the subject is as follows. Chapters I through IV are preparatory: I, structure of the space $\mathfrak{H}$; II, general facts about operators in $\mathfrak{H}$; III, linear manifolds and their projection operators; IV, reducibility of operators. In Chapters V through VIII, the principal work of solving the eigenvalue problem is carried out: V, Cayley transformation; VI, extension elements; VII, deficiency indices; VIII, the extension process. In Chapters IX and X, the exact normal forms for max. H.O. are then established: IX, eigenvalue representation; X, operators without an eigenvalue representation. Finally, Chapters XI and XII contain several general theorems about unbounded operators: XI, semiboundedness; XII, pathology of unboundedness.²⁴ There are also several appendices. Appendix I proves that the function spaces of Introduction I really do satisfy conditions A through E of Chapter I, and hence that our theory is applicable to them.²⁵ In Appendix II we prove a theorem on the solution of the eigenvalue problem for unitary operators (more generally, for “normal” operators). This is a question about bounded operators that can be solved by F. Riesz’s method, but that has never yet been settled within the framework of Hilbert’s theory.²⁶ We need this for Chapter VIII; but since methodologically it belongs to Hilbert’s theory, we prefer to treat it separately in this way. In Appendix III we examine more closely the connection between operators and matrices (cf. Introduction VIII and note 23); in Appendix IV the operator $\bar{R}$ (cf. Introduction VII) is investigated more precisely.

Apart from this, the paper is a logically self-contained whole that can be

²³Using this method the author was able to settle the case of real H.O. (cf. note 20), where the exceptional case indicated in VII had not yet appeared. The process of extending H.O., however, could not be surveyed as it will be, for example, in Chapter VIII; and the method was so nonconstructive that, for instance, the well-ordering theorem had to be invoked (cf. the announcement in *Jahresbericht der Deutschen Mathematiker-Vereinigung* 37, nos. 1–4 (1928), pp. 11–15). Nor could nonreal H.O. be approached in this way. The author does not wish to miss this opportunity to thank E. Schmidt for the interest he has shown in these investigations, especially in the extension of the results to nonreal H.O., and to note that conversations with him repeatedly furnished essential suggestions. In particular, the concept of “hypermaximality” of H.O. (cf. Definition 9), as well as the recognition that this property is necessary and sufficient for the existence of an r.i., is due to E. Schmidt.

²⁴The latter will be sharpened considerably in the author’s paper announced in note 22.

²⁵Together with several considerations from Chapter I, this of course yields a proof of the Fischer–F. Riesz theorem.

²⁶It essentially amounts to the fact that two commuting bounded H.O. admit a common eigenvalue form (or spectral form) in Hilbert’s sense. Cf. Hellinger–Toeplitz, encyclopedia article II, C. 13, p. 1562 ff., where the question is solved for completely continuous H.O.

read without any prior knowledge of the literature on “infinitely many variables.” (Accordingly, the contents of [Chapters I](#) and [III](#), and of [Appendix I](#), apart from their arrangement, are not new.)

I. The Abstract Hilbert Space

We characterize the abstract complex Hilbert space by five properties A through E. It proves convenient to append to some of them several simple theorems that already follow from the particular property alone.

A. $\mathfrak{H}$ is a linear space.

That is: in $\mathfrak{H}$ there are an addition $f + g$, a subtraction $f - g$, a scalar multiplication af (f, g in $\mathfrak{H}$, a a complex number), and an element 0 ; and these obey the rules of ordinary vector algebra.

Definition 1. A subset $\mathfrak{M}$ of $\mathfrak{H}$ is called a linear manifold (abbreviated: *lin. mfd.*) if, together with $f_1, \dots, f_n$, it also contains $a_1 f_1 + \dots + a_n f_n$. If $\mathfrak{M}$ is any subset of $\mathfrak{H}$, then the set of all

$$a_1 f_1 + \dots + a_n f_n \quad (f_1, \dots, f_n \text{ in } \mathfrak{M}; a_1, \dots, a_n \text{ constants})$$

is the smallest linear manifold containing $\mathfrak{M}$; it is called the linear manifold spanned by $\mathfrak{M}$.

Definition 2. The n elements $f_1, \dots, f_n$ are called linearly independent (abbreviated: *lin. indep.*) if

$$a_1 f_1 + \dots + a_n f_n = 0 \quad \implies \quad a_1 = \dots = a_n = 0.$$

The n linear manifolds $\mathfrak{M}_1, \dots, \mathfrak{M}_n$ are called linearly independent if

$$f_1 + \dots + f_n = 0, \quad f_1 \in \mathfrak{M}_1, \dots, f_n \in \mathfrak{M}_n,$$

implies $f_1 = \dots = f_n = 0$.

B. There is in $\mathfrak{H}$ an inner product, analogous to that of vector calculus, which induces a metric.

That is: there is a function (f, g) (f, g in $\mathfrak{H}$, and (f, g) a complex number) having the following properties:

1. $(af, g) = a(f, g)$, 2. $(f_1 + f_2, g) = (f_1, g) + (f_2, g)$,
3. $(f, g) = \overline{(g, f)}$, 4. $(f, f) \geq 0$, and it equals 0 only when $f = 0$.

From 1 and 2, by 3, it follows that

$$1'. (f, ag) = \bar{a}(f, g), \quad 2'. (f, g_1 + g_2) = (f, g_1) + (f, g_2).$$

The formula

$$|f| = \sqrt{(f, f)}$$

defines an “absolute value,” and $|f - g|$ defines the metric (cf. [Theorem 2](#) and note 27).

Theorem 1. *One always has*

$$|(f, g)| \leq |f||g|.$$

Proof. First, by 2, 2', 3, and 4,

$$\begin{aligned} \Re(f, g) &= \left(\frac{f+g}{2}, \frac{f+g}{2} \right) - \left(\frac{f-g}{2}, \frac{f-g}{2} \right) \\ &\leq \left(\frac{f+g}{2}, \frac{f+g}{2} \right) + \left(\frac{f-g}{2}, \frac{f-g}{2} \right) \\ &= \frac{1}{2}((f, f) + (g, g)) = \frac{|f|^2 + |g|^2}{2}. \end{aligned}$$

We replace f, g by $af, a^{-1}g$ ($a > 0$): the left-hand side remains invariant, while the right-hand side has minimum $|f||g|$. We then replace f by ϑf ($|\vartheta| = 1$): the right-hand side remains invariant, while the left-hand side has maximum $|(f, g)|$. Hence

$$|(f, g)| \leq |f||g|.$$

Theorem 2. *One always has*

$$|af| = |a||f|, \quad |f + g| \leq |f| + |g|.$$

Proof. The first assertion follows immediately from 1 and 1', the second from [Theorem 1](#):

$$(f + g, f + g) = (f, f) + (g, g) + 2\Re(f, g),$$

$$|f + g|^2 = |f|^2 + |g|^2 + 2\Re(f, g) \leq |f|^2 + |g|^2 + 2|f||g|,$$

and therefore $|f + g| \leq |f| + |g|$.

Definition 3. Two elements f, g are orthogonal (abbreviated: orth.) if $(f, g) = 0$; two linear manifolds are orthogonal if every element of one is orthogonal to every element of the other. A set $\mathfrak{M}$ is called orthonormal (abbreviated: orthonorm.) if, for f, g in $\mathfrak{M}$,

$$(f, g) = \begin{cases} 1, & f = g, \\ 0, & f \neq g. \end{cases}$$

It is called complete (abbreviated: complete orthonorm.) if it is not a proper subset of another orthonormal set. This plainly says that

$$(f, g) = 0 \quad (f \text{ fixed and } g \text{ arbitrary in } \mathfrak{M})$$

excludes $|f| = 1$, hence also excludes $|f| > 0$, and therefore implies $f = 0$. (As is well known, for several elements or linear manifolds, pairwise orthogonality implies linear independence.)

Definition 4. By [Theorem 2](#), $|f - g|$ is a distance in $\mathfrak{S}$,²⁷ and hence $\mathfrak{S}$ is topologized. A closed (abbreviated: cl.) linear manifold is therefore one that contains its accumulation points. If to the linear manifold spanned by $\mathfrak{M}$ we adjoin all of its accumulation points, we obtain the smallest closed linear manifold containing $\mathfrak{M}$: the closed linear manifold spanned by $\mathfrak{M}$.

²⁷A "distance" $\overline{f, g}$ is, as is well known, required to satisfy the following axioms:

1. $\overline{f, g} \geq 0$, and it is 0 only when $f = g$;
2. $\overline{f, g} = \overline{g, f}$;
3. $\overline{f, h} \leq \overline{f, g} + \overline{g, h}$.

In linear spaces one adds

$$4. \quad \overline{f + h, g + h} = \overline{f, g}, \quad \overline{af, ag} = |a| \overline{f, g}.$$

By [Theorem 2](#), all these properties are possessed by $\overline{f, g} = |f - g|$. Note that, by its bilinearity (1 and 1' in [B](#)) and [Theorem 1](#), (f, g) is continuous.

C. In the metric $|f - g|$, $\mathfrak{H}$ is separable.

That is: a certain countable set is everywhere dense in $\mathfrak{H}$.

D. $\mathfrak{H}$ possesses arbitrarily many (finitely many) linearly independent elements.²⁸

E. $\mathfrak{H}$ is complete.

That is: if a sequence $f_1, f_2, \dots$ in $\mathfrak{H}$ satisfies the Cauchy convergence condition (for every $\varepsilon > 0$ there is an $N = N(\varepsilon)$ such that $m, n \geq N$ implies $|f_m - f_n| \leq \varepsilon$), then it is convergent (there exists an f in $\mathfrak{H}$ such that for every $\varepsilon > 0$ there is an $N = N(\varepsilon)$ for which $n \geq N$ implies $|f_n - f| \leq \varepsilon$).

In [Appendix I](#) we shall show that all the function spaces named in [Introduction I](#) and [II](#) have properties A through E, and hence are special cases of the abstract Hilbert space $\mathfrak{H}$. Here, however, we proceed to prove that $\mathfrak{H}$ is in fact completely characterized by properties A through E: that $\mathfrak{H}$ is isomorphic to ordinary Hilbert space. For this it is necessary to prove several generally known theorems about orthonormal and complete orthonormal systems.

Theorem 3. *Every orthonormal set $\mathfrak{M}$ is at most countable; every complete orthonormal set $\mathfrak{M}$ is countably infinite.*

Remark. *Henceforth we therefore write all orthonormal and complete orthonormal sets as sequences $\varphi_1, \varphi_2, \dots$ (which in the former case may terminate).*

Proof. Let $\mathfrak{M}$ be orthonormal. For any two f, g in $\mathfrak{M}$ with $f \neq g$,

$$|f - g|^2 = |f|^2 + |g|^2 - 2\Re(f, g) = 2, \quad |f - g| = \sqrt{2}.$$

By the separability of $\mathfrak{H}$ (condition [C](#)), $\mathfrak{M}$ is therefore at most countable.

It remains to show that a finite orthonormal system $\mathfrak{M}$ is not complete. Let its elements be $\varphi_1, \dots, \varphi_n$. Among the elements $a_1\varphi_1 + \dots + a_n\varphi_n$ there cannot be $n + 1$ that are linearly independent; hence, by condition [D](#), $\mathfrak{H}$ must contain an element f different from all of them. Then

$$f^* = f - a_1\varphi_1 - \dots - a_n\varphi_n$$

²⁸This excludes the possibility that $\mathfrak{H}$ be a finite-dimensional Euclidean space. Conditions [A](#) and [B](#) were, moreover, modeled on an axiomatic system for vector space given by H. Weyl for finitely many dimensions (*Raum, Zeit, Materie*, 5th ed., Berlin, 1923, §§ 2–4).

is certainly nonzero and is orthogonal to all $\varphi_1, \dots, \varphi_n$ if we set

$$a_k = (f, \varphi_k) \quad (k = 1, \dots, n).$$

By the concluding observation in [Definition 3](#), $\mathfrak{M}$ is therefore incomplete.

Theorem 4. *Let $\varphi_1, \varphi_2, \dots$ be an orthonormal system. Then the series*

$$\sum_{n=1}^{\infty} (f, \varphi_n) \overline{(g, \varphi_n)}$$

is absolutely convergent for all f, g in $\mathfrak{H}$; in particular, when $f = g$,

$$\sum_{n=1}^{\infty} |(f, \varphi_n)|^2 \leq |f|^2.$$

Proof. For $a_n = (f, \varphi_n)$, one has, as is well known,

$$\begin{aligned} \left| f - \sum_{n=1}^N a_n \varphi_n \right|^2 &= |f|^2 - 2\Re \sum_{n=1}^N (f, a_n \varphi_n) + \sum_{m,n=1}^N (a_m \varphi_m, a_n \varphi_n) \\ &= |f|^2 - 2\Re \sum_{n=1}^N \overline{a_n} (f, \varphi_n) + \sum_{m,n=1}^N a_m \overline{a_n} (\varphi_m, \varphi_n) \\ &= |f|^2 - 2 \sum_{n=1}^N |a_n|^2 + \sum_{n=1}^N |a_n|^2 \\ &= |f|^2 - \sum_{n=1}^N |a_n|^2. \end{aligned}$$

Since the left-hand side is always nonnegative,

$$\sum_{n=1}^N |a_n|^2 \leq |f|^2, \quad \sum_{n=1}^N |(f, \varphi_n)|^2 \leq |f|^2.$$

Consequently $\sum_{n=1}^{\infty} |(f, \varphi_n)|^2$ converges, absolutely, and is at most $|f|^2$. Since

$$\left| (f, \varphi_n) \overline{(g, \varphi_n)} \right| \leq \frac{1}{2} |(f, \varphi_n)|^2 + \frac{1}{2} |(g, \varphi_n)|^2,$$

the series $\sum_{n=1}^{\infty} (f, \varphi_n) \overline{(g, \varphi_n)}$ also converges absolutely. This proves everything.

Theorem 5. Let $\varphi_1, \varphi_2, \dots$ be an orthonormal system. The series

$$\sum_{n=1}^{\infty} a_n \varphi_n$$

converges if and only if

$$\sum_{n=1}^{\infty} |a_n|^2$$

converges.

Proof. By condition E, it suffices to show that the Cauchy convergence condition has the same form for the two series. This follows, for $m \geq n$, from

$$\begin{aligned} \left| \sum_{p=1}^m a_p \varphi_p - \sum_{p=1}^n a_p \varphi_p \right|^2 &= \left| \sum_{p=n+1}^m a_p \varphi_p \right|^2 \\ &= \sum_{p,q=n+1}^m (a_p \varphi_p, a_q \varphi_q) \\ &= \sum_{p,q=n+1}^m a_p \overline{a_q} (\varphi_p, \varphi_q) \\ &= \sum_{p=n+1}^m |a_p|^2 = \sum_{p=1}^m |a_p|^2 - \sum_{p=1}^n |a_p|^2. \end{aligned}$$

Corollary. If

$$f = \sum_{n=1}^{\infty} a_n \varphi_n$$

(in the case of convergence), then

$$(f, \varphi_n) = a_n.$$

Proof. For $N \geq n$,

$$\left(\sum_{p=1}^N a_p \varphi_p, \varphi_n \right) = \sum_{p=1}^N a_p (\varphi_p, \varphi_n) = a_n.$$

By continuity of the inner product, we may let $N \rightarrow \infty$, obtaining $(f, \varphi_n) = a_n$.^{T3}

^{T3}Translator's note: The source prints a_m in this final clause; a_n is clearly intended.

Theorem 6. Let $\varphi_1, \varphi_2, \dots$ be an orthonormal system. For every f , the series

$$f' = \sum_{n=1}^{\infty} a_n \varphi_n, \quad a_n = (f, \varphi_n) \quad (n = 1, 2, \dots),$$

converges, and $f - f'$ is orthogonal to all $\varphi_1, \varphi_2, \dots$

Proof. The series converges by [Theorems 4](#) and [5](#); and by the corollary to [Theorem 5](#),

$$(f', \varphi_n) = a_n = (f, \varphi_n), \quad (f - f', \varphi_n) = 0.$$

This proves everything.

Theorem 7. For an orthonormal system $\varphi_1, \varphi_2, \dots$ to be complete, each of the following three conditions is necessary and sufficient:

$\alpha)$ $\varphi_1, \varphi_2, \dots$ span the closed linear manifold $\mathfrak{H}$.

$\beta)$ For every f ,

$$f = \sum_{n=1}^{\infty} a_n \varphi_n, \quad a_n = (f, \varphi_n) \quad (n = 1, 2, \dots).$$

$\gamma)$ For every f, g ,

$$(f, g) = \sum_{n=1}^{\infty} (f, \varphi_n) \overline{(g, \varphi_n)}.$$

Proof. First, completeness implies $\beta)$: for the $f - f'$ of [Theorem 6](#) must vanish. Second, $\beta)$ implies $\alpha)$, because f is an accumulation point of

$$\sum_{n=1}^N a_n \varphi_n,$$

and hence of the linear manifold spanned by $\varphi_1, \varphi_2, \dots$; it is therefore an element of the corresponding closed linear manifold. Third, $\beta)$ implies $\gamma)$, since

$$\sum_{n=1}^N a_n \varphi_n \longrightarrow f, \quad \left| \sum_{n=1}^N a_n \varphi_n \right|^2 \longrightarrow |f|^2 \quad (N \rightarrow \infty),$$

while

$$\left| \sum_{n=1}^N a_n \varphi_n \right|^2 = \sum_{m,n=1}^N a_m \overline{a_n} (\varphi_m, \varphi_n) = \sum_{n=1}^N |a_n|^2.$$

Thus

$$\sum_{n=1}^{\infty} |a_n|^2 = \sum_{n=1}^{\infty} |(f, \varphi_n)|^2 = |f|^2.$$

If here we replace f by $(f + g)/2$ and $(f - g)/2$ and subtract, we obtain the real part of γ); if we replace f, g by if, g , we obtain its imaginary part.

Fourth, α) implies completeness. Otherwise there would be an $f \neq 0$ orthogonal to all $\varphi_1, \varphi_2, \dots$. The entire closed linear manifold spanned by $\varphi_1, \varphi_2, \dots$, that is, $\mathfrak{H}$, would then be orthogonal to f , and hence so would f itself:

$$|f|^2 = (f, f) = 0, \quad f = 0,$$

contrary to the assumption. Fifth, γ) also implies completeness; for the same f , condition γ), with $f = g$, would give

$$|f|^2 = (f, f) = \sum_{n=1}^{\infty} 0^2 = 0, \quad f = 0.$$

Thus we have the logical scheme

$$\begin{aligned} & \text{completeness} \implies \beta), \\ & \beta) \implies \alpha) \implies \text{completeness}, \quad \beta) \implies \gamma) \implies \text{completeness}. \end{aligned}$$

Consequently these four assertions are indeed equivalent.

Theorem 8. *For every sequence $f_1, f_2, \dots$ of arbitrary elements of $\mathfrak{H}$, there is an orthonormal system $\varphi_1, \varphi_2, \dots$ spanning the same linear manifold, and therefore also the same closed linear manifold. Either of the two sequences may already terminate after finitely many terms.*

Remark. *Choose $f_1, f_2, \dots$ everywhere dense, as condition C permits. It then plainly spans the closed linear manifold $\mathfrak{H}$. The same is true of the orthonormal system $\varphi_1, \varphi_2, \dots$ constructed by Theorem 8; hence, by condition α) of Theorem 7, it is complete. This proves the existence of complete orthonormal systems.*

Proof. First, we may replace $f_1, f_2, \dots$ by a subsequence consisting entirely of linearly independent elements and spanning the same linear manifold. Indeed, let g_1 be the first $f_n \neq 0$, let g_2 be the first f_n not of the form $a_1 g_1$ with a_1 arbitrary, let g_3 be the first f_n not of the form $a_1 g_1 + a_2 g_2$ with a_1, a_2 arbitrary, and so on. (If there is no f_n outside the set of all $a_1 g_1 + \dots + a_p g_p$, we stop with g_p .) The sequence $g_1, g_2, \dots$ plainly has the property just required.

We now apply E. Schmidt's "orthogonalization procedure" to $g_1, g_2, \dots$:

$$\begin{aligned} \gamma_1 &= g_1, & \varphi_1 &= \frac{1}{|\gamma_1|} \gamma_1, \\ \gamma_2 &= g_2 - (g_2, \varphi_1) \varphi_1, & \varphi_2 &= \frac{1}{|\gamma_2|} \gamma_2, \\ \gamma_3 &= g_3 - (g_3, \varphi_1) \varphi_1 - (g_3, \varphi_2) \varphi_2, & \varphi_3 &= \frac{1}{|\gamma_3|} \gamma_3, \\ &\dots\dots\dots \end{aligned}$$

The sequence $\varphi_1, \varphi_2, \dots$ plainly spans the same linear manifold as $g_1, g_2, \dots$, and hence as $f_1, f_2, \dots$; it is also immediately seen to form an orthonormal system.

On the basis of these facts, we can now show that $\mathfrak{H}$ can be mapped one-to-one onto ordinary Hilbert space, that is, the set of all sequences of complex numbers $(x_1, x_2, \dots)$ with finite $\sum_{n=1}^{\infty} |x_n|^2$, in such a way that, with $\longleftrightarrow$ denoting the correspondence,

$$\begin{aligned} f \longleftrightarrow (x_1, x_2, \dots) &\implies af \longleftrightarrow (ax_1, ax_2, \dots), \\ \left. \begin{aligned} f &\longleftrightarrow (x_1, x_2, \dots), \\ g &\longleftrightarrow (y_1, y_2, \dots) \end{aligned} \right\} &\implies f + g \longleftrightarrow (x_1 + y_1, x_2 + y_2, \dots), \end{aligned}$$

and, moreover,

$$(f, g) = \sum_{n=1}^{\infty} x_n \overline{y_n}.$$

Indeed, let $\varphi_1, \varphi_2, \dots$ be an arbitrary but fixed complete orthonormal system, and assign to f the sequence $(x_1, x_2, \dots)$ with

$$x_n = (f, \varphi_n) \quad (n = 1, 2, \dots).$$

That $(x_1, x_2, \dots)$ belongs to Hilbert space follows from [Theorem 4](#); that an f exists for every such $(x_1, x_2, \dots)$ follows from [Theorem 5](#) and its corollary. One-to-one correspondence follows from completeness: if $(f, \varphi_n) = (g, \varphi_n)$ for every n , then $f - g$ is orthogonal to every φ_n , and hence is 0. Of the three concluding assertions, the first two are clear, while the third follows from condition γ) of [Theorem 7](#).

Thus it has been shown that $\mathfrak{H}$ is isomorphic to the ordinary Hilbert space, and hence that any two systems $\mathfrak{H}$ satisfying A through E are isomorphic to one another; that is, A through E form a "categorical" system of axioms: they determine all the properties of $\mathfrak{H}$.

II. Operators in Hilbert Space

An operator is a function defined on a subset of $\mathfrak{H}$ (possibly on all of $\mathfrak{H}$) and taking values in $\mathfrak{H}$. We shall first define several especially important kinds of operators.

Definition 5. An operator R is linear (abbreviated: *lin.*) if, whenever $Rf_1, \dots, Rf_n$ are defined, $R(a_1f_1 + \dots + a_nf_n)$ is also defined, and in fact

$$R(a_1f_1 + \dots + a_nf_n) = a_1Rf_1 + \dots + a_nRf_n.$$

An operator R is closed (abbreviated: *cl.*) if, for every sequence $f_1, f_2, \dots$ for which all Rf_n are defined, the convergences

$$f_n \longrightarrow f, \quad Rf_n \longrightarrow f^*$$

imply that Rf is defined and equals f^* .²⁹

Definition 6. An operator R is Hermitian (abbreviated: a H.O.) if, first, its domain spans $\mathfrak{H}$ as a closed linear manifold,³⁰ and, second, whenever Rf and Rg are defined,

$$(f, Rg) = (Rf, g).$$

An operator R is length-preserving if it is closed and linear and, whenever Rf is defined,

$$|f| = |Rf|.$$

We begin by proving two theorems about Hermitian and length-preserving operators. We also recall that an operator S is called an *extension* (abbreviated: *ext.*) of an operator R if, wherever Rf is defined, Sf is also defined and equals Rf ; it is called a *proper extension* (abbreviated: *prop. ext.*) if $R \neq S$, that is, if S is defined at more points than R .

Theorem 9. Let R be a H.O. There exists a linear H.O. $\widehat{R}$ that is an extension of R and of which every H.O. with these same properties is in turn an extension; and there exists a closed linear H.O. $\widetilde{R}$ that is an extension of R (and hence also of $\widehat{R}$) and of which every H.O. with these same properties is in turn an extension.

²⁹This is evidently a continuity-like property; in every “compact” topological space it is in fact equivalent to continuity. In Hilbert space, however, it means much less: for Hermitian operators it is essentially always present; cf. [Theorem 9](#).

³⁰To this extent we weaken condition γ) of [Introduction V](#) (with a view to the application to matrices; cf. [Appendix III](#)): there, density of the domain was required.

Proof. If we could define $\widehat{R}$ as follows: $\widehat{R}f$ is defined whenever

$$f = a_1 f_1 + \cdots + a_n f_n$$

and $Rf_1, \dots, Rf_n$ are defined, and its value is then

$$a_1 Rf_1 + \cdots + a_n Rf_n;$$

and define $\widetilde{R}$ as follows: $\widetilde{R}f$ is defined whenever there exists a sequence $g_1, g_2, \dots$ for which all $\widehat{R}g_n$ are defined and

$$g_n \longrightarrow f, \quad \widehat{R}g_n \longrightarrow f^*,$$

and its value is then f^* —then these $\widehat{R}$ and $\widetilde{R}$ would plainly have all the desired properties. The question is whether these prescriptions are consistent.

They certainly are if

$$a_1 f_1 + \cdots + a_m f_m = b_1 g_1 + \cdots + b_n g_n,$$

where $Rf_1, \dots, Rf_m, Rg_1, \dots, Rg_n$ are defined, implies

$$a_1 Rf_1 + \cdots + a_m Rf_m = b_1 Rg_1 + \cdots + b_n Rg_n,$$

and if

$$f_n \longrightarrow f, \quad g_n \longrightarrow f,$$

where all $\widehat{R}f_n, \widehat{R}g_n$ are defined, together with

$$\widehat{R}f_n \longrightarrow f^*, \quad \widehat{R}g_n \longrightarrow g^*,$$

implies $f^* = g^*$. It plainly suffices to prove the following two assertions:

$$a_1 f_1 + \cdots + a_n f_n = 0 \quad (Rf_1, \dots, Rf_n \text{ defined})$$

implies

$$a_1 Rf_1 + \cdots + a_n Rf_n = 0,$$

and

$$f_n \longrightarrow 0, \quad \widehat{R}f_n \longrightarrow f^* \quad (\widehat{R}f_n \text{ all defined})$$

implies $f^* = 0$.

Suppose first that the premise of the first assertion holds. Whenever Rg is defined,

$$\begin{aligned} (g, a_1 R f_1 + \cdots + a_n R f_n) &= \sum_{p=1}^n (g, a_p R f_p) \\ &= \sum_{p=1}^n (Rg, a_p f_p) \\ &= (Rg, a_1 f_1 + \cdots + a_n f_n) = 0. \end{aligned}$$

Thus all such g are orthogonal to $a_1 R f_1 + \cdots + a_n R f_n$, and so is their closed linear span, $\mathfrak{H}$. Hence $a_1 R f_1 + \cdots + a_n R f_n$ is orthogonal to itself and therefore is 0.

Now consider the premise of the second assertion. Whenever $\widehat{R}g$ is defined, since $\widehat{R}$ is a H.O.,

$$\begin{aligned} (g, f^*) &= \left(g, \lim_{n \rightarrow \infty} \widehat{R} f_n \right) = \lim_{n \rightarrow \infty} (g, \widehat{R} f_n) \\ &= \lim_{n \rightarrow \infty} (\widehat{R} g, f_n) = \left(\widehat{R} g, \lim_{n \rightarrow \infty} f_n \right) = 0. \end{aligned}$$

Thus all such g are orthogonal to f^* , and so is their closed linear span, $\mathfrak{H}$; this again implies $f^* = 0$.

Theorem 10. *Let U be a length-preserving operator. Then its domain and its range are both closed linear manifolds, and U maps them continuously and one-to-one onto one another. Wherever U is defined,*

$$(f, g) = (Uf, Ug).$$

Remark. *If both the domain and the range equal $\mathfrak{H}$, then, in the customary terminology, U is called unitary.*

Proof. Since U is linear, its domain and range are linear manifolds. Moreover,

$$|Uf - Ug| = |U(f - g)| = |f - g|.$$

Thus $Uf = Ug$ implies $f = g$, so the mapping is one-to-one; the continuity of both the mapping and its inverse follows as well. Together with the closedness of U , this implies that its domain and range are both closed sets.

In $|f|^2 = |Uf|^2$, replace f successively by $(f + g)/2$ and $(f - g)/2$ and subtract. This gives

$$\operatorname{Re}(f, g) = \operatorname{Re}(Uf, Ug).$$

Replacing f, g by if, g gives the same result for the imaginary parts. Therefore

$$(f, g) = (Uf, Ug).$$

This proves everything.

Several properties of $\widetilde{R}$ (for R a H.O.) are evident:

$$\widetilde{\widetilde{R}} = \widetilde{R};$$

if S is an extension of R , then $\widetilde{S}$ is an extension of $\widetilde{R}$. We now define maximality for Hermitian operators.

Definition 7. A H.O. R is called maximal (abbreviated: *max.*) if it has no proper extension that is also a H.O.

Since $\widetilde{R}$ is an extension of R , a maximal R must satisfy $R = \widetilde{R}$. Furthermore, let R be a H.O., let S be a maximal H.O., and let S be an extension of R ; then $S = \widetilde{S}$ is also an extension of $\widetilde{R}$. Thus every maximal H.O. is closed and linear, and a H.O. R has exactly the same maximal extensions as the closed linear H.O. $\widetilde{R}$. In later chapters we shall investigate the extension of arbitrary Hermitian operators to maximal ones (cf. also [Introduction VI](#) and [VII](#)); what has just been said shows that throughout this investigation we may restrict ourselves to closed linear operators.

We define further:

Definition 8. Let R be a H.O. The vector f is an extension element of R , and f^* is assigned to f by R , if

$$(f, Rg) = (f^*, g)$$

for every g for which Rg is defined. The pair f, f^* belongs to the $+-$, $0-$, or $--$ -class according as

$$\operatorname{Im}(f^*, f) > 0, \quad = 0, \quad < 0.$$

Theorem 11. The vector f^* is uniquely determined by f (cf. [Definition 8](#)). If Rf is defined, then f is an extension element of the 0 -class, and $f^* = Rf$. The operator R is maximal if and only if there are no further extension elements of the 0 -class.

Proof. If both f^* and f^{**} were assigned to f , then

$$(f^*, g) = (f^{**}, g), \quad (f^* - f^{**}, g) = 0$$

for every g for which Rg is defined, and therefore also for the entire closed linear span of such g , namely $\mathfrak{H}$. It follows once again that $f^* - f^{**} = 0$. The second

assertion is clear. For the third, the sufficiency of the condition is clear, while its necessity follows because any extension element of the 0-class could be used to form a proper extension of R .

The characterization of maximality given by [Theorem 11](#) suggests the following concept.

Definition 9. A H.O. R is called hypermaximal (abbreviated: hypermax.) if it has no extension elements other than those for which R is defined. (That is, if R is maximal and the $+-$ and $--$ -classes are empty.)

The importance of this concept³¹ will appear later; cf., for example, [Theorem 36](#).

We prove one more theorem concerning the continuity of linear operators.

