

Gentzen's Consistency Proofs

┌ GENTZEN'S CONSISTENCY PROOFS ┐

by

M.E. Szabo

A thesis
submitted to the Faculty of
Graduate Studies and Research
in partial fulfilment of the requirements
for the degree of Master of Science

Department of Mathematics
McGill University
Montreal

April 1967

I wish to thank Prof. John Denton
for his kind assistance and encourage-
ment during the preparation of this thesis.

TABLE OF CONTENTS

	Page
Introduction	(i-xxxix)
Gentzen's Consistency Proof for Elementary Number Theory	(1-130)
Section I: Reflections on the Purpose and Possibility of Consistency Proofs	(1-15)
Section II: The Formalization of Elementary Number Theory	(15-49)
Section III: Disputable and Indisputable Forms of Inference in Elementary Number Theory	(50-73)
Section IV: The Consistency Proof	(74-116)
Section V: Reflections on the Consistency Proof	(117-130)
Gentzen's New Version of the Consistency Proof for Elementary Number Theory	(130-188)
Excerpts from a Galley Proof of Gentzen's Initial Draft of the Consistency Proof for Elementary Number Theory	(188-206)
Footnotes	
Bibliography	
Glossary	

INTRODUCTION

The new conceptual power of Cantor's set theory has had a dynamic effect upon creative mathematical thinking and, at the same time, introduced a paradoxical element into mathematics that has shaken the very foundations of this discipline. A consequent thorough re-examination of the nature of mathematics has produced a wealth of ideas and a variety of approaches to foundational studies. Amidst the inevitable clashes between a purely "logical", a partly "philosophical" and a "genuinely mathematical" method of attack, i.e., between the logicians, the intuitionists, and the formalists, Gentzen made it the aim of his mathematical investigations to establish beyond doubt the reliability of large parts of classical mathematics. This undertaking was to be carried out in three distinct stages according to the degree to which the concept of infinity is involved:

The first stage consists of elementary number theory, into which "infinity" enters merely in the form of

(ii)

the stipulation of an infinite domain of objects; the second stage consists of analysis in which the introduction of irrational numbers and infinite series makes it necessary to treat an infinite set as an individual object; the third stage consists of general set theory in which infinite sets of infinite sets are freely admitted.¹⁾

Gentzen believed that elementary number theory, analysis, and a good part of set theory were entirely reliable.²⁾ In this view he was undoubtedly encouraged by Hilbert³⁾, whose assistant he was from 1934 until Hilbert's eventual retirement.

The statement that a branch of mathematics is reliable is a statement about that branch of mathematics⁴⁾ and this leads to the distinction between the theory to be vindicated and a meta-theory (or proof-theory) within which the notion of reliability can be formulated. 'Reliability' here means 'freedom from contradiction' and expresses the fact that once a

1) Gentzen (7), pp 65-66

2) Gentzen (8), p. 204 (7), p. 79

3) Gentzen (8) p. 205

4) Gentzen (5) p. 10

(iii)

theory has been formalized and the concept of a 'proof' defined, it can be shown that a contradictory statement such as ' $1 = 2$ ' cannot be obtained from true statements through the application of the logical rules that are permissible in the theory. The technical term for reliability used in this paper is the term 'consistency', and a proof that ends in a contradiction obtained from true statements by means of the logical rules of the theory will be referred to as an 'inconsistent derivation'.

The present translation contains two published versions of the "consistency" proof for "elementary number theory". They will be referred to by (5) and (10) as listed in the bibliography. In (5) the consistency follows from the non-derivability of the sequent ' $\rightarrow \mathcal{U} \wedge \neg \mathcal{U}$ ',⁵⁾ whilst in (10) it takes the form of the non-derivability of the empty sequent ' $\rightarrow$ ',⁶⁾.

In proving the consistency of a formalized axiomatic theory, the question arises whether the formalization

5) Gentzen (5) p. 74

6) Gentzen (10) p. 145

(iv)

chosen actually encompasses the full range of meaningful statements that can be made in the informal theory. In this respect Gentzen took note of Gödel's Incompleteness Theorem and formulated his rules flexibly enough to allow for whatever adaptation his formalization may require.⁷⁾ Gödel's further result concerning the impossibility of 'internal' consistency proofs in formalized axiomatic theories is overcome by using transfinite induction up to Cantor's first $\aleph_1$ -number as the non-derivable technique in elementary number theory. A conclusive proof of this non-derivability is given in Gentzen's last paper published in 1942.⁸⁾

The next question that arises is what methods of proof can be used to establish the consistency of elementary number theory. As Gentzen points out, there can be no absolute consistency proof.⁹⁾ In the Hilbert tradition, Gentzen decided that the consistency proof must be carried out by means of methods of proof that are unimpeachable and he believed that such methods

7) Gentzen (5) p. 126

8) Gentzen (11) p. 140; Gentzen (10) p. (185)

9) Gentzen (5) p. 13; Gentzen (7) p. 72

(v)

should therefore be constructive, in particular, they should be 'finitist'.¹⁰⁾ This notion is used in the sense of Hilbert and Bernays who, on page 32, Vol. 1, of the Grundlagen der Mathematik, introduce it in the following way:

"...die ausgeführte Betrachtung der Anfangsgründe von Zahlentheorie und Algebra dient dazu, uns das direkte inhaltlich, in Gedanken-Experimenten an anschaulich vorgestellten Objecten¹¹⁾ sich vollziehende und von axiomatischen Annahmen freie Schliessen in seiner Anwendung und Handhabung vorzuführen. Diese Art des Schliessens wollen wir, um einen kurzen Ausdruck zu haben, als das 'finite' Schliessen und ebenso auch die diesem Schliessen zugrunde liegende methodische Einstellung als die 'finite' Einstellung oder den 'finiten' Standpunkt bezeichnen. Im gleichen Sinne wollen wir allemal mit dem Wort 'Finit' zum Ausdruck bringen, dass die betreffende Überlegung, Behauptung oder Definition sich an die Grenzen der grundsatzlichen Vorstellbarkeit von Objekten sowie der grundsatzlichen Ausführbarkeit von Prozessen halt und sich somit im Rahmen konkreter Betrachtung vollzieht."

Whether the natural numbers are actually 'intuitively stipulated objects' in the sense of Hilbert and Bernays is important for the consistency proof and while Gentzen allowed for infinitely many numerals 1, 2, 3, 4, 5,¹²⁾ in (5) he

10) Gentzen (5) p. 7 and p. 50

11) Translator's italics

12) Gentzen (5) p. 18

introduced only the single numeral 1 in (10) together with the successor function.¹³⁾ This change was no doubt intended to bring out the constructive nature of the natural numbers as generated from the axioms of Peano.¹⁴⁾ This interpretation of Gentzen's intention seems reasonable in the light of the principle which Gentzen adopted as a result of a critical examination of Russell's antinomy:

"An infinite totality must not be regarded as actually existing and closed (actual infinity) but only as something becoming which can be extended constructively further and further from something finite (potential infinity)."¹⁵⁾

The difficulties of the antinomies are thus overcome by the rejection the 'actualist' interpretation of infinity. 'Actualist' is a term which the translator has used as a rendering of Gentzen's 'an-sich Auffassung' in order to preserve the philosophical neutrality which Gentzen had adopted. In his address to the Mathematical Congress in Paris in 1937, Gentzen nevertheless suggested that a certain parallel might be drawn between the 'constructivist' and the 'actualist'

13) Gentzen (10) p. 132

14) Gentzen (2) p. 1

15) Gentzen (5) p. 58

views of mathematics and the philosophical schools of 'idealism' and 'realism'.¹⁶⁾ This is the only place in all of his writings that he ever mentions a possible philosophical parallel between mathematical and philosophical theories.

In keeping with the principle adopted, the consistency proof should be entirely 'constructive'. Gentzen considers Brouwer's 'intuitionist' and Hilbert's 'finitist' approach to be two examples of this technique.¹⁷⁾ Yet he sees Brouwer's approach as too radical since it leads to the banishment of the large non-constructive part of analysis which has, for example, stood the test in a variety of applications in physics.¹⁸⁾

Gentzen thus aims at proving the consistency of certain non-finitist branches of mathematics by means of finitist (and therefore) constructivist techniques. In order to achieve this end, paragraphs 10 and 11 of (5) are devoted to a finitist interpretation of the logical connectives $\&$, $\vee$, $\neg$, $\supset$, $\forall$, $\exists$. The finitist interpretation of $\forall x F(x)$ for example, is the following:

16) Gentzen (8) p. 202

17) Gentzen (7) p. 71

18) Gentzen (7) p. 71

(viii)

"If, starting with 1, we substitute x for successive natural numbers, then however far we may proceed in the formation of numbers, in each case a true proposition results.¹⁹⁾ An existential proposition $\exists x F(x)$, in agreement with Hilbert's 'Partialurteil'²⁰⁾ expresses the finitist fact that if the proposition $F(n)$ has been recognized as meaningful and valid for an individual n , we may conclude $\exists x F(x)$.²¹⁾ In keeping with the finitist point of view all predicates and functions occurring in the formalization of elementary number theory must furthermore be 'decidably defined', i.e., a rule or procedure must be specified in each case providing a mechanical test for deciding of a predicate in finitely many steps whether or not it holds of a particular object, and which makes it possible to calculate the value of a function for any arbitrary element in a specified domain. It should be noted, incidentally, that while allowing for the introduction of arbitrary functions in (5),²²⁾ Gentzen restricts himself to only one function, the successor function, in (10).²³⁾ This simplifies the proof and

19) Gentzen (5) p. 60

20) Hilbert and Bernays (1) p. 32

21) Gentzen (5) pp. 62-63

22) Gentzen (5) p. 18

23) Gentzen (10) p. 132. D. (186)

the modifications necessary for the introduction of arbitrary effectively calculable functions are described later.²⁴⁾

The consistency proof falls into five sections:

Section I contains a general introduction to and motivation for consistency proofs;

Section II contains the formalization of elementary number theory as a formal axiomatic theory;

Section III deals with the finitist interpretation of the formalized axiomatic theory and contains a reference to Gödel's discovery that intuitionist and classical elementary number theory are in some sense equivalent.²⁵⁾ It is of historical interest that Gentzen had proved this result independently in 1933, the year Gödel published his result, but that he withdrew his galley proof before final publication when Gödel's result became known;

Section IV constitutes the core of the consistency proof and for this section there exist three versions:

24) Gentzen (10) p. ()

25) Gentzen (5) p. 71

(x)

The first version was never published, since Gentzen withdrew his galley proof when he was accused of making implicit use of the intuitionist fan theorem;²⁶⁾ the second and third versions are the actually published results in which transfinite induction on the ordinal numbers up to ϵ_0 is employed. The first proof might be called the "natural" proof, since the logical calculus used is Gentzen's NK-calculus²⁷⁾ of natural deduction, whereas the second proof, which might be called the "logistic" proof, is based on Gentzen's LK-calculus,²⁸⁾ a Hilbert-type first order predicate calculus. The equivalence of the calculi NK, LK, and a Hilbert-type calculus LDK is actually proved in Part II of the Investigations.²⁹⁾

The New Version of the consistency proof contained in the present translation is actually a new version of Section IV of (5) while excerpts from the Galley Proof contain those parts of Section IV of (5) that were re-written during the proof-reading of (5).

The translator hopes that the inclusion of the

26) Personal communication from Prof. P. Bernays

27) Gentzen (5) footnote (9) and Gentzen (3) I, Section II, pp. 4-8 (tr.) especially (5.3), p.8 (tr.)

28) Gentzen (10) p. 131 and Gentzen (3) I, Section III, pp. 8-10 (tr.)

29) Gentzen (3) I, p. 8 (1.2) (tr.)

Galley Proof in this thesis will help to illustrate Gentzen's genius for assessing the value of criticism advanced against him and will also bring out the remarkable speed with which Gentzen's managed to develop the radically new approach to the consistency problem through transfinite induction while correcting the first draft of his consistency proof. The following correspondences should enable the reader to compare the relevant passages of ^{the} first consistency proof with those of the galley proof:

Articles 14.3-14.63 (pp. ~~188-206~~) correspond to articles 13.93-15.4 (pp. ~~98-116~~).³⁰

Furthermore, pp. ~~204-206~~ : ".... In the transformation of the derivation in Paragraph 12 and in its applications (at 14.441, 14.442 and 14.443)"

corresponds to

p. ~~118~~ : ".... The following functions, in particular, "the ordinal number of the derivation" (15.2).

Also, p. ~~206~~ : ".... Complete induction transfinite proposition"

30) Cf. footnote (20) in Gentzen (5)

corresponds to

"..... Furthermore, propositions
of to ordinal number diminishes"
(15.3), p. 114.

Also, pp. 206 : "..... I hope that these reflections
..... can be further diminished"

corresponds to

the entire section 16.11 on pp. .

Any other changes in the initial draft were purely
editorial and have no bearing on Gentzen's arguments.

At this point a word must be said about the terminology
adopted in translating technical terms. In most cases
the appended glossary will resolve any difficulties
that may arise in this connection; with the following
exceptions: The main difficulties that arose in trans-
lating the present papers were associated with Gentzen's
notions of "Richtigkeit" and "Korrektheit" and those
of "Sinn" and "Bedeutung". Gentzen predicates "Richtig-
keit" of axioms, theorems, propositions, formulae,
sequents, etc. as well as of inferences, inference
figure schemata, and forms of inference. On the
other hand, he also predicates "Korrektheit" of
inferences, inference figure schemata, and forms of

inference as well as of derivations and rules of inference. In conformity with English usage, "Richtigkeit" has been translated by "truth" so that of a given proposition constructed from objects, functions, and predicates, together with the logical connectives, it can be "calculated" whether it is "true" or "false" (p. 51). On the other hand, an idea is lost in the translation of Article 7.3 (p. 53), where Gentzen asserts that it is easily proven that the logical rules of inference are "richtig" in the sense that their application to "richtig" mathematical basic sequents leads to other "richtig" derivable sequents. Here the translator speaks of "correct" logical rules of inference and of "true" sequents. This is justified since the term "richtig" can stand for almost anything from "true", "correct", 31) "well-formed", "well defined" to "formally valid". In fact, in the consistency proof the notion of "truth" can in most cases be taken as synonymous with "formal truth" in the

31) "correct" ordinal number, Gentzen (5) p. 108

sense of the propositional calculus (Article 7.2, p. 52) and of the predicate calculus (Articles 7.3, p. 53, 13.4, p. 81, 1.2, p. 134). Nevertheless, Gentzen does think of a "formally true" sequent as expressing more than the fact that it is true by definition (p. 81) when he speaks of "obviously true" sequents, and he also seems to feel that the number-theoretical axioms represent "immediately obvious" propositions (pp. 48, 62, 125). Such asides are in some sense Gentzen's "private" views and have no bearing on the argument developed in the consistency proof. The last paragraph of the first proof (p. 130) and the remark on p. 172 concerning the definition of the ordinal numbers make Gentzen's attitude to this distinction amply clear. We are also forced to speak of the "intuitive notion of truth" being replaced by the "statability of a reduction procedure". This is a translation of "inhaltlicher Richtigkeitsbegriff" and "Angebbbarkeit einer Reduziervorschrift". The latter represents a technical notion and hence its somewhat un-English translation seems justified.

Incidentally, inferences, forms of inference, inference figure schemata, rules of inference and derivations will always be spoken of as "correct". This allows us to consider a derivation, i.e., a proof, as "correct", if it contains no error in its construction. The view that sequents can be "obviously true" and that number-theoretical axioms are "immediately self-evident" leads us to the notions of "inhaltlich" and "formal". Since there exists no adequate English adjective for rendering "inhaltlich" in any literal way, it has been translated by "intuitive". Other writers have used "concrete" (Kneebone) or coined the new adjective "contensive" (Curry). If we speak of a proposition as being "intuitively true" we will mean that it has a definite "sense" beyond its formal "meaning". The word "sense" is here intended to coincide with the intuitionists notion of "sense" illustrated by the following example:

"We start with the natural numbers 1, 2, 3, etc. They are so familiar to us, that it is difficult to reduce this notion to simpler ones. Yet I

shall try to describe their sense in plain words. In the perception of an object we conceive the notion of an entity by a process of abstracting from the particular qualities of the object. We also recognize the possibility of an indefinite repetition of the conception of entities. In these notions lies the source of the concept of the natural numbers. (L.E.J. Brouwer 1907, p. 3; 1948, p. 1237)³²⁾

Gentzen's "Sinn" is mostly translated by "sense" so that a "senseless" proposition, or a proposition "devoid of sense", will express the fact that it has no "sense" in Brouwer's terminology. By denying an "actualist infinity", for example,³³⁾ we are at the same time robbing propositions such as "Fermat's last theorem is true or is not true" of their "sense".³⁴⁾ Consequently, according to the intuitionist view, such propositions cannot even be asserted. A large part of the consistency proof is thus concerned with ascribing a "finitist sense" to actualist propositions, viz., for every provable proposition in the formalism developed, a reduction rule must be stated and this rule represents the "finitist sense" of the proposition.³⁵⁾ It should be noted that informally the word "Bedeutung" also means "significance" and

32) Heyting (1) p. 13

33) Gentzen (5) p. 58

34) Gentzen (5) p. 57

35) Gentzen (5) p. 129

"consequence" and it has been translated in this way whenever appropriate. An interesting contrast is noticeable between Gentzen's and Frege's notions of the "Sinn" and the "Bedeutung" of a proposition. For the "Bedeutung" of a proposition Frege takes its truth-value³⁶⁾ and for its "Sinn" the thought expressed.³⁷⁾ Thus, "The morning star is a body illuminated by the sun" and "The evening star is a body illuminated by the sun" are two propositions with the same "Bedeutung" but with a different "Sinn". This observation should suffice to show that the notions of "Sinn" and "Bedeutung", or "sense" and "meaning", as used by Gentzen, must not be identified with their counterparts in Frege's famous paper. Gentzen is concerned rather with the "mathematical sense" and "mathematical meaning" of a proposition.

Let us now return to the logical calculi NK and LK used in the two consistency proofs and examine their similarities and differences:

36) Frege, On Sense and Reference, p.496) (Hartman, Philosophy
 37) Frege, On Sense and Reference, p.495) of Recent Times I,
 McGraw-Hill 1967)

The NK calculus, or the classical calculus of natural deduction, makes no explicit use of "sequents". It is indeed based on the idea of introducing assumptions and then applying logical rules to them in order to deduce a proposition from them. But here, in keeping with mathematical practice, the inference proper makes the proposition to be proved independent of the assumptions. In sub-case 2 (4.42) p. 30 of Euclid's proof of the non-existence of a largest prime, for example, it is assumed that there exists an arbitrary number $\underline{d}$ with the property that $\underline{d} > 1 \& \underline{d} \leq \underline{n} + 1$ and it is eventually inferred from this assumption that $\neg \underline{d} \mid \underline{n}! + 1$. It therefore holds without assumption that $(\underline{d} > 1 \& \underline{d} \leq \underline{n} + 1) \supset \neg \underline{d} \mid \underline{n}! + 1$.