Theorem 12. For a linear operator R , each of the following assertions is equivalent to continuity:

$$\alpha) \quad |Rf| \leq C|f|,$$

$$\beta) \quad |(f, Rg)| \leq C|f||g|.$$

If R is a H.O., the following is equivalent as well:

$$\gamma) \quad |(f, Rf)| \leq C|f|^2.^{32}$$

Remark. Following Hilbert, for linear Hermitian operators we also say bounded instead of continuous. We do not investigate continuity for nonlinear operators. (Hilbert's definition is β .)

Proof. By linearity, α) gives

$$|Rf - Rg| = |R(f - g)| \leq C|f - g|,$$

and hence continuity. In β) we need only put $f = Rg$ to obtain

$$|Rg|^2 \leq C|Rg||g|, \quad |Rg| \leq C|g|,$$

which is α). If R is continuous, then, because $R0 = 0$ by linearity, there is an $\varepsilon > 0$ such that $|f| \leq \varepsilon$ implies $|Rf| \leq 1$. Hence, again by linearity,

$$|Rf| \leq \frac{1}{\varepsilon}|f|,$$

³¹It is due to E. Schmidt; cf. [note 23](#). When one restricts oneself to real Hilbert space, the distinction between maximal and hypermaximal does not arise.

³²For a H.O. every (f, Rf) must be real, since

$$(f, Rf) = (Rf, f) = \overline{(f, Rf)}.$$

and therefore

$$|(f, Rg)| \leq |f||Rg| \leq \frac{1}{\varepsilon}|f||g|;$$

thus β) holds. Consequently α) and β) really are equivalent to continuity.

Plainly γ) says less than β), so for a H.O. it is enough to derive β) from γ). This is done as in [Theorem 1](#). Substitute $(f + g)/2$ and $(f - g)/2$ for f and subtract; the result is

$$\operatorname{Re}(f, Rg) \leq C \left(\frac{1}{2}|f|^2 + \frac{1}{2}|g|^2 \right).$$

As in [Theorem 1](#), replace f, g by $af, a^{-1}g$ with $a > 0$ and take the minimum of the right-hand side; then replace f, g by $\vartheta f, g$ with $|\vartheta| = 1$ and take the maximum of the left-hand side. We obtain

$$|(f, Rg)| \leq C|f||g|,$$

that is, β). Here, however, Rf and Rg were assumed to be defined. We can dispense with the definedness of Rf : Rf does not occur in β), and the inequality holds for every f for which Rf is defined, hence on an everywhere-dense set (because R is a linear H.O.); since both sides are continuous, it holds for every f whatsoever.

For a linear H.O. the condition for boundedness is

$$-C|f|^2 \leq (f, Rf) \leq C|f|^2$$

([Theorem 12](#), γ)). We split this condition in the usual way:

Definition 10. A linear H.O. R is called semibounded above, respectively semibounded below, if always

$$(f, Rf) \leq C|f|^2, \quad \text{respectively} \quad (f, Rf) \geq -C|f|^2.$$

(Together, the two conditions are equivalent to boundedness by [Theorem 12](#), γ).) If, in particular,

$$(f, Rf) \geq 0,$$

then R is called definite.

III. Linear Manifolds and Projection Operators

Definition 11. Let $\mathfrak{M}, \mathfrak{N}$ be two linear manifolds, not necessarily closed. Then $\mathfrak{M} + \mathfrak{N}$ and $\mathfrak{M} \times \mathfrak{N}$, respectively the linear manifold spanned by their union and their intersection, are again linear manifolds. If, in particular, $\mathfrak{M}, \mathfrak{N}$ are closed linear manifolds, then $\mathfrak{M} \times \mathfrak{N}$ plainly is one as well. (Whether $\mathfrak{M} + \mathfrak{N}$ is also closed is uncertain; for orthogonal $\mathfrak{M}, \mathfrak{N}$ it follows from [Theorem 17](#).) In what follows, $\mathfrak{M}, \mathfrak{N}$ are to be assumed closed linear manifolds. If $\mathfrak{M}$ is a subset of $\mathfrak{N}$, then their difference $\mathfrak{N} - \mathfrak{M}$, that is, the set of all elements of $\mathfrak{N}$ orthogonal to every element of $\mathfrak{M}$, is also a closed linear manifold. The manifold $\mathfrak{H} - \mathfrak{M}$ is also called the complement of $\mathfrak{M}$.

Theorem 13. Let $\mathfrak{M}$ be a closed linear manifold. Every f can be decomposed in one and only one way as

$$f = g + h, \quad g \in \mathfrak{M}, \quad h \in \mathfrak{H} - \mathfrak{M}.$$

Remark. The vector g is then called the projection of f into $\mathfrak{M}$. The assignment of g to f is plainly an operator defined everywhere; we call it the projection operator of $\mathfrak{M}$ and denote it by $P_{\mathfrak{M}}$ (where $\mathfrak{M}$ is a closed linear manifold).

Proof. Uniqueness is clear. From

$$g_1 + h_1 = g_2 + h_2$$

with $g_1, g_2 \in \mathfrak{M}$ and $h_1, h_2 \in \mathfrak{H} - \mathfrak{M}$, it follows that

$$g_1 - g_2 = h_2 - h_1.$$

This vector belongs both to $\mathfrak{M}$ and to $\mathfrak{H} - \mathfrak{M}$, and is therefore orthogonal to itself and hence 0; thus $g_1 = g_2$ and $h_1 = h_2$. It remains to prove existence.

Since $\mathfrak{H}$ is separable, there is a sequence $f_1, f_2, \dots$ dense in $\mathfrak{M}$,³³ and this sequence therefore spans the closed linear manifold $\mathfrak{M}$. By [Theorem 8](#), there is a normalized orthogonal system $\varphi_1, \varphi_2, \dots$ that does the same. By [Theorem 6](#),

$$g = \sum_{n=1}^{\infty} (f, \varphi_n) \varphi_n$$

converges; together with $\varphi_1, \varphi_2, \dots$, its value belongs to $\mathfrak{M}$. The vector $h = f - g$ is orthogonal to every $\varphi_1, \varphi_2, \dots$, and therefore also to their closed linear span $\mathfrak{M}$; hence $h \in \mathfrak{H} - \mathfrak{M}$. This proves everything.

³³Every subset of a separable topological space is itself separable; see, for example, F. Hausdorff, *Einführung in die Mengenlehre*, Leipzig–Berlin, 1914.

Theorem 14. *The projection operator $P_{\mathfrak{M}}$ is an everywhere-defined maximal H.O., and*

$$P_{\mathfrak{M}}^2 = P_{\mathfrak{M}}.$$

Proof. We know that $P_{\mathfrak{M}}$ is defined everywhere. It is a H.O. if we can prove

$$(f, P_{\mathfrak{M}}g) = (P_{\mathfrak{M}}f, g).$$

Writing

$$f = h + k, \quad g = h' + k',$$

where $h, h' \in \mathfrak{M}$ and $k, k' \in \mathfrak{H} - \mathfrak{M}$, this is the assertion $(f, h') = (h, g)$. But $(k, h') = 0$ and $(h, k') = 0$, so

$$(f, h') = (h + k, h') = (h, h') = (h, h' + k') = (h, g).$$

If $f \in \mathfrak{M}$, the decomposition of [Theorem 13](#) is $f = f + 0$, and hence $P_{\mathfrak{M}}f = f$. Since $P_{\mathfrak{M}}f \in \mathfrak{M}$, it follows that

$$P_{\mathfrak{M}}(P_{\mathfrak{M}}f) = P_{\mathfrak{M}}f,$$

that is, $P_{\mathfrak{M}}^2 = P_{\mathfrak{M}}$. Since $P_{\mathfrak{M}}$ is defined everywhere, it must be maximal.

The maximal Hermitian operators E having the property $E^2 = E$ are, however, of independent interest to us (they correspond to Hilbert's "single forms"); we therefore first investigate them in their own right.

Theorem 15. *For every maximal H.O. E with the property*

$$E^2 = E, \quad {}^{34}$$

there is a closed linear manifold $\mathfrak{M}$ such that

$$P_{\mathfrak{M}} = E.$$

The manifold $\mathfrak{M}$ is the range of E .

Proof. Since $\mathfrak{M}$ is plainly the range of $P_{\mathfrak{M}}$, it suffices to show that $P_{\mathfrak{M}} = E$ when $\mathfrak{M}$ is the range of E .

Because $E^2 = E$, as is easily seen, $\mathfrak{M}$ is the set of all f satisfying $Ef = f$; thus it is a closed linear manifold (since E is maximal, it is also linear and closed).

³⁴That is, if Ef is defined, then $E(Ef)$ is also defined and equals Ef .

Whenever Ef is defined, $f - Ef$ is orthogonal to every element of $\mathfrak{M}$, that is, to every Eg :

$$(Eg, f - Ef) = (g, Ef - E^2f) = 0.$$

Hence $f - Ef \in \mathfrak{S} - \mathfrak{M}$. It follows that $P_{\mathfrak{M}}f = Ef$, so $P_{\mathfrak{M}}$ is an extension of E . Since E is maximal, therefore $P_{\mathfrak{M}} = E$.

Definition 12. An operator satisfying the conditions of [Theorem 15](#) is called a projection operator (abbreviated: P.O.; by [Theorem 15](#) the P.O.'s are identical with the $P_{\mathfrak{M}}$ and correspond one-to-one to the closed linear manifolds $\mathfrak{M}$). A closed linear manifold $\mathfrak{M}$ and its P.O. $P_{\mathfrak{M}}$ are called corresponding. It is immediately seen that the P.O.'s corresponding to the closed linear manifolds (0) ³⁵ and $\mathfrak{S}$ are respectively 0 and 1,³⁶ and furthermore

$$P_{\mathfrak{S}-\mathfrak{M}} = 1 - P_{\mathfrak{M}}.$$

(Thus 0 and 1 are P.O.'s, and if E is one, then so is $1 - E$.)

Theorem 16. A P.O. E is always bounded and definite; in fact,

$$(f, Ef) = |Ef|^2, \quad |Ef| \leq |f|.$$

Proof. We have

$$(f, Ef) = (f, E^2f) = (Ef, Ef) = |Ef|^2,$$

and hence $(f, Ef) \geq 0$. Apply this to the P.O. $1 - E$:

$$(f, (1 - E)f) \geq 0, \quad (f, Ef) \leq (f, f) = |f|^2.$$

Thus $|Ef|^2 \leq |f|^2$, and hence $|Ef| \leq |f|$. This proves the displayed formulas, which in turn imply boundedness and definiteness.

Theorem 17. Let $\mathfrak{M}, \mathfrak{N}$ be two closed linear manifolds, and let E, F be their corresponding P.O.'s. Then EF is a P.O. if and only if

$$EF = FE,$$

that is, if E and F commute; in that case we also call $\mathfrak{M}, \mathfrak{N}$ commuting, and EF is the P.O. of $\mathfrak{M} \times \mathfrak{N}$.

³⁵The set (0) is the set whose sole element is the zero vector of $\mathfrak{S}$. We do not regard the empty set as a linear manifold.

³⁶The symbols 0 and 1 denote two everywhere-defined operators, defined by $0f = 0$ and $1f = f$.

The operator $E + F$ is a P.O. if and only if

$$EF = 0$$

(or, equivalently, if $FE = 0$); this means that all of $\mathfrak{M}$ is orthogonal to all of $\mathfrak{N}$. We then also call E, F orthogonal, and $E + F$ is the P.O. of $\mathfrak{M} + \mathfrak{N}$.

The operator $E - F$ is a P.O. if and only if

$$EF = F$$

(or, equivalently, if $FE = F$); this means that $\mathfrak{N}$ is a subset of $\mathfrak{M}$. We then also call F a part of E , and $E - F$ is the P.O. of $\mathfrak{M} - \mathfrak{N}$.³⁷

Proof. The operator EF is defined everywhere, and hence is a maximal H.O. if it satisfies the symmetry condition

$$(EFf, g) = (f, EFg).$$

Now

$$(f, EFg) = (Ef, Fg) = (FEf, g),$$

and therefore

$$(EFf, g) - (f, EFg) = ((EF - FE)f, g).$$

For this to vanish for every g , one must have $(EF - FE)f = 0$ for every f , and hence $EF - FE = 0$. Thus our condition is already necessary and sufficient for EF to be a maximal H.O. Since it implies

$$(EF)^2 = EF EF = EE FF = E^2 F^2 = EF,$$

it is also necessary and sufficient for EF to be a P.O.

For $E + F$, the H.O. property is clear, as is maximality (the operator is defined everywhere); it therefore suffices to examine when $(E + F)^2 = E + F$. Since

$$(E + F)^2 = E^2 + F^2 + EF + FE = (E + F) + (EF + FE),$$

this condition says that $EF + FE = 0$. But this implies $EF = 0$: multiply on the left by E to obtain

$$EF + EFE = 0,$$

³⁷Since these operators are defined everywhere, we may form EF and $E \pm F$ without hesitation; otherwise their domains would have to be examined more precisely.

then multiply this equality on the right by E to obtain

$$2EFE = 0;$$

together these give $EF = 0$.³⁸ Conversely, if $EF = 0$, then EF is a P.O., and hence $FE = EF = 0$ and $EF + FE = 0$. Thus $EF = 0$ is indeed characteristic; interchanging E and F shows that $FE = 0$ is equivalent.

The operator $E - F$ is a P.O. if and only if

$$1 - (E - F) = (1 - E) + F$$

is one. Since $1 - E$ and F are P.O.'s, this says

$$(1 - E)F = 0, \quad F - EF = 0,$$

or, equivalently,

$$F(1 - E) = 0, \quad F - FE = 0.$$

It remains to prove the assertions concerning $\mathfrak{M}$, $\mathfrak{N}$, $\mathfrak{M} \times \mathfrak{N}$, and $\mathfrak{M} \pm \mathfrak{N}$. First suppose $EF = FE$. Every $EFf = FEf$ belongs to both $\mathfrak{M}$ and $\mathfrak{N}$, hence to $\mathfrak{M} \times \mathfrak{N}$. Conversely, if $f \in \mathfrak{M} \times \mathfrak{N}$, then $f \in \mathfrak{M}$ and $f \in \mathfrak{N}$, so

$$EFf = Ef = f.$$

Thus $\mathfrak{M} \times \mathfrak{N}$ is the range of EF , and EF is its P.O.

Second suppose $EF = 0$ (and hence $FE = 0$). If $f \in \mathfrak{M}$ and $g \in \mathfrak{N}$, then

$$(f, g) = (Ef, Fg) = (f, EFg) = 0;$$

thus $\mathfrak{M}$ is orthogonal to $\mathfrak{N}$, and the converse is equally clear. For every $f, Ef \in \mathfrak{M}$ and $Ff \in \mathfrak{N}$. Hence $(E + F)f \in \mathfrak{M} + \mathfrak{N}$. If $f \in \mathfrak{M} + \mathfrak{N}$, then $f = g + h$ with $g \in \mathfrak{M}$ and $h \in \mathfrak{N}$, and therefore

$$\begin{aligned} (E + F)f &= (E + F)(g + h) \\ &= Eg + Fg + Eh + Fh \\ &= g + 0 + 0 + h = f. \end{aligned}$$

Thus $\mathfrak{M} + \mathfrak{N}$ is the range of $E + F$, and is therefore a closed linear manifold, while $E + F$ is its P.O. by [Theorem 15](#).

³⁸This device is due to E. Schmidt.

Third suppose $EF = F$ (and hence $FE = F$). This says that $EFf = Ff$ for every f , or that $Eg = g$ for every $g \in \mathfrak{N}$. Since this is the defining equation of $\mathfrak{M}$ (cf. the proof of [Theorem 15](#)), it says that $\mathfrak{N}$ is a subset of $\mathfrak{M}$. Since

$$E - F = E - EF = E(1 - F),$$

$E - F$ is the P.O. of

$$\mathfrak{M} \times (\mathfrak{S} - \mathfrak{N}) = \mathfrak{M} - \mathfrak{N}.$$

All our assertions are now proved.

Theorem 18. *The operator E is a part of F if and only if*

$$|Ef| \leq |Ff|$$

for every f . The sum $E_1 + \cdots + E_p$ is a P.O. (where $E_1, \dots, E_p$ are P.O.'s) if and only if

$$E_mE_n = 0 \quad (m \neq n; m, n = 1, \dots, p).$$

Proof. If E is a part of F , then $E = EF$, and hence

$$|Ef| = |EFf| \leq |Ff|.$$

Conversely, suppose this inequality holds for every f . If $f = Ef$, then

$$|Ff| \geq |Ef| = |f|, \quad (f, Ff) = |Ff|^2 \geq |f|^2 = (f, f),$$

and consequently

$$|(1 - F)f|^2 = (f, (1 - F)f) \leq 0, \quad (1 - F)f = 0.$$

Thus $f = Ff$. Therefore, if E, F are the P.O.'s of $\mathfrak{M}, \mathfrak{N}$, then $\mathfrak{M}$ is a part of $\mathfrak{N}$, and hence E is a part of F .

The sufficiency of $E_mE_n = 0$ for $m \neq n$ for the P.O. character of $E_1 + \cdots + E_p$ follows by repeated application of [Theorem 17](#). Its necessity is seen as follows. For $m \neq n$,

$$\begin{aligned} |E_m f|^2 + |E_n f|^2 &\leq |E_1 f|^2 + \cdots + |E_p f|^2 \\ &= (f, E_1 f) + \cdots + (f, E_p f) \\ &= (f, (E_1 + \cdots + E_p)f) \leq |f|^2. \end{aligned}$$

Thus, if $E_n f = f$, then $|E_m f|^2 \leq 0$, so $E_m f = 0$. Taking $f = E_n g$, we obtain $E_m E_n g = 0$ for every g , that is, $E_m E_n = 0$, as asserted.

We conclude with a theorem on the convergence of projection operators.

Theorem 19. *Let $E_1, E_2, \dots$ be a sequence of P.O.'s, and suppose either that E_m is a part of E_n for every $m < n$, or that E_n is a part of E_m for every $m < n$. Then there is a P.O. E such that, for every f ,*

$$E_n f \longrightarrow E f.$$

Proof. If we replace $E_1, E_2, \dots$ by $1 - E_1, 1 - E_2, \dots$, the first case becomes the second,³⁹ so it suffices to consider the first case.

As n increases, $|E_n f|^2$ increases by [Theorem 18](#), while it remains at most $|f|^2$ by [Theorem 16](#). Hence, for every $\varepsilon > 0$, there is an $N = N(\varepsilon)$ such that, for all $m, n \geq N$,

$$||E_n f|^2 - |E_m f|^2| \leq \varepsilon.$$

If also $m < n$, then $E_n - E_m$ is a P.O., and therefore

$$\begin{aligned} |E_n f - E_m f|^2 &= (f, (E_n - E_m)f) \\ &= (f, E_n f) - (f, E_m f) \\ &= |E_n f|^2 - |E_m f|^2 \leq \varepsilon. \end{aligned}$$

Thus the sequence $E_1 f, E_2 f, \dots$ has a limit; denote it by $E f$.

The operator E is defined everywhere. Since

$$(f, E g) = \lim_{n \rightarrow \infty} (f, E_n g) = \lim_{n \rightarrow \infty} (E_n f, g) = (E f, g),$$

it is a maximal H.O. We need only prove that $E^2 = E$. We have

$$\begin{aligned} (f, E^2 g) &= (E f, E g) = \lim_{n \rightarrow \infty} (E_n f, E_n g) \\ &= \lim_{n \rightarrow \infty} (f, E_n^2 g) = \lim_{n \rightarrow \infty} (f, E_n g) = (f, E g). \end{aligned}$$

Thus indeed $E^2 = E$.

³⁹For if E is a part of F , then $1 - F$ is a part of $1 - E$, for example because

$$(1 - E)(1 - F) = 1 - E - F + EF = 1 - F.$$

IV. Reducibility of Operators

For finite-dimensional matrices, reducibility is customarily understood to mean the possibility of applying a unitary transformation to the matrix so as to bring it into a form in which—if, say, it is $N = M_1 + M_2$ dimensional—only entries lying simultaneously in the first M_1 rows and the first M_1 columns, or in the last M_2 rows and the last M_2 columns, may be nonzero. If the matrix is regarded as a linear transformation, this means that the space can be decomposed into two mutually orthogonal linear manifolds of dimensions M_1 and M_2 , each of which is carried by the transformation into a part of itself.

For operators in $\mathfrak{H}$ we formulate the analogue as follows, restricting ourselves to closed linear operators.

Definition 13. *A closed linear operator R is said to be reduced by the closed linear manifold $\mathfrak{M}$ if, whenever Rf is defined, $RP_{\mathfrak{M}}f$ is also defined and*

$$RP_{\mathfrak{M}}f = P_{\mathfrak{M}}Rf. \text{ }^{40}$$

Every R is plainly reduced by (0) and by $\mathfrak{H}$; and if $\mathfrak{M}$ reduces R , then so does $\mathfrak{H} - \mathfrak{M}$, since

$$P_{\mathfrak{H}-\mathfrak{M}} = 1 - P_{\mathfrak{M}}.$$

The operator R is called *irreducible* if it is reduced only by (0) and $\mathfrak{H}$.

If R is reduced by $\mathfrak{M}$, then for every $f \in \mathfrak{M}$, Rf , whenever defined, plainly also belongs to $\mathfrak{M}$, since

$$Rf = RP_{\mathfrak{M}}f = P_{\mathfrak{M}}Rf.$$

Thus R also gives rise to an operator in $\mathfrak{M}$ (that is, a function whose domain and range are subsets of $\mathfrak{M}$). We call this operator the *part of R lying in $\mathfrak{M}$* .

We now prove two theorems on the reduction of operators.

Theorem 20. *Let R be a closed linear operator and $\mathfrak{M}$ a closed linear manifold. In order that $\mathfrak{M}$ reduce R , it is necessary that whenever Rf is defined, $RP_{\mathfrak{M}}f$ also be defined, and that whenever $f \in \mathfrak{M}$, Rf , if defined, also belong to $\mathfrak{M}$. If R is a H.O., these conditions are also sufficient.* ⁴¹

⁴⁰This may also be understood as the commutativity of R with $P_{\mathfrak{M}}$.

⁴¹In finite-dimensional spaces these conditions are also sufficient for unitary operators; in $\mathfrak{H}$ they are not, but that is immaterial here.

Proof. We have already established necessity. It remains to show that the conditions are sufficient for a Hermitian R , that is, to prove

$$RP_{\mathfrak{M}}f = P_{\mathfrak{M}}Rf.$$

By assumption $RP_{\mathfrak{M}}f \in \mathfrak{M}$, so it remains to show that

$$Rf - RP_{\mathfrak{M}}f = R(f - P_{\mathfrak{M}}f)$$

belongs to $\mathfrak{S} - \mathfrak{M}$. Since $f - P_{\mathfrak{M}}f \in \mathfrak{S} - \mathfrak{M}$, it is enough to show that if $g \in \mathfrak{S} - \mathfrak{M}$, then Rg , whenever defined, also belongs to $\mathfrak{S} - \mathfrak{M}$.

Let Rf be defined. Then $RP_{\mathfrak{M}}f \in \mathfrak{M}$, and hence

$$(f, P_{\mathfrak{M}}Rg) = (P_{\mathfrak{M}}f, Rg) = (RP_{\mathfrak{M}}f, g) = 0.$$

The set of such f is everywhere dense; hence $P_{\mathfrak{M}}Rg$ is orthogonal to everything and is therefore 0. Thus $Rg \in \mathfrak{S} - \mathfrak{M}$.

Theorem 21. *Let $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ be a possibly finite sequence of pairwise orthogonal closed linear manifolds which together span the closed linear manifold $\mathfrak{S}$. Let R be a closed linear operator reduced by each of them, and let the parts of R lying in $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ be $R_1, R_2, \dots$, respectively.*

Then R is determined by $R_1, R_2, \dots$ as follows: Rf is defined if and only if

$$f = f_1 + f_2 + \dots, \quad f_p \in \mathfrak{M}_p,$$

all $R_1f_1, R_2f_2, \dots$ are defined, and both series

$$f_1 + f_2 + \dots, \quad R_1f_1 + R_2f_2 + \dots$$

converge (if the sequence $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ terminates, the convergence assumption of course disappears). Its value is then

$$Rf = R_1f_1 + R_2f_2 + \dots.$$

Proof. That R is defined at the points just described and takes the indicated values follows immediately from the definitions of $R_1, R_2, \dots$, together with the fact that R is closed and linear and is reduced by $\mathfrak{M}_1, \mathfrak{M}_2, \dots$. We still have to prove that, if Rf is defined, then f has the stated form.

Put

$$f_p = P_{\mathfrak{M}_p}f \quad (p = 1, 2, \dots).$$

Then $f_p \in \mathfrak{M}_p$, and the $P_{\mathfrak{M}_p}$, like the $\mathfrak{M}_p$, are pairwise orthogonal. Moreover,

$$R_p f_p = R f_p = R P_{\mathfrak{M}_p} f = P_{\mathfrak{M}_p} R f.$$

It follows from the orthogonality of the $P_{\mathfrak{M}_p}$ that

$$P_{\mathfrak{M}_1}, \quad P_{\mathfrak{M}_1} + P_{\mathfrak{M}_2}, \quad P_{\mathfrak{M}_1} + P_{\mathfrak{M}_2} + P_{\mathfrak{M}_3}, \quad \dots$$

form an increasing sequence of P.O.'s, which by [Theorem 19](#) has a P.O. E as its limit. By continuity, every $\mathfrak{M}_p$ is a part of the closed linear manifold $\mathfrak{M}$ corresponding to E . By the assumption on $\mathfrak{M}_1, \mathfrak{M}_2, \dots$, therefore $\mathfrak{M} = \mathfrak{H}$ and $E = 1$.

Thus the series

$$f_1 + f_2 + \dots, \quad R_1 f_1 + R_2 f_2 + \dots$$

converge, with respective sums

$$1f = f, \quad 1Rf = Rf.$$

This proves everything.

V. The Cayley Transform

After these preparations we can take up our actual task, the closer investigation of the various kinds of Hermitian operators. The basis of the following considerations is a device that transfers the Cayley transform of algebra to our setting.

Let R be a closed linear H.O., and let $\mathfrak{A}$ be its domain. The set $\mathfrak{A}$ is a linear manifold and is everywhere dense (and is therefore closed only when $\mathfrak{A} = \mathfrak{H}$, that is, when R is defined everywhere, which is by no means assumed). Consider the two operators $R \pm i1$, which are also defined on $\mathfrak{A}$. A direct calculation gives

$$\begin{aligned} ((R \pm i1)f, (R \pm i1)g) &= (Rf, Rg) \pm (Rf, ig) \pm (if, Rg) + (f, g) \\ &= (Rf, Rg) + (f, g). \end{aligned}$$

In particular, for $f = g$,

$$|(R \pm i1)f| = |(R - i1)f| = \sqrt{|Rf|^2 + |f|^2}.$$

Thus $f \neq 0$ implies both $(R + i1)f \neq 0$ and $(R - i1)f \neq 0$, and by linearity $f \neq g$ implies

$$(R + i1)f \neq (R + i1)g, \quad (R - i1)f \neq (R - i1)g.$$

Suppose $R + i1$ and $R - i1$ map $\mathfrak{A}$ onto $\mathfrak{E}$ and $\mathfrak{F}$, respectively. These mappings among $\mathfrak{A}, \mathfrak{E}, \mathfrak{F}$ are therefore all one-to-one. In particular, denote the mapping from $\mathfrak{E}$ onto $\mathfrak{F}$ by U . Since $R \pm i1$, like R , are closed and linear, the same is true of U ; and the preceding equality gives

$$|Uf| = |f|.$$

Thus U is length-preserving by [Definition 6](#). Since $\mathfrak{E}, \mathfrak{F}$ are its domain and range, both are closed linear manifolds by [Theorem 10](#).

Theorem 22. *Let R be a closed linear H.O. with domain $\mathfrak{A}$, and suppose the operators $R + i1$ and $R - i1$ map $\mathfrak{A}$ onto $\mathfrak{E}$ and $\mathfrak{F}$, respectively. These mappings among $\mathfrak{A}, \mathfrak{E}, \mathfrak{F}$ are one-to-one; in particular, the mapping from $\mathfrak{E}$ onto $\mathfrak{F}$ is a length-preserving operator U . Both $\mathfrak{E}$ and $\mathfrak{F}$ are closed linear manifolds.*

Proof. This was just proved.

Definition 14. *The operators R, U of [Theorem 22](#) are called the Cayley transforms of one another. Here R is always a closed linear H.O., and U is always length-preserving.*⁴²

Theorem 23. *A length-preserving operator U is the Cayley transform of the closed linear H.O. R if and only if*

$$\begin{aligned} \alpha) \quad & \mathfrak{A} \text{ is the set of all } \varphi - U\varphi, \quad \varphi \in \mathfrak{E}, \\ \beta) \quad & R(\varphi - U\varphi) = i(\varphi + U\varphi) \quad \text{for every } \varphi \in \mathfrak{E}. \end{aligned}$$

Proof. Let U be the Cayley transform of R . If $f \in \mathfrak{A}$, then

$$(R + i1)f = \varphi \in \mathfrak{E}, \quad U\varphi = (R - i1)f,$$

and hence

$$\varphi - U\varphi = 2if, \quad f = \frac{\varphi}{2i} - U\frac{\varphi}{2i}.$$

Conversely, if $\varphi \in \mathfrak{E}$, then $(R + i1)f = \varphi$ for some $f \in \mathfrak{A}$, and therefore

$$(R - i1)f = U\varphi, \quad \varphi - U\varphi = 2if.$$

Thus $\varphi - U\varphi \in \mathfrak{A}$. Finally,

$$\varphi + U\varphi = 2Rf,$$

⁴²For the moment we know only that every closed linear H.O. R has exactly one Cayley transform U , which is length-preserving. The scope of the converse will follow from [Theorem 24](#).

and so

$$R(\varphi - U\varphi) = 2iRf = i(\varphi + U\varphi).$$

Thus $\alpha)$ and $\beta)$ hold.

Conversely, suppose U satisfies $\alpha)$ and $\beta)$. If $f \in \mathfrak{A}$, then

$$f = \varphi - U\varphi, \quad \varphi \in \mathfrak{E},$$

by $\alpha)$, and

$$Rf = i(\varphi + U\varphi)$$

by $\beta)$. Hence

$$(R + i1)f = 2i\varphi, \quad (R - i1)f = 2iU\varphi.$$

Thus $(R + i1)f \in \mathfrak{E}$ and

$$U(R + i1)f = (R - i1)f.$$

On the other hand, starting with any $\varphi \in \mathfrak{E}$, we have

$$f = \varphi - U\varphi \in \mathfrak{A}, \quad Rf = i(\varphi + U\varphi),$$

and hence

$$2i\varphi = (R + i1)f, \quad \varphi = (R + i1)\frac{f}{2i}.$$

Thus $\mathfrak{E}$ is precisely the set of all $(R + i1)f$ with $f \in \mathfrak{A}$, and

$$U(R + i1)f = (R - i1)f.$$

Therefore U really is the Cayley transform of R .

Theorem 24. *For a length-preserving operator U there exists a closed linear H.O. R of which it is the Cayley transform if and only if the set of all*

$$\varphi - U\varphi, \quad \varphi \in \mathfrak{E},$$

where $\mathfrak{E}$ is the domain of U , is everywhere dense. In that case R is uniquely determined by U .

Proof. Necessity follows from $\alpha)$ of [Theorem 23](#), since $\mathfrak{A}$ is everywhere dense. Sufficiency and the unique determination of R by U likewise follow from $\alpha)$ and $\beta)$ of [Theorem 23](#), once we know that $\beta)$ does not involve an inconsistency in the definition of R , and that the R so defined is a closed linear H.O.