The intuitive meaning of the logical symbols employed is explained in Article 3.12. In the NK calculus, the structure of the above argument is formalized by the inference figure $\frac{[A]}{\underline{B}} \quad A \supset B$. Here we have of course Gentzen's version of Herbrand's famous

Deduction Theorem. In (5), the above figure takes the form of $\frac{\mathcal{U}, \pi \rightarrow \mathcal{B}}{\pi \rightarrow \mathcal{U} \supset \mathcal{B}}$. The inference figure has

been replaced by an inference figure schema. The rule of substitution can therefore be dispensed with and to the "propositional" variables A and B now correspond the "syntactic" variables $\mathcal{A}$ and $\mathcal{B}$.

This change in notation has the additional advantage of making it clear that in the inference figure $\frac{\mathcal{A} \supset \mathcal{B}}{\mathcal{A} \supset \mathcal{B}}$ the propositions A and B are "mentioned" but not "used", in the now familiar distinction.³⁸⁾ Normally, valid propositions result from valid propositions only if the symbols to be replaced are mentioned but not used in the proposition concerned. An elaborate system of quotation marks can also be employed to make this distinction clear. The use of a symbol is thus often indicated by putting quotation marks around the symbol: In the proposition "7 is a number", the symbol 7 is mentioned but not used, whereas in "7 is an Arabic numeral" it is actually used. Consequently, "7 is a number" and "VII is a number" are true propositions, whereas in the case of "7 is an Arabic numeral" and "VII is an Arabic numeral", the second proposition is false. Gentzen avoids

38) Beth (1) pp. 257-258

(xx)

such difficulties in the consistency proofs by employing inference figure schemata in place of inference figures, i.e., by introducing syntactic variables in place of propositional variables.

The other major difference between the calculus in (5) and the calculus NK is the explicit introduction of sequents in (5). A sequent is an expression of the form $\mathcal{U}_1, \dots, \mathcal{U}_\mu \rightarrow \mathcal{B}$ where $\mathcal{B}$ is a syntactic variable for a proposition that depends on the assumptions $\mathcal{U}_1, \mathcal{U}_2, \dots, \mathcal{U}_\mu$. How such a notion arises "naturally" from the examination of Euclid's classical proof can be seen from ~~at~~ 4.2, subcase 2: The proposition $\neg \underline{d} \mid \underline{a}! + 1$ depends on the assumptions $\forall y[(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n}) \supset \neg \underline{y} \mid \underline{a}! + 1]$ as well as on $\neg(\underline{n} + 1) \mid \underline{a}! + 1$ and $\underline{d} > 1 \ \& \ \underline{d} \leq \underline{n} + 1$.

This dependence is symbolized by writing
$$\forall y[(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n}) \supset \neg \underline{y} \mid \underline{a}! + 1], \neg(\underline{n} + 1) \mid \underline{a}! + 1, \underline{d} > 1 \ \& \ \underline{d} \leq \underline{n} \longrightarrow \neg \underline{d} \mid \underline{a}! + 1$$

i.e., a proposition of the form $\mathcal{U}_1, \mathcal{U}_2, \mathcal{U}_3 \rightarrow \mathcal{B}$
where $\mathcal{U}_1$ stands for $\forall y[(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n}) \supset \neg \underline{y} \mid \underline{a}! + 1]$;
 $\mathcal{U}_2$ stands for $\neg(\underline{n} + 1) \mid \underline{a}! + 1$;

$\mathcal{U}_j$ stands for $j > 1 \ \& \ j \leq n$;
 and $\mathcal{B}$ stands for $\neg j \mid j! + 1$.

$\mathcal{U}_1, \mathcal{U}_2, \mathcal{U}_3$ are called the antecedent formulae,
 and $\mathcal{B}$ is called the succedent formula. It is
 intuitively clear that $\mathcal{B}$ remains a valid proposi-
 tion if we change the order of the assumptions made,
 or if we add a further assumption on which $\mathcal{B}$ does
 not depend, or if in the case where the same assump-
 tion appears more than once, we cancel one of the
 occurrences. This leads naturally to what will be
 called the "structural" rules of inference of Inter-
 change, Thinning, and Contraction. In (5), a thin-
 ning is called an "omission of an antecedent formula"
 and a contraction an "adjunction of an addition antece-
 dent formula". The reason for this terminology is
 precisely the fact that Gentzen considers these
 rules to arise "naturally" and are thus part of
 his "indisputable" forms of inference. This inter-
 pretations of Gentzen's motives seems reasonable,
 especially since he had developed the formal calculus
 of sequents before the consistency proof and had

used the terminology of (10) in that calculus.

It is worth noting that Gentzen seems to have been led to the formulation of his predicate logic in terms of "sequents" by studying a paper by P. Hertz³⁹⁾ entitled "Über Axiomensysteme für beliebige Satzsysteme". In fact, it would appear that a critical study and solution of the problem posed by Hertz has had a remarkable influence on Gentzen's entire methodological thinking. Gentzen's very first publication⁴⁰⁾ makes this amply clear. The "cut", for example, which led eventually to the famous "Hauptsatz",⁴¹⁾ is a generalization a kind of syllogism found in Hertz's paper.⁴²⁾ The "chain rule", too, which is needed in (5) in order to change the vertical arrangement of proofs into a horizontal one for the purpose of assigning measures of complexity to the different proofs in number theory, can be found in Hertz's paper.⁴²⁾

For the purpose of an easier understanding of Gentzen's calculi NK and LK used in (5) and (10) it should be

39) Math. Ann., 89, (1923), Heft 1, 2; 101, (1929)

40) Gentzen (1)

41) Gentzen (3) I, p. 11 (2.5, 2.513); II, pp.2-3,(2.1)(tr.)

42) Math. Ann., 101 (1929) pp. 459, 462, 473 and
Gentzen (1) p. 331

pointed out that sequents can be afforded an intuitive meaning in terms of the logical connectives of the propositional calculus:

$$1) \mathcal{U}_1, \mathcal{U}_2, \dots, \mathcal{U}_\mu \rightarrow \mathcal{B}_1, \mathcal{B}_2, \dots, \mathcal{B}_\nu \quad (\mu, \nu \geq 1)$$

can be expressed by the implication

$$\mathcal{U}_1 \& \mathcal{U}_2 \& \dots \& \mathcal{U}_\mu \supset \mathcal{B}_1 \vee \dots \vee \mathcal{B}_\nu \quad 43)$$

$$2) \rightarrow \mathcal{B}_1, \dots, \mathcal{B}_\nu \quad \text{as} \quad \mathcal{B}_1 \vee \dots \vee \mathcal{B}_\nu$$

$$3) \mathcal{U}_1, \dots, \mathcal{U}_\mu \rightarrow \quad \text{as} \quad \neg (\mathcal{U}_1 \& \dots \& \mathcal{U}_\mu)$$

$$4) \text{ The empty sequent } \rightarrow \quad \text{as} \quad F \quad (\text{the false or any false proposition})$$

The discussion up to this point makes it clear that the introductions and eliminations of the connectives $\&, \vee, \supset, \vee, \exists$ in (5) have analogous counterparts in NK. The main difference between the calculi arises in the treatment of negation due to the difficulties inherent in a finitist interpretation of that connective.⁴⁴⁾ On the other hand, the LK-calculus agrees in its entirety with the calculus employed in (10). The inference of complete induction, which is of central importance in (5) and (10), does not of course appear among the rules of inference formalized in NK and LK.

43) Gentzen (3) II, pp. 6-8 (tr.); (5) 4.56, p.35; p.43, 45-46

44) Gentzen (3) I, pp. 6-8 (tr.)

It must be pointed out that the notion of a "sequent", as it appears above, and as it is used in (10), is a generalization of the "natural" notion of a sequent as it arises in (5). This generalization is achieved by allowing multiple succedent formulae and thus symmetrizing the antecedent and succedent. The difference in (5) and (10) is the same as that exhibited in the calculi LJ and LK of the Investigations.⁴⁵⁾ To the introduction and elimination of a logical connective in (5) there thus corresponds in (10) the introduction of that connective in the succedent and antecedent of the sequent.⁴⁶⁾

For his initial supply of "true" sequents, Gentzen takes certain "logical" and "mathematical" basic sequents. In (5)⁴⁷⁾, a "logical" basic sequent is a sequent of the form $\mathcal{D} \rightarrow \mathcal{D}$, a "mathematical" basic sequent a sequent of the form $\rightarrow \mathcal{E}$. In (10), a "logical" basic sequent is still of the form $\mathcal{D} \rightarrow \mathcal{D}$, but a "mathematical" basic sequent

⁴⁵⁾ Gentzen (3) I, p. 10 (2.3) (tr.)

⁴⁶⁾ Gentzen (10), 1.6, p. 142

⁴⁷⁾ Gentzen (5), p. 38

becomes a sequent consisting entirely of prime formulae (i.e., formulae without logical connectives) and becoming a "true" sequent with every arbitrary substitution of numerical terms for possible occurrences of free variables.⁴⁸⁾ The fact that "logical" and "mathematical" basic sequents are to be taken as "true" follows from the "definitive form" for sequents stated at 13.4 (5) and also the remarks made at 7.3, p. 58 and p. 134. It must be observed quite generally that Gentzen's notion of the "truth" of a sequent is a generalization of the "truth" and "falsity" of an implication as it is customarily defined in propositional logic, just as the sequent itself, as illustrated above, can be regarded as a generalization of the implication

$$A \supset B \quad \text{to} \quad A_1 \& \dots \& A_n \supset B_1 \vee \dots \vee B_m$$

As far as the "truth-content" of number-theoretical propositions is concerned, Gentzen makes his position quite clear in the Investigations, II p. 7 (tr.). He considers propositions of the kind $3=3$, $4 \neq 5 \supset 5=4$, in

48) Gentzen (10), p. 139

general, any arithmetical axiom as "true", as long as every numerical special case is intuitively true.

He states further that

"it is almost self-evident that from such propositions no contradictions are derivable by means of propositional logic. A proof for this would hardly be more than a formal paraphrasing of an intuitively clear situation of fact."⁴⁹⁾

This observation is of course entirely in harmony with the finitist attitude in its literal sense as explained in the excerpt from Hilbert and Bernays.⁵⁰⁾ Gentzen observes further that universally quantified arithmetic axiom formulae are also entirely reliable as long as each numerical special case is intuitively true.⁵¹⁾

What the consistency proof must therefore do is to prescribe a method whereby any arbitrary derived sequent, i.e., a number-theoretical theorem, that is non-contradictory, can be brought into a form in which its truth is intuitively recognizable. Here the notion of the "intuitive truth-content" represents

49) Gentzen (3) II, p. 7 (tr.)

50) Cf. footnote 10)

51) Gentzen (3) II, p. 7 (tr.)

a property that is common to all non-contradictory end-sequents but fails to hold for the sequent $\rightarrow \mathcal{U} \& \neg \mathcal{U}$ in (5) and the empty sequent in (10). This property is furthermore invariant under the reduction of sequents and derivations of sequents to their definitive form (by the definition of a derivation and the assumed reliability of the logical calculi employed).

Let us take a closer look at the procedure whereby Gentzen proves the consistency of formalized elementary number theory. We shall deal with (5) in detail.

The reader will recall that the consistency of the propositional calculus in Hilbert-Ackermann (pp. 32-33) is proved by assigning certain numerical values to proposition variables and showing that the axioms of the calculus always take the value 0 and that this value is preserved by the permissible logical rules of the calculus.

Gentzen in fact follows the method here indicated, although the procedure is of course considerably more complicated. All propositions of elementary

number theory become formulae in the formalization of elementary number theory carried out in (5) (Section II) and each formula is in turn written as a sequent. A sequent is either a logical basic sequent, a mathematical basic sequent, or the end-sequent of a derivation. Thus we must in some way state a rule whereby the "truth" of the end-sequent of a derivation can be "calculated". This is achieved by stating, first, a reduction rule for sequents, i.e., end-sequents of derivations, which reduces the sequent in question to its "definitive" form in finitely many steps so that we can decide by inspection whether or not the sequent is "true". It is then shown further that no contradiction can be derived in the theory formalized if we start from "true" propositions, i.e., logical basic sequents and mathematical basic sequents, and apply to them the rules of the logical calculus developed. In formal language this amounts to showing, as pointed out at the beginning of this monograph, that the sequent $\rightarrow \mathcal{U} \not\vdash \mathcal{U}$ cannot be the end-sequent of a derivation, or, in the case of (10),

that the empty sequent $\longrightarrow$ cannot be derived. The sufficiency of this demonstration can be seen by considering the following consequences of the two statements:

In (5):

$$\begin{array}{c}
 \frac{\frac{\frac{\longrightarrow \mathcal{U} \& \neg \mathcal{U}}{\longrightarrow \mathcal{U}}}{\neg 1=2 \rightarrow \mathcal{U}}}{\longrightarrow \neg \neg 1=2} \\
 \hline
 \longrightarrow 1=2
 \end{array}$$

In (10):

$$\frac{\frac{\frac{\longrightarrow \mathcal{U} \& \neg \mathcal{U}}{\longrightarrow \neg \mathcal{U}}}{\neg 1=2 \rightarrow \neg \mathcal{U}}}{\longrightarrow \neg \neg 1=2}$$

Let us examine the method of reduction developed in (5) somewhat more closely: The reduction rule is stated in the following form:

Sequents, in general, are reduced to the definitive form of 13.4 by first eliminating the connectives $\forall, \exists, >$ and replacing them by $\forall, \&, \neg$. This is permissible since the former connectives can be expressed equivalently by means of the latter without affecting the finitist interpretation of the calculus. Next all free variables are replaced by

(xxx)

arbitrary numerals and all minimal terms by their associated "functional values". Then three cases are distinguished depending on whether the succedent formula is of the form $\forall x F(x)$, or $\mathcal{U} \& \mathcal{B}$ or $\neg \mathcal{U}$. These are then reduced to definitive form. If the succedent formula is false and no antecedent formula is false, three further cases arise depending on whether one of the antecedent formulae has the form $\forall x F(x)$, $\mathcal{U} \& \mathcal{B}$ or $\neg \mathcal{U}$. The definitive form is finally reached in all cases. Special rules for logical and mathematical basic sequents are stated in 13.91 and 13.92, pp. 86-87 (5), so that they too are brought into definitive form. The next step consists in reducing an actually derived sequent to definitive form, i.e., a sequent that is the end-sequent of a derivation. In the case where the end-sequent is already in definitive form, no reduction is defined. In order to be able to state his reduction rule, Gentzen first changes the vertical i.e., tree-like, arrangement of the derivations into a horizontal one by modifying the notion of a derivation and by introducing the

so-called "chain rule". The details of the changes which this entails can easily be seen from the consistency proof proper.⁵²⁾

The crucial step in the argument consists in showing that the reduction procedure leads to the definitive form of a sequent in finitely many steps. This follows easily if no complete induction occurs in the proof. The consistency of elementary number theory without complete induction is, after all, already a consequence of the Hauptsatz and was proved in the Investigations II, pp. 5-7 (tr.) as an illustration of the consequences of that important theorem. As Gentzen points out,⁵³⁾ the special position of complete induction is due to the fact that the number of reduction steps required can become arbitrarily large. In (5) 14.24, the total number of steps required for the reduction of $\pi, \Delta \rightarrow \mathcal{F}(z)$, for example, depended on n (the value of z) and n in turn, depended on choices if z was a free variable. Thus there exists no general bound

52) Gentzen (5) pp. 92-93

53) Gentzen (5) p. 124

for the total number of reduction steps required for the reduction of $\mathcal{T}, \Delta \rightarrow \mathcal{F}(t)$. In (10) the difficulty arises when a proposition with the maximal number of connectives is proved by complete induction⁵⁴⁾ making necessary the special reduction step of a "CJ-reduction". Here the difficulty is not the number of individual complete inductions that may occur in a derivation. They can be fused into a single induction, as Gentzen has shown.⁵⁵⁾ Thus,

"the number of complete inductions occurring in a number-theoretical proof is no measure of the "complexity" of the proof in its meta-mathematical discussion; although this number does have some bearing on this point, it is not ~~the~~ number of inductions but their "degree", i.e., the complexity of the induction proposition, that counts."

In order to show that elementary number theory with complete induction is free from contradiction, we must therefore resort to ranking all possible proofs according to their complexity and the measures here required are the transfinite ordinals. Gentzen puts it this way:⁵⁶⁾

⁵⁴⁾ Gentzen (10) p. 148

⁵⁵⁾ Gentzen (12)

⁵⁶⁾ Gentzen (7) pp. 77-78

(xxxiii)

In carrying out the consistency proof for elementary number theory, one has to consider all conceivable proofs in number theory and to show that in a certain sense, to be defined formally below, each individual proof yields a "correct" result, in effect, no contradiction. The "correctness" of a proof rests on the correctness of certain other, simpler proofs that are contained in it as special cases or as parts. This situation leads us to arrange the proofs in linear order in such a way that those proofs on whose correctness the correctness of another proof depends are made to precede the latter proof in the sequence. This arrangement of the proofs is achieved by assigning to each proof a certain transfinite ordinal number. The proofs preceding a given proof are precisely those whose ordinal numbers precede its ordinal number in the sequence of ordinal numbers. At first sight, the natural numbers might appear to suffice as ordinal numbers for such an arrangement. Yet in actual fact the transfinite ordinal numbers are needed for the following reason: It may happen that the correctness of a proof depends on the correctness of infinitely many simpler proofs. An example: In a proof a proposition is proved by complete induction for all natural numbers. In that case the correctness of this proof obviously depends on the correctness of the infinitely many individual proofs obtained by specializing to a particular natural number. In such cases it is not sufficient to use a natural number as ordinal number, because each natural number is preceded by only finitely many other numbers in the natural ordering. Hence we need the transfinite ordinal numbers in order to represent the arrangement of the proofs according to its complexity.

Furthermore, it now becomes apparent precisely why the inference of transfinite induction is needed as the crucial inference for the consistency proof: With this inference we prove the "correctness" of each individual proof. For the proof no. 1 is trivially correct; and if the correctness of all proofs that precede a particular proof in the arrangement is established, then the proof in question is also

correct, since the arrangement was made in such a way that the correctness of a proof depends on the correctness of specific earlier proofs. From this we can now obviously deduce the correctness of every proof by means of that same transfinite induction and have thus, in particular, established the desired consistency."

The particular form in which the ordinal numbers are defined in (5) and (10) has obviously no bearing on the result. In footnote 21) of (5) Gentzen points out the connection between the two methods of introducing the ordinal numbers. The difference in the two approaches lies again in the greater emphasis on the constructive nature of the elements of Cantor's second number class in (10) versus their "natural" formulation in (5) just as in the case of the difference in the introduction of the natural numbers in the two proofs. The constructive nature of the transfinite ordinals up to ϵ_0 is brought out very clearly in Gentzen (7), where he points out that⁵⁷⁾

THERE ARE REALLY

57) Gentzen (7) p. 76

"there are really only two operations involved through whose repeated application all these numbers are quite automatically generated: 1) given a number, we can form its successor (addition of 1); 2) given an infinite sequence of numbers, we can form a new number ranking after the whole of the sequence (formation of a limit). This procedure may not appear to be constructive since the formation of ω already seems to imply the actualist conception of the completed sequence of the natural numbers. Yet this is not implied; it is quite possible here to interpret the concept of infinity potentially by saying, for instance: The number ω stands in the ordering relation $\omega < \omega$ to every natural number n , however far one may go in forming constructively such ordinal numbers. The infinite sequences that occur in forming the other ordinal numbers must be interpreted constructively in the same way."

Nevertheless, it is precisely at this point that the consistency proof goes beyond the formal framework of elementary number theory and thus escapes the limitations imposed on consistency proofs by Godel's Theorem. As Gentzen remarks:⁵⁸⁾

"The transfinite induction in the consistency proof is now precisely that rule of inference which necessarily, by the theorem of Godel, cannot be shown to be correct by means of the techniques of elementary number theory."⁵⁹⁾

58) Gentzen (7) pp. 78-79

59) Gentzen (11)

It has been emphasized from the outset that all methods of proof used in the consistency proof should be entirely finitist in the sense of Hilbert and Bernays. It now turns out that the notion of "finitist" as quoted earlier is too narrow for the purposes of the consistency proof. The schema of transfinite induction contains a universally quantified premise. Hilbert and Bernays thus recognize the necessity for extending their notion of finitist to this new situation. To quote:⁶⁰⁾

"Wir wollen uns überlegen, wie diese Schlussweise ... als gültig einzusehen ist, und zwar auf eine Art, bei der die Abweichung von unserem bisherigen Verfahren der finiten Beweisführung lediglich darin besteht, dass Allsätze also Prämissen von Sätzen zugelassen werden. Dabei kommen als Prämissen immer nur solche Allsätze vor, die sich nachtraglich auf Grund des Ergebnisses der Überlegung also zutreffend erweisen."

The fact that in the course of a proof by means of this "restricted" transfinite induction, as it is employed in the consistency proof, the quantified propositions are themselves actually verified "on

60) Hilbert and Bernays II, p. 363

the basis of the results of the arguments carried out" would thus seem to bring this inference within the class of intuitionistically acceptable methods of proof. "Finitist" in the wider sense and "intuitionist" in the narrower sense have here in some ways become synonymous, especially since intuitionists tend to accept transfinite induction as long as it ranges no further than Cantor's first ϵ -number. Even E. Borel, one of the constructivist's most staunch supporters, was prepared to accept Cantor's second number class as constructively given.⁶¹⁾ To this Kleene observes that⁶²⁾

"to what extent the Gentzen proof can be accepted as securing classical number theory in the sense of that problem formulation is in the present state of affairs a matter of individual judgment depending on how ready one is to accept induction up to ϵ_0 as a finitary (finitist) method."

It would appear, therefore, that what must turn the scales in favour of Gentzen is precisely the constructive nature of the transfinite ordinal used in the

61) Kneebone (1) p. 246

62) Kleene (1) p. 479

induction. We are here dealing with numbers that can be uniquely displayed and well-ordered and which permit a very natural arithmetical manipulation. Gentzen himself pointed out the categorical difference between the denumerable quantities involved in Cantor's second number class and higher cardinalities when he said that⁶³⁾

"in general set theory, for example, a careful proof-theoretical investigation will eventually confirm the view that all cardinalities that exceed the denumerable ones have, in a very real sense, only an illusory existence and that it would be wisest to do without these concepts."

Nevertheless, Gentzen continued to strive for a consistency proof for classical analysis, considering this branch of mathematics as an "idealization", in the Hilbert sense⁶⁴⁾. He felt that this view of the "second level" of mathematics would eventually restore unity among mathematicians if not among philosophers.⁶⁵⁾ Unfortunately it has not been possible to date to realize the second stage of Hilbert's programme in the way Gentzen had envisaged its solution. We there-

63) Gentzen (7) p. 74

64) Gentzen (9) p. 268

65) Gentzen (9) p. 266 seq.

fore rest our case by quoting Gentzen's own words that summarize the value of the consistency proof for elementary number theory in the best possible way:⁶⁶⁾

"The proof certainly reveals that it is possible to reason consistently "as though" everything in the infinite domain of objects were as actualistically determined as in finite domains. Yet whether and in how far anything "real" corresponds to the actualist sense of a transfinite proposition - apart from what its restricted finitist sense expresses - is a question which the consistency proof does not answer."

66) Gentzen (5) p. 130

THE CONSISTENCY OF ELEMENTARY NUMBER THEORY

By "elementary number theory", I mean the theory of the natural numbers that does not make use of techniques from analysis such as, e.g., irrational numbers or infinite series.

The aim of the present paper is to prove the consistency of elementary number theory or, rather, to reduce the question of consistency to certain general fundamental principles.

How such a consistency proof can be carried out at all and for what reasons it is necessary or at least very desirable to do so will be discussed in Section I.

SECTION I.

REFLECTIONS ON THE PURPOSE AND POSSIBILITY OF CONSISTENCY PROOFS

In paragraph 1, I consider the question why consistency proofs are necessary and, in paragraph 2, how such proofs are possible.⁽¹⁾ In doing so, I shall briefly restate those aspects of the problem, already familiar to many readers, which are of particular relevance to the remainder of this paper.

Paragraph 1

THE ANTI~~NOM~~IES OF SET THEORY AND THEIR SIGNIFICANCE FOR MATHEMATICS AS A WHOLE⁽²⁾

1.1. Mathematics is regarded as the most certain of all the sciences. That it could lead to results which contradict one another seems impossible. This faith in the indubitable certainty of mathematical proofs was sadly shaken around 1900 by the discovery of the "antinomies (or "paradoxes") of set theory". It so happens that in this specialized branch of mathematics contradictions arise in contexts in which no uniquely identifiable mistake in the inferences used can be found.

Particularly instructive is "Russell's Antinomy" which I shall now discuss in detail.

1.2. A set is a collection of arbitrary objects ("elements of the set"). An "empty set", which has no elements at all, is also admitted. We now divide the sets into "sets of the first kind", i.e., sets which contain themselves as an element, and "sets of the second kind", i.e., sets which do not contain themselves as an element.

We now consider the set M which has for its elements the entire collection of the sets of the second kind. Does this set itself belong to the first or the second kind? Both alternatives are absurd: For if the set M belongs to the first kind, i.e., if it contains itself as an element, then this contradicts its definition by which all of its elements were supposed to be sets of the second kind. Suppose, therefore, that the set M belongs to the second kind, i.e., that it does not contain itself as an element. Since it has all sets of the second kind as elements by definition, it must in that case also contain itself as an element and we have thus once again arrived at a

contradiction.

1.3. The result is Russell's antinomy which shows how easily an obvious contradiction can result from a small number of admittedly somewhat subtle inferences.

What is the actual significance of this fact for mathematics as a whole? We may be inclined, at first, to dismiss the entire argument as unmathematical by claiming that the concept of a "set of arbitrary objects" is too vague to count as a mathematical concept.

Yet this objection becomes void if we restrict ourselves to quite specific purely mathematical objects by making the following stipulation for example: The only objects admitted as elements of a "set" are first: Arbitrary natural numbers (1,2,3,4 etc.); second: Arbitrary sets consisting of admissible elements.

Example: The following three elements form an admissible set: First, the number 4; second, the set of all natural numbers; third, the set whose two elements are the number 3 and the set of all natural numbers.

Using this purely mathematical concept of a set, we can then repeat the above (1.2) argument and obtain the same contradiction.

1.4. The fact that we happen to have chosen the natural numbers for our initial objects has obviously no bearing at all on the emergence of the antinomy. It cannot, therefore, be said that a contradiction has been revealed in the domain of the natural numbers; the fault must

be sought rather in the logical inferences employed.

1.5. It is thus natural to go back to look for a definite error in the reasoning that has led to the antinomy. We might, for example, argue that the set $\mathcal{M}$ was defined by referring to the totality of all sets (which was indeed subdivided into sets of the first and second kinds, and where $\mathcal{M}$ was formed with sets of the second kind). The set was then itself added to this totality, which raised the question of whether it belongs to the first or second kind. Such a procedure is circular; it is illicit to define an object by means of a totality and to add it then to that totality so that in some sense it contributes to its own definition ("circulus vitiosus").

We might feel that the correct interpretation of the set $\mathcal{M}$ should rather be the following:

If a definite totality of sets is given then this totality may be subdivided into sets of the first and second kinds. Yet if the sets of the second kind (or alternatively, the first kind) are combined to a new set $\mathcal{M}$, then that set constitutes something completely new and cannot itself be added to that totality.

1.6. The fact that the forms of inference leading to the antinomy seem correct at first sight is based on the idea that the concept of a "set" denotes something "actual" (and the totality of all sets, therefore, constitutes a predetermined closed totality); the critique advanced against this view implies that new sets can be formed only

"constructively" so that a new set depends in its construction on already existing sets.

1.7. If we were to think that the antinomy has thus been explained away quite satisfactorily, we must at once face up to a new difficulty: The form of reasoning (the *circulus vitiosus*) which we have just declared to be inadmissible is already being used in analysis in a quite similar form in the usual proofs of some rather simple theorems, e.g., the theorem: "A function which is continuous on a closed interval and is of different sign at the endpoints has a zero in the interval."

The proof of this result is essentially carried out in the following way: The totality of points in the interval is divided into points of the first and second kinds so that a point is of the first kind if the function has the same sign for all points to its right up to the end of the interval and it belongs to the second kind if this is not the case. The limit point defined by this subdivision is then the required zero. It belongs itself to the points of the interval. Hence we have the "*circulus vitiosus*": The real number concerned is defined by referring to the totality of the real numbers (in an interval) and is then itself added to that totality.

This form of inference is nevertheless considered correct in analysis on the following grounds: The number concerned is, after all, not newly created by the given definition, it already actually exists within the totality of the real numbers and is merely singled out from

this totality by its definition.

Yet exactly the same can be said about the antinomy mentioned above: The set *m* is already actually present in the totality of all sets (defined at 1.3) and is merely singled out by its definition (at 1.2) from this totality.

Considerable differences certainly exist between the forms of inference used to derive the antinomy and those customary in proofs from analysis. Yet we must ask ourselves whether these differences are radical enough to justify a further use of these inferences in analysis--since no contradictions have yet arisen--or whether their similarity with the inferences that have let the antinomies should not prompt us to eliminate these inferences also from analysis. Here the opinions of mathematicians concerned with these questions diverge.

1.8. We can indeed challenge the correctness of other forms of inference customary in mathematics because of certain remote analogies that may be drawn between them and inferences leading to the antinomies. Especially radical in this respect are the "intuitionists" (Brouwer), who even object to forms of inference customary in number theory, not only because these inferences might possibly lead to contradictions, but because the theorems to which they lead have no actual sense and are therefore worthless. I shall come back to this point later in greater detail. (Paragraphs 9-11 and 17.3).

Less radical are the "logicists" (Russell). They draw a line between permissible and non-permissible forms of inference, and the antinomies

turn out to be a consequence of a non-permissible *circulus vitiosus*. At one time the logicians had also disallowed the inference applied in the example from analysis cited above ("ramified theory of types"), yet this inference was later re-admitted.

1.9. Altogether we are left with the following picture:

The contradictions (antinomies) which had occurred in set theory, a specialized branch of mathematics, had given rise to further doubts about the correctness of certain forms of inference customary in the rest of mathematics. Various attempts to draw a line between permissible and non-permissible forms of inference have led to different approaches to the subject.

In order to end this unsatisfactory state of affairs, Hilbert drew up the following programme:

The consistency of the whole of mathematics, in so far as it actually is consistent, is to be proved along exact mathematical lines. This proof is to be carried out by means of forms of inference that are completely unimpeachable ("finitist" forms of inference.)

How such a consistency proof is conceivable at all will be discussed more fully in Paragraph 2.

In the remainder of this paper, I shall then carry out such a consistency proof for elementary number theory. Yet even here we shall meet forms of inference whose closer inspection will give us cause for concern. More about this in Section 3. One point should however be made clear

from the outset; those forms of inference which might possibly be considered disputable occur hardly ever in actual number theoretical proofs, (11.4); we must, therefore, not be misled and because of the great self-evidence of these proofs consider a consistency proof as superfluous.

Paragraph 2

HOW ARE CONSISTENCY PROOFS POSSIBLE?

2.1. General remarks about consistency proofs.

2.11. The consistency of geometries is usually proved by appealing to an arithmetic model. Here the consistency of arithmetic is therefore pre-supposed. In a similar way we can also establish a correspondence between some parts of arithmetic, e.g., the theory of the complex numbers and that of the real numbers.

What remains to be proved ultimately is the consistency of the theory of the natural numbers (elementary number theory) and the theory of the real numbers (analysis) of which the former forms a part; and finally the consistency of set theory as far as that theory is consistent.

2.12. This task is basically different and more difficult than that of reducing the consistency of one theory to that of another theory by mapping the objects of the former theory onto the objects of the latter. Let us look more closely at the situation in the case of the natural numbers:

These numbers can obviously not be mapped onto a simpler domain of objects. Nor are we indeed concerned with the consistency of the

domain of numbers itself, i.e., with the consistency of the basic relationships between the numbers as determined by the "axioms" (e.g., the "Peano Axioms" of number theory). To prove the consistency of these axioms without invoking other equivalent assumptions seems inconceivable. We are concerned rather with the consistency of our logical reasoning about the natural numbers (starting from their axioms) as it occurs in the proofs of number theory. For it is precisely our logical reasoning which in its unrestricted application leads to the antinomy (1.4). Yet such general notions as that of an arbitrary set of sets (1.3) is of course no longer considered to be part of number theory. Elementary number theory comprises merely finite sets (of natural numbers, for example). If infinite sets of natural numbers are included we are already in the domain of the real numbers and hence in analysis. This is the fundamental distinction between elementary number theory and analysis.

From here we reach set theory by extending the concept of a "set" still further.

How can the consistency of arithmetic be actually proved?

2.2 "Proof Theory".

2.21. The assertion that a mathematical theory is consistent constitutes a proposition about the proofs possible in that theory. It says,

after all, that none of these proofs leads to a contradiction. In order to carry out a consistency proof we must therefore make the possible proofs in the theory themselves objects of a new "meta-theory". The theory that has arbitrary mathematical proofs for its objects is called "proof theory" or "meta-mathematics".

2.22 An example of a theorem in proof theory is the "principle of duality" in projective geometry:

It says roughly that from a theorem about points and straight lines (in the plane) another true theorem results if the word "point" is replaced by "straight line" and the word "straight line" by "point". The theorem "for any two distinct straight lines there exists exactly one point coinciding with both straight lines (i.e., lying on them)", for example, has a dual counterpart in the theorem: "For any two distinct points there exists exactly one straight line coinciding with both points (i.e., passing through them)".

The principle of duality is justified thus: The axioms of projective geometry in the plane are such that the dual transformation of an axiom always yields another axiom. If any theorem has therefore been derived from these axioms then a uniform replacement in the proof of the word "point" by "straight line" and of the word "straight line" by "point" thus yields a proof for the dual theorem.

This justification is obviously proof-theoretical since it is about the "proof of a theorem".

(This example also shows that proof theory is capable of advancing mathematics proper.)

2.23. The "formalization" of mathematical proofs.

As the objects of our proof theory we shall take the proofs carried out in mathematics proper. These proofs are customarily expressed in the words of our language. These have the disadvantage that there are many different utterances for the same proposition, that an arbitrariness exists in the order of the words, sometimes even ambiguities. In order to make an exact study of proofs possible it is therefore desirable to begin by giving them a uniform uniquely predetermined form. This is achieved by the "formalization" of the proofs: the words of our language are replaced by definite symbols, the logical forms of inference by formal rules for the formation of new formalized propositions from already proven ones.

In Section II, I shall carry out such a formalization for elementary number theory.

The example of the principle of duality (2.22) shows clearly the difficulties that are inherent in proof theory without a formalization: the linguistic expression of the theorem "for two mutually distinct straight lines there exists exactly one point that coincides with both straight lines" had to be chosen artificially in such a way that the replacement of "point" by "straight line" and vice versa again resulted in a linguistically meaningful theorem. Even in

carrying out the proof of the principle of duality we are left with the feeling that we have not offered a really rigorous proof. In order to make this proof rigorous, we do in fact require an exact formalization of the propositions and proofs (for the domain of projective geometry).

2.3. The forms of inference used in the consistency proof; the theorem of Gödel.

2.31. How can a consistency proof (for elementary number theory, for example) be carried out by means of proof theory?

To begin with, it will have to be made precise what is to be understood by a formalized "number-theoretical proof". Then it must be established that among all such possible "proofs" there can exist none which leads to a "contradiction". (This is a simple property of "proofs" which can be verified immediately for any given "proof".)

Such a consistency proof is once again a mathematical proof in which certain inferences and specific concepts must be used. Their reliability (especially their consistency) must already be pre-supposed. There can be no "absolute consistency proof". A consistency proof can merely reduce the correctness of certain forms of inference to the correctness of other forms of inference.

It is therefore clear that in a consistency proof one can use only forms of inference that count as considerably more secure than the forms of inference of the theory whose consistency is to be proved.

2.32. Of the greatest significance at this point is the following proof-theoretical theorem proved by K. Gödel:⁽³⁾ "It is not possible to prove the consistency of a formally given (demarcated) theory which comprises elementary number theory (nor that of elementary number theory itself) by means of the entire collection of techniques proper to the theory concerned (given that that theory is really consistent)".

From this it follows that the consistency of elementary number theory, for example, cannot be established by means of a part of the methods of proof used in elementary number theory nor indeed by all of these methods. To what extent then is a genuine re-interpretation still possible?

It remains actually quite conceivable that the consistency of elementary number theory can in fact be verified by means of techniques which, in part, no longer belong to elementary number theory, but which can nevertheless be considered to be more reliable than the doubtful components of elementary number theory itself.

2.4. In the following (Sections II-IV) I shall carry out a consistency proof for elementary number theory. In doing so I shall indeed apply techniques of proof which do not belong to elementary number theory (16.2). Several different consistency proofs already exist in the literature ⁽⁴⁾ all of which reach essentially the same point, viz., the verification of the consistency of elementary number theory with the

exclusion of the inference of "complete induction" which, as is well known, constitutes a very important and frequently used form of inference in number theory. The inclusion of complete induction in my proof presents certain difficulties (16.2).

SECTION II

THE FORMALIZATION OF ELEMENTARY NUMBER THEORY

As pointed out at 2.23, it is desirable for a proof-theoretical discussion of a mathematical theory to give that theory a precise formally determined structure. In order to prove the consistency of elementary number theory, I shall therefore begin by carrying out such a formalization of elementary number theory.⁽⁵⁾

This task falls into two parts:

1. The formalization of the propositions occurring in elementary number theory (Paragraph 3).
2. The formalization of the methods of proof used in elementary number theory, i.e., forms of inference and specific concepts (Paragraphs 4 - 6).

Paragraph 3

THE FORMALIZATION OF THE PROPOSITIONS OCCURRING IN ELEMENTARY NUMBER THEORY

3.1. Preparatory Remarks.

3.11. A formalization of mathematical propositions represents nothing fundamentally new even outside of proof theory. It is indeed true to say that mathematics has always undergone a successive formalization, i.e., a replacement of language by mathematical symbols. There are, for example, propositions which are written entirely in symbols, e.g.,

$$(\underline{a} + \underline{b}) \cdot (\underline{a} - \underline{b}) = \underline{a}^2 - \underline{b}^2$$
, in words: "The product of the sum and the difference of the numbers $\underline{a}$ and $\underline{b}$ is equal to the difference of the squares of both numbers".

The proposition "If $\underline{a} = \underline{b}$, then $\underline{b} = \underline{a}$ ", on the other hand, is generally still represented by using words. Completely formalized, it is written: $\underline{a} = \underline{b} \supset \underline{b} = \underline{a}$.