The first point follows from the fact that $\varphi \neq \psi$ implies

$$\varphi - U\varphi \neq \psi - U\psi,$$

or equivalently that $\varphi \neq 0$ implies $\varphi - U\varphi \neq 0$. Indeed, if $\varphi = U\varphi$, then, for every $\chi \in \mathfrak{E}$,

$$(\varphi, \chi) = (U\varphi, U\chi) = (\varphi, U\chi), \quad (\varphi, \chi - U\chi) = 0.$$

But the vectors $\chi - U\chi$ are everywhere dense, so φ is orthogonal to an everywhere-dense set, hence to all of $\mathfrak{E}$, and therefore $\varphi = 0$.

For the second point, the closedness and linearity of R follow from those of U , while the density of the domain of R , namely the set of all $\varphi - U\varphi$, was assumed. It remains only to prove

$$(f, Rg) = (Rf, g),$$

that is, that interchanging f and g in (f, Rg) changes it into its complex conjugate. Put

$$f = \varphi - U\varphi, \quad g = \psi - U\psi.$$

Then

$$\begin{aligned} (f, Rg) &= (\varphi - U\varphi, i(\psi + U\psi)) \\ &= -i\{(\varphi, \psi) - (U\varphi, \psi) + (\varphi, U\psi) - (U\varphi, U\psi)\} \\ &= i\left((U\varphi, \psi) - \overline{(U\psi, \varphi)}\right) \\ &= i(U\varphi, \psi) + i\overline{(U\psi, \varphi)}, \end{aligned}$$

which makes the assertion evident.

Theorem 25. *Let R, S be closed linear Hermitian operators, and let U, V be their Cayley transforms. Then R is an extension, respectively a proper extension, of S if and only if U is an extension, respectively a proper extension, of V . A closed linear manifold $\mathfrak{M}$ reduces R if and only if it reduces U .*

Proof. Each assertion follows immediately from R to U by [Theorem 22](#), and from U to R by [Theorem 23](#).

In the preceding theorems we reduced the problem of deciding the extension relations of R to the much simpler corresponding problem for U . The Cayley transforms of all R are precisely all the U of [Theorem 5](#); here one must note that if R is to be an extension of S , so that U is an extension of a V satisfying [Theorem](#)

5, then U need only be length-preserving.⁴³ Thus, instead of considering the initially rather opaque extensions of R , we need only consider the easily described length-preserving extensions of V , which is itself length-preserving. To orient ourselves in questions of this kind, we prove two simple, geometrically intuitive theorems about length-preserving mappings and closed linear manifolds.

First observe that by the *dimension* of a closed linear manifold $\mathfrak{M}$ we mean the supremum of all finite n for which $\mathfrak{M}$ contains n linearly independent elements; it is therefore one of the numbers

$$0, 1, 2, \dots, \infty.$$

(A substantially more general definition of the dimension of sets in $\mathfrak{S}$ will be given in the next chapter, [Definition 16](#).)

Theorem 26. *Let $\mathfrak{M}$ be a closed linear manifold, and let $\varphi_1, \dots, \varphi_n$ be a normalized orthogonal system spanning it (such a system certainly exists; cf. the proof of [Theorem 13](#); here $n = 0, 1, 2, \dots, \infty$, and in the last case the sequence $\varphi_1, \varphi_2, \dots$ does not terminate). Then $\mathfrak{M}$ is n -dimensional.*

Two closed linear manifolds $\mathfrak{M}, \mathfrak{N}$ can be respectively the domain and the range of a length-preserving operator U if and only if they have the same dimension. Let $\varphi_1, \dots, \varphi_n$ be a normalized orthogonal system spanning $\mathfrak{M}$. Then the normalized orthogonal systems $\psi_1, \dots, \psi_n$ spanning $\mathfrak{N}$ correspond one-to-one to such operators U : to each such system there corresponds exactly one U satisfying

$$U\varphi_1 = \psi_1, \dots, U\varphi_n = \psi_n.$$

Proof. Let $\varphi_1, \dots, \varphi_n$ be a normalized orthogonal system spanning the closed linear manifold $\mathfrak{M}$. Since $\varphi_1, \dots, \varphi_n$ are linearly independent, $\mathfrak{M}$ has at least n dimensions. Thus, when $n = \infty$, we are done. For finite n , we must show that $n + 1$ elements of $\mathfrak{M}$ can never be linearly independent. Indeed, they must have the form

$$f_\mu = a_{\mu 1}\varphi_1 + \dots + a_{\mu n}\varphi_n \quad (\mu = 1, \dots, n + 1),$$

and there must be a linear relation among the $n + 1$ vectors in n dimensions

$$(a_{\mu 1}, \dots, a_{\mu n}) \quad (\mu = 1, \dots, n + 1),$$

and hence also among the f_μ .

⁴³For the vectors $\varphi - U\varphi$ include the vectors $\varphi - V\varphi$, and hence, like the latter, are everywhere dense.

We turn to the second and third assertions. Let $\mathfrak{M}$ be the domain and $\mathfrak{N}$ the range of the length-preserving operator U . If $\varphi_1, \dots, \varphi_n$ is a normalized orthogonal system spanning $\mathfrak{M}$, then, because U preserves length,

$$\psi_1 = U\varphi_1, \dots, \psi_n = U\varphi_n$$

is a normalized orthogonal system spanning $\mathfrak{N}$. Thus $\mathfrak{M}$ and $\mathfrak{N}$ both have n dimensions.

Conversely, let $\mathfrak{M}, \mathfrak{N}$ be two closed linear manifolds of the same dimension, say n . Let

$$\varphi_1, \dots, \varphi_p \quad \text{and} \quad \psi_1, \dots, \psi_q$$

be normalized orthogonal systems spanning $\mathfrak{M}$ and $\mathfrak{N}$, respectively. We must have $p = q = n$. The series

$$\sum_{r=1}^n x_r \varphi_r, \quad \sum_{r=1}^n x_r \psi_r$$

converge under the same circumstances: always when n is finite, and when $n = \infty$ if $\sum_{r=1}^{\infty} |x_r|^2$ is finite, by [Theorem 5](#)—and only then. They therefore exhaust $\mathfrak{M}$ and $\mathfrak{N}$, respectively. (The inner product of $\sum_{r=1}^n x_r \varphi_r$ with φ_s is x_s , so it determines $x_1, \dots, x_n$ uniquely.) We can therefore define an operator U by

$$U \left(\sum_{r=1}^n x_r \varphi_r \right) = \sum_{r=1}^n x_r \psi_r.$$

It has domain $\mathfrak{M}$ and range $\mathfrak{N}$, and is plainly linear. If

$$f = \sum_{r=1}^n x_r \varphi_r,$$

then

$$|f|^2 = |Uf|^2 = \sum_{r=1}^n |x_r|^2, \quad |f| = |Uf|.$$

Thus the operator is continuous, and since its domain is closed, it is itself closed. Hence U is length-preserving.

Conversely, length preservation, and hence in particular closed linearity, together with

$$U\varphi_\mu = \psi_\mu \quad (\mu = 1, \dots, n),$$

implies

$$U \left(\sum_{r=1}^n x_r \varphi_r \right) = \sum_{r=1}^n x_r \psi_r.$$

Our U is therefore uniquely determined by these properties. This proves all the assertions of the theorem.

Theorem 27. *Let U be a length-preserving operator, and let $\mathfrak{E}, \mathfrak{F}$ be its domain and range (thus two closed linear manifolds). All length-preserving extensions of U are obtained as follows.*

Choose two closed linear manifolds $\mathfrak{M}, \mathfrak{N}$ of the same dimension, with

$$\mathfrak{M} \subseteq \mathfrak{H} - \mathfrak{E}, \quad \mathfrak{N} \subseteq \mathfrak{H} - \mathfrak{F},$$

and a length-preserving operator U' having $\mathfrak{M}$ as domain and $\mathfrak{N}$ as range. These determine exactly one length-preserving extension V of U , characterized by the following conditions: its domain and range are, respectively,

$$\mathfrak{E} + \mathfrak{M}, \quad \mathfrak{F} + \mathfrak{N};$$

on $\mathfrak{E}$ one has $Vf = Uf$, and on $\mathfrak{M}$ one has $Vf = U'f$.

By choosing $\mathfrak{M}, \mathfrak{N}, U'$ in all possible ways, one obtains all length-preserving extensions V of U , each exactly once.

Proof. It is clear that every length-preserving extension V of U arises in this way. Let $\mathfrak{R}, \mathfrak{Q}$ be the domain and range of V ; it suffices to put

$$\mathfrak{M} = \mathfrak{R} - \mathfrak{E}, \quad \mathfrak{N} = \mathfrak{Q} - \mathfrak{F},$$

and to let U' be the restriction of V to $\mathfrak{M}$. Then all the conditions are satisfied. It is also clear that our conditions determine at most one V : if both V' and V'' satisfy them, then $V'f = V''f$ throughout $\mathfrak{E}$ and throughout $\mathfrak{M}$, and hence, by linearity, throughout their common domain $\mathfrak{E} + \mathfrak{M}$; thus $V' = V''$. It remains only to show that every system $\mathfrak{M}, \mathfrak{N}, U'$ really determines at least one length-preserving V .

Since $\mathfrak{M}$ and $\mathfrak{N}$ are parts of $\mathfrak{H} - \mathfrak{E}$ and $\mathfrak{H} - \mathfrak{F}$, respectively, they are orthogonal to $\mathfrak{E}$ and $\mathfrak{F}$, respectively. Hence $\mathfrak{M}$ and $\mathfrak{E}$ are linearly independent, and the decomposition

$$f = g + h, \quad g \in \mathfrak{E}, \quad h \in \mathfrak{M},$$

is unique in $\mathfrak{E} + \mathfrak{M}$. Define

$$V(g + h) = Ug + U'h.$$

Since g, h are orthogonal, as are $Ug, U'h$, and since U, U' preserve length,

$$\begin{aligned} |V(g + h)|^2 &= |Ug + U'h|^2 \\ &= |Ug|^2 + |U'h|^2 \\ &= |g|^2 + |h|^2 = |g + h|^2. \end{aligned}$$

The operator V is defined on the closed linear manifold $\mathfrak{E} + \mathfrak{M}$, is plainly linear, and by the preceding equality is continuous; it is therefore closed, and the same equality shows that it is length-preserving. Moreover, as required, it extends both U and U' .

By [Theorems 26](#) and [27](#) we survey all length-preserving extensions of a length-preserving U , and hence all closed linear Hermitian extensions of a closed linear H.O. R ; all of them can be constructed explicitly. For example, R is maximal, that is, it has no proper extension, exactly when U has none—thus when it is impossible in [Theorem 27^{T4}](#) to choose $\mathfrak{M}$ and $\mathfrak{N}$ different from (0) (cf. [note 35](#)). This is plainly the case only if

$$\mathfrak{H} - \mathfrak{E} = (0) \quad \text{or} \quad \mathfrak{H} - \mathfrak{F} = (0),$$

that is, if $\mathfrak{E} = \mathfrak{H}$ or $\mathfrak{F} = \mathfrak{H}$. These questions will be investigated more closely in the next three chapters.

VI. Extension Elements

We now wish to characterize more precisely the closed linear manifolds $\mathfrak{E}, \mathfrak{F}$, which already came into play in the preceding chapter, in relation to the concepts of [Definition 8](#).

For the purposes of this chapter it is convenient to use the terminology of the residue-class calculus of algebra: if $\mathfrak{M}$ is a linear manifold, not necessarily closed, then

$$f = g \pmod{\mathfrak{M}}$$

means that $f - g \in \mathfrak{M}$. The symbols $R, U, \mathfrak{A}, \mathfrak{E}, \mathfrak{F}$ shall retain their previous meanings, as, for example, in [Theorem 22](#).

Theorem 28. *The manifolds $\mathfrak{H} - \mathfrak{E}$ and $\mathfrak{H} - \mathfrak{F}$ are, respectively, the sets of all extension elements f to which R assigns if and $-if$.*

^{T4}*Translator's note:* The source cites Theorem 26 in this sentence; the choice of $\mathfrak{M}, \mathfrak{N}$ is made in Theorem 27.

Proof. By [Theorem 23](#), the assertion that f^* is assigned to f means

$$(f^*, \varphi - U\varphi) = (f, i(\varphi + U\varphi)) = (-if, \varphi + U\varphi).$$

Thus $f^* = if$, respectively $f^* = -if$, means that f (and with it $f^* = \pm if$) is orthogonal to every $\varphi \in \mathfrak{E}$, respectively to every $U\varphi$, that is, to all of $\mathfrak{F}$. Equivalently,

$$f \in \mathfrak{H} - \mathfrak{E}, \quad \text{respectively} \quad f \in \mathfrak{H} - \mathfrak{F}.$$

Consequently all of $\mathfrak{H} - \mathfrak{E}$, respectively all of $\mathfrak{H} - \mathfrak{F}$, apart from 0, belongs to the $+$ -class, respectively the $-$ -class, of extension elements.

Theorem 29. *The set of all extension elements of R is*

$$\mathfrak{A} + (\mathfrak{H} - \mathfrak{E}) + (\mathfrak{H} - \mathfrak{F}).$$

Proof. The extension elements plainly form a linear manifold. Since this manifold must, as we already know, contain

$$\mathfrak{A}, \quad \mathfrak{H} - \mathfrak{E}, \quad \mathfrak{H} - \mathfrak{F},$$

their sum is a part of it; it remains to show that it is no larger.

Let f be an extension element. By the preceding calculation this means

$$(f^*, \varphi - U\varphi) = (-if, \varphi + U\varphi),$$

or

$$(f^* + if, \varphi) = (f^* - if, U\varphi)$$

for every $\varphi \in \mathfrak{E}$, with f^* fixed. Let U^{-1} be the inverse of U , with domain $\mathfrak{F}$ and range $\mathfrak{E}$. Then, wherever these operators are defined, that is, on $\mathfrak{E}$ or $\mathfrak{F}$ as appropriate,

$$P_{\mathfrak{F}}U = U, \quad U^{-1}U = 1, \quad P_{\mathfrak{E}}U^{-1} = U^{-1}, \quad UU^{-1} = 1.$$

Now

$$\begin{aligned} (f^* - if, U\varphi) &= (f^* - if, P_{\mathfrak{F}}U\varphi) \\ &= (P_{\mathfrak{F}}f^* - iP_{\mathfrak{F}}f, U\varphi) \\ &= (UU^{-1}P_{\mathfrak{F}}f^* - iUU^{-1}P_{\mathfrak{F}}f, U\varphi) \\ &= (U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f, \varphi). \end{aligned}$$

Therefore

$$(\{f^* + if\} - \{U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f\}, \varphi) = 0.$$

Since this holds for every $\varphi \in \mathfrak{E}$,

$$P_{\mathfrak{E}}(\{f^* + if\} - \{U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f\}) = 0,$$

and hence

$$\{P_{\mathfrak{E}}f^* + iP_{\mathfrak{E}}f\} - \{U^{-1}P_{\mathfrak{F}}f^* - iU^{-1}P_{\mathfrak{F}}f\} = 0.$$

Equivalently,

$$-U^{-1}P_{\mathfrak{F}}f^* + P_{\mathfrak{E}}f^* = -i(U^{-1}P_{\mathfrak{F}}f + P_{\mathfrak{E}}f),$$

or, after applying $-U$,

$$P_{\mathfrak{F}}f^* - UP_{\mathfrak{E}}f^* = i(P_{\mathfrak{F}}f + UP_{\mathfrak{E}}f).$$

In general,

$$g - P_{\mathfrak{M}}g = P_{\mathfrak{H}-\mathfrak{M}}g = 0 \pmod{\mathfrak{H} - \mathfrak{M}}.$$

It follows that

$$P_{\mathfrak{E}}g = P_{\mathfrak{F}}g \pmod{(\mathfrak{H} - \mathfrak{E}) + (\mathfrak{H} - \mathfrak{F})}.$$

Hence the preceding equality gives

$$P_{\mathfrak{E}}f^* - UP_{\mathfrak{E}}f^* = i(P_{\mathfrak{E}}f + UP_{\mathfrak{E}}f) \pmod{(\mathfrak{H} - \mathfrak{E}) + (\mathfrak{H} - \mathfrak{F})}.$$

The left-hand side belongs, by its very form, to $\mathfrak{A}$; therefore the right-hand side belongs to

$$\mathfrak{A} + (\mathfrak{H} - \mathfrak{E}) + (\mathfrak{H} - \mathfrak{F}).$$

Likewise,

$$P_{\mathfrak{E}}f - UP_{\mathfrak{E}}f$$

belongs by its very form to $\mathfrak{A}$. These two facts imply that $P_{\mathfrak{E}}f$ belongs to the same sum. Finally,

$$f = P_{\mathfrak{E}}f \pmod{\mathfrak{H} - \mathfrak{E}},$$

so the same is true of f . This proves everything.

Theorem 30. *The linear manifolds*

$$\mathfrak{A}, \quad \mathfrak{H} - \mathfrak{E}, \quad \mathfrak{H} - \mathfrak{F}$$

are linearly independent.

Proof. We must infer from

$$f + \varphi + \psi = 0, \quad f \in \mathfrak{A}, \quad \varphi \in \mathfrak{H} - \mathfrak{E}, \quad \psi \in \mathfrak{H} - \mathfrak{F},$$

that $f = \varphi = \psi = 0$. Equivalently, it suffices to prove that if

$$\varphi \in \mathfrak{H} - \mathfrak{E}, \quad \psi \in \mathfrak{H} - \mathfrak{F}, \quad \varphi + \psi \in \mathfrak{A},$$

then $\varphi = \psi = 0$. Assume these latter premises.

By what has already been said, $\varphi, \psi, \varphi + \psi$ are extension elements to which R assigns, respectively,

$$i\varphi, \quad -i\psi, \quad R(\varphi + \psi),$$

the last of which is defined. But $i(\varphi - \psi)$ is plainly also assigned to $\varphi + \psi$, and hence

$$i(\varphi - \psi) = R(\varphi + \psi).$$

Since $i\varphi$ and $-i\psi$ are assigned to φ and ψ , respectively, we must have

$$\begin{aligned} (i\varphi, \varphi + \psi) &= (\varphi, R(\varphi + \psi)) = (\varphi, i(\varphi - \psi)), \\ 2i(\varphi, \varphi) &= 0, \quad \varphi = 0, \end{aligned}$$

and

$$\begin{aligned} (-i\psi, \varphi + \psi) &= (\psi, R(\varphi + \psi)) = (\psi, i(\varphi - \psi)), \\ -2i(\psi, \psi) &= 0, \quad \psi = 0. \end{aligned}$$

Thus every extension element f can be decomposed in one and only one way as

$$f = f^0 + f^+ + f^-, \quad f^0 \in \mathfrak{A}, \quad f^+ \in \mathfrak{H} - \mathfrak{E}, \quad f^- \in \mathfrak{H} - \mathfrak{F};$$

conversely, every such f is an extension element. We call f^0, f^+, f^- , respectively, the 0-, +-, and --components of f . Since $Rf^0, if^+, -if^-$ are assigned to f^0, f^+, f^- , respectively, R assigns to

$$f = f^0 + f^+ + f^-$$

the vector

$$f^* = Rf^0 + if^+ - if^-.$$

This decomposition is important because it permits a more precise characterization of the 0-, +-, and --classes.

Theorem 31. *An extension element f belongs to the 0-, +-, or --class according as*

$$|f^+| = |f^-|, \quad |f^+| > |f^-|, \quad |f^+| < |f^-|.$$

Proof. Since

$$f = f^0 + f^+ + f^-, \quad f^* = Rf^0 + if^+ - if^-,$$

and since if^+ and $-if^-$ are assigned to f^+ and f^- , respectively, we have

$$\begin{aligned} (f^*, f) &= (Rf^0 + if^+ - if^-, f^0 + f^+ + f^-) \\ &= (Rf^0, f^0) + \{(Rf^0, f^+) + (if^+, f^0)\} \\ &\quad + \{(Rf^0, f^-) + (-if^-, f^0)\} \\ &\quad + \{(if^+, f^-) + (-if^-, f^+)\} \\ &\quad + \{(if^+, f^+) + (-if^-, f^-)\} \\ &= (f^0, Rf^0) + \{(f^0, if^+) + (if^+, f^0)\} \\ &\quad + \{(f^0, -if^-) + (-if^-, f^0)\} \\ &\quad + \{(f^+, -if^-) + (-if^-, f^+)\} \\ &\quad + i\{|f^+|^2 - |f^-|^2\}. \end{aligned}$$

The first term and the first three brace-enclosed expressions are plainly real, whereas the last brace-enclosed contribution is purely imaginary. It follows that

$$\operatorname{Im}(f^*, f) = |f^+|^2 - |f^-|^2,$$

which proves the assertion.

VII. The Deficiency Indices

Definition 15. *Let $R, U, \mathfrak{A}, \mathfrak{E}, \mathfrak{F}$ have their previous meanings, as in [Theorem 22](#). Let m, n be the dimensions of*

$$\mathfrak{H} - \mathfrak{E}, \quad \mathfrak{H} - \mathfrak{F},$$

respectively. We call m, n the deficiency indices of R . Thus

$$m, n = 0, 1, 2, \dots, \infty.$$

These numbers will play a decisive role in questions of maximality and hypermaximality. There is a close relation between them and the three classes of

extension elements, which we shall now establish. First we generalize the notion of dimension in $\S$.

Definition 16. Let $\mathfrak{M}$ be a linear manifold, not necessarily closed, and let $\mathfrak{U}$ be a subset of $\S$, which need not even be a linear manifold. By the dimension of $\mathfrak{U}$ modulo $\mathfrak{M}$ we mean the supremum of all finite n for which there exist n linearly independent elements whose entire linear manifold,⁴⁴ possibly with the exception of 0,⁴⁵ lies in $\mathfrak{U}$ and outside $\mathfrak{M}$. Thus n is one of the numbers

$$0, 1, 2, \dots, \infty.$$

In what follows, $\mathfrak{U}$ will not in fact be entirely arbitrary, but will always have the following properties. First, it is a set of residue classes modulo $\mathfrak{M}$: if

$$f = g \pmod{\mathfrak{M}},$$

then either neither or both of f, g belong to $\mathfrak{U}$. Second, whenever $f \in \mathfrak{U}$, then $af \in \mathfrak{U}$ for every $a \neq 0$. It is immediately seen that, among such sets $\mathfrak{U}$, only the empty set and $\mathfrak{M}$ have dimension 0. It is also clear that the dimension of a closed linear manifold $\mathfrak{U}$ modulo (0) agrees with the dimension defined in [Chapter V](#).

The $+$ -, $-$ -, and 0-classes are sets of this kind modulo $\mathfrak{A}$, as are unions of any of them. Since the first two do not contain 0, whereas the third does, in the zero-dimensional case they must be empty and equal to $\mathfrak{A}$, respectively.

Theorem 32. The $+$ -class, and also its union with the 0-class, is m -dimensional modulo $\mathfrak{A}$; the $-$ -class, and also its union with the 0-class, is n -dimensional modulo $\mathfrak{A}$; the 0-class is $\min(m, n)$ -dimensional modulo $\mathfrak{A}$.

Proof. Let $p \leq m$ be finite, as m may be ∞ , and let $f_1, \dots, f_p$ be linearly independent elements of the m -dimensional manifold $\S - \mathfrak{E}$. Their entire linear manifold lies in $\S - \mathfrak{E}$, and hence, except for 0, in the $+$ -class, and therefore automatically outside $\mathfrak{A}$. Thus the $+$ -class has at least m dimensions modulo $\mathfrak{A}$. To prove the first two assertions, it remains to show that its union with the 0-class has at most m dimensions.

For $m = \infty$ this is trivial. If m is finite, we must show that, whenever $f_1, \dots, f_{m+1}$ are linearly independent extension elements, some nonzero combination

$$a_1 f_1 + \dots + a_{m+1} f_{m+1} \neq 0$$

belongs either to $\mathfrak{A}$ or to the $-$ -class.

⁴⁴Since n is finite, this linear manifold is automatically closed.

⁴⁵This exception will prove convenient below.

Write

$$f_\mu = f_\mu^0 + f_\mu^+ + f_\mu^- \quad (\mu = 1, \dots, m+1).$$

The vectors f_μ^+ all lie in the m -dimensional manifold $\mathfrak{S} - \mathfrak{E}$, so they are not linearly independent. Choose $a_1, \dots, a_{m+1}$, not all zero, such that

$$a_1 f_1^+ + \dots + a_{m+1} f_{m+1}^+ = 0.$$

Then

$$a_1 f_1 + \dots + a_{m+1} f_{m+1} = g^0 + g^-, \quad g^0 \in \mathfrak{A}, \quad g^- \in \mathfrak{S} - \mathfrak{F}.$$

For $g^- = 0$ this vector belongs to $\mathfrak{A}$; for $g^- \neq 0$, it belongs to the $--$ -class by [Theorem 31](#). This proves the first two assertions. The next two follow in the same way by interchanging $+$ and $-$.

It remains to prove the final assertion. Since the dimension of the 0-class modulo $\mathfrak{A}$ is at most both m and n , it suffices to show that it is at least $\min(m, n)$. Thus, for every finite

$$p \leq \min(m, n),$$

we must exhibit linearly independent vectors $f_1, \dots, f_p$ such that every nonzero combination

$$a_1 f_1 + \dots + a_p f_p$$

belongs to the 0-class but not to $\mathfrak{A}$.

There are p linearly independent elements in $\mathfrak{S} - \mathfrak{E}$. We may “orthogonalize” them immediately (cf. [Theorem 8](#)), and hence may assume them normalized and orthogonal:

$$\varphi_1, \dots, \varphi_p.$$

Likewise, let

$$\psi_1, \dots, \psi_p$$

be normalized orthogonal elements of $\mathfrak{S} - \mathfrak{F}$. Put

$$f_1 = \varphi_1 + \psi_1, \dots, f_p = \varphi_p + \psi_p.$$

Then, unless $a_1 = \dots = a_p = 0$,

$$\begin{aligned} a_1 f_1 + \dots + a_p f_p &= (a_1 \varphi_1 + \dots + a_p \varphi_p) \\ &\quad + (a_1 \psi_1 + \dots + a_p \psi_p). \end{aligned}$$

By [Theorem 30](#) this vector certainly does not belong to $\mathfrak{A}$, and in particular is nonzero. Moreover,

$$|a_1\varphi_1 + \cdots + a_p\varphi_p|^2 = |a_1\psi_1 + \cdots + a_p\psi_p|^2 = |a_1|^2 + \cdots + |a_p|^2.$$

Hence it belongs to the 0-class by [Theorem 31](#). Thus $f_1, \dots, f_p$ are linearly independent and have all the required properties.

Theorem 33. *The operator R is maximal if and only if $m = 0$ or $n = 0$, and hypermaximal if and only if $m = n = 0$. Equivalently, R is maximal if and only if $\mathfrak{E} = \mathfrak{S}$ or $\mathfrak{F} = \mathfrak{S}$, and hypermaximal if and only if*

$$\mathfrak{E} = \mathfrak{F} = \mathfrak{S}.$$

Proof. By [Theorem 11](#), maximality means that the 0-class equals $\mathfrak{A}$; by [Definition 9](#), hypermaximality means in addition that the +- and --classes are empty. By the remark preceding [Theorem 32](#), the first condition means that the 0-class is zero-dimensional, while the additional condition means that the +- and --classes are zero-dimensional. Finally, by [Theorem 32](#), these conditions say

$$\min(m, n) = 0,$$

that is, $m = 0$ or $n = 0$, and, respectively, $m = 0$ and $n = 0$.

The second formulation follows immediately from the first.

The contrast between maximality and hypermaximality is now clear, as is the role of m, n . We add the following observation concerning [Theorem 32](#). If R is replaced by

$$aR + b1 \quad (a, b \text{ real, } a \neq 0),$$

which is again a closed linear H.O. with the same domain as R , then $\mathfrak{E}, \mathfrak{F}$ change in an opaque way, and at first we cannot see what happens to m, n . But $\mathfrak{A}$, the set of extension elements, and the 0-, +-, and --classes plainly do not change at all when $a > 0$; when $a < 0$, only the last two classes are interchanged. It follows from [Theorem 32](#) that m, n likewise remain unchanged or are merely interchanged.⁴⁶

⁴⁶Let $\bar{R}$ be the following operator: $\bar{R}f$ is defined for every extension element f of R , and its value is the assigned vector f^* . The operator $\bar{R}$ is plainly closed and linear and is an extension of R . (If R is hypermaximal, then plainly $\bar{R} = R$; otherwise $\bar{R}$ is not even a H.O., since $m > 0$ or $n > 0$, and hence the +- or the --class is nonempty, while in these classes

$$(\bar{R}f, f) = (f^*, f)$$

VIII. The Extension Process

Already at the end of [Chapter VI](#) we obtained an overview of all possible extensions of a closed linear H.O. R ; we can now decide when and how an extension to a maximal or hypermaximal operator is possible.

Theorem 34. *A closed linear H.O. R with deficiency indices m, n can be extended to one with deficiency indices m', n' if and only if there exists*

$$p \in \{0, 1, 2, \dots, \infty\}$$

such that

$$m' + p = m, \quad n' + p = n.$$

Proof. Consider the extension process of [Theorem 27](#), with $U, \mathfrak{U}, \mathfrak{E}, \mathfrak{F}$ having their usual meanings and $\mathfrak{M}, \mathfrak{N}$ as in that theorem. The manifolds

$$\mathfrak{S} - \mathfrak{E}, \quad \mathfrak{S} - \mathfrak{F}$$

is not real.) By [Theorem 28](#), $\mathfrak{S} - \mathfrak{E}$ and $\mathfrak{S} - \mathfrak{F}$ are the sets of solutions of

$$\overline{R}f = if, \quad \text{respectively} \quad \overline{R}f = -if.$$

Thus m and n are the maximum numbers of linearly independent solutions of these equations.

By the preceding observation, the maximum number of linearly independent solutions of

$$a\overline{R}f + bf = if$$

is m , respectively n , according as $a > 0$, respectively $a < 0$. This equation says

$$\overline{R}f = \rho f, \quad \rho = -\frac{b}{a} + \frac{1}{a}i.$$

We may therefore say: let ρ be nonreal. For the equation

$$\overline{R}f = \rho f$$

the maximum number of linearly independent solutions is m when $\text{Im } \rho > 0$, and n when $\text{Im } \rho < 0$. (For real ρ , the maximum number is, as usual, the “multiplicity of the point eigenvalue” ρ . When ρ is nonreal, the equation $Rf = \rho f$ must have no solution other than $f = 0$, since otherwise

$$(Rf, f) = \rho(f, f)$$

would not be real.)

These facts are related to certain results of Carleman; cf. the citation in [note 4](#).

are m - and n -dimensional, respectively. The question is whether they have closed linear submanifolds $\mathfrak{M}, \mathfrak{N}$ of the same dimension, say p , such that

$$(\mathfrak{S} - \mathfrak{E}) - \mathfrak{M}, \quad (\mathfrak{S} - \mathfrak{F}) - \mathfrak{N}$$

are m' - and n' -dimensional, respectively.

More generally, the question is this: when does a k -dimensional closed linear manifold $\mathfrak{R}$ have a p -dimensional part, that is, a closed linear manifold $\mathfrak{P}$, such that $\mathfrak{R} - \mathfrak{P}$ has k' dimensions? We shall see that this happens exactly when

$$k = k' + p.$$

Substituting

$$\mathfrak{R} = \mathfrak{S} - \mathfrak{E}, \quad \mathfrak{P} = \mathfrak{M}, \quad k = m, \quad k' = m',$$

and then

$$\mathfrak{R} = \mathfrak{S} - \mathfrak{F}, \quad \mathfrak{P} = \mathfrak{N}, \quad k = n, \quad k' = n',$$

gives the theorem.

The condition is necessary. Indeed, suppose $\mathfrak{P}$ and $\mathfrak{R} - \mathfrak{P}$ are p - and k' -dimensional, respectively. They are then spanned, as closed linear manifolds, by normalized orthogonal systems

$$\varphi_1, \dots, \varphi_p, \quad \psi_1, \dots, \psi_{k'},$$

respectively. Hence

$$\mathfrak{R} = \mathfrak{P} + (\mathfrak{R} - \mathfrak{P})$$

is spanned by

$$\varphi_1, \dots, \varphi_p, \psi_1, \dots, \psi_{k'}.$$

These systems are mutually orthogonal because they belong to $\mathfrak{P}$ and $\mathfrak{R} - \mathfrak{P}$. Thus $\mathfrak{R}$ has dimension $p + k'$, and therefore $k = k' + p$. Here either k' or p may also be ∞ .

The condition is also sufficient. Suppose $\mathfrak{R}$ has dimension $k = k' + p$. It is then spanned by a normalized orthogonal system

$$\varphi_1, \dots, \varphi_p, \psi_1, \dots, \psi_{k'},$$

again allowing k' or p to be ∞ . Let $\mathfrak{P}$ be the closed linear manifold spanned by $\varphi_1, \dots, \varphi_p$. Then $\mathfrak{R} - \mathfrak{P}$ is spanned by $\psi_1, \dots, \psi_{k'}$, and the proof is complete.

Theorem 35. *Let R be a closed linear H.O. with deficiency indices m, n . The operator R can always be extended to a maximal H.O. If $m \neq n$, none of its extensions is hypermaximal. If $m = n$ and these numbers are finite, every maximal extension is hypermaximal; if $m = n = \infty$, then among the maximal extensions there are both hypermaximal ones and ones that are not hypermaximal.*

Proof. By [Theorem 33](#), maximality means that at least one deficiency index is 0, hypermaximality means that both are 0, and maximality without hypermaximality means that exactly one is 0. Applying [Theorem 34](#) gives precisely the assertions of the theorem.

Corollary to Theorem 35. *If R is not already maximal, then its maximal extensions are possible in continuum many different ways; and, whenever such extensions exist, the same is true separately of its hypermaximal extensions and of its maximal but nonhypermaximal extensions.*

Proof. By [Theorem 27](#), our extension method amounted to enlarging $\mathfrak{E}$ and $\mathfrak{F}$ by two closed linear manifolds $\mathfrak{M}, \mathfrak{N}$, which, when R is not maximal, had positive dimension.^{T5} But a length-preserving mapping U from $\mathfrak{M}$ onto $\mathfrak{N}$ also enters, and it can be chosen in continuum many ways: by [Theorem 26](#), every normalized orthogonal system

$$\psi_1, \dots, \psi_p$$

spanning $\mathfrak{N}$ gives one such U , and there are continuum many of these systems, since, for example, ψ_1 may be replaced by any $\alpha\psi_1$ with $|\alpha| = 1$.⁴⁷

If R is merely an H.O., then, as we know, its maximal and hypermaximal extensions are the same as those of the closed linear H.O. $\widetilde{R}$, to which the results just obtained are applicable. In particular, if R has only one maximal extension, then $\widetilde{R}$ itself must already be maximal; the same holds for hypermaximal extensions.

We now turn to a closer investigation of hypermaximal H.O. ([Chapter IX](#)), and then of maximal but nonhypermaximal H.O. ([Chapter X](#)).

^{T5}*Translator's note:* The source prints $\mathfrak{S} - \mathfrak{E}$ and $\mathfrak{S} - \mathfrak{F}$ here; Theorem 27 shows that $\mathfrak{E}$ and $\mathfrak{F}$ are meant.

⁴⁷There cannot be more than continuum many extensions, since there are only continuum many closed operators in $\mathfrak{S}$, as follows easily from the separability of $\mathfrak{S}$.

IX. Hypermaximal Operators and the Eigenvalue Representation

If R is hypermaximal, then

$$\mathfrak{E} = \mathfrak{F} = \mathfrak{H};$$

that is, its Cayley transform U is unitary: it is meaningful everywhere and has an inverse U^{-1} that is likewise unitary and meaningful everywhere. Thus the hypermaximal H.O. R correspond one-to-one to the unitary operators U satisfying the condition of [Theorem 24](#), namely, those for which the set of all $\varphi - U\varphi$ is everywhere dense.

Already in the proof of [Theorem 24](#) we observed that, if the $\varphi - U\varphi$ are everywhere dense, then $U\varphi = \varphi$ implies $\varphi = 0$. Here the converse also holds. For if the $\varphi - U\varphi$ are not everywhere dense, then they span a closed linear manifold $\mathfrak{M} \neq \mathfrak{H}$ (they themselves already form a linear manifold), and hence $\mathfrak{H} - \mathfrak{M} \neq (0)$. Let $f \neq 0$ be an element of $\mathfrak{H} - \mathfrak{M}$. Every $\varphi - U\varphi$ belongs to $\mathfrak{M}$ and is therefore orthogonal to f :

$$\begin{aligned} (f, \varphi - U\varphi) &= (f, \varphi) - (f, U\varphi) \\ &= (Uf, U\varphi) - (f, U\varphi) = (Uf - f, U\varphi) \end{aligned}$$

must vanish for every φ ; consequently $Uf - f = 0$, or $Uf = f$.

The unitary operators in question can therefore also be characterized as those for which $Uf = f$ has no solution other than 0 (the number 1 is not a “point eigenvalue”). Thus we control the hypermaximal H.O. R once we know these U . This, however, is a problem concerning bounded operators (U preserves length and is therefore continuous), and it can be solved by the methods of Hilbert’s theory of bounded forms, with an addition that also belongs to that circle of ideas. We carry out the necessary considerations, partly following F. Riesz’s simple method, in [Appendix II](#), and state only the result here.

Definition 17. Let $a < x < b$ be an open interval (below, a, b will once be 0, 1 and once be $-\infty, \infty$). By a resolution of the identity (abbreviated: r.i.) we mean a family of P.O.’s $E(\lambda)$, with λ ranging over all of (a, b) , having the following properties:

- $\alpha)$ If $\lambda \leq \mu$, then $E(\lambda)$ is a part of $E(\mu)$.
- $\beta)$ If $\lambda \geq \lambda_0$ and $\lambda \rightarrow \lambda_0$, then, for every f ,

$$E(\lambda)f \longrightarrow E(\lambda_0)f.$$

γ) If $\lambda \rightarrow a$, respectively $\lambda \rightarrow b$, then

$$E(\lambda)f \rightarrow 0, \quad \text{respectively} \quad E(\lambda)f \rightarrow f.$$

If now $(a, b) = (0, 1)$ and $E(\rho)$ is an r.i., then for every f there is exactly one f^* such that, for every g ,

$$(f^*, g) = \int_0^1 e^{2\pi i \rho} d(E(\rho)f, g)$$

holds.⁴⁸ The operator U defined by $Uf = f^*$ is unitary, and $Uf = f$ has no solution other than 0. Conversely, every U with these properties is produced in exactly one way by an r.i. $E(\rho)$ through the procedure just sketched (cf., as stated, [Appendix II](#)).

We now prove:

Theorem 36. Let $(a, b) = (-\infty, \infty)$, and let $F(\lambda)$ be an r.i. If the integral

$$\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2$$

is finite,⁴⁹ then there is exactly one f^* such that, for every g ,

$$(f^*, g) = \int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g),$$

where the integral on the right is absolutely convergent. We define an operator R for these f by $Rf = f^*$.

This R is a hypermaximal H.O., and every hypermaximal H.O. R can be produced in this way by means of exactly one r.i.

Proof. The r.i. $E(\rho)$ for the interval $(0, 1)$ and the r.i. $F(\lambda)$ for the interval $(-\infty, \infty)$ correspond one-to-one by the following substitutions.⁵⁰

$$E(\rho) \mapsto F(\lambda) = E\left(-\frac{1}{\pi} \operatorname{arccot} \lambda\right),$$

⁴⁸This is a Stieltjes integral, and it is absolutely convergent. We write ρ in place of λ .

⁴⁹As λ increases, $F(\lambda)$ increases, and hence so does $|F(\lambda)f|^2$. Moreover, λ^2 is never negative. By its nature, therefore, this integral either converges and is finite or diverges properly to $+\infty$.

⁵⁰The substitutions

$$\rho = -\frac{1}{\pi} \operatorname{arccot} \lambda, \quad \lambda = -\cot \pi \rho$$

put the intervals $0 < \rho < 1$ and $-\infty < \lambda < \infty$ into one-to-one correspondence. For arccot , the value between $-\pi$ and 0 is always to be taken.

$$F(\lambda) \mapsto E(\rho) = F(-\cot \pi \rho).$$

It therefore suffices to show the following. If $E(\rho)$ is the r.i. of a unitary operator U (cf. the discussion preceding the theorem), then $F(\lambda)$ really does produce an operator R by the prescription of the theorem, and this operator is the Cayley transform of U . Everything we know about the relation between Cayley transforms, and everything stated before the theorem about hypermaximal H.O., unitary operators, and r.i., then proves all the assertions.

Let, therefore, $F(\lambda)$ be an r.i., let

$$E(\rho) = F(-\cot \pi \rho),$$

and let U be the corresponding unitary operator. We first show that, when

$$\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2 = C^2$$

is finite, the f^* of the theorem can actually be constructed.

Let $\lambda < \lambda' < \mu$. Since $F(\lambda)$ is a part of $F(\mu)$, $F(\mu) - F(\lambda)$ is a P.O., and therefore

$$\begin{aligned} |((F(\mu) - F(\lambda))f, g)| &= |((F(\mu) - F(\lambda))f, (F(\mu) - F(\lambda))g)| \\ &\leq |(F(\mu) - F(\lambda))f| |(F(\mu) - F(\lambda))g|. \end{aligned}$$

Moreover, in general,

$$\begin{aligned} |(F(\mu) - F(\lambda))h|^2 &= (h, (F(\mu) - F(\lambda))h) \\ &= (h, F(\mu)h) - (h, F(\lambda)h) \\ &= |F(\mu)h|^2 - |F(\lambda)h|^2. \end{aligned}$$

Consequently,^{T6}

$$\begin{aligned} |((F(\mu) - F(\lambda))f, g)| &\leq \sqrt{|F(\mu)f|^2 - |F(\lambda)f|^2} \sqrt{|F(\mu)g|^2 - |F(\lambda)g|^2}, \\ |\lambda' \{(F(\mu)f, g) - (F(\lambda)f, g)\}| &\leq \sqrt{\lambda'^2 \{|F(\mu)f|^2 - |F(\lambda)f|^2\}} \sqrt{|F(\mu)g|^2 - |F(\lambda)g|^2}. \end{aligned}$$

Now the two integrals

$$\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2 = C^2, \quad \int_{-\infty}^{\infty} d|F(\lambda)g|^2 = |F(\lambda)g|^2 \Big|_{\lambda=-\infty}^{\lambda=\infty} = |g|^2$$

^{T6}*Translator's note:* The source writes $E(\lambda)$ rather than $F(\lambda)$ throughout this intermediate estimate.

We use F consistently with the statement of the theorem and with the surrounding proof.

converge absolutely. Hence the familiar inequality

$$\sqrt{a_1 b_1} + \cdots + \sqrt{a_p b_p} \leq \sqrt{(a_1 + \cdots + a_p)(b_1 + \cdots + b_p)} \quad (a_1, \dots, a_p, b_1, \dots, b_p \geq 0)$$

has, first, the consequence that

$$\int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g)$$

also converges absolutely, and, second, that⁵¹

$$\left| \int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g) \right| \leq \sqrt{\left(\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2 \right) \left(\int_{-\infty}^{\infty} d|F(\lambda)g|^2 \right)} = C|g|.$$

For a moment, with f fixed and g variable, denote this integral by $L(g)$. Then $L(g)$ has the properties

$$\begin{aligned} L(a_1 g_1 + \cdots + a_n g_n) &= \overline{a_1} L(g_1) + \cdots + \overline{a_n} L(g_n), \\ |L(g)| &\leq C|g|; \end{aligned}$$

it is conjugate-linear and continuous. By a theorem of F. Riesz, there is therefore exactly one f^* such that

$$(f^*, g) = L(g)$$

always holds.⁵² This is precisely the assertion concerning f^* .

The operator R can therefore actually be constructed. Let R^* be the Cayley transform of U ; it remains to prove that $R = R^*$. It is enough to show that R is an

⁵¹It suffices to recall the definition of the Stieltjes integral. The inequality just used is merely a modified form of Schwarz's inequality.

⁵²The theorem is proved as follows. Since all the (f^*, g) are prescribed, there cannot be two distinct f^* , so it suffices to prove existence. Let $\varphi_1, \varphi_2, \dots$ be a complete orthonormal system, and put

$$L(\varphi_n) = a_n \quad (n = 1, 2, \dots).$$

By Theorem 7, β),

$$g = \sum_{n=1}^{\infty} (g, \varphi_n) \varphi_n;$$

hence, by the conjugate-linearity and continuity of L ,

$$L(g) = \sum_{n=1}^{\infty} \overline{(g, \varphi_n)} a_n,$$

extension of R^* . For whenever Rf and Rg have meaning, interchanging f and g changes (Rf, g) into its complex conjugate, by the definition in [Theorem 36](#); hence

$$(Rf, g) = (f, Rg).$$

Thus R is an H.O. (its domain contains, by assumption, that of R^* and is therefore everywhere dense) and an extension of the maximal R^* , which is maximal by the unitarity of U ; consequently $R = R^*$.

Suppose, then, that R^*f has meaning. We must prove that Rf has meaning and equals R^*f ; equivalently, that

$$\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2$$

is finite and that, for every g ,

$$(R^*f, g) = (Rf, g) = \int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g).$$

and this series must converge. If $\sum_{n=1}^{\infty} |a_n|^2$ is finite, then

$$f^* = \sum_{n=1}^{\infty} a_n \varphi_n$$

converges by [Theorem 5](#), and $(f^*, \varphi_n) = a_n$. It then follows from [Theorem 7](#), γ), that

$$L(g) = \sum_{n=1}^{\infty} (f^*, \varphi_n) \overline{(g, \varphi_n)} = (f^*, g),$$

which is the assertion.

It remains to prove that $\sum_{n=1}^{\infty} |a_n|^2$ is finite. Since $|L(g)| \leq C|g|$, taking

$$g = x_1 \varphi_1 + \cdots + x_m \varphi_m$$

gives

$$\left| \sum_{n=1}^m a_n \overline{x_n} \right| \leq C \sqrt{\sum_{n=1}^m |x_n|^2}.$$

Set $x_n = a_n$ for $n = 1, \dots, m$. Then

$$\sum_{n=1}^m |a_n|^2 \leq C^2.$$

Thus $\sum_{n=1}^{\infty} |a_n|^2$ is finite, and is at most C^2 .

The fact that R^*f has meaning says, moreover, that

$$f = \varphi - U\varphi;$$

this will be our starting point.

Under the transformation $\lambda = -\cot \pi\rho$, the two integrals

$$\int_{-\infty}^{\infty} \lambda^2 d|F(\lambda)f|^2, \quad \int_{-\infty}^{\infty} \lambda d(F(\lambda)f, g)$$

may be written as

$$\int_0^1 \cot^2 \pi\rho d|E(\rho)f|^2, \quad \int_0^1 -\cot \pi\rho d(E(\rho)f, g).$$

To substitute $f = \varphi - U\varphi$ here, we must first calculate several integrals:

$$\begin{aligned} (\varphi \pm U\varphi, g) &= \int_0^1 d(E(\rho)\varphi, g) \pm \int_0^1 e^{2\pi i\rho} d(E(\rho)\varphi, g) \\ &= \int_0^1 (1 \pm e^{2\pi i\rho}) d(E(\rho)\varphi, g). \end{aligned}$$

It follows that⁵³

$$\begin{aligned} (\varphi - U\varphi, E(\rho)g) &= \int_0^1 (1 - e^{2\pi i\rho'}) d(E(\rho')\varphi, E(\rho)g) \\ &= \int_0^1 (1 - e^{2\pi i\rho'}) d(E(\rho)E(\rho')\varphi, g) \\ &= \int_0^1 (1 - e^{2\pi i\rho'}) d(E(\min(\rho, \rho'))\varphi, g) \\ &= \int_0^\rho (1 - e^{2\pi i\rho'}) d(E(\rho')\varphi, g). \end{aligned}$$

Therefore, using complex conjugation in the second line⁵⁴ and the general Stieltjes

⁵³For $\rho' > \rho$, the expression

$$(E(\min(\rho, \rho'))\varphi, g)$$

is constant, and hence its Stieltjes differential, and therefore the corresponding integral, is 0.

⁵⁴Take the complex conjugate of the preceding formula with $g = \varphi$.

formula in the next step⁵⁵, we obtain

$$\begin{aligned}
 |E(\rho)(\varphi - U\varphi)|^2 &= (\varphi - U\varphi, E(\rho)(\varphi - U\varphi)) \\
 &= \int_0^\rho (1 - e^{2\pi i \rho'}) d(E(\rho')\varphi, \varphi - U\varphi) \\
 &= \int_0^\rho (1 - e^{2\pi i \rho'}) d \left[\int_0^{\rho'} (1 - e^{-2\pi i \rho''}) d(E(\rho'')\varphi, \varphi) \right] \\
 &= \int_0^\rho (1 - e^{2\pi i \rho'})(1 - e^{-2\pi i \rho'}) d(E(\rho')\varphi, \varphi) \\
 &= \int_0^\rho 4 \sin^2(\pi \rho') d |E(\rho')\varphi|^2.
 \end{aligned}$$

It follows, using the formula of note 55 once more, that

$$\begin{aligned}
 \int_0^1 \cot^2(\pi \rho) d |E(\rho)(\varphi - U\varphi)|^2 &= \int_0^1 \cot^2(\pi \rho) d \left[\int_0^\rho 4 \sin^2(\pi \rho') d |E(\rho')\varphi|^2 \right] \\
 &= \int_0^1 \cot^2(\pi \rho) 4 \sin^2(\pi \rho) d |E(\rho)\varphi|^2 \\
 &= \int_0^1 4 \cos^2(\pi \rho) d |E(\rho)\varphi|^2.
 \end{aligned}$$

This is finite, since it is bounded above by the finite quantity

$$\int_0^1 4 d |E(\rho)\varphi|^2 = 4|\varphi|^2.$$

Furthermore,

$$\begin{aligned}
 \int_0^1 -\cot(\pi \rho) d(E(\rho)(\varphi - U\varphi), g) &= \int_0^1 -\cot(\pi \rho) d \left[\int_0^\rho (1 - e^{2\pi i \rho'}) d(E(\rho')\varphi, g) \right] \\
 &= \int_0^1 -\cot(\pi \rho)(1 - e^{2\pi i \rho}) d(E(\rho)\varphi, g) \\
 &= \int_0^1 i(1 + e^{2\pi i \rho}) d(E(\rho)\varphi, g)
 \end{aligned}$$

⁵⁵The general formula for Stieltjes integrals is

$$\int_A^B p(\rho) d \left[\int_A^\rho q(\rho') d r(\rho') \right] = \int_A^B p(\rho) q(\rho) d r(\rho).$$

$$\begin{aligned}
&= (i(\varphi + U\varphi), g) \\
&= (R^*(\varphi - U\varphi), g).
\end{aligned}$$

This proves everything required.

We have obtained for hypermaximal H.O. the normal form analogous to Hilbert's spectral representation or eigenvalue representation for bounded forms (cf. the citation in note 7). In doing so, however, we have already exhausted all eigenvalue representations for these operators, so there certainly can be none for maximal but nonhypermaximal H.O. (Admittedly, we have not yet constructed an example of such an H.O.) We therefore refrain from further investigation of hypermaximal H.O., which would offer scarcely any new points of view, and turn to maximal but nonhypermaximal H.O.

X. Maximal but Nonhypermaximal Operators

Such an H.O. R , with deficiency indices m, n , is characterized by either

$$m = 0, \quad n > 0, \quad \text{or} \quad m > 0, \quad n = 0.$$

Since replacing R by $-R$ interchanges m, n , it is enough to consider $m = 0, n > 0$. Form $U, \mathfrak{E}, \mathfrak{F}$. Then

$$\mathfrak{E} = \mathfrak{H}, \quad \mathfrak{G}_0 = \mathfrak{H} - \mathfrak{F}$$

is n -dimensional, and U maps $\mathfrak{H}$ isometrically onto $\mathfrak{F}$.

Let $\mathfrak{G}_\nu$ be the image of $\mathfrak{G}_0$ under U^ν ($\nu = 0, 1, 2, \dots$); it too is a closed linear manifold. For $\nu > 0$, $\mathfrak{G}_\nu$ is a part of $\mathfrak{F}$ and is therefore orthogonal to $\mathfrak{G}_0$. Applying the mapping U^k ($k = 0, 1, 2, \dots$), the images $\mathfrak{G}_{k+\nu}$ and $\mathfrak{G}_k$ are orthogonal, since (f, g) is invariant under this mapping; cf. [Theorem 10](#). Thus, in general,

$$\mathfrak{G}_k \perp \mathfrak{G}_l \quad (k, l = 0, 1, 2, \dots; k \neq l).$$

Let $\mathfrak{G}_0, \mathfrak{G}_1, \mathfrak{G}_2, \dots$ span the closed linear manifold $\mathfrak{M}$, and put

$$\mathfrak{H} - \mathfrak{M} = \mathfrak{N}.$$

We already know that U maps $\mathfrak{G}_\nu$ one-to-one and isometrically onto $\mathfrak{G}_{\nu+1}$. We now show that it likewise maps $\mathfrak{N}$ onto itself. As is easily seen, $\mathfrak{N}$ is the set of all elements of $\mathfrak{H}$ orthogonal to every $\mathfrak{G}_\nu$. Because (f, g) is invariant, its image is the

set of all elements in the image of $\mathfrak{H}$ that are orthogonal to every image of the $\mathfrak{G}_v$. These are the elements of

$$\mathfrak{F} = \mathfrak{H} - \mathfrak{G}_0$$

orthogonal to every $\mathfrak{G}_{v+1}$, hence the elements of $\mathfrak{H}$ orthogonal to $\mathfrak{G}_0$ and to all $\mathfrak{G}_{v+1}$. This is indeed $\mathfrak{N}$ again.

Let $\varphi_1^0, \dots, \varphi_p^0$ be an orthonormal system spanning the closed linear manifold $\mathfrak{G}_0$, where possibly $p = \infty$, and put

$$U^v \varphi_q^0 = \varphi_q^v \quad (v = 0, 1, 2, \dots; q = 1, \dots, p).$$

Then $\varphi_1^v, \dots, \varphi_p^v$ is also an orthonormal system and spans $\mathfrak{G}_v$. For $k \neq l$, φ_q^k lies in $\mathfrak{G}_k$ and $\varphi_{q''}^l$ in $\mathfrak{G}_l$, and the two are therefore orthogonal. Thus all

$$\varphi_q^v \quad (v = 0, 1, 2, \dots; q = 1, \dots, p)$$

form an orthonormal system; moreover, they span the same closed linear manifold as the $\mathfrak{G}_v$, namely $\mathfrak{M}$.

For each $q = 1, \dots, p$, let

$$\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$$

span the closed linear manifold $\mathfrak{M}_q$. Then $\mathfrak{M}_1, \dots, \mathfrak{M}_p$ span $\mathfrak{M}$, and $\mathfrak{M}_q$ is a part of $\mathfrak{M}$. Since

$$U \varphi_q^v = \varphi_q^{v+1},$$

U maps $\mathfrak{M}_q$ into itself. In summary:

$$\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$$

span $\mathfrak{H}$, just as $\mathfrak{M}, \mathfrak{N}$ do, and U maps them respectively into $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$, that is, into subsets of themselves. Moreover, $\mathfrak{M}_1, \dots, \mathfrak{M}_p$, as parts of $\mathfrak{M}$, are orthogonal to $\mathfrak{N}$ and also to one another, because their spanning vectors $\varphi_{q'}^k, \varphi_{q''}^l$ have $q' \neq q''$. It follows that each of

$$\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$$

reduces U , and therefore, by [Theorem 25](#), also its Cayley transform R .

Consider the parts of R and U lying in $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$. Denote them, respectively, by

$$R_1, \dots, R_p, S, \quad \text{and} \quad U_1, \dots, U_p, V.$$

By [Theorem 21](#), R is characterized by $R_1, \dots, R_p, S$. Now all the closed linear manifolds $\mathfrak{M}_1, \dots, \mathfrak{M}_p$ are infinite-dimensional, since $\mathfrak{M}_q$ is spanned by the orthonormal system

$$\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$$

Thus they are all Hilbert spaces, meaning that they satisfy conditions A through E of [Chapter I](#), with the definitions of $af, f + g, (f, g)$ unchanged,⁵⁶ and hence all the concepts developed for $\mathfrak{H}$ may be used in them without qualification. Thus R_q is an H.O. in $\mathfrak{M}_q$. In particular, its domain is everywhere dense in $\mathfrak{M}_q$, since it consists of the projections of the everywhere-dense domain of R in $\mathfrak{H}$. The operator U_q preserves length and is the Cayley transform of R_q .

Three cases are possible in $\mathfrak{N}$. First, $\mathfrak{N}$ may be infinite-dimensional, and hence a Hilbert space; then S is again an H.O., V is unitary, since it maps $\mathfrak{N}$ onto $\mathfrak{N}$, and S is its Cayley transform, hence S is hypermaximal. Second, $\mathfrak{N}$ may be N -dimensional, where $N = 1, 2, \dots$, and hence an N -dimensional complex Euclidean space (cf. [note 56](#)); then S is an H.O. in it, that is, an ordinary finite-dimensional Hermitian operator. Third, $\mathfrak{N}$ may be zero-dimensional, that is, $\mathfrak{N} = (0)$; then it need not be considered at all.

Thus nothing new can occur in $\mathfrak{N}$ and S , and we therefore turn to $\mathfrak{M}_q, R_q, U_q$ ($q = 1, \dots, p$). They all behave alike: R_q is the Cayley transform of U_q , and in $\mathfrak{M}_q$ there is a complete orthonormal system

$$\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$$

such that

$$U_q \varphi_q^v = U \varphi_q^v = \varphi_q^{v+1}.$$

Thus

$$U_q \left(\sum_{v=0}^{\infty} x_v \varphi_q^v \right) = \sum_{v=0}^{\infty} x_v \varphi_q^{v+1}, \quad \sum_{v=0}^{\infty} |x_v|^2 < \infty.$$

This completely determines U_q , and hence R_q , once $\varphi_q^0, \varphi_q^1, \varphi_q^2, \dots$ are regarded as known. Using the isomorphism of $\mathfrak{M}_q$ with $\mathfrak{H}$, we may also say that R_q must

⁵⁶Every closed linear manifold plainly satisfies A through C and E, while D is equivalent to infinite dimensionality. A finite-dimensional closed linear manifold consists only of 0 when it has dimension 0. If it has $N = 1, 2, \dots$ dimensions, it is spanned by an orthonormal system $\varphi_1, \dots, \varphi_N$ and therefore consists of all

$$x_1 \varphi_1 + \dots + x_N \varphi_N;$$

that is, it is the N -dimensional complex Euclidean space of all $x_1, \dots, x_N$.

be isomorphic to the following operator $\bar{R}$ in $\mathfrak{H}$: $\bar{R}$ is the Cayley transform of $\bar{U}$, where, for a fixed complete orthonormal system $\psi_0, \psi_1, \psi_2, \dots$ in $\mathfrak{H}$,

$$\bar{U} \left(\sum_{v=0}^{\infty} x_v \psi_v \right) = \sum_{v=0}^{\infty} x_v \psi_{v+1}, \quad \sum_{v=0}^{\infty} |x_v|^2 < \infty.$$

It remains first to show that $\bar{R}$ really exists, that is, that the evidently length-preserving $\bar{U}$ satisfies the condition of [Theorem 24](#). If it does, then, since $\bar{U}$ is defined on all of $\mathfrak{H}$, so that $\mathfrak{E} = \mathfrak{H}$ and $m = 0$, $\bar{R}$ is a maximal H.O. The range $\mathfrak{F}$ of $\bar{U}$ is the set of all

$$\sum_{v=0}^{\infty} x_v \psi_{v+1},$$

that is, the closed linear manifold spanned by $\psi_1, \psi_2, \dots$. Hence $\mathfrak{H} - \mathfrak{F}$ is the set of all $a\psi_0$, and is one-dimensional.^{T7} Consequently $\bar{R}$ has deficiency indices 0, 1; it is the first example of a maximal but nonhypermaximal H.O.

It remains to prove its existence, that is, the everywhere-denseness of the set of all $\varphi - \bar{U}\varphi$. Since this set is a linear manifold, it suffices to show that all elements of a complete orthonormal system are its accumulation points, for example all ψ_v . Now

$$\begin{aligned} & \left(\psi_v + \frac{l-1}{l} \psi_{v+1} + \dots + \frac{1}{l} \psi_{v+l-1} \right) - \bar{U} \left(\psi_v + \frac{l-1}{l} \psi_{v+1} + \dots + \frac{1}{l} \psi_{v+l-1} \right) \\ &= \left(\psi_v + \frac{l-1}{l} \psi_{v+1} + \dots + \frac{1}{l} \psi_{v+l-1} \right) - \left(\psi_{v+1} + \frac{l-1}{l} \psi_{v+2} + \dots + \frac{1}{l} \psi_{v+l} \right) \\ &= \psi_v - \frac{1}{l} (\psi_{v+1} + \dots + \psi_{v+l}). \end{aligned}$$

This differs from ψ_v by

$$\frac{1}{l} (\psi_{v+1} + \dots + \psi_{v+l}),$$

whose squared norm and norm are, respectively,

$$\left| \frac{1}{l} (\psi_{v+1} + \dots + \psi_{v+l}) \right|^2 = \frac{1}{l}, \quad \left| \frac{1}{l} (\psi_{v+1} + \dots + \psi_{v+l}) \right| = \frac{1}{\sqrt{l}}.$$

With a suitable l , this is arbitrarily small. Everything is therefore proved.

^{T7}*Translator's note:* The source prints $a\psi_1$ here; the orthogonal complement of the span of $\psi_1, \psi_2, \dots$ is plainly the span of ψ_0 .

We formulate the result:

Theorem 37. *Let $\psi_0, \psi_1, \psi_2, \dots$ be a complete orthonormal system. There is then exactly one H.O. $\bar{R}$ whose Cayley transform $\bar{U}$ is defined by*

$$\bar{U} \left(\sum_{\nu=0}^{\infty} x_{\nu} \psi_{\nu} \right) = \sum_{\nu=0}^{\infty} x_{\nu} \psi_{\nu+1}, \quad \sum_{\nu=0}^{\infty} |x_{\nu}|^2 < \infty.$$

The operator $\bar{R}$ has deficiency indices $0, 1$ and is therefore maximal but not hypermaximal.

Proof. This was just proved.

Below we shall derive several further important general properties of $\bar{R}$; Part VI of the “Remarks” is devoted to its explicit representation.^{T8} First we return to our earlier developments and prove:

Theorem 38. *Let R be a maximal but nonhypermaximal H.O.; thus its deficiency indices are either*

$$0, n \quad (n = 1, 2, \dots, \infty)$$

or

$$m, 0 \quad (m = 1, 2, \dots, \infty).$$

Put $p = n$, respectively $p = m$. Then $\mathfrak{H}$ can be decomposed into $p + 1$ closed linear manifolds

$$\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N},$$

which are pairwise orthogonal and together span the closed linear manifold $\mathfrak{H}$; when $p = \infty$, the sequence $\mathfrak{M}_1, \mathfrak{M}_2, \dots$ does not terminate. They have the following properties:

α) Each of $\mathfrak{M}_1, \dots, \mathfrak{M}_p, \mathfrak{N}$ reduces R ; denote the parts of R lying in them by

$$R_1, \dots, R_p, S,$$

respectively.