3.12. The linguistic expression "If $\mathcal{U}$ holds, then $\mathcal{B}$ holds", formally written as $\mathcal{U} \supset \mathcal{B}$, is an example of the logical composition of propositions for the purpose of forming a new proposition. Further compositions of propositions are constructed with the symbols $\&$, $\vee$, $\neg$, $\forall$ and $\exists$ with the following meanings: $\mathcal{U} \& \mathcal{B}$ means " $\mathcal{U}$ holds and $\mathcal{B}$ holds", $\mathcal{U} \vee \mathcal{B}$: " $\mathcal{U}$ holds or $\mathcal{B}$ holds" (i.e., at least one of the two propositions holds), $\neg \mathcal{U}$: " $\mathcal{U}$ does not hold".
 $\forall x \mathcal{U}(x)$: " $\mathcal{U}(x)$ holds for all x ", $\exists x \mathcal{U}(x)$: "There is a x , so that $\mathcal{U}(x)$ holds".

3.13. As an example we shall consider "Goldbach's Conjecture" ("Every even natural number can be represented as the sum of two prime numbers"),

which can be formally written as:

$$\forall x \{ 2|x \supset \exists y \exists z [y+z=x \wedge (\text{Prime } y + \text{Prime } z)] \}.$$

Here $\text{Prim } \underline{a}$ stands for "a is a prime number"; $\underline{a}|\underline{b}$, as usual, for

"a is a divisor of b". All variables shall refer only to the natural

(= positive whole) numbers.

3.14. The symbols $=$, Prime and $|$ are "predicate symbols"; once its argument places have been filled by numbers, such a symbol constitutes a proposition. The symbol $+$ is a "function symbol"; once its argument places have been filled by numbers, it represents another number.

The formal counterpart of a proposition is generally called a "formula".

(Just as in mathematics, for example, $(\underline{a} + \underline{b}) \cdot (\underline{a} - \underline{b}) = \underline{a}^2 - \underline{b}^2$ is called a "formula", although in a special sense.)

After these remarks I shall now give a precise characterization of those formal expressions which are to be admitted into our formalized number theory for the purpose of representing propositions.

3.2. Precise definition of a formula. (6)

3.21. The following kinds of symbol will serve for the formation of formulae:

3.211. Symbols for individual natural numbers: 1,2,3,4,5,6,7,8,9, 10,11,12, , briefly called "numerals". (No symbols for other numbers will be needed.)

3.212. Variables for natural numbers: These I divide into free and bound variables (vid.seq.). Any other symbol that has not yet been used may serve as a variable; yet it must be stated in each case whether such a symbol is to be a free or a bound variable.

3.213. Symbols for individual functions, briefly called "function symbols": $+$, $\cdot$, and others as needed (cf. 6.1).

3.214. Symbols for individual predicates, briefly called "predicate symbols": $=$, $<$, Prime , $|$ and others as needed (cf. 6.1).

3.215. Symbols for the logical composition (logical connectives)* of propositions: $\wedge$, $\vee$, $\supset$, $\neg$, $\forall$, $\exists$.

3.22. Definition of a term (formal expression for a - individual or indeterminate - number):

3.221. Numerals (3.211) and free variables (3.212) are terms.

* Translator's addition.

3.222. If s and t are terms, then so are $s + t$ and $s \cdot t$;
other terms may be formed analogously by means of further function
symbols that may have been introduced (3.213).

3.223. No expressions other than those formed in accordance with
3.221 and 3.222 are terms.

3.224. Example of a term: $[(a + 2)^3 \cdot t] + 4$; where a and b are
free variables.

Brackets serve as usual the purpose of avoiding ambiguities in con-
nection with the grouping of the individual symbols.

3.23. I now define the notion of a formula (formal counterpart of a
number-theoretical proposition):

3.231. A predicate symbol (3.214) whose "argument places" are filled
by arbitrary terms (3.22) yields a formula.

Example: $(2 + a) \cdot 4 < t$

3.232. If $\mathcal{U}$ is a formula, then so is $\neg \mathcal{U}$. If $\mathcal{U}$ and $\mathcal{B}$ are
formulae, then so are $\mathcal{U} \wedge \mathcal{B}$, $\mathcal{U} \vee \mathcal{B}$, and $\mathcal{U} \supset \mathcal{B}$.

3.233. From a formula results another formula if all free occurrences
of a variable in the former formula are replaced by a not yet occurring

bound variable x , and if the entire formula is at the same time prefixed by $\forall x$ or $\exists x$.

3.234. No expressions other than those formed in accordance with 3.231, 3.232 and 3.233 are formulae.

3.24. As in the case of terms, brackets must be used to display unambiguously how a formula has been constructed in accordance with 3.232 and 3.233.

Examples of Formulae: Cf. 3.13, 3.11, 3.231.

The intuitive sense of a formula follows from the remarks in 3.1. It should be observed that a formula with free variables constitutes an "indefinite" proposition which becomes a "definite" proposition only if all free variables in it are replaced by terms without free variables, e.g. numerals. (7)

A minimal term is a term consisting of one function symbol with numerals in the argument places, e.g.: $1 + 3$.

A minimal formula is a formula consisting of one predicate symbol with numerals in the argument places, e.g.: $4 = 12$.

A transfinite formula is a formula containing at least one $\forall$ or $\exists$ symbol.

3.25. German and Greek letters will be used as "syntactic variables", i.e., as variables for our proof-theoretical considerations about number

theory.

3.3. The question arises whether our concept of formula is wide enough for the representation of all propositions occurring in elementary number theory.

Strictly speaking the answer is no. There certainly are propositions in elementary number theory (examples will follow) for which no immediate formal representation exists in terms of the methods formalized. Yet such propositions may safely be disregarded as long as equivalent propositions exist in each case which are representable in our formalism.

For this a number of important examples:

3.31. As the objects of number theory I have taken into account only the natural numbers. Yet the rest of the integers as well as, occasionally, the fractions are of course also needed in number theory. It is not difficult however, to reinterpret all propositions about integers and fractions as propositions about the natural numbers by observing that the negative integers can be made to correspond to pairs of positive integers and the fractions to pairs of integers. (An example: $\frac{a}{b} = \frac{c}{d}$ is interpreted as $a.d = c.b$.) Even in the case where finite sets of natural numbers or of integers or fractions are included among the objects of number theory (e.g., the "complete systems of residues") it is still possible to reinterpret all propositions as propositions

about the natural numbers, although in this case such interpretations are considerably more complicated. The same holds for propositions in which diophantine equations etc. are taken as objects.

Here I do not intend to discuss these methods of re-interpretation further; they present no fundamental difficulties (especially for the consistency proof) and anyone who concerns himself somewhat more closely with these matters will easily see their feasibility (cf. also 17.2).

If infinite sets of natural numbers, integers, or fractions are admitted, such a reinterpretation is in general no longer possible precisely because we are here already dealing with objects from analysis (cf. 2.12). It is, after all, customary to define the real numbers themselves as certain infinite sets of rational numbers.

3.32. Functions and predicates occur in number theory in a variety of forms. In defining a formula I have taken account of this fact by admitting at 3.213 and 3.214 "further symbols as needed". Further details about the introduction of arbitrary functions and predicates follow in Paragraph 6.

3.33. As far as the logical compositions of propositions are finally concerned, the following are for example customary utterances:

"The proposition $\mathcal{U}$ holds if and only if the proposition holds."

This composition of propositions is of course represented as follows:

$(\mathcal{B} \supset \mathcal{U}) \ \& \ (\mathcal{U} \supset \mathcal{B})$.

"There exists exactly one number x , for which the proposition

$U(x)$ holds." For this we write: $\exists x[U(x) \wedge \forall y(U(y) \supset y = x)]$,

with obviously the same meaning. (Suitable bound variables are to be chosen for x and y ; where $U(y)$ is the expression resulting from $U(x)$ by the replacement of x by y .)

"There are infinitely many numbers x for which the proposition $U(x)$ holds." This simply means that "For every number there exists a number greater than the former for which U holds."; and in this form the proposition is representable in our formalism.

"The sum total of numbers x for which the proposition $U(x)$ holds, is n." This proposition - in which n is left indeterminate - can be represented in our formalism only in a considerably paraphrased form, possibly as follows: We include the finite sets of natural numbers among the objects of the theory and paraphrase the above proposition thus: "There exists a set of natural numbers whose sum total of elements is n and for which it further holds: for each one of its elements the proposition U holds and every number for which U holds belongs to the set." Here "sum total" is a function, "belongs" a predicate, and both must be defined in advance. The concept of a finite set can finally be paraphrased again according to 3.31.

There exists of course a variety of other linguistic expressions all of which can be reduced to immediately formalizable utterances.

()

3.34. I shall return to the question of the completeness of the formalism in a quite general sense after the consistency proof has been carried out (17.1).

Paragraph 4

EXAMPLE OF A PROOF FROM ELEMENTARY NUMBER THEORY

4.1. I now proceed to the formalization of the methods of proof used in elementary number theory. I.e., I shall have to list as completely as possible all forms of inference and methods of forming concepts used in proofs of elementary number theory and assign to them a formally fixed form which avoids all the ambiguities of their linguistic representation.

Only if a precise formal definition can then be given of what is meant by an elementary number-theoretical "proof" can we begin with the proof theory of elementary number theory.

I shall begin by giving an example of a number theoretical proof in this paragraph, and shall classify the individual forms of inference according to definite criteria by means of examples from this proof. In Paragraph 5, I shall then give a precise general formulation to these forms of inference.

Finally, in Paragraph 6, I shall discuss the methods of forming concepts and the here relevant number-theoretical "axioms".

4.2. As an example of a proof from elementary number theory, I shall choose Euclid's well known proof of the theorem: "There are infinitely many prime numbers."

I shall first carry out the proof in words in a version which has been adapted somewhat to the purpose in hand.

In the following (throughout Paragraph 4) I use the letters $\underline{a}, \underline{b}_1, \underline{b}_2, \underline{c}, \underline{d}, \underline{l}, \underline{m}, \underline{n}$, as free variables, the letters $\underline{p}, \underline{q}$ as bound variables (for natural numbers).

The theorem to be proved can be formulated more precisely as follows:
 "For every natural number there exists a larger one which is a prime number."

Suppose now that $\underline{a}$ is an arbitrary natural number. We must therefore show that there exists a prime number which is larger than $\underline{a}$. We consider the number $\underline{a}! + 1$. If it is a prime number then it already validates our assertion. If it is not a prime number then it has a divisor $\underline{b}_1$, (excluding 1 and itself). This divisor is larger than $\underline{a}$ for no number from 2 to $\underline{a}$ can divide $\underline{a}! + 1$, since any such division leaves a remainder of 1. If $\underline{b}_1$ is a prime number it validates our assertion. If it is not a prime number then it too has a divisor $\underline{b}_2$ other than 1 and itself. This number also divides $\underline{a}! + 1$ since it divides $\underline{b}_1$. Hence $\underline{b}_2$ is also larger than $\underline{a}$. By continual repetition of this process we obtain a sequence of numbers: $\underline{a}! + 1, \underline{b}_1, \underline{b}_2, \dots$ whose terms become smaller and smaller. Hence the sequence must terminate at some point, i.e., its last number is a prime number which divides $\underline{a}! + 1$ and is larger than $\underline{a}$. Hence the existence of a prime number which is larger than $\underline{a}$ has been verified. Since $\underline{a}$ was a quite arbitrary natural number it follows: for every natural number there exists a larger one which is a prime number. Q.E.D.

4.3. In the proof I have pre-supposed various simple theorems as already known. These can be reduced to still simpler facts by further

proofs, yet this is unimportant for our present purpose since we are interested, above all, in the inferences which occur in the various steps of the above proof.

Here we must keep in mind that through practice we are accustomed to carrying out entire sequences of proof at once without being conscious of each individual inference contained in that step. In order to single out the actual elementary inferences I shall therefore go through Euclid's proof once again and bring to light all individual inferences contained in some parts of the proof. At the same time, I shall formalize according to Paragraph 3 the various propositions as they occur.

4.4. Detailed Analysis of Euclid's Proof.

The proof contains a somewhat disguised "complete induction" (cf. the place: "by continual repetition of this process....."). The usual normal form of the inference by complete induction is this:

The validity of a proposition is proved for the number 1; then it is shown that if the proposition holds for an arbitrary natural number $\underline{n}$ it also holds for $\underline{n} + 1$; hence this proposition holds for any arbitrary natural number.

It will also be convenient to reduce to this normal form the disguised complete induction which here occurs; to do this I shall choose the following proposition as the "induction proposition", formulated for a number $\underline{m}$: "Either there exists a prime number among the numbers from 1 to $\underline{m}$ which is greater than $\underline{a}$ or none of these numbers,

except 1, divides $\underline{a}! + 1$. Formally:

$$\{\exists \underline{z} [\underline{z} \leq \underline{m} \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})]\} \vee \forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{m}) \supset \neg \underline{y} | (\underline{a}! + 1)].$$

The proof now runs as follows:

4.41. The induction proposition must first be proved for $\underline{m} = 1$.

Here its second part is satisfied automatically since there is obviously no number which is larger than 1 and smaller than or equal to 1. Explicitly: for an arbitrary $\underline{c}$ it holds that

$\neg(\underline{c} > 1 \ \& \ \underline{c} \leq 1)$; this we assume as given. Then it also holds that

$(\underline{c} > 1 \ \& \ \underline{c} \leq 1) \supset \neg \underline{c} | (\underline{a}! + 1)$, and, since $\underline{c}$ was arbitrary,

$$\forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq 1) \supset \neg \underline{y} | (\underline{a}! + 1)] \quad . \quad \text{From this}$$

follows, in accordance with the meaning of $\vee$ (3.12), the entire

induction proposition for $\underline{m} = 1$, viz:

$$\{\exists \underline{z} [\underline{z} \leq 1 \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})]\} \vee \forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq 1) \supset \neg \underline{y} | (\underline{a}! + 1)] .$$

4.42. Next comes the "induction step", i.e.: we assume that the

induction proposition has been proved for an arbitrary number $\underline{n}$, so

that

$$\{\exists \underline{z} [\underline{z} \leq \underline{n} \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})]\} \vee \forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n}) \supset \neg \underline{y} | (\underline{a}! + 1)].$$

is valid and is now to be proven to be valid for $\underline{n} + 1$. This is done

as follows:

On the basis of the induction assumption two cases are possible:

1. $\exists \underline{z} [\underline{z} \leq \underline{n} \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})]$

2. $\forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n}) \supset \neg \underline{y} | (\underline{a}! + 1)]$

In the first case it follows without difficulty that $\exists \underline{z} [\underline{z} \leq \underline{n} + 1 \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})]$. I shall not discuss this further. In this case therefore the induction assumption has already been proved for $\underline{n} + 1$, viz.,

$$\{ \exists \underline{z} [\underline{z} \leq \underline{n} + 1 \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})] \} \vee \forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n} + 1) \supset \neg \underline{y} \mid (\underline{a}! + 1)] .$$

Let us now look at the second case:

$$\forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n}) \supset \neg \underline{y} \mid (\underline{a}! + 1)] .$$

It holds that $(\underline{n} + 1) \mid (\underline{a}! + 1) \vee \neg (\underline{n} + 1) \mid (\underline{a}! + 1)$. We can thus distinguish two sub-cases:

First sub-case: $(\underline{n} + 1) \mid (\underline{a}! + 1)$. From this it follows that

$$\text{Prime}(\underline{n} + 1) \ \& \ (\underline{n} + 1) > \underline{a} , \text{ which I shall briefly show since}$$

the only forms of inference here used are those for which we already have examples in the remaining parts of the proof:

$\underline{n} + 1$ is a prime number; for if it had a divisor other than 1 and itself, it would be smaller than $\underline{n} + 1$, and would then also divide

contradicting our assumption that $\forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq \underline{n}) \supset \neg \underline{y} \mid (\underline{a}! + 1)]$.

. $\underline{n} + 1$ is furthermore larger than $\underline{a}$; for the numbers from

2 to $\underline{a}$ do not divide $\underline{a}! + 1$ since such a division always leaves a

remainder of 1. Hence it holds in fact that $\text{Prime}(\underline{n} + 1) \ \& \ (\underline{n} + 1) > \underline{a}$

; also $\underline{n} + 1 \leq \underline{n} + 1$, hence it holds that $\underline{n} + 1 \leq \underline{n} + 1 \ \&$

$\text{Prime}(\underline{n} + 1) \ \& \ (\underline{n} + 1) > \underline{a}$ and consequently also that

$$\exists \underline{z} [\underline{z} \leq \underline{n} + 1 \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})] ,$$

and thus

$$\{\exists z [z \leq n+1 \& (\text{Prime } z \& z > a)]\} \vee \forall y [(y > 1 \& y \leq n+1) \supset \neg y | (a!+1)].$$

Second sub-case: $\neg(n+1) | (a!+1)$. Suppose that d is an arbitrary number with a property that $d > 1 \& d \leq n+1$. From $d \leq n+1$ follows $d \leq n \vee d = n+1$, which is to be taken as given.

Suppose first that $d \leq n$; it also holds that $\forall y [(y > 1 \& y \leq n) \supset \neg y | (a!+1)]$, hence in particular that $(d > 1 \& d \leq n) \supset \neg d | (a!+1)$. From $d > 1$, together with $d \leq n$, it follows that $d > 1 \& d \leq n$, and together with the preceding therefore $\neg d | (a!+1)$.

If, however, $d = n+1$, then, because of $\neg(n+1) | (a!+1)$, it also follows that $\neg d | (a!+1)$.

Thus it holds in general that $\neg d | (a!+1)$, a consequence of the assumption $d > 1 \& d \leq n+1$. Hence we can write $(d > 1 \& d \leq n+1) \supset \neg d | (a!+1)$, and further, since d was an arbitrary number,

$$\forall y [(y > 1 \& y \leq n+1) \supset \neg y | (a!+1)],$$

and thus once again

$$\{\exists z [z \leq n+1 \& (\text{Prime } z \& z > a)]\} \vee \forall y [(y > 1 \& y \leq n+1) \supset \neg y | (a!+1)],$$

We have therefore in all cases obtained the induction proposition for $n+1$, and this completes the induction step.

4.43. The proof is now quickly completed:

From the complete induction follows the validity of the induction proposition for arbitrary numbers. We require it only for the number $\underline{a}! + 1$:

$$\{\exists z [z \leq \underline{a}! + 1 \ \& \ (\text{Prime } z \ \& \ z > \underline{a})]\} \\ \vee \forall y [(y > 1 \ \& \ y \leq \underline{a}! + 1) \supset \neg y | (\underline{a}! + 1)].$$

From the second case it follows in particular that

$$(\underline{a}! + 1 > 1 \ \& \ \underline{a}! + 1 \leq \underline{a}! + 1) \supset \neg (\underline{a}! + 1) | (\underline{a}! + 1).$$

Yet it holds that $\underline{a}! + 1 > 1 \ \& \ \underline{a}! + 1 \leq \underline{a}! + 1$, which we assume as given; hence it follows that $\neg (\underline{a}! + 1) | (\underline{a}! + 1)$. On the other hand it holds of course that $(\underline{a}! + 1) | (\underline{a}! + 1)$, we have thus obtained a contradiction, i.e., the second case cannot possibly occur; formally:

$$\neg \forall y [(y > 1 \ \& \ y \leq \underline{a}! + 1) \supset \neg y | (\underline{a}! + 1)].$$

Only the first case remains, i.e.: $\exists z [z \leq \underline{a}! + 1 \ \& \ \text{Prime } z \ \& \ z > \underline{a}]$.