β) Each $\mathfrak{M}_q$ ($q = 1, 2, \dots, p$) is infinite-dimensional and is a Hilbert space, that is, isomorphic to $\mathfrak{H}$. In the case of deficiency indices $0, n$, all R_q are isomorphic to $\bar{R}$; in the case of deficiency indices $m, 0$, all R_q are isomorphic to $-\bar{R}$, where $\bar{R}$ is the operator of [Theorem 37](#).

^{T8}Translator’s note: The “Remarks” cited here are not included in the present article, so there is no internal cross-reference target.

γ) There are three possibilities for $\mathfrak{N}$. First, it may be infinite-dimensional and a Hilbert space, hence isomorphic to $\mathfrak{H}$; then S is a hypermaximal H.O. Second, it may be N -dimensional, where $N = 1, 2, \dots$, and hence an N -dimensional complex Euclidean space; then S is an ordinary finite-dimensional Hermitian operator. Third, it may be zero-dimensional, that is, $\mathfrak{N} = (0)$; then it and S may be omitted.

These $R_1, \dots, R_p, S$ determine R by [Theorem 21](#).

Proof. For $m = 0, n > 0$ this was proved in the first part of this chapter; one need only note [Theorem 37](#). For $m > 0, n = 0$, it follows by replacing R by $-R$, thereby interchanging m, n .

In finite-dimensional spaces every Hermitian operator can be put into diagonal form. As long as the space has more than one dimension, the operator is therefore reducible. In $\mathfrak{H}$, hypermaximal H.O. are likewise reducible, as follows.

Let R be hypermaximal, and let $F(\lambda)$ be its r.i. from [Theorem 36](#). Suppose first that there is a λ such that

$$F(\lambda) \neq 0, 1.$$

Then the closed linear manifold $\mathfrak{M}$ whose P.O. is

$$P_{\mathfrak{M}} = F(\lambda)$$

is neither (0) nor $\mathfrak{H}$. This $\mathfrak{M}$ reduces R . First,

$$\begin{aligned} \int_{-\infty}^{\infty} \lambda'^2 d|F(\lambda')F(\lambda)f|^2 &= \int_{-\infty}^{\infty} \lambda'^2 d|F(\min(\lambda', \lambda))f|^2 \\ &= \int_{-\infty}^{\lambda} \lambda'^2 d|F(\lambda')f|^2, \end{aligned}$$

by the observation of note [53](#); and this is finite whenever

$$\int_{-\infty}^{\infty} \lambda'^2 d|F(\lambda')f|^2$$

is finite. Second, in that case,^{T9}

$$\begin{aligned} (F(\lambda)Rf, g) &= (Rf, F(\lambda)g) \\ &= \int_{-\infty}^{\infty} \lambda' d(F(\lambda')f, F(\lambda)g) \end{aligned}$$

^{T9}*Translator's note:* From these displays through the following scalar case, the source switches from the previously declared F back to E . We retain F consistently.

$$\begin{aligned}
&= \int_{-\infty}^{\infty} \lambda' d(F(\lambda)F(\lambda')f, g) \\
&= \int_{-\infty}^{\infty} \lambda' d(F(\lambda')F(\lambda)f, g) \\
&= (RF(\lambda)f, g)
\end{aligned}$$

for every g . Thus

$$F(\lambda)Rf = RF(\lambda)f.$$

Suppose, on the other hand, that every $F(\lambda)$ equals either 0 or 1. Then, by α) of [Definition 17](#), there is a λ_0 such that the former holds for $\lambda < \lambda_0$ and the latter for $\lambda > \lambda_0$. The number λ_0 is finite by γ), and at $\lambda = \lambda_0$ the latter holds by β).^{T10} Therefore

$$(Rf, g) = \int_{-\infty}^{\infty} \lambda' d(F(\lambda')f, g) = \lambda_0(f, g), \quad Rf = \lambda_0 f.$$

Thus $\lambda_0 1$ is an extension of R , and hence, by maximality,

$$R = \lambda_0 1.$$

But this operator is plainly reduced by every closed linear manifold.

Reducibility is therefore a vestige of transformability to diagonal form, and as such is most closely connected with the possibility of an eigenvalue representation. Irreducibility of a matrix already signifies the occurrence of exceptional cases in elementary-divisor theory, something excluded by Hermitian character in finite dimensions and, as we have seen, also in § for hypermaximal operators. From this point of view, the following property of $\bar{R}$ is particularly noteworthy.

Theorem 39. *A maximal H.O. is irreducible if and only if, for a suitable choice of the complete orthonormal system $\psi_0, \psi_1, \psi_2, \dots$ in [Theorem 37](#), it coincides with $\bar{R}$ or with $-\bar{R}$. Both alternatives can never occur for the same H.O.*

Proof. Let R be maximal and irreducible. By the preceding discussion it is certainly not hypermaximal, so [Theorem 38](#) applies. Since the $\mathfrak{M}_1$ occurring there reduces R and is nonzero, it must equal $\mathfrak{H}$. Thus $p = 1$ and $\mathfrak{R}$ is absent. Hence $R = R_1$, and by β) of that theorem it is $\bar{R}$ or $-\bar{R}$, with a suitable choice of $\psi_0, \psi_1, \psi_2, \dots$. In the former case the deficiency indices of R are 0, 1, while in the latter they are 1, 0; the two alternatives are therefore incompatible.

^{T10}*Translator's note:* The source interchanges the references to β) and γ) in this sentence; the references are corrected here.

It remains to show that $\overline{R}$, and hence also $-\overline{R}$, is irreducible; equivalently, by [Theorem 25](#), that $\overline{U}$ is irreducible. Suppose that $\overline{U}$ is reduced by $\mathfrak{M}$, and hence also by

$$\mathfrak{N} = \mathfrak{S} - \mathfrak{M}.$$

The everywhere-defined $\overline{U}$ maps each of these manifolds into a subset of itself. Put

$$\mathfrak{M}' = \overline{U}(\mathfrak{M}), \quad \mathfrak{N}' = \overline{U}(\mathfrak{N}).$$

If $\mathfrak{M}' = \mathfrak{M}$ and $\mathfrak{N}' = \mathfrak{N}$, then the range of $\overline{U}$ would contain both $\mathfrak{M}$ and $\mathfrak{N}$, and hence also

$$\mathfrak{M} + \mathfrak{N} = \mathfrak{S},$$

which is not the case. Thus $\mathfrak{M}'$ or $\mathfrak{N}'$ is a proper subset of $\mathfrak{M}$ or $\mathfrak{N}$, respectively, and

$$\mathfrak{M} - \mathfrak{M}' \neq (0) \quad \text{or} \quad \mathfrak{N} - \mathfrak{N}' \neq (0).$$

Let $f \neq 0$ belong to $\mathfrak{M} - \mathfrak{M}'$, respectively to $\mathfrak{N} - \mathfrak{N}'$. It is then orthogonal to $\mathfrak{M}'$, respectively $\mathfrak{N}'$. Being an element of $\mathfrak{M}$, respectively $\mathfrak{N}$, it is also orthogonal to $\mathfrak{N}$ and $\mathfrak{N}'$, respectively to $\mathfrak{M}$ and $\mathfrak{M}'$. It is therefore orthogonal to the entire range of $\overline{U}$,

$$\mathfrak{M}' + \mathfrak{N}'.$$

That is, it is orthogonal to every $\psi_1, \psi_2, \dots$, and hence

$$f = a\psi_0 \quad (a \neq 0).$$

Consequently ψ_0 also belongs to $\mathfrak{M} - \mathfrak{M}'$, respectively to $\mathfrak{N} - \mathfrak{N}'$, and thus to $\mathfrak{M}$, respectively $\mathfrak{N}$. But $\mathfrak{M}$ and $\mathfrak{N}$ contain their $\overline{U}$ -images; whichever of them contains ψ_0 therefore also contains $\psi_1, \psi_2, \dots$, and hence all of $\mathfrak{S}$. Thus either

$$\mathfrak{M} = \mathfrak{S},$$

or

$$\mathfrak{S} - \mathfrak{M} = \mathfrak{N} = \mathfrak{S}, \quad \mathfrak{M} = (0).$$

This proves everything.

Together with $\overline{R}$, the operator

$$a\overline{R} + b1 \quad (a, b \text{ real, } a \neq 0)$$

is also maximal and irreducible, and hence, up to the choice of coordinate system, equals either $\overline{R}$ or $-\overline{R}$. For $a > 0$ its deficiency indices are 0, 1, so in that case $\overline{R}$ is the possibility; for $a < 0$ they are 1, 0, so in that case $-\overline{R}$ is the possibility.

Finally, we mention that the decomposition of the H.O. R in [Theorem 38](#) admits a kind of converse, which makes it possible to construct closed linear Hermitian operators R having any prescribed deficiency indices m, n . Since $m = n = 0$ is realized by every hypermaximal H.O. (for example, by 1), we may assume $m + n > 0$.

There then exist $m + n$ infinite-dimensional closed linear manifolds

$$\mathfrak{M}_1, \dots, \mathfrak{M}_m, \mathfrak{N}_1, \dots, \mathfrak{N}_n$$

which are pairwise orthogonal and together span $\mathfrak{H}$ as a closed linear manifold.⁵⁷ Each is isomorphic to $\mathfrak{H}$, so we can realize $-\overline{R}$ in the first m of them and $\overline{R}$ in the remaining n ; denote the resulting operators by

$$R_1, \dots, R_m, S_1, \dots, S_n.$$

It is easily verified that, when combined according to the prescription of [Theorem 21](#), these yield a closed linear H.O. R , and that R has deficiency indices m, n .

XI. Semiboundedness

We shall give several properties that every nonhypermaximal (but maximal) H.O. must lack. (Most of the results of the following theorem will soon also be obtained by another route.)

Theorem 40. *A maximal but nonhypermaximal H.O. can be neither semibounded above nor semibounded below (and hence, a fortiori, cannot be bounded); it can neither be defined everywhere nor take every value; and it cannot be real in any complete normalized orthogonal system (cf. [note 20](#)).*

Proof. For such an operator there is a closed linear manifold $\mathfrak{M}$ reducing it such that the part lying in $\mathfrak{M}$ is $\overline{R}$ or $-\overline{R}$, by [Theorem 38](#). If it possessed any of the

⁵⁷Let $\varphi_1, \varphi_2, \dots$ be a complete normalized orthogonal system. Divide it into $m + n$ subsequences

$$\psi_1^1, \psi_2^1, \dots; \dots; \psi_1^m, \psi_2^m, \dots; \chi_1^1, \chi_2^1, \dots; \dots; \chi_1^n, \chi_2^n, \dots$$

(either m or n may also equal ∞), and let $\mathfrak{M}_1, \dots, \mathfrak{M}_m, \mathfrak{N}_1, \dots, \mathfrak{N}_n$ be the closed linear manifolds spanned by these respective subsequences.

first four properties just named, then $\overline{R}$ or $-\overline{R}$ would have to possess the same property. But $\overline{R}$ (and hence also $-\overline{R}$) has none of them, as will be shown in its detailed discussion in [Appendix IV](#), § 4.

As for the last property: if an operator is real, in the sense of [note 20](#), in a complete normalized orthogonal system, then for it there is no distinction between i and $-i$, and hence its deficiency indices are equal. For a maximal but nonhypermaximal H.O., however, they are unequal; cf. [Theorem 33](#).

We now turn to the systematic investigation of semibounded operators.

Theorem 41. *A H.O. R that takes every value is hypermaximal. It takes each value exactly once.*

Proof. Let f be an extension element of R , and let f^* be assigned to it. By assumption there is a g for which Rg is defined and $Rg = f^*$. For every h for which Rh is defined,

$$(f, Rh) = (f^*, h) = (Rg, h) = (g, Rh).$$

Since Rh takes every value, it follows that $f = g$; hence Rf is defined. Thus R is hypermaximal. At the same time we have proved that $Rf = Rg$ implies $f = g$, by putting $f^* = Rf$.

Theorem 42. *A linear H.O. R satisfying*

$$(f, Rf) \geq |f|^2$$

for every f in its domain can be extended to a definite hypermaximal H.O.

Proof. First form the closed linear H.O. $\widetilde{R}$. Plainly it too satisfies

$$(f, \widetilde{R}f) \geq |f|^2.$$

Let $\mathfrak{M}$ be the range of $\widetilde{R}$; it is plainly a linear manifold. We have

$$|f|^2 \leq (f, \widetilde{R}f) \leq |f| |\widetilde{R}f|, \quad |\widetilde{R}f| \geq |f|,$$

and

$$|\widetilde{R}f - \widetilde{R}g| = |\widetilde{R}(f - g)| \geq |f - g|.$$

Thus $\widetilde{R}f = \widetilde{R}g$ implies $f = g$, so $\widetilde{R}$ has an inverse $\widetilde{R}^{-1}$. This inverse is plainly also closed and linear, is defined on $\mathfrak{M}$, and is continuous because

$$|f - g| \geq |\widetilde{R}^{-1}f - \widetilde{R}^{-1}g|.$$

Consequently $\mathfrak{M}$ is closed, and we may form the closed linear manifold

$$\mathfrak{N} = \mathfrak{L} - \mathfrak{M}.$$

Let $f \in \mathfrak{N}$. Whenever $\widetilde{R}g$ is defined, it belongs to $\mathfrak{M}$, and hence

$$(f, \widetilde{R}g) = 0.$$

Thus f is an extension element to which 0 is assigned. If $\widetilde{R}f$ is defined, therefore, $\widetilde{R}f = 0$, and then $f = 0$. Hence the domain $\mathfrak{A}$ of $\widetilde{R}$, which is a linear manifold, and $\mathfrak{N}$ are linearly independent.

We may therefore define a new operator S as follows: Sf is defined when

$$f = g + h, \quad g \in \mathfrak{A}, \quad h \in \mathfrak{N},$$

the decomposition being unique, and its value is

$$Sf = \widetilde{R}g.$$

The operator S is an extension of $\widetilde{R}$ and is a H.O. Indeed, if

$$f = g + h, \quad f' = g' + h',$$

with $g, g' \in \mathfrak{A}$ and $h, h' \in \mathfrak{N}$, then

$$\begin{aligned} (f, Sf') &= (g + h, \widetilde{R}g') = (g, \widetilde{R}g') \\ &= (\widetilde{R}g, g') = (\widetilde{R}g, g' + h') = (Sf, f'). \end{aligned}$$

The manifold $\mathfrak{N}$ reduces S . By definition, Sf is always defined on $\mathfrak{N}$ and equals 0 there; and since every Sf equals some $\widetilde{R}g$ and therefore lies in $\mathfrak{M}$, we also have

$$P_{\mathfrak{N}}Sf = 0 = SP_{\mathfrak{N}}f.$$

Hence $\mathfrak{M} = \mathfrak{L} - \mathfrak{N}$ also reduces S . Denote the parts of S lying in $\mathfrak{M}, \mathfrak{N}$ by S' and S'' , respectively.

The range of S is $\mathfrak{M}$. Since

$$Sf = P_{\mathfrak{M}}Sf = SP_{\mathfrak{M}}f,$$

S already takes each of its values in $\mathfrak{M}$, so S' also has range $\mathfrak{M}$. As the one-to-one linear image of the infinite-dimensional domain $\mathfrak{A}$ (which is infinite-dimensional

because it is everywhere dense), $\mathfrak{M}$ is likewise infinite-dimensional and hence is a Hilbert space. We may therefore apply [Theorem 41](#): S' is hypermaximal, while S'' is in fact 0.

It follows immediately that S itself must be hypermaximal. Moreover, it is definite. Indeed, if

$$f = g + h, \quad g \in \mathfrak{A}, \quad h \in \mathfrak{N},$$

then

$$(f, Sf) = (g + h, \tilde{R}g) = (g, \tilde{R}g) \geq |g|^2 \geq 0.$$

This proves everything.

Theorem 43. *Suppose a linear H.O. R satisfies either*

$$(f, Rf) \geq C|f|^2$$

for every f in its domain, or

$$(f, Rf) \leq C|f|^2$$

for every such f ; these are precisely the conditions of semiboundedness. Then R can be extended to a hypermaximal H.O. S satisfying, respectively,

$$(f, Sf) \geq C'|f|^2 \quad \text{or} \quad (f, Sf) \leq C'|f|^2.$$

Here C' may be chosen to be any number less than C , respectively greater than C .⁵⁸

Proof. For $C' = 0$, $C = 1$, and the $\geq$ sign, this is [Theorem 42](#). All other cases reduce to this one by considering, in place of R , S ,

$$\frac{1}{C - C'}R - \frac{C'}{C - C'}1, \quad \frac{1}{C - C'}S - \frac{C'}{C - C'}1.$$

XII. Pathology of Unbounded Operators

We begin with several auxiliary theorems.

Theorem 44. *If R is a linear H.O. that is not semibounded above, then, for every $n = 1, 2, \dots$ and every C , there exists an n -dimensional linear manifold on which Rf is always defined and throughout which*

$$(f, Rf) \geq C|f|^2.$$

⁵⁸The theorem probably remains true for $C' = C$, but no proof of this is available.

Proof. Let $f_1, \dots, f_k$ be arbitrary elements for which $Rf_1, \dots, Rf_k$ are defined. We first show that there must be a φ such that

$$\varphi \perp f_1, \dots, f_k, \quad |\varphi| = 1, \quad R\varphi \text{ is defined}, \quad (\varphi, R\varphi) \geq D.$$

Suppose the contrary. Since only the linear manifold spanned by $f_1, \dots, f_k$ matters, we may “orthogonalize” them (cf. [Theorem 8](#)), replacing them by a normalized orthogonal system

$$\varphi_1, \dots, \varphi_l \quad (l \leq k).$$

Let f be arbitrary, subject only to Rf being defined. By adjoining f and orthogonalizing, we can write

$$f = \sum_{\mu=1}^l a_\mu \varphi_\mu + b\varphi,$$

where

$$|\varphi| = 1, \quad \varphi \perp \varphi_1, \dots, \varphi_l.$$

Since $Rf, R\varphi_1, \dots, R\varphi_l$ are defined, $R\varphi$ is also defined; by our contrary assumption,

$$(\varphi, R\varphi) \leq D.$$

Now

$$|f|^2 = \sum_{\mu=1}^l |a_\mu|^2 + |b|^2,$$

and

$$\begin{aligned} (f, Rf) &= \sum_{\mu, \nu=1}^l a_\mu \overline{a_\nu} (\varphi_\mu, R\varphi_\nu) \\ &\quad + 2 \operatorname{Re} \sum_{\mu=1}^l b \overline{a_\mu} (\varphi, R\varphi_\mu) + |b|^2 (\varphi, R\varphi) \\ &\leq \sum_{\mu, \nu=1}^l |a_\mu| |a_\nu| |\varphi_\mu| |R\varphi_\nu| \\ &\quad + 2 \sum_{\mu=1}^l |b| |a_\mu| |\varphi| |R\varphi_\mu| + |b|^2 D. \end{aligned}$$

If

$$D' = \max_{\mu=1,\dots,l} (|R\varphi_\mu|, D),$$

then

$$\begin{aligned} (f, Rf) &\leq D' \left(\sum_{\mu=1}^l |a_\mu| + |b| \right)^2 \\ &\leq (l+1)D' \left(\sum_{\mu=1}^l |a_\mu|^2 + |b|^2 \right) = (l+1)D'|f|^2. \end{aligned}$$

Thus R would be semibounded above, contrary to the hypothesis.

We go one step further. Let $g_1, \dots, g_k$ now be entirely arbitrary. Then there exists a φ for which $R\varphi$ is defined and

$$|\varphi| = 1, \quad (\varphi, R\varphi) \geq D, \quad |(\varphi, g_\mu)| \leq \varepsilon \quad (\mu = 1, \dots, k; \varepsilon > 0).$$

(Thus it is not quite exactly orthogonal.) Indeed, the vectors f for which Rf is defined are everywhere dense. Choose $f_1, \dots, f_k$ with

$$|g_\mu - f_\mu| \leq \varepsilon \quad (\mu = 1, \dots, k),$$

and choose φ as in the result just proved. Then

$$|\varphi| = 1, \quad (\varphi, R\varphi) \geq D,$$

and

$$|(\varphi, g_\mu)| = |(\varphi, g_\mu - f_\mu)| \leq |\varphi| |g_\mu - f_\mu| \leq \varepsilon.$$

We can now prove the assertion itself. Choose successively:

$$\begin{aligned} \varphi_1, \quad R\varphi_1 \text{ defined}, \quad |\varphi_1| = 1, \quad (\varphi_1, R\varphi_1) \geq D; \\ \varphi_2, \quad R\varphi_2 \text{ defined}, \quad |\varphi_2| = 1, \quad (\varphi_2, R\varphi_2) \geq D, \quad |(R\varphi_1, \varphi_2)| \leq 1; \\ \vdots \\ \varphi_n, \quad R\varphi_n \text{ defined}, \quad |\varphi_n| = 1, \quad (\varphi_n, R\varphi_n) \geq D, \\ |(R\varphi_1, \varphi_n)| \leq 1, \dots, |(R\varphi_{n-1}, \varphi_n)| \leq 1. \end{aligned}$$

Thus all $R\varphi_\mu$ are defined,

$$|\varphi_\mu| = 1, \quad (\varphi_\mu, R\varphi_\mu) \geq D,$$

and

$$|(R\varphi_\mu, \varphi_\nu)| \leq 1 \quad (\mu < \nu),$$

hence, by the Hermitian property, for every $\mu \neq \nu$. It follows that

$$\begin{aligned} \left| \sum_{\mu=1}^n a_\mu \varphi_\mu \right|^2 &= \sum_{\mu, \nu=1}^n a_\mu \overline{a_\nu} (\varphi_\mu, \varphi_\nu) \\ &\leq \sum_{\mu, \nu=1}^n |a_\mu| |a_\nu| |\varphi_\mu| |\varphi_\nu| \\ &\leq \left(\sum_{\mu=1}^n |a_\mu| \right)^2 \leq n \sum_{\mu=1}^n |a_\mu|^2, \end{aligned}$$

while

$$\begin{aligned} \left(\sum_{\mu=1}^n a_\mu \varphi_\mu, R \left\{ \sum_{\mu=1}^n a_\mu \varphi_\mu \right\} \right) &= \sum_{\mu, \nu=1}^n a_\mu \overline{a_\nu} (\varphi_\mu, R\varphi_\nu) \\ &= \sum_{\mu=1}^n |a_\mu|^2 (\varphi_\mu, R\varphi_\mu) + \sum_{\substack{\mu, \nu=1 \\ \mu \neq \nu}}^n a_\mu \overline{a_\nu} (\varphi_\mu, R\varphi_\nu) \\ &\geq D \sum_{\mu=1}^n |a_\mu|^2 - \sum_{\substack{\mu, \nu=1 \\ \mu \neq \nu}}^n |a_\mu| |a_\nu| \\ &= (D+1) \sum_{\mu=1}^n |a_\mu|^2 - \left(\sum_{\mu=1}^n |a_\mu| \right)^2 \\ &\geq (D+1-n) \sum_{\mu=1}^n |a_\mu|^2. \end{aligned}$$

If $D > n-1$, the second inequality shows that

$$\sum_{\mu=1}^n a_\mu \varphi_\mu = 0$$

implies

$$\sum_{\mu=1}^n |a_\mu|^2 = 0, \quad a_1 = \cdots = a_n = 0.$$

Thus $\varphi_1, \dots, \varphi_n$ are linearly independent, and the linear manifold of all $\sum_{\mu=1}^n a_\mu \varphi_\mu$ is n -dimensional. Moreover, the two inequalities give

$$\left(\sum_{\mu=1}^n a_\mu \varphi_\mu, R \left\{ \sum_{\mu=1}^n a_\mu \varphi_\mu \right\} \right) \geq \frac{D+1-n}{n} \left| \sum_{\mu=1}^n a_\mu \varphi_\mu \right|^2.$$

Taking also

$$D \geq nC + n - 1,$$

we therefore have throughout this linear manifold

$$(f, Rf) \geq C|f|^2.$$

This proves the theorem.

Theorem 45. *Let R be a linear H.O. that is not semibounded above, let $f_1, \dots, f_k$ be arbitrary elements, and let C be any number. Then there exists a φ for which $R\varphi$ is defined such that*

$$\varphi \perp f_1, \dots, f_k, \quad |\varphi| = 1, \quad (\varphi, R\varphi) \geq C.$$

*Proof.*⁵⁹ Choose the linear manifold in [Theorem 44](#) to be $(k+1)$ -dimensional. In it the equations

$$(f_1, \varphi) = 0, \dots, (f_k, \varphi) = 0,$$

which represent at most k independent linear conditions, have a solution $f \neq 0$. Multiplication by a scalar allows us to arrange $|f| = 1$. This f has all the required properties.

Theorem 46. *Every linear H.O. R that is not semibounded above is an extension of a definite linear H.O. S .*

Proof. It suffices to exhibit an everywhere-dense sequence $f_1, f_2, \dots$ such that all Rf_μ are defined and

$$\left(\sum_{\mu=1}^n x_\mu f_\mu, R \left\{ \sum_{\mu=1}^n x_\mu f_\mu \right\} \right) \geq 0 \quad (n = 1, 2, \dots).$$

Indeed, the operator S defined on the linear manifold spanned by $f_1, f_2, \dots$ —an everywhere-dense, and hence nonclosed, manifold—and agreeing there with R is plainly a linear definite H.O., and R is its extension.

⁵⁹When $Rf_1, \dots, Rf_k$ were defined, we already obtained this at the beginning of the proof of [Theorem 44](#); but we must free ourselves from that assumption.

Let $g_1, g_2, \dots$ be an everywhere-dense sequence such that all Rg_μ are defined. Thus $g_1, g_2, \dots$ may be taken to be a subsequence dense in the domain $\mathfrak{A}$ of R ; the domain itself is everywhere dense, and such a sequence exists because $\mathfrak{A}$, as a subset of $\mathfrak{H}$, is separable; cf. [note 33](#). Let $C_1, C_2, \dots$ be a sequence of numbers to be specified later.

Using [Theorem 45](#), choose a sequence $\varphi_1, \varphi_2, \dots$ with the following properties:

$$\begin{aligned} R\varphi_1 \text{ defined, } |\varphi_1| &= 1, \quad (\varphi_1, R\varphi_1) \geq C_1; \\ R\varphi_2 \text{ defined, } |\varphi_2| &= 1, \quad (\varphi_2, R\varphi_2) \geq C_2, \quad (R\varphi_1, \varphi_2) = 0; \\ R\varphi_3 \text{ defined, } |\varphi_3| &= 1, \quad (\varphi_3, R\varphi_3) \geq C_3, \quad (R\varphi_1, \varphi_3) = (R\varphi_2, \varphi_3) = 0; \\ &\vdots \end{aligned}$$

Thus all $R\varphi_\mu$ are defined,

$$|\varphi_\mu| = 1, \quad (\varphi_\mu, R\varphi_\mu) \geq C_\mu, \quad (R\varphi_\mu, \varphi_\nu) = 0 \quad (\mu < \nu),$$

and therefore the last equality holds whenever $\mu \neq \nu$. Put

$$f_1 = g_1 + \varphi_1, \quad f_2 = g_2 + \frac{1}{2}\varphi_2, \quad f_3 = g_3 + \frac{1}{3}\varphi_3, \quad \dots,$$

and claim that, with a suitable choice of $C_1, C_2, \dots$, the sequence $f_1, f_2, \dots$ has all the desired properties.

First, it is everywhere dense:

$$|f_n - g_n| = \left| \frac{1}{n}\varphi_n \right| = \frac{1}{n} \longrightarrow 0.$$

For every f there is a subsequence $g_{N_1}, g_{N_2}, \dots$ of $g_1, g_2, \dots$ such that $g_{N_j} \rightarrow f$, and then also $f_{N_j} \rightarrow f$.