Suppose that $\underline{b}$ is such a number so that $\underline{b} \leq \underline{a}! + 1 \ \& \ (\text{Prime } \underline{b} \ \& \ \underline{b} > \underline{a})$ holds.

Then it holds in particular that $\text{Prime } \underline{b} \ \& \ \underline{b} > \underline{a}$, from which $\exists z [\text{Prime } z \ \& \ z > \underline{a}]$ follows. Yet $\underline{a}$ was a quite arbitrary natural number, hence this result holds for all natural numbers, i.e., $\forall y \exists z (\text{Prime } z \ \& \ z > y)$.

This is the final result of Euclid's proof.

4.5. Classification of the individual forms of inference by reference to examples from Euclid's Proof.

Let us now focus our attention on the individual inferences occurring in the above proof. Here the following classification almost suggests itself:

For every logical connective $\&$, $\vee$, $\supset$, $\neg$, $\forall$, and $\exists$ there exists certain associated forms of inference. These may be further divided into forms of inference by which the connective concerned is introduced and other forms of inference by which the same connective is eliminated from a proposition. As examples for each individual case I shall give an inference from Euclid's Proof:

4.51. A $\forall$ -introduction occurs at the end of the proof, viz: after $\exists z [\text{Prime } z \ \& \ z > a]$ was proved for number $\underline{a}$, it was inferred that $\forall y \exists z (\text{Prime } z \ \& \ z > y)$

A $\forall$ -elimination took place at 4.42, subcase 2, where from $\forall y [(y > 1 \ \& \ y \leq n) \supset \neg y | (a! + 1)]$ it was inferred that $(\underline{d} > 1 \ \& \ \underline{d} \leq \underline{n}) \supset \neg \underline{d} | (\underline{a}! + 1)$

4.52. A $\&$ -introduction (from 4.42, 2 subcase): the two propositions and $\underline{d} \leq \underline{n}$ together yielded the proposition

$$\underline{d} > 1 \ \& \ \underline{d} \leq \underline{n}.$$

A $\&$ -elimination (from 4.43):

From $\underline{e} \leq \underline{a}! + 1 \ \& \ (\text{Prime } \underline{e} \ \& \ \underline{e} > \underline{a})$ it was inferred that $\text{Prime } \underline{e} \ \& \ \underline{e} > \underline{a}$.

4.53. A $\exists$ -introduction (from 4.43):

From $\text{Prime } \underline{e} \ \& \ \underline{e} > \underline{a}$ it was inferred that $\exists \underline{z} (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})$

A $\exists$ -elimination (from 4.43):

The proposition

$$\exists \underline{z} [\underline{z} \leq \underline{a}! + 1 \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})]$$

was valid. From it was inferred that $\underline{e} \leq \underline{a}! + 1 \ \& \ (\text{Prime } \underline{e} \ \& \ \underline{e} > \underline{a})$,

where $\underline{e}$ stood for any one of the numbers which existed by virtue of the previous proposition.

4.54. A $\forall$ -introduction (from 4.41):

$$\text{From } \forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq 1) \supset \neg \underline{y} | (\underline{a}! + 1)]$$

it was inferred that

$$\{\exists \underline{z} [\underline{z} \leq 1 \ \& \ (\text{Prime } \underline{z} \ \& \ \underline{z} > \underline{a})]\} \\ \vee \forall \underline{y} [(\underline{y} > 1 \ \& \ \underline{y} \leq 1) \supset \neg \underline{y} | (\underline{a}! + 1)].$$

A $\forall$ -elimination (from 4.42): the proposition

$$\{\exists \underline{z} [\underline{z} \leq \underline{n} \wedge (\text{Prime } \underline{z} \wedge \underline{z} > \underline{a})]\} \vee \forall \underline{y} [(\underline{y} > 1 \wedge \underline{y} \leq \underline{n}) \supset \neg \underline{y} | (\underline{a}! + 1)].$$

was valid. From it resulted the distinction of cases:

$$\text{Case (i) : } \exists \underline{z} [\underline{z} \leq \underline{n} \wedge (\text{Prime } \underline{z} \wedge \underline{z} > \underline{a})];$$

$$\text{Case (ii): } \forall \underline{y} [(\underline{y} > 1 \wedge \underline{y} \leq \underline{n}) \supset \neg \underline{y} | (\underline{a}! + 1)].$$

This distinction of cases was terminated by the fact that the same proposition

$$\{\exists \underline{z} [\underline{z} \leq \underline{n} + 1 \wedge (\text{Prime } \underline{z} \wedge \underline{z} > \underline{a})]\} \vee \forall \underline{y} [(\underline{y} > 1 \wedge \underline{y} \leq \underline{n} + 1) \supset \neg \underline{y} | (\underline{a}! + 1)]$$

could eventually be inferred in both cases.

4.55. AD-introduction (from 4.42, subcase 2):

Starting with the assumption $\underline{d} \leq 1 \wedge \underline{d} \leq \underline{n} + 1$ we reached the result: $\neg \underline{d} | \underline{a}! + 1$. Hence $(\underline{d} > 1 \wedge \underline{d} \leq \underline{n} + 1) \supset \neg \underline{d} | (\underline{a}! + 1)$ was valid.

AD-elimination (from 4.42, subcase 2):

From $\underline{d} \leq 1 \wedge \underline{d} \leq \underline{n}$ and $(\underline{d} > 1 \wedge \underline{d} \leq \underline{n}) \supset \neg \underline{d} | (\underline{a}! + 1)$ it was inferred that $\neg \underline{d} | (\underline{a}! + 1)$.

4.56. For negation ($\neg$) the situation is not quite as simple; for here there exist several distinct forms of inference and these cannot be divided clearly into $\neg$ -introductions and $\neg$ -eliminations. I shall come back to this later (5.26). Here I shall cite only a single important example from Euclid's Proof, viz. a "reductio ad absurdum" - inference (from 4.43): $\neg \forall \underline{y} [(\underline{y} > 1 \wedge \underline{y} \leq \underline{a}! + 1) \supset \neg \underline{y} | (\underline{a}! + 1)]$

was inferred from the fact that the assumption $\forall y [(y > 1 \wedge y \leq n+1) \supset \neg y | (a!+1)]$ led to a contradiction, viz., to the proposition $\neg(a!+1) | (a!+1)$, whereas $(a!+1) | (a!+1)$ is indeed provable.

Paragraph 5

THE FORMALIZATION OF THE FORMS OF INFERENCE OCCURRING IN ELEMENTARY NUMBER THEORY

5.1. Preliminary Remarks.

My next task is to formulate the different kinds of forms of inference, which have been introduced by means of the above examples, in their most general form.

The determination of the individual forms of inference is not entirely unique. Yet the sub-division into introductions and eliminations of the individual logical connectives which I have chosen seems to me especially lucid and natural.

What, then, does the general form of a form of inference look like?

E.g., as the general form of the $\wedge$ - elimination one would be inclined to put simply the following: if a proposition of the form $U \wedge B$ is proven (where U and B are arbitrary formulae), then U (or B) is also valid.

Yet we must still keep in mind the following: the structure of a mathematical proof does not in general consist merely of a passing from valid propositions to other valid propositions by the application of the inferences. It happens, rather, that a proposition is often assumed as

valid and further propositions are deduced from it whose validity therefore depends on the validity of this assumption. Examples from Euclid's proof: The "reductio" (4.56), the $\supset$ -introduction (4.55), the induction step in the complete induction (4.42).

In order to describe completely the meaning of any proposition occurring in a proof we must therefore state in each case upon which of the assumptions that may have been made the proposition in question depends.

I therefore make it a rule that, together with every (formalized) proposition $\mathcal{B}$ occurring in a formalized proof the (formalized) assumptions $\mathcal{U}_1, \dots, \mathcal{U}_\mu$ upon which the proposition depends must also be listed in the following form:

$$\mathcal{U}_1, \mathcal{U}_2, \dots, \mathcal{U}_\mu \longrightarrow \mathcal{B}$$

which reads: From the assumptions $\mathcal{U}_1, \dots, \mathcal{U}_\mu$ follows $\mathcal{B}$. Such an expression I call a "sequent". If there are no assumptions, we write $\rightarrow \mathcal{B}$.

An example from Euclid's proof: The proposition $\neg d \mid (a!+1)$ from 4.42, sub-case 2, must, in order to display its dependence on assumptions, be represented by the following sequent:

$$\forall y [(y > 1 \ \& \ y \leq n) \supset \neg y \mid (a!+1)], \neg (n+1) \mid (a!+1), \\ d > 1 \ \& \ d \leq n+1 \longrightarrow \neg d \mid (a!+1).$$

Since every proposition of the original proof is now represented by a

sequent in the formalized proof we can formulate the forms of inference directly for sequents.

Our earlier example, the $\&$ - elimination, would now have to be formulated thus: "If the sequent $G_1, \dots, G_\mu \rightarrow \mathcal{A} \& \mathcal{B}$ is proven ($\mu \geq 0$), then $G_1, \dots, G_\mu \rightarrow \mathcal{A}$ or $G_1, \dots, G_\mu \rightarrow \mathcal{B}$ is also valid."

In the following, general schemata for the remaining forms of inference will be given in the same way.

5.2. Precise General Formulation of the Individual Forms of Inference.

5.21. Definition of a sequent ⁽⁹⁾ (formal expression for the meaning of a proposition in a proof together with its dependence on possible assumptions):

A sequent is an expression of the form:

$$\mathcal{A}_1, \mathcal{A}_2, \dots, \mathcal{A}_\mu \longrightarrow \mathcal{B}$$

$(\mathcal{A}_1, \mathcal{A}_2, \dots, \mathcal{A}_\mu)$

where arbitrary formulae (3.23) may take the place of $\mathcal{A}$ and $\mathcal{B}$. The

formulae $\mathcal{A}_1, \mathcal{A}_2, \dots, \mathcal{A}_\mu$ I call antecedent formulae of the sequent. ~~and $\mathcal{B}$ the succedent formula of the sequent.~~ It is permissible that no antecedent formulae occur, then the sequent has the form: $\longrightarrow \mathcal{B}$ yet there must always be a succedent formula.

5.22. Definition of a derivation (formal counterpart of a proof):

A derivation consists of a number of consecutive sequents of which each is either a "basic sequent" or has resulted from certain earlier sequents by a "structural transformation" or by the application of a "rule of inference". The definition of the various concepts will follow presently.

The last sequent of a derivation contains no antecedent formulae, its succedent formula is called the end-formula of the derivation. (It represents the proposition proved by the proof.)

5.23. Definition of a basic sequent:

I distinguish between "logical" and "mathematical" basic sequents.

A logical basic sequent is a sequent of the form $D \rightarrow D$, where D can be any arbitrary formula. (Such a sequent occurs in the formalization of a proof if and when an assumption D is made in the proof.)

A mathematical basic sequent is a sequent of the form $\rightarrow \mathcal{A}$, where the formula $\mathcal{A}$ represents a "mathematical axiom". What is to be understood by a number-theoretical "axiom", in particular, will be explained in Paragraph 6.

5.24. Definition of a structural transformation:

The following kinds of transformation of a sequent are called structural transformations (because they affect only the structure of a sequent, independently of the meaning of the individual formulae):

5.241. Interchange of two antecedent formulae;

5.242. Omission of an antecedent formula equal to another antecedent formula;

5.243. Adjunction of an arbitrary formula to the antecedent formulae;

5.244. Replacement of a bound variable within a formula throughout the scope of a $\forall$ -or $\exists$ -symbol by another bound variable not yet occurring in the formula.

Transformations according to 5.241, 5.242 and 5.244 obviously leave the meaning of the sequent unchanged since it makes no difference to the meaning of the sequent in what order the assumptions are listed or whether one and the same assumption is listed more than once or, finally, what symbol is used for the bound variable. All possibilities of transformation mentioned are thus of a purely formal nature and intuitively of no consequence; they must be stated explicitly only because of the special character of our formalization.

A structural transformation according to 5.243 means that to a proposition we may adjoin an arbitrary assumption upon which, besides other possible assumptions, it is to depend. At first this may seem somewhat strange; yet if a proposition is true, for example, we are forced to admit that in that case it is also valid on the basis of an arbitrary assumption. (If we were to stipulate that this may be asserted only in cases where a "factual dependence" exists, considerable difficulties would arise because of the possibility of proofs in which only an apparent use of an assumption is made.)

5.25. Definition of a rule of inference (formal counterpart of a form of inference):

Altogether we require thirteen rules of inference.

5.250. The German and Greek letters used here have the following meanings:

For $\mathcal{U}$, $\mathcal{B}$ and $\mathcal{C}$ may stand arbitrary formulae; for $\forall x F(x)$ or $\exists x F(x)$ an arbitrary formula of this form, then $F(a)$ or $F(t)$ stands for that formula which has resulted from $F(x)$ by the replacement of the bound variable x by an arbitrary free variable a or an arbitrary term t ; for $\mathcal{T}$, Δ , $\mathcal{H}$ may be put arbitrary, possibly empty, sequences of formulae, separated by commas (as antecedent formulae of the sequent concerned).

Now the individual rules of inference:

5.251. $\&$ - introduction: from the sequents $\mathcal{T} \rightarrow \mathcal{U}$ and $\Delta \rightarrow \mathcal{B}$ follows the sequent $\mathcal{T}, \Delta \rightarrow \mathcal{U} \& \mathcal{B}$.

$\&$ - elimination: from $\mathcal{T} \rightarrow \mathcal{U} \& \mathcal{B}$ follows the $\mathcal{T} \rightarrow \mathcal{U}$ or $\mathcal{T} \rightarrow \mathcal{B}$.

$\vee$ - introduction: from $\mathcal{T} \rightarrow \mathcal{U}$ follows $\mathcal{T} \rightarrow \mathcal{U} \vee \mathcal{B}$ or $\mathcal{T} \rightarrow \mathcal{B} \vee \mathcal{U}$.

$\vee$ - elimination: from $\mathcal{T} \rightarrow \mathcal{U} \vee \mathcal{B}$ and $\mathcal{U}, \Delta \rightarrow \mathcal{C}$ and $\mathcal{B}, \Delta \rightarrow \mathcal{C}$ follows $\mathcal{T}, \Delta, \mathcal{H} \rightarrow \mathcal{C}$.

$\forall$ - introduction: from $\mathcal{T} \rightarrow F(a)$ follows $\mathcal{T} \rightarrow \forall x F(x)$, as long as the free variable a does not occur in $\mathcal{T}$ and $\forall x F(x)$.

$\forall$ - elimination: from $\mathcal{T} \rightarrow \forall x F(x)$ follows $\mathcal{T} \rightarrow F(z)$.

$\exists$ - introduction: from $\mathcal{T} \rightarrow F(z)$ follows $\mathcal{T} \rightarrow \exists x F(x)$.

$\exists$ - elimination: from $\mathcal{T} \rightarrow \exists x F(x)$ and $F(a), \Delta \rightarrow \mathcal{C}$ follows $\mathcal{T}, \Delta \rightarrow \mathcal{C}$, as long as the free variable a does not occur in $\mathcal{T}, \Delta, \mathcal{C}$ and $\exists x F(x)$.

$\supset$ - introduction: from $\mathcal{U}, \mathcal{T} \rightarrow \mathcal{B}$ follows $\mathcal{T} \rightarrow \mathcal{U} \supset \mathcal{B}$

$\supset$ - elimination: from $\mathcal{T} \rightarrow \mathcal{U}$ and $\Delta \rightarrow \mathcal{U} \supset \mathcal{B}$ follows $\mathcal{T}, \Delta \rightarrow \mathcal{B}$.

5.252. The "reductio" rule: from $\mathcal{U}, \mathcal{T} \rightarrow \mathcal{B}$ and $\mathcal{U}, \Delta \rightarrow \neg \mathcal{B}$ follows $\mathcal{T}, \Delta \rightarrow \neg \mathcal{U}$.

"Elimination of the double negation": from $\mathcal{T} \rightarrow \neg \neg \mathcal{U}$ follows $\mathcal{T} \rightarrow \mathcal{U}$.

5.253. The "complete induction" rule: from $\mathcal{T} \rightarrow F(1)$ and $F(a), \Delta \rightarrow F(a+1)$ follows $\mathcal{T}, \Delta \rightarrow F(z)$, as long as the free variable a does not occur in $\mathcal{T}, \Delta, F(1)$ and $F(z)$.

5.26. Some remarks about the rules of inference.

In general the formulation of the individual rules of inference will be understood by referring to the appropriate examples of inferences (4.5). Several points should however be explained:

The $\mathcal{T}, \Delta$ and $\textcircled{H}$ are required since in the most general case we must allow for arbitrarily many assumptions.

The formulation of the $\supset$ -introduction, the "reductio", and of complete induction in which the appropriate assumptions are written down at

the same time practically suggests itself, whereas it seems probably somewhat artificial in the case of the $\forall$ - and $\exists$ -elimination, if these rules are compared with the corresponding examples of inferences (4.5). Yet the formulation is smoothest if in the distinction of cases ($\forall$ -elimination) the two possibilities that result are simply regarded as assumptions which become redundant as soon as the same result ($\mathcal{C}$) has been obtained from both; in the case of the $\exists$ -elimination the situation is similar: the proposition $\mathcal{F}(a)$ inferred from $\exists x \mathcal{F}(x)$ is an assumption only to the extent to which it is assumed of the variable a occurring in it that it represents any one of the numbers with the property $\mathcal{F}$ which exist according to $\exists x \mathcal{F}(x)$. This assumption becomes redundant as soon as a result ($\mathcal{C}$) has been deduced from it in which the variable a no longer occurs.

This leads me at once to a further point requiring some elaboration: it concerns the restrictions on free variables imposed in the case of the $\forall$ -introduction, the $\exists$ -elimination, and complete induction.

In each case the restriction says that in all formulae involved in the rule of inference (including the assumption formulae) the free variable a belonging to the rule of inference may occur only in the formula $\mathcal{F}(a)$ or $\mathcal{F}(a+1)$. It is easily seen by means of examples that this requirement is necessary in general and actually quite obvious; in the case of mathematical proofs it is fulfilled automatically. (By its very purpose the variable a is naturally out of place in the remaining formulae.)

The following must be said about the rules of inference for negation: as already mentioned at 4.56 the choice of elementary forms of inference is here more arbitrary than in the case of the other logical connectives. I should like to mention the following simple alternative rules of inference that might have been adopted:

From $\mathcal{U}, \Gamma \rightarrow \mathcal{B}$ and $\neg \mathcal{U}, \Delta \rightarrow \mathcal{B}$ follows $\Gamma, \Delta \rightarrow \mathcal{B}$.

From $\Gamma \rightarrow \mathcal{U} \vee \mathcal{B}$ and $\Delta \rightarrow \neg \mathcal{B}$ follows $\Gamma, \Delta \rightarrow \mathcal{U}$

(Example at 4.43).

From $\Gamma \rightarrow \neg \mathcal{B}$ and $\mathcal{U}, \Delta \rightarrow \mathcal{B}$ follows $\Gamma, \Delta \rightarrow \neg \mathcal{U}$.

From $\Gamma \rightarrow \neg \mathcal{U}$ follows $\Gamma \rightarrow \mathcal{U} \supset \mathcal{B}$ (Example at 4.41).

From $\Gamma \rightarrow \mathcal{U}$ and $\Delta \rightarrow \neg \mathcal{U}$ follows $\Gamma, \Delta \rightarrow \mathcal{B}$.

As logical basic sequents for the $\rightarrow$ -connective we could also have taken the following:

$\rightarrow \mathcal{U} \vee \neg \mathcal{U}$, "Law of the Excluded Middle". (Example at 4.42); $\rightarrow (\mathcal{U} \& \neg \mathcal{U})$, "Law of Contradiction".

Yet the two rules of inference which I have chosen (5.252) are sufficient; the remaining rules and the basic sequents here listed are already contained in them (if the rules of inference for the other logical connectives are included); this may be verified without any essential difficulties.

5.3. Are our rules of inference actually sufficient for the representation of all inferences that occur in elementary number theory?

5.31. The completeness of the purely logical rules of inference, i.e., the rules belonging to the connectives $\&, \vee, \supset, \neg, \forall, \exists$, has already been proved elsewhere⁽¹⁰⁾; (completeness here means that all correct inferences of the same type are representable by the stated rules).

To these forms of inference we must now, for the purpose of elementary number theory, add "complete induction". Here the question of the completeness of the rules of inference becomes a rather difficult problem; I shall return to it after the consistency proof (17.1) has been carried out. At this point I should merely like to observe the following:

It may be considered as fairly certain that all inferences occurring in the usual number-theoretical proofs are representable in our system as long as they make no use of techniques from analysis. The same may also be said of the frequently used "intuitive" inferences, even if this is not immediately obvious from looking at them.

In order to verify this generally each individual proof would of course have to be examined separately and this would be extremely laborious.

5.32. I shall contend myself with a number of particularly important examples:

Complete induction occurs frequently in certain modified forms which are reducible to our normal form as follows:

5.321. First the "descending" complete induction, which runs as follows:

From $\mathcal{T} \rightarrow \mathcal{F}(1)$ and $\mathcal{F}(a+1), \Delta \rightarrow \mathcal{F}(a)$ follows $\mathcal{T}, \Delta \rightarrow \mathcal{F}(1)$.

Again a must not occur in $\mathcal{T}, \Delta, \mathcal{F}(1)$ and $\mathcal{F}(1)$.

It is transformed thus: $\neg \mathcal{F}(a) \rightarrow \neg \mathcal{F}(a)$ is a basic sequent.

From it follows (5.243) $\mathcal{F}(a+1), \neg \mathcal{F}(a) \rightarrow \neg \mathcal{F}(a)$, and together with $\mathcal{F}(a+1), \Delta \rightarrow \mathcal{F}(a)$ follows by "reductio" (5.252) $\Delta, \neg \mathcal{F}(a) \rightarrow \neg \mathcal{F}(a+1)$.

which equals (5.241) $\neg \mathcal{F}(a), \Delta \rightarrow \neg \mathcal{F}(a+1)$. If we then include

the basic sequent $\neg \mathcal{F}(1) \rightarrow \neg \mathcal{F}(1)$, we can apply the rule of

complete induction in the prescribed form 5.253, with $\neg \mathcal{F}$ as the

induction proposition, and obtain $\neg \mathcal{F}(1), \Delta \rightarrow \neg \mathcal{F}(1)$. By

including $\mathcal{T} \rightarrow \mathcal{F}(1)$, and thus also obtaining (5.243) $\neg \mathcal{F}(1), \mathcal{T} \rightarrow \mathcal{F}(1)$

as valid, we deduce $\mathcal{T}, \Delta \rightarrow \neg \neg \mathcal{F}(1)$ by "reductio",

and from it, by "elimination of the double negation" (5.252):

$$\mathcal{T}, \Delta \rightarrow \mathcal{F}(1)$$

5.322. A further example consists of the following modified complete induction:

From $\mathcal{T} \rightarrow \mathcal{F}(1)$ and $\forall x [x \leq a \supset \mathcal{F}(x)], \Delta \rightarrow \mathcal{F}(a+1)$

follows $\mathcal{T}, \Delta \rightarrow \mathcal{F}(1)$. Again a must not occur in $\mathcal{T}, \Delta \rightarrow \mathcal{F}(1)$

and $\mathcal{F}(1)$; x designates a bound variable not occurring in $\mathcal{F}(1)$.

This induction is easily turned into a normal complete induction (5.253)

with the following induction proposition (stated for an arbitrary number $\underline{m}$): $\forall x [x \leq \underline{m} \supset F(x)]$, in words possibly: "For all numbers from 1 to $\underline{m}$ F holds".

5.323. The corresponding "descending" form runs: From $\mathcal{T} \rightarrow F(\underline{t})$ and $F(n+1), \Delta \rightarrow Fx [x \leq a \wedge F(x)]$ follows $\mathcal{T}, \Delta \rightarrow F(1)$.

This form can be reduced to the normal form of complete induction in the same way as the two previous examples.

The induction in Euclid's Proof was originally of this kind (4.2) and was then reduced to its normal form (4.4).

Paragraph 6

SPECIFIC CONCEPTS AND AXIOMS

IN ELEMENTARY NUMBER THEORY

6.1. In a proof there may also occur "specific concepts" in addition to the actual inferences; these are introductions of new objects, functions, or predicates.

What kinds of specific concepts are in practice used in number theory?

The introduction of new objects such as negative numbers etc. has already been discussed at 3.31, and it was pointed out that these objects are basically dispensable.

The introduction of a new function or a predicate usually takes the form of a verbal "definition" of these concepts.

Examples:

The function $\underline{a}^{\underline{b}}$ is defined as "the number $\underline{a}$, taken $\underline{b}$ times as factor".

The function $\underline{a}!$ is defined as "the product of the numbers from 1 to $\underline{a}$ ".

The number $(\underline{a}, \underline{b})$ is defined as "the greatest common divisor of $\underline{a}$ and $\underline{b}$ ".

The predicate " $\underline{a}$ is a perfect number" means the same as "the number $\underline{a}$ is equal to the sum of its proper divisors".

The predicate $\underline{a} \neq \underline{b}$ means the same as $\neg (\underline{a} = \underline{b})$.

The predicate $\underline{a} \mid \underline{b}$ means the same as $\exists \underline{z} (\underline{a} \cdot \underline{z} = \underline{b})$.

The function $(\frac{\underline{a}}{\underline{b}})$, the "Legendre Symbol" is defined for the case where $\underline{b}$ is an odd prime number as follows: $(\frac{\underline{a}}{\underline{b}}) = 0$ if $\underline{b} \mid \underline{a}$ holds, if $\neg \underline{b} \mid \underline{a}$ holds, then $(\frac{\underline{a}}{\underline{b}}) = 1$, if the number $\underline{a}$ is a quadratic residue mod $\underline{b}$, $(\frac{\underline{a}}{\underline{b}}) = -1$, if $\underline{a}$ is not a quadratic residue mod $\underline{b}$.

The function $ak(\underline{a}, \underline{b}, \underline{c})$, the "Ackermann function", a function significant for certain questions of proof theory, may be defined thus⁽¹¹⁾

("recursively"): $ak(\underline{a}, \underline{b}, 0) = \underline{a} + \underline{b}$
 $ak(\underline{a}, \underline{b}, 1) = \underline{a} \cdot \underline{b}$
 $ak(\underline{a}, \underline{b}, 2) = \underline{a}^{\underline{b}}$
 and further for $\underline{c} \geq 2$:

$ak(\underline{a}, 0, \underline{c} + 1) = \underline{a}$
 $ak(\underline{a}, \underline{b} + 1, \underline{c} + 1) = ak(\underline{a}, ak(\underline{a}, \underline{b}, \underline{c} + 1), \underline{c})$.

I shall not set up general formal schemata for these and other methods of forming concepts. It will turn out that even without such schemata these concepts may be incorporated wholesale in the consistency proof. The same holds for the "axioms" about which I should now like to say a few words.

6.2. In number-theoretical proofs we start from certain simple, immediately obvious propositions for which no further proof is offered. These are the "axioms". They are closely related to the specific concepts in so far as these axioms state basic facts about the predicates and functions occurring in them. Actually, a new concept may be formally introduced by merely stating a number of axioms about it ("implicit definition"). An example: The function $(\underline{a}, \underline{b})$ is completely characterizable by the axioms:

$$\forall \underline{x} \forall \underline{y} [(\underline{x}, \underline{y}) | \underline{x} \neq (\underline{x}, \underline{y}) | \underline{y}] \text{ and } \forall \underline{x} \forall \underline{y} \exists \underline{z} [(\underline{z} | \underline{x} \neq \underline{z} | \underline{y}) \wedge \underline{z} > (\underline{x}, \underline{y})].$$

The choice of the axioms is not determined uniquely. We may aim at making do with as few simple axioms as possible.⁽¹²⁾ Yet in actually working number-theoretical proofs we usually stipulate a larger number of axioms without concern over redundancy, independence, etc. For my consistency proof it is fairly immaterial which axioms are chosen. As in the case of the specific concepts, I shall contend myself with the statement of several examples from which it can be seen what kinds of proposition qualify as axioms:

Some axioms for the predicate $=$ and the function $+$, formalized:

$$\begin{array}{ll} \forall \underline{x} (\underline{x} = \underline{x}) & \forall \underline{x} \neg (\underline{x} + 1 = \underline{x}) \\ \forall \underline{x} \forall \underline{y} (\underline{x} = \underline{y} \supset \underline{y} = \underline{x}) & \forall \underline{x} \forall \underline{y} (\underline{x} + \underline{y} = \underline{y} + \underline{x}) \\ \forall \underline{x} \forall \underline{y} \forall \underline{z} [(\underline{x} = \underline{y} \wedge \underline{y} = \underline{z}) \supset \underline{x} = \underline{z}] & \forall \underline{x} \forall \underline{y} \forall \underline{z} [(\underline{x} + \underline{y}) + \underline{z} = \underline{x} + (\underline{y} + \underline{z})]. \end{array}$$

6.3. The concept of "the....such that".

The following special notion is still worth mentioning:

If a proposition of form

$$\forall x_1, \forall x_2, \dots, \forall x_v, \exists \eta \{ F(x_1, x_2, \dots, x_v, \eta) \wedge \forall \zeta [F(x_1, x_2, \dots, x_v, \zeta) \supset \zeta = \eta] \}$$

in words possibly: "For every combination of numbers $x_1, \dots, x_v$

there exists one and only one number η such that $F(x_1, \dots, x_v, \eta)$

holds", has been proved, then a function may be introduced which

represents precisely this value (η) in its dependence on the combination of numbers ($x_1, \dots, x_v$) ("The... such that").

Formally: For this function one might use the expression (written

for the arguments $a_1, \dots, a_v$): $L_\eta F(a_1, \dots, a_v, \eta)$;

for this expression the following then holds:

$$\forall x_1, \dots, \forall x_v, F(x_1, \dots, x_v, L_\eta F(x_1, \dots, x_v, \eta)) .$$

The x 's may also be empty in which case the L -symbol represents a single number.

Such specific concepts which are not generally needed in practical elementary number theory or which can be replaced by "definitions"

of the kind mentioned above (6.1) are insignificant for the question of consistency since they may always be eliminated from a derivation. (13)

SECTION III
DISPUTABLE AND INDISPUTABLE FORMS OF INFERENCE
(14)
IN ELEMENTARY NUMBER THEORY

The task of the consistency proof will be (2.31) to justify the disputable forms of inference (including specific concepts and axioms) on the basis of indisputable inferences. For a proper understanding of my consistency proof for elementary number theory, which follows in Section IV, we shall therefore have to examine precisely what forms of inference and other techniques of proof from elementary number theory are indeed disputable and which others can be accepted as undoubtedly correct. An unequivocal separation is not possible (cf. 1.8); but we can certainly produce arguments which will make the admissibility of some methods of proof very plausible, whereas a corresponding justification fails for other methods in cases where there exists a remote analogy to the fallacies occurring in the antinomies of set theory which make these techniques suspect.

We shall now develop such arguments by first considering the mathematical theory with a finite domain of objects (Paragraph 7) and by then discussing the peculiarities and difficulties arising from the generalization to an infinite domain of objects (Paragraph 8-11).

Paragraph 7

MATHEMATICS WITH FINITE DOMAINS OF OBJECTS

7.1. The mathematical treatment of a finite domain of objects may be as follows:

The objects of the domain are enumerated; in doing so each object receives a definite designation which applies to no other object.

A function or a predicate is defined thus: Suppose the number of argument places is ν . For every possible enumeration of ν objects from the domain of objects, it is determined whether the object is the associated functional value or, in the case of predicates: whether the predicate does or does not hold for this combination of objects.

We could also permit functions and predicates to remain undefined for some combinations of objects, this represents an unimportant complication.

Since there are always only finitely many enumerations of ν objects, every function and every predicate may be completely defined by such a "definition table".

7.2. For every definite proposition (3.24) which has been constructed in accordance with 3.22, 3.23, from the given objects, functions, and predicates together with the logical connectives, it can furthermore be "calculated" according to the following formal rule whether the proposition is true or false:

The proposition is represented by a formula without free variables.

If it contains the symbol $\forall$ then the term $\forall x F(x)$ concerned is replaced by $[\dots [F(q_1) \& F(q_2)] \& F(q_3) \& \dots] \& F(q_s)]$

where $q_1, \dots, q_s$ represents the entire collection of objects of the domain. The same is done for every $\forall$ that occurs, and each $\exists$ is replaced by a corresponding expression with $\forall$ instead of $\&$.

Then every term that occurs is "calculated" on the basis of the definition tables for the functions occurring in it, i.e., the term is replaced by the object symbol which represents its "value". If several function symbols are nested then the calculation is out step by step by working from the inside to the outside.

Next we determine of every occurring minimal formula (3.24) on the basis of the definition table of the predicate concerned whether the formula represents a true or a false proposition. Then follows the determination of the truth or falsity of those parts of the formula which are made up of arbitrary logical connectives; this is done step by step from the inside according to the following instructions:

$U \& B$ is true, if U and B are both true, otherwise false.

$U \vee B$ is true, if U is true, and also if B is true; it is false only if U and B are both false. $U \supset B$ is false if

U is true and B is false; in every other case $U \supset B$ is true. $\neg U$ is true if U is false, but false if U is true.

The entire procedure follows at once from the actual sense which

we associate with the formal symbols. For us it is important only to realize that in a theory with a finite domain of objects every well-defined proposition is decidable, i.e., that it can be determined by a definite procedure in finitely many steps whether the proposition is true or false.

7.3. It is easily proved that the logical rules of inference (5.2.), applied to this theory, are correct in the sense that their application to "true" mathematical basic sequents leads to "true" derivable sequents. Here the concept of the "truth" of a sequent is to be determined formally in agreement with its intuitive sense as follows: a sequent without free variables is false if all antecedent formulae are true and the succedent formula is false; in every other case it is true. A sequent with free variables is true if every arbitrary replacement of object symbols yields a true sequent.

A verification of this statement would mean no more than a confirmation of the fact that we have indeed chosen our formal rules of inference in such a way that they are in harmony with the intuitive sense of the logical connectives.

7.4. It should still be noted that in practice the above method of introducing objects, functions, and predicates and of "evaluating" the propositions is rarely used in mathematical theories with finite domains of objects; for a large number of objects this would become far too lengthy. In such cases the methods used are rather like those applied in the case of an infinite domain of objects described below.

Paragraph 8

DECIDABLE CONCEPTS AND PROPOSITIONS

IN AN INFINITE DOMAIN OF OBJECTS

8.1. What becomes different if we wish to develop the theory with an infinite domain of objects such as the natural numbers, for example?

8.11. It is then no longer possible to enumerate the objects in order to designate them, since there are infinitely many of them.

The place of an enumeration is taken by a construction rule of the following kind: 1 designates a natural number. Further $1 + 1$, $1 + 1 + 1$, generally: From an expression representing a natural number an expression for a further natural number is obtained by adjoining $+ 1$. (The symbols 2, 3, 4, etc. may be introduced afterwards as abbreviations for $1 + 1$, $1 + 1 + 1$, $1 + 1 + 1 + 1$, etc.; this is of secondary importance.)

This rule which must be expressed in finitely many words generates the infinite number sequence because it contains the possibility of continuing this constructive process through a repetitive procedure. ("Potential infinity".)

8.12. Nor can functions and predicates, as in the case of a finite domain, be defined by an enumeration of all individual values. If we wanted to give a definition table for a number-theoretical function with one argument, for example, we would have to state successively its value for the arguments 1, 2, 3, 4, etc., hence for infinitely many

values. This is impossible. Instead we prescribe a calculation rule; e.g., for the function $2 \cdot a : 2 \cdot 1$ is 2; $2 \cdot (b + 1)$ is equal to $(2 \cdot b) + 2$. This rule makes it possible to calculate the associated functional value uniquely one by one for each natural number.

Generally, a function or a predicate is considered to be decidably defined if a decision procedure is given for it, i.e.: for every given enumeration of natural numbers it must be possible to calculate uniquely the associated functional value by means of this procedure or, in the case of predicates, it must be decidable uniquely whether the predicate concerned holds or does not hold for this collection of numbers.

For all examples of definitions of functions and predicates given at 6.1 such decision procedures can be stated. In the case of specific concepts formed according to 6.3 this may at times no longer be possible. By eliminating these specific concepts we have transferred the doubts associated with them to the logical forms of inference; these will be further discussed below (Paragraphs 9 - 11).

8.2. Let us now consider the propositions in the theory with the infinite domain of objects of the natural numbers.

Of every given definite proposition in which the connectives "all" and "there is" do not occur it can be decided, as in the case of a finite domain, whether it is true or false. The procedure is the same as at 7.2. Instead of being determined by a definition table, the values of the terms of a proposition as well as the truth or falsity of the minimal formulae are now determined by the appropriate

decision rule for the functions or predicates concerned.

The application of the logical rules of inference to propositions of this kind can also be shown to be admissible in the same way as in the case of a finite domain.

It should still be mentioned that a corresponding result also holds for propositions in which the connectives "all" and "there is" refer only to finitely many numbers. Such propositions can be decided in the way described, $\forall$ and $\exists$ must be replaced by $\forall$ and $\exists$ as at 7.2, and the appropriate forms of inference, i.e., the $\forall$ - and $\exists$ -forms of inference (5.251) as well as complete induction (5.253) can also be shown to be admissible in the same way as long as the domain of the - free and bound - variables that occur is limited to the numbers from 1 to a fixed number n .

Paragraph 9

THE "ACTUALIST" INTERPRETATION OF

TRANSFINITE PROPOSITIONS (15)

9.1. Let us now turn to the essentially transfinite propositions, i.e., propositions in which the connectives "all" or "there is" refer to the totality of all natural numbers. Here we are confronted by a fundamentally new state of affairs.

First we must note that the decision rule which is applicable in the case of a finite domain (7.2, 8.2) does not transfer to such transfinite propositions.

In the case of a proposition about all natural numbers, for example, we would have to test infinitely many individual cases, which is impossible. No decision rule for arbitrary transfinite propositions is known and it is doubtful whether such a rule can ever be given. If there were such a rule then it could for example be decided of the thus far unproven last theorem of Fermat (as well as of Goldbach's Conjecture, etc.) by calculation whether it is true or false.

What sense then can be attributed to a proposition whose truth cannot be verified?