We now compute (f, Rf) for the linear combinations of $f_1, f_2, \dots$. For every n ,

$$\begin{aligned} &\left(\sum_{\mu=1}^n a_\mu f_\mu, R \left\{ \sum_{\mu=1}^n a_\mu f_\mu \right\} \right) \\ &= \left(\sum_{\mu=1}^n a_\mu g_\mu + \sum_{\mu=1}^n \frac{a_\mu}{\mu} \varphi_\mu, \sum_{\mu=1}^n a_\mu Rg_\mu + \sum_{\mu=1}^n \frac{a_\mu}{\mu} R\varphi_\mu \right) \\ &= \sum_{\mu, \nu=1}^n a_\mu \overline{a_\nu} (g_\mu, Rg_\nu) + 2 \operatorname{Re} \sum_{\mu, \nu=1}^n \frac{a_\mu \overline{a_\nu}}{\mu} (\varphi_\mu, Rg_\nu) \end{aligned}$$

$$+ \sum_{\mu=1}^n \frac{|a_\mu|^2}{\mu^2} (\varphi_\mu, R\varphi_\mu) + \sum_{\substack{\mu, \nu=1 \\ \mu \neq \nu}}^n \frac{a_\mu \bar{a}_\nu}{\mu\nu} (\varphi_\mu, R\varphi_\nu).$$

Consequently,

$$\begin{aligned} & \left(\sum_{\mu=1}^n a_\mu f_\mu, R \left\{ \sum_{\mu=1}^n a_\mu f_\mu \right\} \right) \\ & \geq - \sum_{\mu, \nu=1}^n |a_\mu| |a_\nu| |g_\mu| |Rg_\nu| - 2 \sum_{\mu, \nu=1}^n |a_\mu| |a_\nu| \frac{|\varphi_\mu| |Rg_\nu|}{\mu} + \sum_{\mu=1}^n |a_\mu|^2 \frac{C_\mu}{\mu^2} \\ & = \sum_{\mu=1}^n \frac{C_\mu}{\mu^2} |a_\mu|^2 - \sum_{\mu, \nu=1}^n \left(|g_\mu| |Rg_\nu| + \frac{2|Rg_\nu|}{\mu} \right) |a_\mu| |a_\nu|. \end{aligned}$$

Put

$$\bar{A}_{\mu\nu} = |g_\mu| |Rg_\nu| + \frac{2|Rg_\nu|}{\mu}.$$

This depends on $g_1, g_2, \dots$, but not on $f_1, f_2, \dots$ or on $\varphi_1, \varphi_2, \dots$, that is, not on $C_1, C_2, \dots$.^{T11} We estimate the subtracted term:

$$\begin{aligned} \sum_{\mu, \nu=1}^n \bar{A}_{\mu\nu} |a_\mu| |a_\nu| &= \sum_{\mu=1}^n \bar{A}_{\mu\mu} |a_\mu|^2 \\ &+ \sum_{\substack{\mu, \nu=1 \\ \mu < \nu}}^n (\bar{A}_{\mu\nu} + \bar{A}_{\nu\mu}) |a_\mu| |a_\nu| \\ &\leq \sum_{\mu=1}^n \bar{A}_{\mu\mu} |a_\mu|^2 \\ &+ \sum_{\substack{\mu, \nu=1 \\ \mu < \nu}}^n \left[2^{\nu-\mu-1} (\bar{A}_{\mu\nu} + \bar{A}_{\nu\mu})^2 |a_\mu|^2 + \frac{1}{2^{\nu-\mu}} |a_\nu|^2 \right] \\ &= \sum_{\mu=1}^n \left\{ \bar{A}_{\mu\mu} + \sum_{\rho=1}^{\mu-1} 2^{\mu-\rho-1} (\bar{A}_{\mu\rho} + \bar{A}_{\rho\mu})^2 + \sum_{\rho=\mu+1}^n \frac{1}{2^{\rho-\mu}} \right\} |a_\mu|^2 \end{aligned}$$

^{T11}Translator's note: The source changes the factor $2/\mu$ in the expansion to $2/\nu$ in the following bound and then drops the factor 2 in the equality and in the definition of $\bar{A}_{\mu\nu}$. The three occurrences have been made consistent with the preceding expansion.

$$\leq \sum_{\mu=1}^n \left\{ \bar{A}_{\mu\mu} + 1 + \sum_{\sigma=0}^{\mu-2} 2^\sigma (\bar{A}_{\mu, \mu-\sigma-1} + \bar{A}_{\mu-\sigma-1, \mu})^2 \right\} |a_\mu|^2.$$

Denote the expression in braces by $\bar{B}_\mu$. Then

$$\left(\sum_{\mu=1}^n a_\mu f_\mu, R \left\{ \sum_{\mu=1}^n a_\mu f_\mu \right\} \right) \geq \sum_{\mu=1}^n \left(\frac{C_\mu}{\mu^2} - \bar{B}_\mu \right) |a_\mu|^2.$$

It therefore suffices to set

$$C_\mu = \mu^2 \bar{B}_\mu,$$

and the proof is complete.^{60,T12}

⁶⁰For the purposes of the work announced in [note 22](#), we record here a certain sharpening of this estimate, although we do not need it here. We have

$$\begin{aligned}
 \left| \sum_{\mu=1}^n a_{\mu} f_{\mu} \right|^2 &= \left| \sum_{\mu=1}^n a_{\mu} g_{\mu} + \sum_{\mu=1}^n \frac{a_{\mu}}{\mu} \varphi_{\mu} \right|^2 \\
 &= \sum_{\mu, \nu=1}^n a_{\mu} \overline{a_{\nu}} (g_{\mu}, g_{\nu}) + 2 \operatorname{Re} \sum_{\mu, \nu=1}^n \frac{a_{\mu} \overline{a_{\nu}}}{\mu} (\varphi_{\mu}, g_{\nu}) \\
 &\quad + \sum_{\mu, \nu=1}^n \frac{a_{\mu} \overline{a_{\nu}}}{\mu \nu} (\varphi_{\mu}, \varphi_{\nu}) \\
 &\leq \sum_{\mu, \nu=1}^n \left\{ |g_{\mu}| |g_{\nu}| + \frac{2|\varphi_{\mu}| |g_{\nu}|}{\mu} + \frac{|\varphi_{\mu}| |\varphi_{\nu}|}{\mu \nu} \right\} |a_{\mu}| |a_{\nu}| \\
 &= \sum_{\mu, \nu=1}^n \left(|g_{\mu}| |g_{\nu}| + \frac{2|g_{\nu}|}{\mu} + \frac{1}{\mu \nu} \right) |a_{\mu}| |a_{\nu}|.
 \end{aligned}$$

Denote the expression in parentheses by $\overline{C}_{\mu\nu}$. As in the text,

$$\sum_{\mu, \nu=1}^n \overline{C}_{\mu\nu} |a_{\mu}| |a_{\nu}| \leq \sum_{\mu=1}^n \left\{ \overline{C}_{\mu\mu} + 1 + \sum_{\sigma=0}^{\mu-2} 2^{\sigma} (\overline{C}_{\mu, \mu-\sigma-1} + \overline{C}_{\mu-\sigma-1, \mu})^2 \right\} |a_{\mu}|^2.$$

Call this brace-enclosed expression $\overline{D}_{\mu}$. Then

$$\left| \sum_{\mu=1}^n a_{\mu} f_{\mu} \right|^2 \leq \sum_{\mu=1}^n \overline{D}_{\mu} |a_{\mu}|^2.$$

If we now put

$$C_{\mu} = \mu^2 (\overline{B}_{\mu} + \mu \overline{D}_{\mu}),$$

then, whenever $a_1 = \dots = a_{p-1} = 0$ and $n \geq p$,

$$\begin{aligned}
 &\left(\sum_{\mu=p}^n a_{\mu} f_{\mu}, R \left\{ \sum_{\mu=p}^n a_{\mu} f_{\mu} \right\} \right) - p \left| \sum_{\mu=p}^n a_{\mu} f_{\mu} \right|^2 \\
 &\geq \sum_{\mu=p}^n \left(\frac{C_{\mu}}{\mu^2} - \overline{B}_{\mu} - p \overline{D}_{\mu} \right) |a_{\mu}|^2 \geq 0.
 \end{aligned}$$

Thus, on the linear manifold spanned by

$$f_p, f_{p+1}, f_{p+2}, \dots,$$

We have already obtained a rather “pathological” result, which is easily sharpened further.

Theorem 47. *Let R be a linear H.O. that is not semibounded above or, respectively, not semibounded below. Then there is a hypermaximal H.O. S satisfying throughout*

$$(f, Sf) \geq C|f|^2, \quad \text{respectively} \quad (f, Sf) \leq C|f|^2,$$

such that R and S are both extensions of the same linear H.O. Here C may be chosen arbitrarily.

Proof. For $C = -1$ and failure of semiboundedness above, this follows from [Theorem 46](#) and [Theorem 43](#), in which one takes $C = 0$ and $C' = -1$. Every other value of C follows from this case by considering, in place of R, S ,

$$R - (C + 1)1, \quad S - (C + 1)1,$$

or, respectively,

$$-R + (C - 1)1, \quad -S + (C - 1)1.$$

The strangest consequences for unbounded Hermitian operators follow from this. The relation of a H.O. to its matrix, and of that matrix to its sections, is thereby qualitatively different from the bounded case. In particular, the “section spectra” of unbounded matrices have very little to do with the spectrum of the matrix itself and do not, in any sense, “approximate” it as they do in Hilbert’s theory. We shall say something about the operator–matrix relation in [Appendix III](#); the detailed discussion of the properties of unbounded matrices will be given in the work announced in [note 22](#).

We give several further characteristic properties of unbounded Hermitian operators. In contrast to [Theorems 41](#) through [47](#), we now restrict ourselves to closed linear operators.

Theorem 48. *Let R be a closed linear H.O. Then the following four assertions are*

which is again everywhere dense and hence nonclosed, one even has

$$(f, Rf) \geq p|f|^2$$

for every $p = 1, 2, \dots$

^{T12}*Translator’s note:* In note 60 the source drops the factor 2 from the bound on the mixed term, although it is present in the preceding exact expansion. The two occurrences have been restored here; consequently the definitions of $\bar{C}_{\mu\nu}$ and $\bar{D}_\mu$ use the corrected bound.

equivalent:

- $\alpha)$ R is bounded,
- $\beta)$ R is bounded and hypermaximal,
- $\gamma)$ R is defined everywhere,
- $\delta)$ there is no closed linear H.O. of which R is a proper extension.

61

Proof. Suppose first that R is not semibounded above. Choose C so that

$$(f, Rf) < C|f|^2$$

occurs, and then, by [Theorem 47](#), choose a hypermaximal H.O. S such that

$$(f, Sf) \geq C|f|^2$$

always holds. Plainly S cannot be an extension of R . Now R, S are extensions of a linear H.O. T , and, since they are closed and linear, they are also extensions of the closed linear H.O. $\tilde{T}$. Since S is not an extension of R , we must have $\tilde{T} \neq R$; thus $\delta)$ is not satisfied.

If R is not semibounded below, then $\delta)$ fails for $-R$, and hence also for R . Thus unboundedness implies the negation of $\delta)$; equivalently, $\delta)$ implies $\alpha)$.

Assertion $\alpha)$ implies $\gamma)$: for R is closed and continuous, so its domain is closed; but the domain is everywhere dense, and hence must equal $\mathfrak{H}$. Assertion $\gamma)$ implies $\delta)$: if R were a proper extension of a closed linear H.O. S , then S would not be maximal and so, by the corollary to [Theorem 35](#), could be extended to several maximal Hermitian operators, for example to an $R' \neq R$. Whenever $R'f$ and Sg are defined,

$$(R'f, g) = (f, R'g) = (f, Sg) = (f, Rg) = (Rf, g).$$

Since the set of such g is everywhere dense, $R'f = Rf$. Thus R would be an extension of R' , with $R' \neq R$, even though R' was to be maximal.

We have the logical scheme

$$\alpha) \implies \gamma) \implies \delta) \implies \alpha),$$

so $\alpha), \gamma), \delta)$ are logically equivalent. Assertion $\beta)$ implies $\alpha)$, while $\beta)$ follows from $\alpha)$ and $\gamma)$; hence it too is equivalent to them.

⁶¹The equivalence of $\alpha)$ and $\gamma)$ is the well-known theorem of Toeplitz, *Mathematische Annalen* **69** (1911), p. 321.

Theorem 49. *Let R be an unbounded closed linear H.O. Then it is an extension of a closed linear H.O. that can be extended to continuum many different maximal Hermitian operators.*

Proof. By [Theorem 48](#), δ), R must be a proper extension of a closed linear H.O. S . This S is not maximal and therefore, by the corollary to [Theorem 35](#), has the required property.

Appendix I. Function Spaces

We shall prove here the assertion announced in [Chapter I](#): all the function spaces mentioned in [Introduction I](#) and [II](#) are abstract Hilbert spaces, that is, they satisfy conditions [A](#) through [E](#) of [Chapter I](#). For the “ordinary Hilbert space” of all sequences

$$(x_1, x_2, \dots) \quad \text{with} \quad \sum_{n=1}^{\infty} |x_n|^2 < \infty,$$

this proof is unnecessary as soon as at least one system satisfying [A](#) through [E](#) is known: every such system is, by the concluding considerations of [Chapter I](#), isomorphic to the ordinary Hilbert space, which must therefore also possess properties [A](#) through [E](#). Hence we need consider only the function spaces of [Introduction I](#).

Let Ω be one of the metric spaces named there, or indeed any more general structure of which we require only the following. On Ω there are to be notions of “measurability,” a “measure,” and an “integral,” denoted

$$\int_{\Omega} f(P) dv,$$

in the Lebesgue sense. We shall not spell out precisely what this “Lebesgue sense” is to mean; it would be easy to describe, but somewhat lengthy. The considerations below will in any event reveal which properties of these notions matter. We also assume that Ω is separable.

The function space $\mathfrak{H}'$ consists of all measurable complex-valued functions $f(P)$ on Ω for which

$$\int_{\Omega} |f(P)|^2 dv$$

is finite; two functions differing only on a set of measure 0 are not regarded as

distinct. The meanings of af and $f + g$ are clear, and we put

$$(f, g) = \int_{\Omega} f(P) \overline{g(P)} dv.$$

These definitions are permissible. If $f, g \in \mathfrak{H}'$, so that

$$\int_{\Omega} |f(P)|^2 dv, \quad \int_{\Omega} |g(P)|^2 dv$$

are finite, then af and $f + g$ also belong to $\mathfrak{H}'$. Indeed,

$$\int_{\Omega} |af(P)|^2 dv$$

is plainly finite, while

$$\int_{\Omega} |f(P) + g(P)|^2 dv$$

is finite because

$$|f(P) + g(P)|^2 \leq 2|f(P)|^2 + 2|g(P)|^2.$$

Moreover,

$$\int_{\Omega} f(P) \overline{g(P)} dv$$

is meaningful, and indeed absolutely convergent, because

$$|f(P) \overline{g(P)}| \leq \frac{1}{2}|f(P)|^2 + \frac{1}{2}|g(P)|^2.$$

We now verify conditions A through E in order.

A, B. These are plainly satisfied.

C. We must find a sequence

$$f_1(P), f_2(P), \dots$$

which is everywhere dense with respect to the distance

$$|f - g| = \sqrt{\int_{\Omega} |f(P) - g(P)|^2 dv}.$$

The measure of Ω may be infinite, but in any case Ω can be exhausted by an increasing sequence

$$\Pi_1, \Pi_2, \dots$$

of subsets of finite measure. If f is any function in $\mathfrak{S}'$, define f_N , for $N = 1, 2, \dots$, by

$$f_N(P) = \begin{cases} f(P), & P \in \Pi_N \text{ and } |f(P)| \leq N, \\ 0, & \text{otherwise.} \end{cases}$$

For every P ,

$$f_N(P) \longrightarrow f(P),$$

and all the integrals

$$\int_{\Omega} |f_N(P)|^2 dv$$

are bounded above by $\int_{\Omega} |f(P)|^2 dv$. Hence

$$\int_{\Omega} |f_N(P) - f(P)|^2 dv \longrightarrow 0;$$

that is, $f_N \rightarrow f$ in our distance. Thus the functions which are nonzero only on a set of finite measure and which are also bounded are everywhere dense. It therefore suffices to find a sequence dense among them.

Let f be a function of this class. Let M be the measure of the set on which it is nonzero, and let A be the supremum of its absolute value. In the closed disk

$$\{z \in \mathbb{C} : |z| \leq A\},$$

choose finitely many Gaussian rational numbers $\rho_1, \dots, \rho_k$ (numbers in $\mathbb{Q} + i\mathbb{Q}$) forming an ε -net, where $\varepsilon > 0$ is otherwise arbitrary.^{T13} Put

$$f_{\varepsilon}(P) = \begin{cases} \rho_h, & f(P) \neq 0 \text{ and } \rho_h \text{ is nearest to } f(P) \text{ among } \rho_1, \dots, \rho_k, \\ 0, & \text{otherwise.} \end{cases}$$

It follows immediately that

$$\int_{\Omega} |f_{\varepsilon}(P) - f(P)|^2 dv \leq M\varepsilon^2.$$

Thus $f_{\varepsilon} \rightarrow f$ as $\varepsilon \rightarrow 0$. Consequently, those functions which take only Gaussian-rational values, take only finitely many distinct values, and for which the set

^{T13}*Translator's note:* The source partitions the real interval $[-A, A]$ using rational numbers, although $\mathfrak{S}'$ consists of complex-valued functions. We instead use a finite ε -net of Gaussian rationals in the disk $|z| \leq A$; the resulting bound below is $M\varepsilon^2$.

corresponding to each nonzero value has finite measure are also everywhere dense. The desired sequence need only be dense among these functions.

If Σ is any subset of Ω of finite measure, let f_Σ be the function equal to 1 on Σ and 0 elsewhere. The preceding class is simply the class of all functions

$$\rho_1 f_{\Sigma_1} + \cdots + \rho_k f_{\Sigma_k},$$

where $\rho_1, \dots, \rho_k$ are Gaussian-rational numbers. If we can find a sequence of sets

$$\Sigma^{(1)}, \Sigma^{(2)}, \dots$$

which is everywhere dense among all sets of finite measure—the distance between two sets Σ', Σ'' being the distance between their functions $f_{\Sigma'}, f_{\Sigma''}$, that is, the measure of their “difference set”

$$(\Sigma' + \Sigma'') - (\Sigma' \cdot \Sigma'')^{62}$$

—then the functions

$$\rho_1 f_{\Sigma^{(n_1)}} + \cdots + \rho_k f_{\Sigma^{(n_k)}}$$

with

$$\begin{aligned} k = 1, 2, \dots, \quad \rho_1, \dots, \rho_k &\in \mathbb{Q} + i\mathbb{Q}, \\ n_1, \dots, n_k &= 1, 2, \dots, \end{aligned}$$

would be everywhere dense in $\mathfrak{S}'$, and we would have reached our goal. It remains to find such a sequence of sets.

Let Σ be a set of finite measure M . As is well known, M is the infimum of the measures of all open point sets containing Σ . Let Σ' be such a set with measure at most $M + \varepsilon$. Then the “difference set” $\Sigma' - \Sigma$ has measure at most ε . Thus the open sets are everywhere dense.

Since Ω is separable, there is in Ω a sequence of open sets

$$\Lambda_1, \Lambda_2, \dots$$

forming an “equivalent neighborhood system,” in the sense that every open set Σ' is the union of all the Λ_n contained in it—say, of

$$\Lambda_{n_1}, \Lambda_{n_2}, \dots$$

⁶²Here $+$, $-$, and $\cdot$ denote union, set difference, and intersection, respectively.

Let the measure of Σ' be M' . Then the measures of

$$\Lambda_{n_1}, \quad \Lambda_{n_1} + \Lambda_{n_2}, \quad \Lambda_{n_1} + \Lambda_{n_2} + \Lambda_{n_3}, \quad \dots$$

converge to M' . If, for example, the measure of

$\Lambda_{n_1} + \dots + \Lambda_{n_k}$ is already at least $M' - \varepsilon$, then the “difference set”

$$\Sigma' - (\Lambda_{n_1} + \dots + \Lambda_{n_k})$$

has measure at most ε . Thus the sets

$$\Lambda_{n_1} + \dots + \Lambda_{n_k} \quad (k = 1, 2, \dots; n_1, \dots, n_k = 1, 2, \dots)$$

are likewise everywhere dense,⁶³ and there are only countably many of them. This brings us to our goal.

D. Let $k = 1, 2, \dots$. Choose k pairwise disjoint sets $K_1, \dots, K_k$ in Ω , each of finite, nonzero measure. Let $f_h(P) = 1$ on K_h and $f_h(P) = 0$ elsewhere ($h = 1, \dots, k$). Then $f_1, \dots, f_k$ are plainly linearly independent (and belong to $\mathfrak{S}'$).

E. Let $f_1, f_2, \dots$ be a sequence of functions satisfying Cauchy's convergence criterion: for every $\varepsilon > 0$ there is an $N = N(\varepsilon)$ such that, for $m, n \geq N$,

$$\int_{\Omega} |f_m(P) - f_n(P)|^2 dv \leq \varepsilon.$$

Following a proof of Weyl, choose

$$N_1 < N_2 < N_3 < \dots, \quad N_p \geq N\left(\frac{1}{8^p}\right).$$

Then

$$\int_{\Omega} |f_{N_{p+1}}(P) - f_{N_p}(P)|^2 dv \leq \frac{1}{8^p}.$$

Consequently, the set on which

$$|f_{N_{p+1}}(P) - f_{N_p}(P)| \geq \frac{1}{2^p}$$

has measure at most 2^{-p} . Hence

$$|f_{N_{p+1}}(P) - f_{N_p}(P)| \leq \frac{1}{2^p}, \quad |f_{N_{p+2}}(P) - f_{N_{p+1}}(P)| \leq \frac{1}{2^{p+1}}, \dots$$

⁶³Naturally, we take only those of finite measure.

everywhere except on a set of measure at most

$$\frac{1}{2^p} + \frac{1}{2^{p+1}} + \cdots = \frac{1}{2^{p-1}}.$$

Where these inequalities hold, the sequence $f_{N_p}(P), f_{N_{p+1}}(P), \dots$ has a limit. It therefore has a limit everywhere except on a set of measure at most 2^{1-p} . Since this is true for every p , the exceptional set has measure zero.

Define

$$f(P) = \lim_{p \rightarrow \infty} f_{N_p}(P)$$

where this limit exists, that is, everywhere except on a null set, and put $f(P) = 0$ otherwise. If $m \geq N(\varepsilon)$, then, for all sufficiently large p (namely, as soon as $N_p \geq N(\varepsilon)$),

$$\int_{\Omega} |f_m(P) - f_{N_p}(P)|^2 dv \leq \varepsilon.$$

We may therefore let $p \rightarrow \infty$, obtaining

$$\int_{\Omega} |f_m(P) - f(P)|^2 dv \leq \varepsilon.$$

Thus $f_m - f$, and hence f , belongs to $\mathfrak{H}'$, and secondly $f_m \rightarrow f$. This proves everything asserted.

Appendix II. The Eigenvalue Problem for Unitary Operators

1.

Throughout what follows we consider everywhere-defined, bounded (continuous) linear operators R on $\mathfrak{H}$; thus [Theorem 12\(α\), \(β\)](#) applies to them. Hermitian character, however, will not always be assumed. For fixed f , the functional

$$L(g) = (f, Rg)$$

has the properties

$$L(a_1 g_1 + \cdots + a_n g_n) = \overline{a_1} L(g_1) + \cdots + \overline{a_n} L(g_n), \quad |L(g)| \leq C |f| |g| = D |g|.$$

The theorem of F. Riesz (cf. [Note 52](#)) is therefore applicable: $L(g) = (f^*, g)$. We denote the uniquely determined f^* by R^*f . Then R^* is an everywhere-defined and plainly linear operator, and

$$(f, Rg) = (R^*f, g), \quad (Rf, g) = (f, R^*g).$$

The first equation is the definition; the second follows by interchanging f, g and taking complex conjugates. In particular,

$$|R^*f|^2 = (R^*f, R^*f) = (f, RR^*f) \leq C|f||R^*f|, \quad |R^*f| \leq C|f|,$$

so R^* too is bounded.⁶⁴

One verifies at once that

$$\begin{aligned} R^{**} &= R, & (R \pm S)^* &= R^* \pm S^*, \\ (RS)^* &= S^*R^*, & (aR)^* &= \bar{a}R^*, \\ 0^* &= 0, & 1^* &= 1. \end{aligned}$$

⁶⁵ For bounded Hermitian operators, $R = R^*$ is characteristic. Every unitary operator U belongs to our class (it is bounded because $|Uf| = |f|$), and within this class it is characterized by

$$UU^* = U^*U = 1, \quad \text{that is,} \quad U^{-1} = U^*.$$

⁶⁶

2.

Now suppose that R is Hermitian. We first attack its eigenvalue problem by the method of F. Riesz. Let

$$p(x) = \sum_{v=0}^n a_v x^v$$

⁶⁴For matrices, this star means taking the conjugate transpose; for integral-operator kernels it has the analogous meaning, and for differential operators it means taking the adjoint.

⁶⁵The operators in question are everywhere defined, so the operations $+$, $-$, and multiplication are meaningful without further qualification; cf. [Note 37](#).

⁶⁶If U is unitary, it has an inverse U^{-1} , and

$$(f, U^{-1}g) = (Uf, UU^{-1}g) = (Uf, g) = (f, U^*g)$$

for all f, g ; hence $U^{-1} = U^*$. Conversely, if $U^{-1} = U^*$, then $(Uf, Ug) = (f, U^*Ug) = (f, g)$, so that $|Uf| = |f|$. Thus U is isometric; since it is everywhere defined and assumes every value (U^{-1} is everywhere defined), it is unitary.

be a polynomial with real coefficients.^{T14} Then

$$p(R) = \sum_{\nu=0}^n a_{\nu} R^{\nu}, \quad R^0 = 1,$$

is a Hermitian operator of our class. Addition, subtraction, and multiplication of these $p(R)$ are performed exactly as for the polynomials $p(x)$; in particular, all of them commute. Choose C such that $|Rf| \leq C|f|$ for every f .

Theorem 1*. *If $p(x) \geq 0$ for $-C \leq x \leq C$, then $p(R)$ is positive.*

Proof. Let $\mathfrak{M}$ be any finite-dimensional, and hence closed, linear manifold, spanned for example by the orthonormal system $\varphi_1, \dots, \varphi_k$. The operator

$$R' = P_{\mathfrak{M}} R P_{\mathfrak{M}}$$

is Hermitian and satisfies

$$|R'f| = |P_{\mathfrak{M}} R P_{\mathfrak{M}} f| \leq |R P_{\mathfrak{M}} f| \leq C |P_{\mathfrak{M}} f| \leq C |f|.$$

It is therefore of our class and is plainly reduced by $\mathfrak{M}$ (its part in $\mathfrak{S} - \mathfrak{M}$ is zero). Its part in $\mathfrak{M}$ is a k -dimensional Hermitian operator, say with eigenvalues $\lambda_1, \dots, \lambda_k$. Since $|R'f| \leq C|f|$, all $|\lambda_j| \leq C$. The operator $p(R')$ is also reduced by $\mathfrak{M}$, and its part there has the eigenvalues $p(\lambda_1), \dots, p(\lambda_k)$, all nonnegative. Thus, if $f \in \mathfrak{M}$,

$$(f, p(R')f) \geq 0.$$

Now let f be arbitrary, and take for $\mathfrak{M}$ the closed linear manifold spanned by $f, Rf, \dots, R^n f$. Then $f \in \mathfrak{M}$, $R'f = Rf, \dots, R'^n f = R^n f$, and therefore $p(R')f = p(R)f$. Hence $(f, p(R)f) \geq 0$. Since f was arbitrary, $p(R)$ is positive.

Theorem 2*. *If $|p(x)| \leq \varepsilon$ for $-C \leq x \leq C$, then*

$$|p(R)f| \leq \varepsilon |f|.$$

Proof. **Theorem 1*** applies to the polynomials $\varepsilon \pm p(x)$; hence the operators $\varepsilon 1 \pm p(R)$ are positive, and

$$(f, \{\varepsilon 1 \pm p(R)\}f) \geq 0, \quad |(f, p(R)f)| \leq \varepsilon |f|^2.$$

^{T14}*Translator's note:* The source begins this sum at $\nu = 1$, although the immediately following definition of $p(R)$ begins at $\nu = 0$ and explicitly includes $R^0 = 1$. The lower limit has therefore been normalized to 0 here.

By [Theorem 12](#) (the equivalence of γ and α), this means precisely that $|p(R)f| \leq \varepsilon|f|$.

Theorem 3*. *Let $P(x)$ be continuous on $-C \leq x \leq C$, and let $p_1(x), p_2(x), \dots$ be polynomials converging there uniformly to $P(x)$. Then $p_n(R)f$ has a limit f^* depending only on R, f , and $P(x)$.*

Proof. This follows immediately from [Theorem 2*](#).

We therefore define $P(R)f = f^*$. This gives an everywhere-defined linear Hermitian operator, because the $p_n(R)$ have those properties; it is bounded by [Theorem 48](#)(α), (γ), or directly because the $p_n(x)$, and hence by [Theorem 2*](#) the $p_n(R)$, are uniformly bounded.

Theorem 4*. *The $P(R)$ are added, subtracted, and multiplied just as the $P(x)$ are. Theorems 1* and 2* hold for them as well; and when $P(x)$ is a polynomial, the old and new definitions agree.*

Proof. The last assertion follows by taking $p_1(x) = p_2(x) = \dots = P(x)$. The remaining assertions hold for the polynomials $p(x)$, and hence, by continuity, also for $P(x)$.

3.

We now depart from F. Riesz's route.

Theorem 5*. *Let $\mathfrak{M}_R$ be the set of all f for which $Rf = 0$, and let $\mathfrak{N}_R$ be the set of all Rf and their limit points. Then*

$$\mathfrak{M}_R = \mathfrak{S} - \mathfrak{N}_R.$$

Proof. The vectors Rf form a linear manifold. It is enough to show that $\mathfrak{M}_R$ consists of all g orthogonal to this manifold. Orthogonality to it means

$$(g, Rf) = 0 \quad \text{for every } f,$$

that is, $(Rg, f) = 0$, and hence $Rg = 0$.

Theorem 6*. *If R, S are positive, then $\mathfrak{M}_{R+S}$ is a part of $\mathfrak{M}_R$, and $\mathfrak{N}_R$ is a part of $\mathfrak{N}_{R+S}$.*

Proof. From $(R + S)g = 0$ it follows that

$$0 \leq (g, Rg) \leq (g, (R + S)g) = 0,$$

and hence $(g, Rg) = 0$. For positive R one has

$$|(f, Rg)|^2 \leq (f, Rf)(g, Rg).$$

This is proved just like the corresponding inequality for $R = 1$ in [Theorem 1](#). Thus $(g, Rg) = 0$ implies $(f, Rg) = 0$ for every f , and therefore $Rg = 0$. The preceding relation says that $\mathfrak{M}_{R+S}$ is a part of $\mathfrak{M}_R$; [Theorem 5*](#) then shows that $\mathfrak{N}_R$ is a part of $\mathfrak{N}_{R+S}$.

We define

$$f_a(x) = \text{Max}(x - a, 0), \quad \mathfrak{M}(R, a) = \mathfrak{M}_{f_a(R)}, \quad \mathfrak{N}(R, a) = \mathfrak{N}_{f_a(R)}.$$

Theorem 7*. *Throughout $\mathfrak{M}(R, a)$, respectively $\mathfrak{N}(R, a)$, one has*

$$(f, Rf) \leq a|f|^2, \quad \text{respectively} \quad (f, Rf) \geq a|f|^2.$$

Proof. Put $g_a(x) = \text{Max}(-(x - a), 0)$. Since

$$f_a(x)g_a(x) = 0, \quad f_a(x) - g_a(x) = x - a,$$

we have, with $S' = f_a(R)$ and $S'' = g_a(R)$,

$$S'S'' = S''S' = 0, \quad S' - S'' = R - a1$$

by [Theorem 4*](#). Both S' and S'' are positive. In $\mathfrak{M}(R, a) = \mathfrak{M}_{S'}$ we have $S'f = 0$, so

$$Rf = af - S''f, \quad (f, Rf) = a|f|^2 - (f, S''f) \leq a|f|^2.$$

In $\mathfrak{N}(R, a) = \mathfrak{N}_{S'}$, by [Theorem 5*](#) and $S''S' = 0$, we lie in $\mathfrak{M}_{S''}$. Thus $S''f = 0$, and

$$Rf = af + S'f, \quad (f, Rf) = a|f|^2 + (f, S'f) \geq a|f|^2.$$

Theorem 8*. *If $a \leq b$, then $\mathfrak{M}(R, a)$ is a part of $\mathfrak{M}(R, b)$, and $\mathfrak{N}(R, b)$ is a part of $\mathfrak{N}(R, a)$. Moreover, for $a < -C$, respectively $a > C$,*

$$\mathfrak{M}(R, a) = (0), \quad \mathfrak{M}(R, a) = \mathfrak{H}, \quad \mathfrak{N}(R, a) = \mathfrak{H}, \quad \mathfrak{N}(R, a) = (0),$$

respectively.

Proof. Since $f_b(x) \geq 0$ and $f_a(x) - f_b(x) \geq 0$, the first assertion follows from [Theorem 6*](#), applied to $f_b(R)$ and $f_a(R) - f_b(R)$.^{T15} For the second, if a nonzero f belonged, when $a < -C$, to $\mathfrak{M}(R, a)$, or, when $a > C$, to $\mathfrak{N}(R, a)$, then

$$(f, Rf) \leq a|f|^2 < -C|f|^2, \quad \text{respectively} \quad (f, Rf) \geq a|f|^2 > C|f|^2,$$

^{T15}*Translator's note:* The source prints $f_b - f_a \geq 0$ and applies the theorem to f_a and $f_b - f_a$. Since $a \leq b$ gives $f_a \geq f_b$, both signs have been corrected.

contrary to the definition of C in [Theorem 12](#)(γ). Hence the corresponding manifold is (0) , and its orthogonal complement is $\mathfrak{S} - (0) = \mathfrak{S}$.

Theorem 9*. *Quite generally,*

$$(Rf, g) = \int a \, d(P_{\mathfrak{M}(R,a)}f, g),$$

where the integral may be taken between any two limits smaller than $-C$ and greater than C .

Proof. For $f = g$, the assertion follows from [Theorems 7*](#) and [8*](#).⁶⁷ Substituting $(f + g)/2$ and $(f - g)/2$ and subtracting gives the equation for the real parts in

⁶⁷The operator R commutes with $f_a(R)$, by [Theorem 4*](#), and is therefore reduced by $\mathfrak{M}(R, a)$: from $f_a(R)f = 0$ follows $f_a(R)Rf = Rf_a(R)f = 0$, and one applies [Theorem 20](#). Thus R commutes with $P_{\mathfrak{M}(R,a)}$.