9.2. The traditional view is this: it is "actually" pre-determined whether a transfinite proposition such as for example, Fermat's last theorem is "true" or "false" independently of whether we know or shall ever know which of the two is the case. Every transfinite proposition is thought of as having a definite actual sense; in particular, the sense of a $\forall$ -proposition is thought to be this: "For every single one of the infinitely many natural numbers the proposition concerned holds"; the sense of a $\exists$ -proposition: "In the infinite totality of the natural numbers there somewhere exists a number for which the proposition concerned holds".

From this interpretation is inferred further that for transfinite propositions the same logical forms of inference are valid as for the finite case since the "actualist" sense of the logical connectives in transfinite propositions corresponds exactly to that in the finite case.

9.3. At this point there now exists ample cause for criticism as long as one has decided to draw the utmost consequences from the insights gained in considering the antinomies of set theory. This I will now do and shall, as a result of a critical examination of Russell's antinomy (1.6), lay down the following principle:

An infinite totality must not be regarded as actually existing and closed (actual infinity) but only as something becoming which can be extended constructively further and further from something finite (potential infinity).

9.4. The constructive methods for the introduction of objects, functions, and predicates stated in Paragraph 8 are in line with this principle. They were explicitly based on the idea of a gradual progression in the number sequence, starting at the beginning, and not on the idea of a completed totality of all natural numbers. The same holds true for the propositions discussed at 8.2, since they also refer to only finitely many objects and not yet to an infinite totality.

9.5. The "actualist" interpretation of transfinite propositions described at 9.2, however, is no longer compatible with this principle, for it is based on the idea of the closed infinite number sequence.

At the same time the view that the logical forms of inference can simply be transferred from finite to infinite domains of objects must be rejected.

I remind the reader of a similar although more trivial case of an

inadmissible generalization from the finite to the infinite, viz., the well-known fallacy: "Every (finite) set of natural numbers contains a largest number; hence the (infinite) set of all natural numbers contains a largest number." This argument leads to contradictions since it does not in fact hold true.

9.6. Having rejected the actualist interpretation of transfinite propositions we are still left with the possibility of ascribing a "finitist" sense to such propositions, i.e., of interpreting them in each case as expressions for definite finitely characterizable states of affairs.

Once this view has been adopted the relevant logical forms of inference must be examined for their compatibility with this interpretation of the propositions.

Such an examination will be carried out in Paragraph 10 below for an extensive portion of the transfinite propositions and their associated forms of inference. In Paragraph 11, I shall discuss the remaining propositional forms and their forms of inference; there our method will meet with difficulties and the significance of the intuitionist (1.8) delimitation between permissible and non-permissible forms of inference within number theory will become apparent; another still stricter delimitation will also turn out to be defensible.

Paragraph 10

FINITIST INTERPRETATION OF THE CONNECTIVES

 $\forall$, $\exists$, $\neg$ AND $\vee$ IN

TRANSFINITE PROPOSITIONS

I imagine first a number theory whose propositions refer to only finitely many numbers. To it I shall adjoin step by step certain types of transfinite propositions.

10.1. The $\forall$ -connective.

10.11. We shall begin with the simplest form of a transfinite proposition: $\forall x \mathcal{F}(x)$, where $\mathcal{F}$ shall not yet contain a $\forall$ or $\exists$, so that the truth of $\mathcal{F}(x)$ is verifiable for each individual number substituted for x (8.2).

True propositions of this form are for example:

$$\forall x (2 \mid x \vee \neg 2 \mid x) ; \quad \forall x (x = x)$$

Such propositions will undoubtedly be regarded as meaningful and true. After all, one need not associate the idea of a closed infinite number of individual propositions with this $\forall$, since its sense can be given a "finitist" interpretation as follows: "if, starting with 1 , we substitute for x successive natural numbers, then however far we may progress in the formation of numbers, a true proposition results in each case."

10.12. This interpretation may be generalized to the case where $\mathcal{F}$ is an arbitrary proposition to which a finitist sense has already been ascribed: $\forall x \mathcal{F}(x)$ may be asserted meaningfully if $\mathcal{F}(x)$ represents a meaningful and true proposition for arbitrary successive replacements of x by numbers.

10.13. The forms of inference associated with the $\forall$ -connective, the $\forall$ -introduction, and the $\forall$ -elimination (5.251), are in harmony with this interpretation: A $\forall$ is introduced if a proof is available that on the basis of certain assumptions ($\mathcal{T}$) - transfinite assumptions are still completely meaningless and out of the question for the time being - $\mathcal{F}(a)$ is true, and from this is inferred that on the basis of the same assumptions $\forall x \mathcal{F}(x)$ holds. This is in order, for if an arbitrary number n is given then it may be substituted for a - in the whole proof - and a proof for $\mathcal{F}(n)$ results (on the same assumptions $\mathcal{T}$ which, by virtue of the restriction on variables for the $\forall$ -introduction, do not contain a and have thus obviously remained unaffected by this substitution). In the case of the $\forall$ -elimination $\mathcal{T} \rightarrow \mathcal{F}(t)$ is deduced from $\mathcal{T} \rightarrow \forall x \mathcal{F}(x)$. Once possible occurrences of free variables have been replaced by numbers, the term t represents a definite number n ; in keeping with its finitist sense the proposition $\forall x \mathcal{F}(x)$ also guarantees that $\mathcal{F}(n)$ holds; hence this form of inference is also acceptable.

10.14. The usual number-theoretical axioms may be formulated in such a way that they follow from propositions without $\forall$ or $\exists$ by a number of $\forall$ -inferences ranging over the entire proposition (cf. 6.2.).

The conclusion that, in terms of the finitist interpretation of the $\forall$, and on the basis of the decidable definitions of the functions and predicates occurring in them these axioms are true is of such self-evidence that it requires no further investigation.

It actually seems hardly possible that this conclusion could be reduced to something basically simpler.

10.2. The $\&$ -connective.

A transfinite proposition of the form $\mathcal{U} \& \mathcal{B}$ is meaningful and may be asserted if $\mathcal{U}$ and $\mathcal{B}$ have already been recognized as meaningful and valid propositions. The rules for the $\&$ -introduction and $\&$ -elimination are obviously in harmony with this interpretation. Here, as above, transfinite assumptions ($\mathcal{T}$, Δ) are excluded for the time being.

10.3. The $\exists$ -connective.

The reader may so far have the impression that the "finitist interpretation" attributes to transfinite propositions really only the same sense as that usually associated with such propositions. That this is not the case will emerge from the following discussion of the $\exists$ and $\forall$ (cf. 10.6).

What sense shall we allow a proposition of the form $\exists x \mathcal{F}(x)$?

The actualist in interpretation: "somewhere in the infinite number sequence there exists a number with the property $\mathcal{F}$ " is for us devoid of sense. Yet if the proposition $\mathcal{F}(n)$ has been recognized as meaningful and valid for a definite number n , we wish to be

able to conclude ($\exists$ -introduction): $\exists x F(x)$. There are no objections to this; the proposition $\exists x F(x)$ now constitutes only a weakening of the proposition $F(n)$ ("Partialaussage" for Hilbert, "Urteilsabstrakt" for Weyl) in that it now attests merely that we have found a number n with the property F although this number itself is no longer mentioned. With this, $\exists x F(x)$ has a finitist sense.

If, instead of being introduced in $F(n)$, the $\exists$ is introduced in a proposition $F(t)$ containing an arbitrary term t , then nothing has essentially changed. For if the free variables occurring in $F(t)$ are replaced by definite numbers (which free variables, after all, stand for) then t becomes a definite calculable number n on the basis of the decidable definitions of functions. If a

$\exists$ -introduction is accompanied by the occurrence of non-transfinite assumptions (τ) the situation is not essentially altered.

Consider now how other propositions can be inferred by the elimination of the $\exists$ from a proven proposition of the form $\exists x F(x)$ on the basis of the finitist sense of that proposition. In contrast with the situation in the case of $\forall$ and $\&$ it is obviously not possible to reclaim the proposition $F(n)$ from $\exists x F(x)$ which had provided the justification for the assertion of $\exists x F(x)$, precisely because the value of n is no longer apparent from $\exists x F(x)$. Yet we may proceed as follows: we conclude $F(a)$, where a is a free variable taking the place of the number n whose value need not be known at this time. If we then succeed in deducing from $F(a)$ a certain proposition $\mathcal{C}$ no longer containing a ,

then this proposition is valid. We have thus a $\exists$ -elimination in accordance with 5.251.

This is the first rule so far in which an associated assumption, viz., $F(a)$, occurs. This assumption can be transfinite. Although we have previously not granted a sense to transfinite propositions as assumptions but only as proven propositions, we can here say: the fact that $\exists x F(x)$ has been proved and is meaningful means that a number n must have been known and is reconstructible on the basis of the proof of $\exists x F(x)$ so that $F(n)$ also represents a meaningful true proposition. Here the assumption $F(a)$ is not regarded as an arbitrary assumption but as the true proposition $F(n)$, where a merely denotes the number n . The proof of $\mathcal{C}$ from the assumption $F(a)$ thus no longer appears as hypothetical but as an ordinary direct proof; and precisely this is its sense.

10.4. The $\vee$ -connective has an easy analogy to $\exists$, as did $\&$ to $\forall$: a transfinite proposition of the form $\mathcal{U} \vee \mathcal{B}$ is meaningful and may be asserted if one of the propositions $\mathcal{U}$ and $\mathcal{B}$ has been recognized as meaningful and valid. The rule of the $\vee$ -introduction corresponds completely to this interpretation. A $\vee$ -elimination is carried out thus: if $\mathcal{U} \vee \mathcal{B}$ is given and if the same proposition $\mathcal{C}$ follows from the assumption $\mathcal{U}$ as well as from the assumption $\mathcal{B}$, then $\mathcal{C}$ holds. This is in order since $\mathcal{U} \vee \mathcal{B}$ entails that either $\mathcal{U}$ or $\mathcal{B}$ has at some point been recognized as valid. In this way a proof for $\mathcal{C}$ from $\mathcal{U}$, or a proof for $\mathcal{C}$ from $\mathcal{B}$, can be made independent of the assumption $\mathcal{U}$, or $\mathcal{B}$, as was done in the case of the $\exists$ -elimination, and we obtain a direct

proof. The second proof becomes redundant and it is thus immaterial whether it has a sense or not.

10.5. At this point it should be explained briefly how the inference of complete induction is at once compatible with the finitist interpretation: Suppose that $F(1)$ is a meaningful valid proposition. The term t in the conclusion $F(t)$ represents a definite number n once possible occurrences of free variables have been replaced by numbers. By replacing a successively by the numbers 1, 2, 3, up to $n - 1$ in the proof of $F(a+1)$ from $F(a)$ we have formed a direct proof, starting from the valid proposition $F(1)$ via $F(2)$, $F(3)$, etc. up to $F(n)$, so that finally, $F(n)$ is now a valid, meaningful proposition.

This may sound trivial; what is essential is that the assumption $F(a)$ which may have been devoid of sense (if it was transfinite) has been afforded a sense by the possibility of transforming the relevant portion of the proof involved into a direct proof in which no longer functions as an assumption.

10.6. The finitist interpretation given to the connectives $\vee$ and $\exists$ differs from the actualist interpretation not only conceptually but also in its practical consequences as the following examples show:

The proposition "Fermat's Last Theorem is either true or not true" is true according to the actualist interpretation. However, according to the finitist interpretation of $\vee$, this proposition cannot be

asserted. For here it would be required that one of the two propositions has already been recognized as valid. This, however, is so far not the case.

A corresponding example containing a $\exists$ is the proposition

$$\exists x \{ [\forall y \forall z \forall u \forall r (x > 2 \supset y^x + z^x = u^x)] \}$$

$$\vee [\exists y \exists z \exists u (x > 2 \ \& \ y^x + z^x = u^x)] \}$$

in words possibly: "There exists a number x so that either Fermat's Theorem is true or there exists a counter-example with the exponent x ". This proposition is true according to the actualist interpretation but may not be asserted according to the finitist interpretation of the $\exists$ since at this time no such number x is known.

Consequently neither of these two propositions is provable by the forms of inference discussed so far since it was possible to attribute a finitist sense to these forms of inference; the additional forms of inference associated with the $\neg$ are needed for this purpose (cf. 11.2).

10.7. The finitist interpretation of transfinite propositions containing the connectives $\vee$, $\&$, $\exists$ and $\forall$ attempted in these paragraphs and the justification of the associated forms of inference involved is in many respects incomplete; the meaning of propositions in which a number of such connectives occur in nested form, in particular, would still have to be discussed in greater detail. I shall not do this since I am here concerned only with establishing fundamentals.

A purely formal consistency proof for this part of number theory could be developed later on the basis of these considerations. Yet such a proof would be of little value since it itself would have to make use of transfinite propositions and the same associated forms of inference which it is intended to "justify". Such a proof would therefore not represent an appeal to more elementary facts, although it would still confirm the finitist character of the formalized rules of inference. Yet we would have to have a clear idea beforehand of what is to be considered finitist (in order to be able to carry out the consistency proof proper with finitist methods of proof).

Paragraph 11

THE CONNECTIVES $\supset$ AND $\neg$ IN

TRANSFINITE PROPOSITIONS : THE INTUITIONIST VIEW

11.1. The $\supset$ -connective.

We now intend to include transfinite propositions containing the connective $\supset$.

What does $\mathcal{U} \supset \mathcal{B}$ mean? Suppose, for example, that there exists a proof in which the proposition $\mathcal{B}$ is proved on the basis of the assumption $\mathcal{U}$ by means of inferences that have already been recognized as permissible. From this we infer by $\supset$ -introduction:

$\mathcal{U} \supset \mathcal{B}$. This proposition is merely intended to express the fact that a proof is available which permits a proof of the proposition $\mathcal{B}$ from the proposition $\mathcal{U}$ once the proposition $\mathcal{U}$ is proven.

The inference of the $\supset$ -elimination is in harmony with this

interpretation: here $\mathcal{B}$ is inferred from $\mathcal{A}$ and $\mathcal{A} \supset \mathcal{B}$;
 this is in order, since $\mathcal{A} \supset \mathcal{B}$ indicates precisely the existence
 of a proof for $\mathcal{B}$ in the case where $\mathcal{A}$ is already proven.

In interpreting $\mathcal{A} \supset \mathcal{B}$ in this way I have presupposed that the
 available proof of $\mathcal{B}$ from the assumption $\mathcal{A}$ contains merely
inferences already recognized as permissible. Yet such a proof
 could itself contain other $\supset$ -inferences and then our interpretation
breaks down. For it would be circular to justify the $\supset$ -inferences
 on the basis of a $\supset$ -interpretation which itself already involves
 the presupposition of the admissibility of the same form of inference.
 The $\supset$ -inferences which occur in the proof would in that case have
 to be justified beforehand; yet this has its difficulties especially
 if the assumption $\mathcal{A}$ has itself the form $\mathcal{C} \supset \mathcal{D}$; if this
 happens we have actually no proof for $\mathcal{D}$ from $\mathcal{C}$ on the basis of
 which we could ascribe a sense to $\mathcal{C} \supset \mathcal{D}$.

In order to cope with this difficulty a more complicated interpretation
 rule would certainly have to be formulated. This constitutes one of
 the principal objectives of the consistency proof which follows in
 Section IV.

11.2. The $\neg$ -connective presents even greater obstacles to a
 finitist interpretation than the $\supset$. Transfinite propositions were
 actually always interpreted in such a way that they could in each
 case be regarded as something that had previously been recognized as
 valid. In its actualist interpretation $\neg \mathcal{A}$ does not however
 express the fact that something holds but rather purely negatively

that something, viz., the proposition $\mathcal{U}$, does not hold.

The following positive interpretation seems nonetheless possible:

$\neg \mathcal{U}$ is to be regarded as meaningful and true if a proof exists to the effect that a falsehood is certain to follow from the assumption of the validity of $\mathcal{U}$. And here the $\neg$ -connective is re-interpreted in terms of the $\supset$ -connective since $\neg \mathcal{U}$ can certainly be defined as equivalent with $\mathcal{U} \supset 1=2$, for example. The inference of "reductio" is in harmony with this interpretation, as may be shown quite formally: From $\mathcal{U}, \neg \rightarrow \mathcal{B}$ and $\mathcal{U}, \Delta \rightarrow \mathcal{B} \supset 1=2$ we wish to derive $\neg, \Delta \rightarrow \mathcal{U} \supset 1=2$. This is done as follows: By $\supset$ -elimination we obtain $\mathcal{U}, \neg, \mathcal{U}, \Delta \rightarrow 1=2$ hence (5.242) $\mathcal{U}, \neg, \Delta \rightarrow 1=2$, and from this by $\supset$ -introduction, $\neg, \Delta \rightarrow \mathcal{U} \supset 1=2$. This completes the reduction of the "reductio" to the $\supset$ -forms of inference.

It should be noted that in this re-interpretation of the $\neg$ in terms of the $\supset$ all doubts associated with the $\supset$ naturally carry over to the $\neg$ -connective to a corresponding degree.

Now there actually arises a further difficulty: The "elimination of the double negation" cannot at all be shown to agree with the given $\neg$ -interpretation. There is no compelling reason why the validity of $(\mathcal{U} \supset 1=2) \supset 1=2$ should follow from the validity of $\mathcal{U}$.

This form of inference conflicts in fact quite categorically with the remaining forms of inference. In the case of the logical

connectives $\forall$, $\&$, $\exists$, $\vee$ and $\supset$ we had in each case an introduction and an elimination inference corresponding to one another in a certain way. (cf. the discussion in Paragraphs 10 and 11.1). In the case of the $\neg$ -connective the "reductio" inference can be regarded both as an introduction (of $\neg$ in $\neg\mathcal{U}$) and an elimination (of $\neg$ in $\neg\mathcal{B}$); the "elimination of the double negation", however, represents an additional $\neg$ -elimination which does not correspond to the $\neg$ -introduction by "reductio". Double negation renders possible indirect proofs of positive propositions ($\mathcal{U}$) from their contraries by means of contradiction, in cases where a positive proof of the same proposition may be completely inaccessible. In this way we can for example prove the two propositions containing $\forall$ and $\exists$ given as examples at 10.6, whereas under the finitist interpretation of $\forall$ and $\exists$ these propositions may not even be asserted.

From this it follows that there is no way at all of including the inference of the "elimination of the double negation" in a finitist interpretation of the kind chosen for $\forall$ and $\exists$.

11.3. Here the intuitionists draw the line in number theory by disallowing the inference of the "elimination of the double negation" for transfinite propositions $\mathcal{U}$. This delimitation is often also effected by disallowing the "law of the excluded middle", $\mathcal{U} \vee \neg\mathcal{U}$, for transfinite $\mathcal{U}$; this comes to the same thing.⁽¹⁶⁾

The "finitist interpretation" of the ^{logical}~~propositional~~ connectives $\forall$, $\&$, $\exists$ and $\vee$ in transfinite propositions described in

Paragraph 10 agrees essentially with the interpretation of the intuitionists. Yet they allow a more general use of the $\supset$ -connective; the $\neg$ -connective is interpreted as at 11.2 by reducing it to $\supset$, and to this corresponds the expression " $\mathcal{U}$ is absurd" in place of " $\mathcal{U}$ does not hold" for $\neg \mathcal{U}$.