Let

$$x_0 < x_1 < \cdots < x_p, \quad x_0 < -C, \quad x_p > C.$$

Then

$$\begin{aligned} (Rf, f) &= \sum_{v=1}^p (R(P_{\mathfrak{M}(R,x_v)} - P_{\mathfrak{M}(R,x_{v-1})})f, f) \\ &= \sum_{v=1}^p (RP_{\mathfrak{M}(R,x_v) - \mathfrak{M}(R,x_{v-1})}f, f) \\ &= \sum_{v=1}^p (RP_{\mathfrak{M}(R,x_v)\mathfrak{M}(R,x_{v-1})}f, f) \\ &= \sum_{v=1}^p (RP_{\mathfrak{M}(R,x_v)\mathfrak{M}(R,x_{v-1})}f, P_{\mathfrak{M}(R,x_v)\mathfrak{M}(R,x_{v-1})}f). \end{aligned}$$

The last step uses the commutativity just noted. By [Theorem 7*](#), each summand lies between

$$x_{v-1}|P_{\mathfrak{M}(R,x_v) - \mathfrak{M}(R,x_{v-1})}f|^2 \quad \text{and} \quad x_v|P_{\mathfrak{M}(R,x_v) - \mathfrak{M}(R,x_{v-1})}f|^2,$$

and

$$|P_{\mathfrak{M}(R,x_v) - \mathfrak{M}(R,x_{v-1})}f|^2 = |P_{\mathfrak{M}(R,x_v)}f|^2 - |P_{\mathfrak{M}(R,x_{v-1})}f|^2.$$

Consequently,

$$\begin{aligned} (Rf, f) &\leq \sum_{v=1}^p x_v (|P_{\mathfrak{M}(R,x_v)}f|^2 - |P_{\mathfrak{M}(R,x_{v-1})}f|^2), \\ (Rf, f) &\geq \sum_{v=1}^p x_{v-1} (|P_{\mathfrak{M}(R,x_v)}f|^2 - |P_{\mathfrak{M}(R,x_{v-1})}f|^2). \end{aligned}$$

general; replacing f, g by if, g then gives the imaginary parts.

Theorem 9* is essentially Hilbert's result for bounded Hermitian operators. We proceed to two Hermitian operators R, S of our class.

Theorem 10*. *If R, S commute, then so do*

$$P_{\mathfrak{M}(R,a)} \quad \text{and} \quad P_{\mathfrak{M}(S,b)}.$$

Proof. Since R, S commute, so do all $P(R), Q(S)$: this is immediate for polynomials and follows by continuity for general continuous functions. In particular, $f_a(R)$ commutes with $f_b(S)$. Hence $\mathfrak{M}(R, a)$ reduces $f_b(S)$, by the argument in the preceding footnote, and so does $\mathfrak{N}(R, a) = \mathfrak{H} - \mathfrak{M}(R, a)$. Thus $\mathfrak{M}(S, b)$ is obtained by first taking, within $\mathfrak{M}(R, a)$, the f for which $f_b(S)f = 0$, then doing the same within $\mathfrak{N}(R, a)$, and finally adding the two closed linear manifolds. The commutativity of the two projection operators follows at once.

4.

We now turn to the so-called normal operators A of our class, characterized by

$$AA^* = A^*A.$$

The Hermitian operators

$$R = \frac{1}{2}(A + A^*), \quad S = \frac{1}{2i}(A - A^*)$$

then commute, and

$$A = R + iS, \quad A^* = R - iS.$$

Form the projection operators

$$E(A, a, b) = P_{\mathfrak{M}(R,a)}P_{\mathfrak{M}(S,b)}.$$

Theorem 11*. *The operator A and the $E(A, a, b)$ satisfy the following conditions:*⁶⁸

These are the upper and lower Stieltjes sums. By the definition of the Stieltjes integral, they yield

$$(Rf, f) = \int a \, d|P_{\mathfrak{M}(R,a)}f|^2 = \int a \, d(P_{\mathfrak{M}(R,a)}f, f).$$

⁶⁸Cf. [Introduction IX](#) and [Note 58](#).

α The operators $E(A, a, b)$ and $E(A, a', b')$ commute, and

$$E(A, a, b)E(A, a', b') = E(A, \min(a, a'), \min(b, b')).$$

β For sufficiently small a or b , $E(A, a, b) = 0$. For sufficiently large a , respectively b , it is independent of a , respectively b ; and for sufficiently large a, b , it equals 1.

Moreover,

$$\begin{aligned} (Af, g) &= \iint (a + ib) d(E(A, a, b)f, g), \\ (A^*f, g) &= \iint (a - ib) d(E(A, a, b)f, g), \end{aligned}$$

where the double integrals extend over a sufficiently large rectangle.

Proof. By [Theorems 9*](#) and [10*](#),

$$(Rf, g) = \iint a d(E(A, a, b)f, g), \quad (Sf, g) = \iint b d(E(A, a, b)f, g).$$

Everything now follows immediately from [Theorem 8*](#), [Theorem 9*](#), [Theorem 10*](#), and the definition of $E(A, a, b)$.

Unitary operators U are normal, since $UU^* = U^*U = 1$, so [Theorem 11*](#) applies to them. Their $E(U, a, b)$, however, have further properties, which we now derive.

First, quite generally,

$$\begin{aligned} (Af, E(A, a, b)g) &= \iint (a' + ib') d(E(A, a', b')f, E(A, a, b)g) \\ &= \iint (a' + ib') d(E(A, a, b)E(A, a', b')f, g) \\ &= \iint (a' + ib') d(E(A, \min(a, a'), \min(b, b'))f, g) \\ &= \int_{-\infty}^a \int_{-\infty}^b (a' + ib') d(E(A, a', b')f, g), \end{aligned}$$

and

$$\begin{aligned} (Af, Ag) &= \iint (a - ib) d(Af, E(A, a, b)g) \\ &= \iint (a - ib) d\left\{ \int_{-\infty}^a \int_{-\infty}^b (a' + ib') d(E(A, a', b')f, g) \right\} \end{aligned}$$

$$\begin{aligned}
&= \iint (a - ib)(a + ib) d(E(A, a, b)f, g) \\
&= \iint (a^2 + b^2) d(E(A, a, b)f, g),
\end{aligned}$$

where the interchange in the second calculation is as in [Note 55](#). Also

$$(f, g) = \iint d(E(A, a, b)f, g).$$

For unitary A , therefore, $(f, g) = (Af, Ag)$ gives

$$\iint (1 - a^2 - b^2) d(E(A, a, b)f, g) = 0,$$

and, on putting $f = g$,

$$\iint (1 - a^2 - b^2) d|E(A, a, b)f|^2 = 0.$$

5.

If $\mathcal{R}$ is a rectangle

$$a' \leq x \leq a'', \quad b' \leq y \leq b''$$

with vertices $(a', b'), (a'', b'), (a', b''), (a'', b'')$, define the projection operator^{T16}

$$\begin{aligned}
E(A, \mathcal{R}) &= E(A, a'', b'') - E(A, a'', b') - E(A, a', b'') + E(A, a', b') \\
&= P_{\mathfrak{N}(R, a')} (1 - P_{\mathfrak{N}(R, a'')}) P_{\mathfrak{N}(S, b')} (1 - P_{\mathfrak{N}(S, b'')}) \\
&= P_{\mathfrak{N}(R, a') \mathfrak{M}(R, a'') \mathfrak{N}(S, b') \mathfrak{M}(S, b'')}.
\end{aligned}$$

Suppose that $\mathcal{R}$ has no point in common with the unit circle. Then throughout $\mathcal{R}$ either $1 - a^2 - b^2 \geq \varepsilon$, or $1 - a^2 - b^2 \leq -\varepsilon$, for some $\varepsilon > 0$. For every f such that $E(A, \mathcal{R})f = f$ (that is, f lies in $\mathfrak{N}(R, a') \mathfrak{M}(R, a'') \mathfrak{N}(S, b') \mathfrak{M}(S, b'')$)⁶⁹, we therefore have

$$\left| \iint_{\mathcal{R}} (1 - a^2 - b^2) d|E(A, a, b)f|^2 \right| = \left| \iint (1 - a^2 - b^2) d|E(A, a, b)f|^2 \right|$$

^{T16}Translator's note: The source reverses $\mathfrak{M}$ and $\mathfrak{N}$ in this rectangle projection, the following parenthesis, and Note 69. Since $E(A, a, b)$ is increasing in each endpoint, the interval factors are $\mathfrak{N}(R, a') \mathfrak{M}(R, a'')$ and $\mathfrak{N}(S, b') \mathfrak{M}(S, b'')$, as corrected here.

⁶⁹Since f lies in both $\mathfrak{N}(R, a')$ and $\mathfrak{M}(R, a'')$, $P_{\mathfrak{M}(R, a)}f$ equals 0 for $a \leq a'$ and f for $a \geq a''$; thus $E(A, a, b)f$ is independent of a outside this interval. The same is true of b outside $b' \leq b \leq b''$. Accordingly, the integral over the whole plane may be identified with the integral over $\mathcal{R}$.

$$\geq \varepsilon \left| \iint d|E(A, a, b)f|^2 \right| = \varepsilon |f|^2.$$

Thus $f = 0$, and $E(A, \mathcal{R}) = 0$.

If $\mathcal{R}, \mathcal{S}$ are disjoint rectangles (their boundaries may meet),^{70, T17} then

$$E(A, \mathcal{R})E(A, \mathcal{S}) = 0.$$

If two rectangles $\mathcal{R}, \mathcal{S}$ together make exactly one rectangle $\mathcal{T}$,⁷¹ then

$$E(A, \mathcal{R}) + E(A, \mathcal{S}) = E(A, \mathcal{T}).$$

Repeated use of these facts yields the following. If $\mathcal{U}$ is a union of finitely many rectangles $\mathcal{R}_1, \dots, \mathcal{R}_n$, then

$$E(A, \mathcal{U}) = E(A, \mathcal{R}_1) + \dots + E(A, \mathcal{R}_n)$$

is a projection operator; it is independent of the rectangular decomposition of $\mathcal{U}$. If $\mathcal{U}, \mathcal{V}$ are two such sets having at most boundary points in common, then

$$E(A, \mathcal{U})E(A, \mathcal{V}) = 0, \quad E(A, \mathcal{U} + \mathcal{V}) = E(A, \mathcal{U}) + E(A, \mathcal{V}).$$

Moreover, if $\mathcal{U}$ does not meet the unit circle, then $E(A, \mathcal{U}) = 0$. It follows easily that if $\mathcal{U}, \mathcal{V}$ have the same points in common with the unit circle, and neither has a corner on that circle, then

$$E(A, \mathcal{U}) = E(A, \mathcal{V}).$$

We can now, so to speak, transform the Stieltjes double integral to polar coordinates.

⁷⁰Write, for example,

$$\mathcal{R} : a' \leq x \leq a'', b' \leq y \leq b'', \quad \mathcal{S} : c' \leq x \leq c'', d' \leq y \leq d''.$$

One of $a'' \leq c', c'' \leq a', b'' \leq d'$, or $d'' \leq b'$ must hold. In the first case $\mathfrak{M}(R, a'')$ is a part of $\mathfrak{M}(R, c')$, so $\mathfrak{M}(R, a'')$ is orthogonal to $\mathfrak{N}(R, c')$. Hence $P_{\mathfrak{M}(R, a'')}P_{\mathfrak{N}(R, c')} = 0$, and a fortiori $E(A, \mathcal{R})E(A, \mathcal{S}) = 0$. The other three cases are identical.

^{T17}*Translator's note:* In Note 70 the source instead states $\mathfrak{N}(R, a'') \perp \mathfrak{M}(R, c')$ and writes the corresponding product of projections. The complementary factors have been corrected here; the corrected orthogonality is the one needed for the conclusion.

⁷¹Either $a'' = c'$ or $c'' = a'$, while simultaneously $b' = d'$ and $b'' = d''$; or $b'' = d'$ or $d'' = b'$, while simultaneously $a' = c'$ and $a'' = c''$. The assertion follows directly from the definition of $E(A, \mathcal{R})$.

6.

Let $0 \leq \rho < \sigma \leq 1$. There certainly exist finite unions of rectangles $\mathcal{U}$ whose intersection with the unit circle is exactly the arc from $2\pi\rho$ to $2\pi\sigma$, and none of whose corners lies on the unit circle. For all such $\mathcal{U}$, $E(A, \mathcal{U})$ is the same projection operator; denote it by $F(\rho, \sigma)$. For $0 \leq \rho < \sigma < \tau \leq 1$,

$$F(\rho, \sigma) + F(\sigma, \tau) = F(\rho, \tau).$$

Putting $F(\rho) = F(0, \rho)$, we therefore have

$$F(\rho, \sigma) = F(\sigma) - F(\rho).$$

Since plainly $F(0, 1) = 1$, it follows that $F(1) = 1$. The $F(\rho)$ are projection operators, and because $F(\sigma) - F(\rho)$ is a projection operator for $\rho \leq \sigma$, $F(\rho)$ is a part of $F(\sigma)$. Define the hitherto undefined $F(0)$ to be 0.

Partition the integration region in

$$(Af, f) = \iint (a + ib) d|E(A, a, b)f|^2$$

(a large rectangle) into rectangles $\mathcal{R}_1, \dots, \mathcal{R}_k$ as follows. Each has diameter at most $\varepsilon > 0$, and none has a corner on the unit circle. Order them so that $\mathcal{R}_1, \dots, \mathcal{R}_h$, where $h \leq k$, meet the circle in the respective arcs

$$2\pi\rho_0, 2\pi\rho_1; \dots; 2\pi\rho_{h-1}, 2\pi\rho_h, \quad 0 = \rho_0 < \rho_1 < \dots < \rho_h = 1.$$

Choose a point ξ_v from each $\mathcal{R}_v$. Then

$$(Af, f) = \sum_{v=1}^k \iint_{\mathcal{R}_v} (a + ib) d|E(A, a, b)f|^2$$

differs by an arbitrarily small amount⁷² from

$$\sum_{v=1}^k \xi_v \iint_{\mathcal{R}_v} d|E(A, a, b)f|^2 = \sum_{v=1}^k \xi_v (E(A, \mathcal{R}_v)f, f)$$

⁷²For every rectangle $\mathcal{S}$, $\iint_{\mathcal{S}} d|E(A, a, b)f|^2 \geq 0$, since it equals $|E(A, \mathcal{S})f|^2$. Thus the difference is at most

$$\varepsilon \iint d|E(A, a, b)f|^2 = \varepsilon |f|^2.$$

$$\begin{aligned}
&= \sum_{v=1}^h \xi_v |F(\rho_{v-1}, \rho_v) f|^2 \\
&= \sum_{v=1}^h \xi_v (|F(\rho_v) f|^2 - |F(\rho_{v-1}) f|^2).
\end{aligned}$$

We may choose

$$\xi_v = e^{2\pi i \rho'_v}, \quad \rho_{v-1} < \rho'_v < \rho_v \quad (v = 1, \dots, h),$$

so the last expression becomes

$$\sum_{v=1}^h e^{2\pi i \rho'_v} (|F(\rho_v) f|^2 - |F(\rho_{v-1}) f|^2).$$

The points $\rho_0, \rho_1, \dots, \rho_h$ are at most ε apart. Hence

$$(Af, f) = \int_0^1 e^{2\pi i \rho} d|F(\rho) f|^2.$$

73

Theorem 12*. For every unitary operator U there is a family of projection operators $E(\lambda)$, $0 \leq \lambda \leq 1$, with the following properties:

α If $\lambda \leq \mu$, then $E(\lambda)$ is a part of $E(\mu)$; moreover $E(0) = 0$ and $E(1) = 1$.

β If $\lambda \geq \lambda_0$ and $\lambda \rightarrow \lambda_0$, then, for every f ,

$$E(\lambda)f \longrightarrow E(\lambda_0)f.$$

For all f, g ,

$$(Uf, g) = \int_0^1 e^{2\pi i \lambda} d(E(\lambda)f, g).$$

Proof. Consider the $F(\lambda)$ just constructed. They have property α , but may fail to have β . Let

$$0 \leq \lambda_0 < 1, \quad \lambda_1 > \lambda_2 > \dots > \lambda_0, \quad \lambda_n \rightarrow \lambda_0.$$

The $F(\lambda_n)$ form a decreasing sequence of projection operators. By [Theorem 19](#), there is a projection operator E such that

$$F(\lambda_n)f \longrightarrow Ef$$

⁷³The integral on the right converges; for example, prescribe $\rho_0, \rho_1, \dots, \rho_h$ and adapt the rectangular partition to them.

for every f . This E depends only on λ_0 : for any two such sequences, suitable subsequences can be interlaced to give another sequence of the same kind, for which the convergence must also hold. Denote this E by $E(\lambda_0)$.

Since $F(\lambda_0)$ is a part of every $F(\lambda_1), F(\lambda_2), \dots$, it is a part of $E(\lambda_0)$. If $\lambda_0 < \mu_0$, choose $\lambda_1 = \mu_0$; then $E(\lambda_0)$ is a part of $F(\mu_0) = F(\lambda_1)$, since it is a part of all the $F(\lambda_n)$. Thus

$$F(\lambda) \text{ is a part of } E(\lambda), \quad E(\lambda) \text{ is a part of } F(\mu) \quad (\lambda < \mu),$$

and consequently $E(\lambda)$ is a part of $E(\mu)$ for $\lambda \leq \mu$.

Now let $\lambda > \lambda_0, \lambda \rightarrow \lambda_0$. The projection $E(\lambda) - E(\lambda_0)$ is a part of $F(\lambda) - E(\lambda_0)$. Hence

$$|E(\lambda)f - E(\lambda_0)f| \leq |F(\lambda)f - E(\lambda_0)f|,$$

and the right-hand side tends to zero by the definition of $E(\lambda_0)$. Thus the $E(\lambda)$ satisfy β and also α , apart from the requirements $E(0) = 0, E(1) = 1$. We enforce those by definition; β may then fail at $\lambda = 0$. Denote the old value of $E(0)$ by $\bar{E}$.

For $F(\lambda)$ we had

$$(Uf, f) = \int_0^1 e^{2\pi i \lambda} d|F(\lambda)f|^2.$$

The functions $|F(\lambda)f|^2$ and $|E(\lambda)f|^2$ are monotone in λ , and, for $\lambda < \mu$,

$$|F(\lambda)f|^2 \leq |E(\lambda)f|^2 \leq |F(\mu)f|^2.$$

It follows readily that the same integral relation holds with $E(\lambda)$. This is the asserted final relation when $f = g$. Substitution of $(f + g)/2$ and $(f - g)/2$, followed by subtraction, gives its real part for arbitrary f, g ; replacing f, g by if, g gives its imaginary part.

It remains to repair the possible failure of β at $\lambda = 0$. For $0 < \lambda < 1$, replace $E(\lambda)$ by $E(\lambda) - \bar{E}$. Everything else remains unchanged and this last point is put in order.

7.

Theorem 12* was the principal difficulty. Some easy uniqueness theorems now follow.

Theorem 13*. *If projection operators $E(\lambda)$ are chosen in accordance with conditions α, β of Theorem 12*, then there is exactly one unitary operator U related to them as in that theorem.*

Proof. There is at most one such U , since [Theorem 12*](#) determines every (Uf, g) , and hence every Uf . It remains to prove existence. The argument used in the proof of [Theorem 36](#), and $|e^{2\pi i\lambda}| = 1$, give

$$\left| \int_0^1 e^{2\pi i\lambda} d(E(\lambda)f, g) \right| \leq |f||g|.$$

Call the integral on the left $L(g)$, with f fixed and g variable. Then

$$L(a_1g_1 + \cdots + a_ng_n) = \overline{a_1}L(g_1) + \cdots + \overline{a_n}L(g_n), \quad |L(g)| \leq |f||g| = C|g|.$$

F. Riesz's theorem applies (cf. [Note 53](#)): there is exactly one f^* such that $L(g) = (f^*, g)$. Define U by $Uf = f^*$. This is plainly an everywhere-defined linear operator; only its unitarity remains to be proved.

First $|(Uf, g)| \leq |f||g|$, so U is bounded and U^* exists. We have

$$(Uf, g) = \int_0^1 e^{2\pi i\lambda} d(E(\lambda)f, g), \quad (U^*f, g) = \int_0^1 e^{-2\pi i\lambda} d(E(\lambda)f, g).$$

The first equation is the definition; the second follows by interchanging f, g and taking complex conjugates. Just as at the end of [Section 4](#), these formulas give

$$(Uf, Ug) = (f, g), \quad (U^*f, U^*g) = (f, g).$$

Consequently

$$(U^*Uf, g) = (UU^*f, g) = (f, g)$$

for all f, g , so $U^*U = UU^* = 1$; that is, U is unitary.

Theorem 14*. *For a given unitary operator U , only one family $E(\lambda)$ of projection operators can be related to it as in [Theorem 12*](#).*

Proof. The formula

$$(U^n f, g) = \int_0^1 e^{2\pi in\lambda} d(E(\lambda)f, g)$$

is trivial for $n = 0$, is the definition for $n = 1$, and follows from $U^{-1} = U^*$ for $n = -1$. If it has been proved for m, n , then it also holds for $m - n$, since the method used at the end of [Section 4](#) gives

$$(U^m f, U^n g) = \int_0^1 e^{2(m-n)\pi i\lambda} d(E(\lambda)f, g),$$

whereas

$$(U^m f, U^n g) = ((U^n)^* U^m f, g) = (U^{*n} U^m f, g) = (U^{-n} U^m f, g) = (U^{m-n} f, g).$$

Thus the formula holds for every $n = 0, \pm 1, \pm 2, \dots$

If $F(\lambda)$ is a second family with the same properties, then the same equation holds with F in place of E . Hence

$$\int_0^1 p(\lambda) d((E(\lambda)f, g) - (F(\lambda)f, g)) = 0$$

first for every $p(\lambda) = e^{2n\pi i\lambda}$, and then for their linear combinations, the trigonometric polynomials. By continuity it holds for every continuous $p(\lambda)$ with $p(0) = p(1)$. Because of condition β in [Theorem 12*](#), it remains valid for functions having finitely many jumps, provided the value at a jump is always taken on its right. Put, for $0 \leq \lambda_0 < 1$,

$$p(\lambda) = \begin{cases} 1, & 0 < \lambda \leq \lambda_0, \\ 0, & \text{otherwise.} \end{cases}$$

Then

$$\int_0^{\lambda_0} d((E(\lambda)f, g) - (F(\lambda)f, g)) = (E(\lambda_0)f, g) - (F(\lambda_0)f, g) = 0.$$

Since this holds for all f, g , $E(\lambda_0) = F(\lambda_0)$. Here $\lambda_0 \neq 1$; at $\lambda_0 = 1$ the assertion follows from condition α of [Theorem 12*](#).

8.

Theorems 12*, 13*, 14* show that the correspondence in Theorem 12* between a unitary operator U and a family of projection operators $E(\lambda)$ satisfying α, β is one-to-one, and exhausts all U and all $E(\lambda)$. To be a resolution of the identity, these $E(\lambda)$ still require

$$E(\lambda)f \longrightarrow f \quad (\lambda < 1, \lambda \rightarrow 1);$$

cf. [Definition 17](#). We shall therefore have proved everything announced in [Chapter IX](#) once we show that this property is equivalent to the assertion that $\varphi = U\varphi$ occurs only for $\varphi = 0$.

First suppose $\varphi = U\varphi$, $\varphi \neq 0$. Then

$$\begin{aligned}\Re((\varphi, \varphi) - (U\varphi, \varphi)) &= \Re \int_0^1 (1 - e^{2\pi i\lambda}) d(E(\lambda)\varphi, \varphi) \\ &= \int_0^1 2 \sin^2(\pi\lambda) d|E(\lambda)\varphi|^2.\end{aligned}$$

The left-hand side vanishes, while, for every $\varepsilon > 0$, the right-hand side is at least

$$\begin{aligned}\int_\varepsilon^{1-\varepsilon} 2 \sin^2(\pi\lambda) d|E(\lambda)\varphi|^2 &\geq 2 \sin^2(\pi\varepsilon) \int_\varepsilon^{1-\varepsilon} d|E(\lambda)\varphi|^2 \\ &= 2 \sin^2(\pi\varepsilon) (|E(1-\varepsilon)\varphi|^2 - |E(\varepsilon)\varphi|^2).\end{aligned}$$

For $\varepsilon \leq \frac{1}{2}$, the last expression is nonnegative, and therefore zero:

$$|E(1-\varepsilon)\varphi| = |E(\varepsilon)\varphi|.$$

As $\varepsilon \rightarrow 0$, the right-hand side tends to zero, and hence so does the left. Thus

$$E(1-\varepsilon)\varphi \rightarrow 0 \neq \varphi.$$

Conversely, suppose $E(\lambda)f \not\rightarrow f$ as $\lambda \rightarrow 1$. By [Theorem 19](#), there is a projection operator E such that $E(\lambda)g \rightarrow Eg$ for every g ; in particular $Ef \neq f$. For $f' = f - Ef$, therefore,

$$f' \neq 0, \quad E(\lambda)f' \rightarrow Ef' = 0.$$

Now $|E(\lambda)f'|$ is nonnegative and monotonically nondecreasing in λ . Since it tends to zero as $\lambda \rightarrow 1$, it must vanish identically. Hence

$$E(\lambda)f' = 0 \quad (\lambda < 1), \quad E(1)f' = f'.$$

It follows that

$$(Uf', g) = \int_0^1 e^{2\pi i\lambda} d(E(\lambda)f', g) = (f', g)$$

for every g , and therefore $Uf' = f'$. This completes the proof.

Appendix III. Operators and Matrices

1.

It is desirable to translate our results on Hermitian operators into the customary terminology of matrices. We shall see that, in the unbounded case, the terminology

of operators is considerably more practical. In the bounded case, by Hilbert's results, the two amount to the same thing. We shall also see that closed linear Hermitian operators correspond to Hermitian matrices whose row sums of absolute squares are all finite, though not necessarily bounded. These are the so-called *square-summable matrices*.

The more detailed theory of these matrices is to be developed in the work announced in [Note 22](#). Here we give only what is needed for orientation and for the purposes of [Appendix IV](#).

A matrix

$$A = \{a_{\mu\nu}\}, \quad \mu, \nu = 1, 2, \dots,$$

will be called *Hermitian square-summable* (abbreviated H. square-summable) if

$$a_{\mu\nu} = \overline{a_{\nu\mu}}, \quad \sum_{\nu=1}^{\infty} |a_{\mu\nu}|^2 < \infty \quad (\mu = 1, 2, \dots).$$

⁷⁴ If $\varphi_1, \varphi_2, \dots$ is a complete normalized orthogonal system, define an operator R on the vectors φ_μ , and only on those vectors, by

$$R\varphi_\mu = \sum_{\nu=1}^{\infty} a_{\mu\nu}\varphi_\nu;$$

the series converges. The closed linear manifold spanned by $\varphi_1, \varphi_2, \dots$ is $\mathfrak{H}$, and

$$(\varphi_\mu, R\varphi_\nu) = \overline{a_{\nu\mu}} = a_{\mu\nu} = (R\varphi_\mu, \varphi_\nu).$$

Thus R is Hermitian. We call R the *elementary Hermitian operator* of $A, \{\varphi_1, \varphi_2, \dots\}$, and call its linear extension $\widehat{R}$, and its closed linear extension $\widetilde{R}$, the *linear* and *closed linear Hermitian operators*, respectively, of $A, \{\varphi_1, \varphi_2, \dots\}$.

The operator $\widehat{R}f$ is defined precisely for finite sums $f = \sum_{\mu=1}^n x_\mu \varphi_\mu$, and

$$\widehat{R}f = \sum_{\mu=1}^n x_\mu \left(\sum_{\nu=1}^{\infty} a_{\mu\nu} \varphi_\nu \right) = \sum_{\nu=1}^{\infty} \left(\sum_{\mu=1}^n a_{\mu\nu} x_\mu \right) \varphi_\nu.$$

The operator $\widetilde{R}$ is not so easy to characterize directly, but its extension elements are. For f to be an extension element, with f^* assigned to it, one must have

$$(f^*, \varphi_\nu) = (f, R\varphi_\nu) = \sum_{\mu=1}^{\infty} (f, \varphi_\mu) (\varphi_\mu, R\varphi_\nu) = \sum_{\mu=1}^{\infty} a_{\mu\nu} (f, \varphi_\mu), \quad \nu = 1, 2, \dots,$$

⁷⁴Then every series $\sum_{\rho=1}^{\infty} a_{\mu\rho} a_{\rho\nu}$ converges absolutely, so the matrix A^2 can be formed.

since R and $\widetilde{R}$ have the same extension elements. Write

$$(f, \varphi_\mu) = x_\mu, \quad f = \sum_{\mu=1}^{\infty} x_\mu \varphi_\mu.$$

There is an f^* satisfying

$$(f^*, \varphi_\nu) = \sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu$$

if and only if

$$\sum_{\nu=1}^{\infty} \left| \sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \right|^2 < \infty,$$

and in that case it is

$$f^* = \sum_{\nu=1}^{\infty} \left(\sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \right) \varphi_\nu.$$

Thus

$$f = \sum_{\mu=1}^{\infty} x_\mu \varphi_\mu, \quad \sum_{\mu=1}^{\infty} |x_\mu|^2 < \infty,$$

is an extension element if and only if, for

$$y_\nu = \sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu, \tag{75}$$

one has $\sum_{\nu=1}^{\infty} |y_\nu|^2 < \infty$; then

$$f^* = \sum_{\nu=1}^{\infty} y_\nu \varphi_\nu.$$

⁷⁵ If, in fact, $\widehat{R}f$ is defined, then f is certainly such an extension element. Furthermore,

$$(f^*, f) = \sum_{\nu=1}^{\infty} y_\nu \overline{x_\nu} = \sum_{\nu=1}^{\infty} \left(\sum_{\mu=1}^{\infty} a_{\mu\nu} x_\mu \right) \overline{x_\nu},$$

⁷⁵Both $\sum_{\mu} |a_{\mu\nu}|^2$ and $\sum_{\mu} |x_\mu|^2$ are finite, so the series $\sum_{\mu} a_{\mu\nu} x_\mu$ converge absolutely.

$$(f, f^*) = \sum_{\mu=1}^{\infty} x_{\mu} \overline{y_{\mu}} = \sum_{\mu=1}^{\infty} x_{\mu} \left(\sum_{\nu=1}^{\infty} a_{\mu\nu} \overline{x_{\nu}} \right).$$

Thus the 0-class is characterized by the interchangeability of the two infinite summations.

2.

We now show that, for every closed linear Hermitian operator R , there are an H. square-summable matrix $A = \{a_{\mu\nu}\}$ and a complete normalized orthogonal system $\varphi_1, \varphi_2, \dots$ such that R is the closed linear Hermitian operator of $A, \{\varphi_1, \varphi_2, \dots\}$. It is enough to specify the φ_{ν} , for

$$a_{\mu\nu} = (R\varphi_{\mu}, \varphi_{\nu})$$

then determines A . If every $R\varphi_{\mu}$ is defined, A is H. square-summable: $a_{\mu\nu} = \overline{a_{\nu\mu}}$ is clear, and

$$\sum_{\nu=1}^{\infty} |a_{\mu\nu}|^2 = \sum_{\nu=1}^{\infty} |(R\varphi_{\mu}, \varphi_{\nu})|^2 = |R\varphi_{\mu}|^2 < \infty.$$

Let S be the elementary operator of $A, \{\varphi_1, \varphi_2, \dots\}$. Then

$$S\varphi_{\mu} = \sum_{\nu=1}^{\infty} a_{\mu\nu} \varphi_{\nu} = \sum_{\nu=1}^{\infty} (R\varphi_{\mu}, \varphi_{\nu}) \varphi_{\nu} = R\varphi_{\mu};$$

thus S is the restriction of R to $\varphi_1, \varphi_2, \dots$. We must choose this system so that $\tilde{S} = R$. In other words, for every f for which Rf is defined, there must be a sequence $f_1, f_2, \dots$ of finite linear combinations of $\varphi_1, \varphi_2, \dots$ such that

$$f_n \rightarrow f, \quad Rf_n \rightarrow Rf,$$

and of course every $R\varphi_n$ must be defined.

This is achieved if there is a sequence $g_1, g_2, \dots$, with every Rg_n defined, such that for each f in the domain of R a subsequence $g_{N_1}, g_{N_2}, \dots$ satisfies

$$g_{N_n} \rightarrow f, \quad Rg_{N_n} \rightarrow Rf.$$

This sequence is automatically everywhere dense. It is then enough to orthogonalize $g_1, g_2, \dots$, as in [Theorem 8](#), to obtain $\varphi_1, \varphi_2, \dots$.

Such a sequence is constructed as follows. Let $h_1, h_2, \dots$ be everywhere dense but otherwise arbitrary. For every triple of positive integers m, n, k for which there is an f in the domain of R satisfying

$$|h_m - f| \leq \frac{1}{k}, \quad |h_n - Rf| \leq \frac{1}{k},$$

choose one such f , denoted $f_{m,n,k}$. Written as a single sequence, these $f_{m,n,k}$ have all the required properties. Indeed, given f, Rf and $\varepsilon > 0$, choose k with $2/k \leq \varepsilon$, and choose h_m, h_n such that

$$|h_m - f| \leq \frac{1}{k}, \quad |h_n - Rf| \leq \frac{1}{k}.$$

Then $f_{m,n,k}$ is defined, and

$$|h_m - f_{m,n,k}| \leq \frac{1}{k}, \quad |h_n - Rf_{m,n,k}| \leq \frac{1}{k}.$$

Consequently

$$|f - f_{m,n,k}| \leq \varepsilon, \quad |Rf - Rf_{m,n,k}| \leq \varepsilon.$$

Since ε was arbitrary, the construction is complete.

3.

For a given R , there may be several pairs

$$A = \{a_{\mu\nu}\}, \{\varphi_1, \varphi_2, \dots\}, \quad B = \{b_{\mu\nu}\}, \{\psi_1, \psi_2, \dots\},$$

whose closed linear Hermitian operator is R ; both A and B are to be H. square-summable. The two complete normalized orthogonal systems determine a unitary infinite matrix

$$\{u_{\mu\nu}\}, \quad u_{\mu\nu} = (\varphi_\mu, \psi_\nu).$$

Indeed,

$$\begin{aligned} \sum_{\rho=1}^{\infty} u_{\mu\rho} \overline{u_{\nu\rho}} &= \sum_{\rho=1}^{\infty} (\varphi_\mu, \psi_\rho)(\psi_\rho, \varphi_\nu) = (\varphi_\mu, \varphi_\nu) = \delta_{\mu\nu}, \\ \sum_{\rho=1}^{\infty} \overline{u_{\rho\mu}} u_{\rho\nu} &= \sum_{\rho=1}^{\infty} (\psi_\mu, \varphi_\rho)(\varphi_\rho, \psi_\nu) = (\psi_\mu, \psi_\nu) = \delta_{\mu\nu}. \end{aligned} \tag{U1}$$

Moreover,

$$\begin{aligned}\varphi_\mu &= \sum_{v=1}^{\infty} (\varphi_\mu, \psi_v) \psi_v = \sum_{v=1}^{\infty} u_{\mu v} \psi_v, \\ \psi_\mu &= \sum_{v=1}^{\infty} (\psi_\mu, \varphi_v) \varphi_v = \sum_{v=1}^{\infty} \overline{u_{v\mu}} \varphi_v.\end{aligned}\tag{U2}$$

The usual transformation formulas also hold for $a_{\mu\nu}, b_{\mu\nu}$:

$$\begin{aligned}a_{\mu\nu} &= (\varphi_\mu, R\varphi_\nu) \\ &= \sum_{\rho=1}^{\infty} (\varphi_\mu, \psi_\rho) (\psi_\rho, R\varphi_\nu) \\ &= \sum_{\rho=1}^{\infty} (\varphi_\mu, \psi_\rho) (R\psi_\rho, \varphi_\nu) \\ &= \sum_{\rho=1}^{\infty} (\varphi_\mu, \psi_\rho) \left[\sum_{\sigma=1}^{\infty} (R\psi_\rho, \psi_\sigma) (\psi_\sigma, \varphi_\nu) \right] \\ &= \sum_{\rho=1}^{\infty} \left(\sum_{\sigma=1}^{\infty} b_{\rho\sigma} u_{\mu\rho} \overline{u_{\nu\sigma}} \right).\end{aligned}$$

Interchanging μ, ν , and ρ, σ , and taking complex conjugates, using $a_{\mu\nu} = \overline{a_{\nu\mu}}$ and $b_{\rho\sigma} = \overline{b_{\sigma\rho}}$, gives

$$a_{\mu\nu} = \sum_{\rho=1}^{\infty} \left(\sum_{\sigma=1}^{\infty} b_{\rho\sigma} u_{\mu\rho} \overline{u_{\nu\sigma}} \right) = \sum_{\sigma=1}^{\infty} \left(\sum_{\rho=1}^{\infty} b_{\rho\sigma} u_{\mu\rho} \overline{u_{\nu\sigma}} \right).\tag{U3}$$

Similarly, the inverse formulas are^{T18}

$$b_{\mu\nu} = \sum_{\rho=1}^{\infty} \left(\sum_{\sigma=1}^{\infty} a_{\rho\sigma} \overline{u_{\rho\mu}} u_{\sigma\nu} \right) = \sum_{\sigma=1}^{\infty} \left(\sum_{\rho=1}^{\infty} a_{\rho\sigma} \overline{u_{\rho\mu}} u_{\sigma\nu} \right).\tag{U4}$$

Thus far the customary theory of unitary transformations of Hermitian matrices remains valid, but no farther. Given an H. square-summable matrix $A = \{a_{\mu\nu}\}$,

^{T18}*Translator's note:* In the second iterated sum of (U4), the source prints $b_{\rho\sigma}$. The inverse transformation requires $a_{\rho\sigma}$, as used in both sums here.

a complete normalized orthogonal system $\varphi_1, \varphi_2, \dots$, and an arbitrary unitary transformation matrix $\{u_{\mu\nu}\}$ satisfying (U₁), an attempt to apply (U₂)–(U₄) need not initially produce any convergent series at all. Even if all the convergences in (U₃) and (U₄) occur, the closed linear Hermitian operator of A , $\{\varphi_1, \varphi_2, \dots\}$ may change under the transformation. We shall not pursue these relations here; they will be discussed fully elsewhere (cf. Note 22).

4.

It remains to mention how to determine the closed linear manifolds $\mathfrak{E}$, $\mathfrak{F}$ associated with the closed linear Hermitian operator $\widetilde{R}$ of A , $\{\varphi_1, \varphi_2, \dots\}$, where R is its elementary Hermitian operator; cf. Theorem 22. $\mathfrak{E}$, respectively $\mathfrak{F}$, is the set of all $(\widetilde{R} + i1)f$, respectively $(\widetilde{R} - i1)f$. Closedness of $\widetilde{R}$, together with the properties of $\widetilde{R} \pm i1$ established in Theorem 22, readily shows that this is the set of limit points of $(R \pm i1)f$.⁷⁶ This, in turn, is the linear manifold spanned by the $(R \pm i1)f$, hence by the $(R \pm i1)\varphi_\mu$. Accordingly, $\mathfrak{E}$, respectively $\mathfrak{F}$, is the closed linear manifold spanned by

$$(R \pm i1)\varphi_\mu = \sum_{\nu=1}^{\infty} a_{\mu\nu}\varphi_\nu \pm i\varphi_\mu.$$

Appendix IV. The Operator $\overline{R}$

1.

The operator $\overline{R}$ and its Cayley transform $\overline{U}$ were defined in Chapter X, Theorem 37, as follows. Let $\varphi_0, \varphi_1, \varphi_2, \dots$ be a complete normalized orthogonal system and put

$$\overline{U} \left(\sum_{\nu=0}^{\infty} x_\nu \varphi_\nu \right) = \sum_{\nu=0}^{\infty} x_\nu \varphi_{\nu+1}, \quad \sum_{\nu=1}^{\infty} |x_\nu|^2 < \infty.$$

Thus $\overline{R}f$ is defined for all

$$f = \varphi - \overline{U}\varphi = x_0\varphi_0 + (x_1 - x_0)\varphi_1 + (x_2 - x_1)\varphi_2 + \dots,$$

⁷⁶If a sequence $(\widetilde{R} \pm i1)f$ converges, then so does $U^{\pm 1}(\widetilde{R} \pm i1)f = (\widetilde{R} \mp i1)f$; hence f and $\widetilde{R}f$ converge. The converse is immediate: if f and $\widetilde{R}f$ converge, then so does $(\widetilde{R} \pm i1)f$.

and then

$$\bar{R}f = i(\varphi + \bar{U}\varphi) = ix_0\varphi_0 + i(x_0 + x_1)\varphi_1 + i(x_1 + x_2)\varphi_2 + \cdots.$$

Equivalently,

$$\bar{R} \left(\sum_{v=0}^{\infty} y_v \varphi_v \right)$$

is defined when

$$|y_0|^2 + |y_0 + y_1|^2 + |y_0 + y_1 + y_2|^2 + \cdots$$

converges; for necessarily

$$x_0 = y_0, \quad x_1 = y_0 + y_1, \quad x_2 = y_0 + y_1 + y_2, \dots$$

Its value is then

$$iy_0\varphi_0 + (2iy_0 + iy_1)\varphi_1 + (2iy_0 + 2iy_1 + iy_2)\varphi_2 + \cdots.$$

In this way we obtain a kind of matrix for $\bar{R}$:

$$a_{\mu\nu} = \begin{cases} 0, & \mu > \nu, \\ i, & \mu = \nu, \\ 2i, & \mu < \nu. \end{cases}$$

It is not, of course, a matrix in the sense of [Appendix III](#): it is neither Hermitian nor square-summable. This is unsurprising, since all $\bar{R}\varphi_\mu$, $\mu = 0, 1, 2, \dots$, are undefined.

To give a genuine matrix for $\bar{R}$, we must first find a complete normalized orthogonal system $\psi_1, \psi_2, \dots$ for which every $\bar{R}\psi_\mu$ is defined. Put, for $m, k = 0, 1, 2, \dots$,

$$\begin{aligned} \omega_{m,k} = & \varphi_{m2^{k+1}} + \varphi_{m2^{k+1}+1} + \cdots + \varphi_{(2m+1)2^k-1} \\ & - \varphi_{(2m+1)2^k} - \varphi_{(2m+1)2^{k+1}} - \cdots - \varphi_{(m+1)2^{k+1}-1}. \end{aligned}$$

Our criterion shows that $\bar{R}(\sum_{v=0}^{\infty} x_v \varphi_v)$ is certainly defined if only finitely many x_v are nonzero and $\sum_{v=0}^{\infty} x_v = 0$; this is the case here.

The equality

$$(\omega_{m,k}, \omega_{n,l}) = 0$$

certainly holds when the two expressions have no φ_ν -term in common. It also holds when all the φ_ν occurring in one occur in the other, all with coefficient +1, or all with coefficient -1. Except when $m = n, k = l$, one of these alternatives always occurs; on the other hand,

$$|\omega_{m,k}|^2 = 2^{k+1}.$$

Consequently

$$\psi_{m,k} = 2^{-(k+1)/2} \omega_{m,k}$$

form a normalized orthogonal system, and all $\bar{R}\psi_{m,k}$ are defined.

This system is complete. Indeed, let

$$f = \sum_{\nu=0}^{\infty} x_\nu \varphi_\nu$$

be orthogonal to all $\psi_{m,k}$, equivalently to all $\omega_{m,k}$. This says

$$\begin{aligned} & x_{m2^{k+1}} + x_{m2^{k+1}+1} + \cdots + x_{(2m+1)2^k-1} \\ & - x_{(2m+1)2^k} - x_{(2m+1)2^k+1} - \cdots - x_{(m+1)2^{k+1}-1} = 0. \end{aligned}$$

For $k = 0$, this yields

$$x_0 - x_1 = 0, \quad x_2 - x_3 = 0, \quad x_4 - x_5 = 0, \quad x_6 - x_7 = 0, \dots;$$

that is,

$$x_0 = x_1, \quad x_2 = x_3, \quad x_4 = x_5, \quad x_6 = x_7, \dots$$

For $k = 1$, it gives

$$x_0 + x_1 - x_2 - x_3 = 0, \quad x_4 + x_5 - x_6 - x_7 = 0, \dots,$$

so

$$2x_0 = 2x_2, \quad 2x_4 = 2x_6, \dots,$$

and therefore

$$x_0 = x_1 = x_2 = x_3, \quad x_4 = x_5 = x_6 = x_7, \dots$$

For $k = 2$, it gives

$$x_0 + x_1 + x_2 + x_3 - x_4 - x_5 - x_6 - x_7 = 0,$$

and hence $4x_0 = 4x_4$, and so forth. Thus, altogether,

$$x_0 = x_1 = x_2 = \cdots.$$

But $\sum_{\nu=0}^{\infty} |x_\nu|^2$ can be finite only if all these numbers vanish. Hence $f = 0$, proving completeness.

2.

Define the matrix $\{b_{m,k;n,l}\}$ by

$$b_{m,k;n,l} = (\psi_{m,k}, \bar{R}\psi_{n,l}) = 2^{-(k+l)/2-1}(\omega_{m,k}, \bar{R}\omega_{n,l}).$$

It is H. square-summable, and its elementary Hermitian operator S is the restriction of $\bar{R}$ to the $\psi_{m,k}$. Consequently $\bar{R}$ is an extension of S , and, being closed and linear, an extension of $\bar{S}$. If $\bar{S}$ is maximal, then necessarily

$$\bar{R} = \bar{S};$$

that is, $\bar{R}$ is the closed linear Hermitian operator of $\{b_{m,k;n,l}\}, \{\psi_{m,k}\}$.

We prove maximality by showing that the manifold $\mathfrak{E}$ of $\bar{S}$ is all of $\mathfrak{H}$. By [Appendix III.4](#), the vectors $\bar{R}\psi_{m,k} + i\psi_{m,k}$, or equivalently $\bar{R}\omega_{m,k} + i\omega_{m,k}$, must span $\mathfrak{H}$ as a closed linear manifold. Thus it is enough to show that if

$$f = \sum_{v=0}^{\infty} x_v \varphi_v$$

is orthogonal to all $\bar{R}\omega_{m,k} + i\omega_{m,k}$, then $f = 0$.

The formulas in the preceding section give

$$\begin{aligned} \bar{R}\omega_{m,k} &= i\varphi_{m2^{k+1}} + 3i\varphi_{m2^{k+1}+1} + \cdots + (2^{k+1} - 1)i\varphi_{(2m+1)2^k-1} \\ &\quad + (2^{k+1} - 1)i\varphi_{(2m+1)2^k} + (2^{k+1} - 3)i\varphi_{(2m+1)2^{k+1}} + \cdots \\ &\quad + 3i\varphi_{(m+1)2^{k+1}-2} + i\varphi_{(m+1)2^{k+1}-1}, \end{aligned}$$

and therefore

$$\begin{aligned} \bar{R}\omega_{m,k} + i\omega_{m,k} &= 2i\varphi_{m2^{k+1}} + 4i\varphi_{m2^{k+1}+1} + \cdots + 2^{k+1}i\varphi_{(2m+1)2^k-1} \\ &\quad + (2^{k+1} - 2)i\varphi_{(2m+1)2^k} + (2^{k+1} - 4)i\varphi_{(2m+1)2^{k+1}} + \cdots \\ &\quad + 2i\varphi_{(m+1)2^{k+1}-2}. \end{aligned}$$

Orthogonality of f to these vectors means

$$2i\bar{x}_{m2^{k+1}} + 4i\bar{x}_{m2^{k+1}+1} + \cdots + 2^{k+1}i\bar{x}_{(2m+1)2^k-1} + \cdots + 2i\bar{x}_{(m+1)2^{k+1}-2} = 0.$$

Thus

$$x_0 = 0, \quad x_2 = 0, \quad x_4 = 0, \dots;$$

$$x_0 + 2x_1 + x_2 = 0, \quad x_4 + 2x_5 + x_6 = 0, \dots;$$

$$x_0 + 2x_1 + 3x_2 + 4x_3 + 3x_4 + 2x_5 + x_6 = 0, \dots,$$

and so forth. It follows successively that

$$x_0 = x_1 = x_2 = \dots = 0,$$

and hence $f = 0$.

It remains to calculate

$$b_{m,k;n,l} = 2^{-(k+l)/2-1}(\omega_{m,k}, \bar{R}\omega_{n,l}),$$

which is immediate because both vectors have been expressed in terms of the φ_v . If $m \neq n$ and $k = l$, there are no common φ_v -terms, so the result is zero; for $m = n, k = l$, direct calculation again gives zero. Let $k > l$. If n lies outside

$$m2^{k-l}, \dots, (m+1)2^{k-l} - 1,$$

there are again no common φ_v -terms, and the result is zero. If n lies in this interval, then

$$n = m2^{k-l} + \rho \quad \text{or} \quad n = (m+1)2^{k-l} - \rho - 1, \quad \rho = 0, 1, \dots, 2^{k-l-1} - 1.$$

In the two cases, respectively,

$$\begin{aligned} (\omega_{m,k}, \bar{R}\omega_{n,l}) &= \sum_{\alpha=\rho 2^{l+1}}^{(2\rho+1)2^l-1} [-i(2\alpha+1)] - \sum_{\alpha=(2\rho+1)2^l}^{(\rho+1)2^{l+1}-1} [-i(2\alpha+1)] \\ &= -i 2^l(-2 \cdot 2^l) = i 2^{2l+1}, \end{aligned}$$

or

$$\begin{aligned} (\omega_{m,k}, \bar{R}\omega_{n,l}) &= \sum_{\alpha=(2\rho+1)2^l}^{(\rho+1)2^{l+1}-1} [-i(2\alpha+1)] - \sum_{\alpha=\rho 2^{l+1}}^{(2\rho+1)2^l-1} [-i(2\alpha+1)] \\ &= -i 2^l(2 \cdot 2^l) = -i 2^{2l+1}. \end{aligned}$$

For $k < l$, Hermitian symmetry and interchange of m, n , and k, l , give the corresponding values

$$0, \quad -i 2^{2k+1}, \quad i 2^{2k+1}.$$

After the normalizing factor in $b_{m,k;n,l}$, the values

$$0, \quad \pm i 2^{2k+1}, \quad \pm i 2^{2l+1}$$

become

$$0, \quad \pm i 2^{(3k-l)/2}, \quad \pm i 2^{(3l-k)/2},$$

respectively.

We reduce the double index $m, k = 0, 1, 2, \dots$ to the single index $p = 1, 2, \dots$ by

$$p = (2m + 1)2^k.$$

The result may then be stated as follows.

Theorem 1.** *In the complete normalized orthogonal system $\psi_1, \psi_2, \dots$ described above, there is an H . square-summable matrix $B = \{b_{pq}\}$ such that $\bar{B}$ is the closed linear Hermitian operator of $B, \{\psi_1, \psi_2, \dots\}$. Its entries $b_{pq}, p, q = 1, 2, \dots$, are defined as follows.*

Partition the sequence of positive integers into cells. The interval

$$m2^{k+1} + 1, \dots, (m + 1)2^{k+1}$$

is a cell; the intervals

$$m2^{k+1} + 1, \dots, (2m + 1)2^k$$

and

$$(2m + 1)2^k + 1, \dots, (m + 1)2^{k+1}$$

are its first and second halves, respectively; $(2m + 1)2^k$ is its center and k its degree. Every positive integer p is the center of exactly one cell, called the cell of p .

If neither of the cells of p, q is a proper part of the other, then

$$b_{pq} = 0.$$

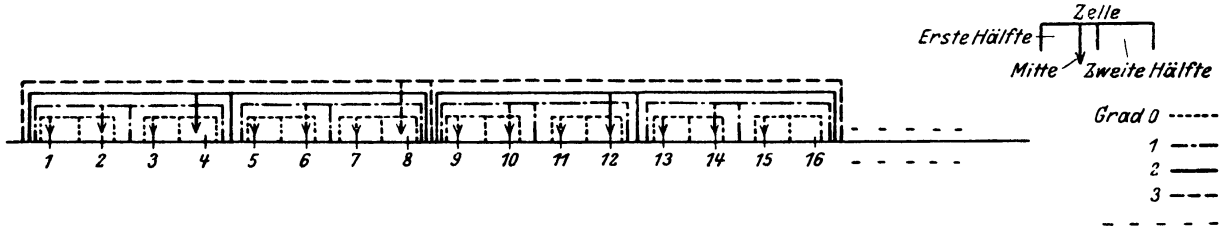
If one is a proper part of the other, and k, l are the degrees of the smaller and larger cells, respectively, then

$$b_{pq} = \pm i 2^{(3k-l)/2}.$$

The sign is determined by the following table:

	it lies in the first half	it lies in the second half
p 's cell is smaller	—	+
q 's cell is smaller	+	—

Thus B is i times a skew-symmetric matrix.^{T19}



3.

We obtain a better view of $\overline{R}$ by realizing it in function spaces. Let Ω be the interval $0, 1$, and form its function space: the set of all Lebesgue-measurable functions $f(x)$ on $0 < x < 1$ for which

$$\int_0^1 |f(x)|^2 dx < \infty.$$

By [Appendix I](#), this is a Hilbert space isomorphic to $\mathfrak{H}$; denote it by $\mathfrak{H}^*$.

The functions

$$e^{2n\pi ix}, \quad n = 0, \pm 1, \pm 2, \dots,$$

form a complete normalized orthogonal system in $\mathfrak{H}^*$. Let $\mathfrak{I}$ be the closed linear manifold spanned by

$$e^{2n\pi ix}, \quad n = 0, 1, 2, \dots$$

It is infinite-dimensional and therefore itself a Hilbert space; in $\mathfrak{I}$ we shall realize $\overline{R}$.

The operator

$$Uf(x) = e^{2\pi ix} f(x)$$

is isometric, indeed unitary, on $\mathfrak{H}^*$, and maps $\mathfrak{I}$ into a part of itself. Put

$$\varphi_v(x) = e^{2v\pi ix}, \quad v = 0, 1, 2, \dots$$

Then $\varphi_0, \varphi_1, \varphi_2, \dots$ is a complete normalized orthogonal system in $\mathfrak{I}$, and

$$U\varphi_v = \varphi_{v+1}.$$

^{T19}Translator's note: In the reproduced source diagram, the German labels mean: *Zelle*, "cell"; *Erste Hälfte*, "first half"; *Mitte*, "center"; *Zweite Hälfte*, "second half"; and *Grad*, "degree." Thus *Grad* 0, ..., *Grad* 3 denote degrees 0 through 3.

Thus U agrees on $\mathfrak{I}$ with $\overline{U}$, and $\overline{R}$ is its Cayley transform. Hence $\overline{R}f(x)$, for $f \in \mathfrak{I}$, is defined when

$$f(x) = \varphi(x) - U\varphi(x) = (1 - e^{2\pi ix})\varphi(x), \quad \varphi \in \mathfrak{I},$$

and then

$$\overline{R}f(x) = i(\varphi(x) + U\varphi(x)) = i(1 + e^{2\pi ix})\varphi(x).$$

Equivalently, it is defined when both

$$f(x), \quad \frac{f(x)}{1 - e^{2\pi ix}}$$

belong to $\mathfrak{I}$, and in that case

$$\overline{R}f(x) = i \frac{1 + e^{2\pi ix}}{1 - e^{2\pi ix}} f(x) = -\cot(\pi x) f(x).$$

Let $f(x) \in \mathfrak{I}$. For $f(x)/(1 - e^{2\pi ix})$ to belong to $\mathfrak{I}$, its absolute square must be integrable and it must be orthogonal to every

$$e^{2n\pi ix}, \quad n = -1, -2, \dots$$

This implies

$$\int_0^1 \cot^2(\pi x) |f(x)|^2 dx < \infty,$$

since $\overline{R}f(x) = -\cot(\pi x)f(x)$. We now show that this condition conversely implies

$$\frac{f(x)}{1 - e^{2\pi ix}} \in \mathfrak{I} \quad (f \in \mathfrak{I}).$$

Indeed,

$$\left| \frac{1}{1 - e^{2\pi ix}} \right|^2 = \left| \frac{\cot(\pi x) + i}{2i} \right|^2 \leq \frac{1}{2} \cot^2(\pi x) + \frac{1}{2}.$$

Thus

$$\int_0^1 \left| \frac{f(x)}{1 - e^{2\pi ix}} \right|^2 dx < \infty,$$

and the functions

$$\left| \frac{f(x)}{1 - \rho e^{2\pi ix}} \right|^2, \quad 0 < \rho < 1,$$

are absolutely uniformly integrable. Consequently, as $0 < \rho < 1$, $\rho \rightarrow 1$,

$$\int_0^1 \frac{f(x)}{1 - \rho e^{2\pi i x}} e^{-2n\pi i x} dx \longrightarrow \int_0^1 \frac{f(x)}{1 - e^{2\pi i x}} e^{-2n\pi i x} dx.$$

We must show that the right-hand side vanishes for $n = -1, -2, \dots$; it is enough to prove this for the left-hand side, for every $0 < \rho < 1$. Now

$$\frac{1}{1 - \rho e^{2\pi i x}} = \sum_{v=0}^{\infty} \rho^v e^{2v\pi i x},$$

and, for fixed ρ , the series converges uniformly in x . Therefore

$$\begin{aligned} \sum_{v=0}^N \rho^v \int_0^1 f(x) e^{-2(n-v)\pi i x} dx &= \int_0^1 f(x) \left(\sum_{v=0}^N \rho^v e^{2v\pi i x} \right) e^{-2n\pi i x} dx \\ &\longrightarrow \int_0^1 \frac{f(x)}{1 - \rho e^{2\pi i x}} e^{-2n\pi i x} dx \end{aligned}$$

as $N \rightarrow \infty$. Since $n = -1, -2, \dots$ and $v = 0, 1, 2, \dots$, one has $n - v = -1, -2, \dots$. Because $f \in \mathfrak{F}$, the left-hand side is always zero; hence the right-hand side vanishes as well.

We can now prove the following.

Theorem 2.** *Let $\mathfrak{H}^*$ be the Hilbert function space just described, and let $\mathfrak{F}_k^+$, respectively $\mathfrak{F}_k^-$, be the closed linear manifolds spanned by all*

$$e^{2n\pi i x} \quad \text{with } n \geq k, \quad \text{respectively with } n \leq k,$$

where $n, k = 0, \pm 1, \pm 2, \dots$. The $\mathfrak{F}_k^\pm$ are infinite-dimensional, and hence are again Hilbert spaces.

Define R on $\mathfrak{H}^*$ as follows: if $f(x)$ and $-\cot(\pi x)f(x)$ belong to $\mathfrak{H}^*$, that is, if

$$\int_0^1 |f(x)|^2 dx < \infty, \quad \int_0^1 \cot^2(\pi x) |f(x)|^2 dx < \infty,$$

then $Rf(x)$ is defined and

$$Rf(x) = -\cot(\pi x)f(x).$$

If $f(x)$ lies in one of the $\mathfrak{F}_k^\pm$ and $Rf(x)$ is defined, then $Rf(x)$ lies in the same $\mathfrak{F}_k^\pm$. Thus R may be restricted on each $\mathfrak{F}_k^\pm$ to a Hermitian operator there.⁷⁷ On every $\mathfrak{F}_k^+$ this gives $\overline{R}$, and on every $\mathfrak{F}_k^-$ it gives $-\overline{R}$.

⁷⁷This is by no means the reducibility of R by $\mathfrak{F}_k^\pm$: whenever Rf is defined, $RP_{\mathfrak{F}_k^\pm}f$ need not also be defined.

Proof. For $\mathfrak{I}_0^+ = \mathfrak{I}$, this has already been proved. The unitary map

$$f(x) \mapsto e^{2k\pi ix} f(x)$$

takes $(\mathfrak{H}^*, R, \mathfrak{I}_0^+)$ to $(\mathfrak{H}^*, R, \mathfrak{I}_k^+)$. The unitary map

$$f(x) \mapsto f(1-x)$$

takes $(\mathfrak{H}^*, R, \mathfrak{I}_{-k}^+)$ to $(\mathfrak{H}^*, -R, \mathfrak{I}_k^-)$. Hence the assertions hold for all $\mathfrak{I}_k^\pm$.

It is also easy to show that this R is a hypermaximal Hermitian operator on $\mathfrak{H}^*$.

4.

To every

$$f(x) = \sum_{n=0}^{\infty} a_n e^{2n\pi ix}, \quad \sum_{n=0}^{\infty} |a_n|^2 < \infty,$$

in $\mathfrak{I}_0^+$,⁷⁸ we may associate the function

$$F(z) = \sum_{n=0}^{\infty} a_n z^n$$

analytic in the unit disk; it is analytic there because $\sum |a_n|^2 < \infty$, and hence $a_n \rightarrow 0$. The circle Parseval identity gives^{T20}

$$\sum_{n=0}^{\infty} |a_n|^2 = \lim_{r \rightarrow 1} \sum_{n=0}^{\infty} r^{2n} |a_n|^2 = \lim_{r \rightarrow 1} \frac{1}{2\pi i} \int_{|z|=r} |F(z)|^2 \frac{dz}{z}.$$

Thus the functions $F(z)$ analytic in the unit disk for which

$$\frac{1}{2\pi i} \int_{|z|=r} |F(z)|^2 \frac{dz}{z}$$

⁷⁸The series $\sum_{n=0}^{\infty} a_n e^{2n\pi ix}$ converges, in the metric used throughout this paper, “in the mean”:

$$\int_0^1 \left| f(x) - \sum_{n=0}^N a_n e^{2n\pi ix} \right|^2 dx \longrightarrow 0 \quad (N \rightarrow \infty),$$

but not necessarily pointwise.

^{T20}*Translator’s note:* The source prints r^n and omits the normalizing factor $1/(2\pi i)$ in the contour integral. Both have been corrected to the standard Parseval identity.

remains bounded as $r \rightarrow 1$ form a Hilbert space. Since

$$-\cot(\pi x) = i \frac{e^{2\pi i x} + 1}{e^{2\pi i x} - 1},$$

$\bar{R}$ is represented there by

$$\bar{R}F(z) = i \frac{z+1}{z-1} F(z),$$

provided the result again belongs to the stated class.

Another noteworthy realization of $\bar{R}$ is obtained by considering functions $f(x)$ defined on $0 < x < \infty$ with

$$\int_0^\infty |f(x)|^2 dx < \infty.$$

It then corresponds to the differential operator

$$i \frac{d}{dx}$$

with the boundary condition $f(0) = 0$. We shall not pursue this realization here.⁷⁹

It remains to verify the properties of $\bar{R}$ used in [Theorem 40](#). First, $\bar{R}$ is not everywhere defined and does not assume every value. Realize it, by [Theorem 2**](#), in $\mathfrak{I}_0^+$. Since

$$\int_0^1 \cot^2(\pi x) dx = \int_0^1 \tan^2(\pi x) dx = \infty,$$

$\bar{R}$ is undefined for $f(x) = 1$, and never assumes that value.

Second, $\bar{R}$ is bounded neither below nor above. Realize it as in [Theorem 1**](#) and put

$$f = \psi_{2^\nu} \pm i\psi_{2^{\nu+1}}.$$

Then

$$(f, \bar{R}f) = \mp 2^{\nu+\frac{1}{2}}, \quad |f|^2 = 2.$$

Thus, for any C , values with

$$(f, \bar{R}f) > C|f|^2 \quad \text{or} \quad (f, \bar{R}f) < C|f|^2$$

are obtained by choosing $2^{\nu-\frac{1}{2}} > |C|$.

⁷⁹It will be discussed in the work mentioned in [Note 22](#).