The "elimination of the double negation" undoubtedly stands in definite contrast to the remaining forms of inference to such a degree that it might quite reasonably be disallowed. In fact, I consider a still more radical critique, especially of the general use of the $\supset$ (11.1), as equally well justified.

A THEOREM BY GÖDEL ABOUT THE EQUIVALENCE OF INTUITIONIST AND THE WHOLE OF ELEMENTARY NUMBER THEORY

As was first proven by K. Gödel ⁽¹⁷⁾, it is possible to eliminate the inference of the "elimination of the double negation" involving a transfinite $\mathcal{U}$ from any given elementary number-theoretical proof by a special interpretation of transfinite propositions so that every proof of this kind becomes intuitionistically acceptable.

In this way, the whole of actualist number theory becomes reduced to intuitionist number theory. In particular, the former is consistent if the latter is.

The interpretation involved takes the following form: the logical connectives $\&$, $\vee$, $\supset$ and $\neg$ are assigned their intuitionist sense. Not so in the case of $\forall$ and $\exists$; $\mathcal{U} \vee \mathcal{B}$ is interpreted as $\neg [(\neg \mathcal{U}) \& \neg \mathcal{B}]$, $\exists x \mathcal{F}(x)$

as $\neg \forall x \neg F(x)$. The reason for this interpretation is that the $\forall$ and $\exists$ cannot here be assigned their intuitionist meaning since the examples of propositions stated at 10.6 are provable in actualist number theory but not in intuitionist number theory. Yet if $\forall$ and $\exists$ are replaced in these examples by $\forall$, $\exists$ and $\neg$ in the way described then propositions result which are also intuitionistically provable.

In my consistency proof the "elimination of the double negation" actually presents no essential difficulties (13.93).

11.4. The forms of inference which we have not been able to justify so far by means of a finitist interpretation and which are therefore disputable for the time being occur very rarely in proofs carried out in practical number theory. It follows from our discussion that such inferences are principally the "elimination of the double negation" (and the "law of the excluded middle") applied to transfinite propositions as well as the use of transfinite propositions containing nested $\supset$ -and $\neg$ -connectives.

Transfinite propositions of a more complicated structure hardly ever occur in practice. In Euclid's proof presented in Paragraph 4, for example, the only essentially transfinite propositions are the two propositions occurring at the end: $\exists z (\text{Prime } z \wedge z > a)$ and $\forall y \exists z [\text{Prime } z \wedge z > y]$. The whole proof is entirely finitist. The other transfinite propositions which occur in it, i.e., those containing $\forall$ or $\exists$, are such that their bound variables range only over a finite segment of the number sequence.

As an example of a more difficult proof I have looked through Rev. Zeller's proof of the "law of quadratic reciprocity"⁽¹⁸⁾ and here I have also been unable to find a "disputable inference".

We are indeed justified in having the impression of an unquestionable correctness in the case of this and similar proofs. In these proofs we tend automatically to look more for a finitist than an actualist interpretation of the transfinite propositions.

The task of the consistency proof for elementary number theory is thus more a justification of theoretically possible rather than actually occurring inferences.

SECTION IV

THE CONSISTENCY PROOF

I shall now prove the consistency of elementary number theory as a whole as formalized in Section II.

In carrying out this consistency proof we must make certain, as was pointed out in 2.31, that the inferences and specific concepts used in the proof itself are indisputable or at least considerably more reliable than the doubtful forms of inference of elementary number theory. It follows from our discussion in Section III that this requirement can be regarded as met if the methods of proof used are "finitist" (in the sense of Paragraphs 9 - 11). The extent of our success in this direction will be examined more closely in Section V (16.1).

Paragraphs 13 to 15 contain the core of the consistency proof whereas Paragraph 12 is concerned with some relatively simple preliminaries.

Paragraph 12

THE ELIMINATION OF THE SYMBOLS $\forall$, $\exists$,

AND $\supset$ FROM A GIVEN DERIVATION.

Take any number-theoretical derivation (5.22) as given. It is to be shown that it is consistent, i.e., that its end-formula cannot have the form $\mathcal{U} \wedge \neg \mathcal{U}$.

We begin by stating a rule for a transformation of the given derivation. As a result of this transformation the connectives $\vee$, $\exists$ and $\supset$ will no longer occur in the derivation.

12.1. In actualist logic, which is what we are in effect dealing with in unrestricted number theory, the different logical connectives can be represented by other connectives in various ways. By means of three connectives, viz., $\neg$, any one of the three connectives $\&$, $\vee$ and $\supset$ as well as any one of the two connectives $\forall$ and $\exists$, all others may be expressed. I shall make use of this fact to facilitate the consistency proof and shall retain the symbols $\&$, $\forall$ and $\neg$ and express $\vee$, $\exists$ and $\supset$ in terms of these.

This does not mean that the ambiguities (11.1) associated with the are thus conjured away, they stay with us in an equivalent form in the $\neg$.

The replacement takes the form:

For $\mathcal{U} \vee \mathcal{B}$ we put $\neg ((\neg \mathcal{U}) \& \neg \mathcal{B})$
 For $\mathcal{U} \supset \mathcal{B}$ we put $\neg (\mathcal{U} \& \neg \mathcal{B})$
 For $\exists x \mathcal{F}(x)$ we put $\neg \forall x \neg \mathcal{F}(x)$

All $\vee$ -, $\exists$ - and $\supset$ -symbols occurring in the derivation are replaced in this way. The order in which this is done is obviously immaterial.

12.2. We must now examine to what extent the given derivation has

remained correct after these replacements and, where this is not the case, modify the derivation accordingly. That such a modification is possible is very plausible since the new formulations for the $\vee$, $\exists$ and $\supset$ are indeed equivalent to the original ones in the actualist interpretation. The precise formal verification is consequently not difficult:

Logical basis sequents (5.23) have been turned into other logical basic sequents.

The same holds true for mathematical basic sequents as long as we presuppose that a mathematical axiom in which the $\vee$ -, $\exists$ - and

$\supset$ -connectives occur becomes another mathematical axiom after the replacement of these connectives by $\neg$, $\&$ and $\vee$. This requirement is easily met: we simply formulate all axioms in advance without the use of $\vee$, $\exists$ and $\supset$.

Structural transformations (5.24) and application instances of the rules of inference (5.25) have obviously remained correct as long as we are not dealing with one of the rules associated with the connectives

$\vee$, $\exists$ and $\supset$. The latter rules must be replaced by applying other rules of inference in accordance with the following instructions:

A $\vee$ -introduction "from $T \rightarrow \mathcal{U}$ follows $T \rightarrow \mathcal{U} \vee \mathcal{B}$ " after the replacement takes the form: "From $T^* \rightarrow \mathcal{U}^*$ follows $T^* \rightarrow \neg((\neg \mathcal{U}^*) \& \neg \mathcal{B}^*)$ ". $\mathcal{U}^*$ designates the formula

which has resulted from $\mathcal{U}$ by replacement; $\mathcal{B}^*$ and $\mathcal{T}^*$ are to be understood in the same way.

In words, the new version which uses the forms of inference for $\&$ and $\neg$ reads as follows: $\mathcal{U}^*$ holds on the assumptions $\mathcal{T}^*$. If $(\neg \mathcal{U}^*) \& \neg \mathcal{B}^*$ were to hold, then so would $\neg \mathcal{U}^*$ in particular, and this cannot be the case since it contradicts $\mathcal{U}^*$, i.e., $\neg((\neg \mathcal{U}^*) \& \neg \mathcal{B}^*)$ holds on the assumptions $\mathcal{T}^*$.

To this corresponds the following formal instruction: the appropriate place in the derivation is to be transformed thus: $(\neg \mathcal{U}^*) \& \neg \mathcal{B}^* \longrightarrow (\neg \mathcal{U}^*) \& \neg \mathcal{B}^*$ is a basic sequent; by $\&$ -elimination we obtain $(\neg \mathcal{U}^*) \& \neg \mathcal{B}^* \longrightarrow \neg \mathcal{U}^*$, this together with the sequent $(\neg \mathcal{U}^*) \& \neg \mathcal{B}^*, \mathcal{T}^* \longrightarrow \mathcal{U}^*$, obtained from $\mathcal{T}^* \longrightarrow \mathcal{U}^*$ by means of 5.243 by "reductio", yields $\mathcal{T}^* \longrightarrow \neg((\neg \mathcal{U}^*) \& \neg \mathcal{B}^*)$.

The other form of the $\vee$ -introduction is dealt with in the same way.

A $\vee$ -elimination has the following form after the replacement: "From $\mathcal{T}^* \longrightarrow \neg((\neg \mathcal{U}^*) \& \neg \mathcal{B}^*)$ and $\mathcal{U}^*, \Delta^* \longrightarrow \mathcal{C}^*$ and $\mathcal{B}^*, \Theta^* \longrightarrow \mathcal{C}^*$ follows $\mathcal{T}^*, \Delta^*, \Theta^* \longrightarrow \mathcal{C}^*$ ". This is transformed thus: $\neg \mathcal{C}^* \longrightarrow \neg \mathcal{C}^*$ yields $\mathcal{U}^*, \neg \mathcal{C}^* \longrightarrow \neg \mathcal{C}^*$, this together with $\mathcal{U}^*, \Delta^* \longrightarrow \mathcal{C}^*$ by "reductio" $\Delta^*, \neg \mathcal{C}^* \longrightarrow \neg \mathcal{U}^*$; similarly $\mathcal{B}^*, \neg \mathcal{C}^* \longrightarrow \neg \mathcal{C}^*$ together with $\mathcal{B}^*, \Theta^* \longrightarrow \mathcal{C}^*$ yields the sequent $\Theta^*, \neg \mathcal{C}^* \longrightarrow \neg \mathcal{B}^*$; taking both results together we obtain $\Delta^*, \neg \mathcal{C}^*, \Theta^*, \neg \mathcal{C}^* \longrightarrow (\neg \mathcal{U}^*) \& \neg \mathcal{B}^*$ by $\&$ -introduction,

hence (5.242, 5.241) $\neg \mathcal{C}^*, \Delta^*, \mathcal{Q}^* \rightarrow (\neg \mathcal{U}^*) \& \neg \mathcal{B}^*$; from $\mathcal{T}^* \rightarrow \neg ((\neg \mathcal{U}^*) \& \neg \mathcal{B}^*)$ follows $\neg \mathcal{C}^*, \mathcal{T}^* \rightarrow \neg ((\neg \mathcal{U}^*) \& \neg \mathcal{B}^*)$, thus, by "reductio", we obtain $\Delta^*, \mathcal{Q}^*, \mathcal{T}^* \rightarrow \neg \neg \mathcal{C}^*$, and, finally, by "elimination of the double negation" $\Delta^*, \mathcal{Q}^*, \mathcal{T}^* \rightarrow \mathcal{C}^*$, hence (5.241) $\mathcal{T}^*, \Delta^*, \mathcal{Q}^* \rightarrow \mathcal{C}^*$.

A $\mathcal{J}$ -introduction or $\mathcal{J}$ -elimination is dealt with analogously to the $\mathcal{V}$ -introduction or $\mathcal{V}$ -elimination; a $\mathcal{V}$ -elimination takes the place of a $\&$ -elimination or a $\mathcal{V}$ -introduction the place of a $\&$ -introduction in the appropriate place of the derivation. The details are straightforward.

A $\supset$ -introduction after the replacement takes the form: "From $\mathcal{U}^*, \mathcal{T}^* \rightarrow \mathcal{B}^*$ follows $\mathcal{T}^* \rightarrow \neg (\mathcal{U}^* \& \neg \mathcal{B}^*)$ ".

This is transformed thus: $\mathcal{U}^* \& \mathcal{B}^* \rightarrow \mathcal{U}^* \& \mathcal{B}^*$ yields $\mathcal{U}^* \& \neg \mathcal{B}^* \rightarrow \mathcal{U}^*$ as well as $\mathcal{U}^* \& \neg \mathcal{B}^* \rightarrow \neg \mathcal{B}^*$, hence also $\mathcal{U}^*, \mathcal{U}^* \& \neg \mathcal{B}^* \rightarrow \neg \mathcal{B}^*$; this together with $\mathcal{U}^*, \mathcal{T}^* \rightarrow \mathcal{B}^*$ yields $\mathcal{T}^*, \mathcal{U}^* \& \neg \mathcal{B}^* \rightarrow \neg \mathcal{U}^*$, hence $\mathcal{U}^* \& \neg \mathcal{B}^*, \mathcal{T}^* \rightarrow \neg \mathcal{U}^*$. By including $\mathcal{U}^* \& \neg \mathcal{B}^* \rightarrow \mathcal{U}^*$ we obtain $\mathcal{T}^* \rightarrow \neg (\mathcal{U}^* \& \neg \mathcal{B}^*)$.

A $\supset$ -elimination after the replacement takes the form: "From $\mathcal{T}^* \rightarrow \mathcal{U}^*$ and $\Delta^* \rightarrow \neg (\mathcal{U}^* \& \neg \mathcal{B}^*)$ follows $\mathcal{T}^*, \Delta^* \rightarrow \mathcal{B}^*$ ". This is transformed thus: $\mathcal{T}^* \rightarrow \mathcal{U}^*$ and $\neg \mathcal{B}^* \rightarrow \neg \mathcal{B}^*$ yield $\mathcal{T}^*, \neg \mathcal{B}^* \rightarrow \mathcal{U}^* \& \neg \mathcal{B}^*$, hence $\neg \mathcal{B}^*, \mathcal{T}^* \rightarrow \mathcal{U}^* \& \neg \mathcal{B}^*$; by including

$\neg B^*, \Delta^* \longrightarrow \neg (A^* \& \neg B^*)$ we obtain $\neg^* \Delta^* \longrightarrow \neg^* \neg B^*$,
and from this $\neg^* \Delta^* \longrightarrow B^*$.

12.3. We have thus succeeded in transforming the given derivation into a derivation in which the symbols $\vee$, $\&$ and $\supset$ no longer occur. It should be observed that the end-formula of the derivation has undergone a change only if it contained a $\vee$, $\&$ or $\supset$.

12.4. It is worth noting that according to what was said at 11.3 the given derivation is now already essentially an intuitionistically admissible number-theoretical derivation; for wherever the "elimination of the double negation" is still used it could be replaced by other rules of inference.

Paragraph 13

THE REDUCTION OF SEQUENTS

The notion of the "statability of a reduction rule" for a sequent, to be defined below, will serve as a formal replacement of the intuitive concept of truth; it provides us with a special finitist interpretation of propositions and takes the place of their actualist interpretation (cf. Paragraphs 9 - 11).

In a sequent in which the connectives $\vee$, $\&$ and $\supset$ no longer occur, an individual reduction step can be carried out in the following way (13.11 to 13.53):

13.11. Suppose that the sequent contains at least one free variable. In that case we replace every occurrence of this free variable by one and the same arbitrarily chosen numeral.

13.12. Suppose that the sequent contains no free variables and that somewhere in one of its formulae a minimal term (3.24) occurs (e.g., as part of a longer term). In that case we replace minimal term by its associated "functional value", i.e., by that numeral which represents its value for the given numbers as arguments by virtue of the definition of the function concerned (cf. 8.12).

Thus I am now requiring of the functions that they are decidably defined in the sense of 8.12.

13.21. Suppose that the sequent contains no free variables and no minimal terms and that its succedent formula (5.21) has the form

$$\forall x F(x) \quad . \quad \text{In that case we replace it by a formula} \\ F(n) \quad , \text{ i.e., by a formula which results from } F(x)$$

by the substitution of an arbitrarily chosen numeral n for the variable x .

13.22. Suppose that the sequent contains no free variables and no minimal terms and that its succedent formula has the form $A \& B$. In that case we replace it by the formula A or by the formula B , as we please.

13.23. Suppose that the sequent contains no free variables and no

minimal terms and that its succedent formula has the form $\neg \mathcal{U}$.

In that case we replace it by the formula $1 = 2$,⁽¹⁹⁾ and at the same time adjoin the formula $\mathcal{U}$ (in the last place) to the antecedent formulae of the sequent (cf. 11.2).

13.3. If none of the possibilities listed above applies then the succedent formula of the sequent must be a minimal formula (3.24).

I am now requiring of predicates, as was done for functions above, that they are decidably defined in the sense of 8.12.

We can consequently decide of a given minimal formula on the basis of the definition of the predicate concerned whether it represents a true or false proposition.

13.4. Suppose that the sequent contains no free variables and no minimal terms and that its succedent formula is a true minimal formula; or: that the succedent formula is a false minimal formula (e.g., $1 = 2$) and that one of its antecedent formulae is also a false minimal formula.

For such an obviously true sequent (cf. 7.3.) no reduction step is defined.

13.5. Suppose that the sequent contains no free variables and no minimal terms; that its succedent formula is a false minimal formula; and that none of its antecedent formulae are false minimal formulae. In that case the following three different kinds of reduction step are permissible (counterpart to 13.2):

13.51. Suppose that an antecedent formula has the form $\forall x F(x)$. To it we adjoin an antecedent formula $F(n)$, i.e., a formula which results from $F(x)$ by the substitution of a numeral n for the variable x . In doing so we may either retain or omit the formula $\forall x F(x)$.

13.52. Suppose that an antecedent formula has the form $A \wedge B$. In that case we adjoin to it either the formula A or the formula B . In doing so we may omit or also retain the formula $A \wedge B$.

13.53. Suppose that an antecedent formula has the form $\neg A$. We replace it by the succedent formula A . In doing so we may either omit or retain the formula $\neg A$.

13.6. A reduction rule for a sequent in which the connectives $\vee$, $\exists$ and $\supset$ do not occur is a rule which renders possible in each case the "reduction" of a sequent in finitely many individual reduction steps (in accordance with 13.11 to 13.53) to one of the correct definitive forms (13.4) regardless of how we may choose the numeral n involved, or which of the two formulae A and B (in the case of 13.22) we may choose when carrying out a reduction step in which there exists a "choice", i.e., one of the steps described at 13.11, 13.21 and 13.22.

13.7. If several possibilities arise in any other reduction step (e.g., in the case of 13.5) no choice exists since we shall require the reduction rule to be such that it determines what kind of reduction

step is to take place. Also e.g., what numeral n is to be used when adjoining an antecedent formula $F(n)$ and whether the affected formula $\forall x F(x)$ is to be omitted or not.

13.8. Illustration of the Reduction Concept.

13.81. The Reduction of True Sequents containing no Variables.

In order to illustrate the reduction concept I shall begin by showing that for sequents without variables and without the symbols $\forall$, $\exists$ and $\supset$, the concept of the statability of a reduction rule coincides with the concept of truth in the sense of a calculation procedure (7.2, 7.3):

Such a "true" sequent is to be reduced to its definitive form according to the following rule: First, all terms that may occur are to be replaced by their "numerical values" (13.12). If the definitive form (13.4) has not yet been reached, a reduction step is to be carried out by which the sequent is transformed into another "true" sequent in which fewer logical connectives occur than before. This is always possible. For reductions according to 13.22 and 13.23 certainly fulfil this requirement. In the case of 13.5 the following reduction step among the different ones possible is to be applied:

If a false antecedent formula of the form $u \& v$ occurs, then either u or v must be false; in that case the formula $u \& v$ is replaced by u or v . If a false antecedent formula of the form

$\neg \mathcal{U}$ occurs it is omitted and the succedent formula is replaced by $\mathcal{U}$.

Each one of the given reduction steps obviously leads to another true sequent, in particular to one with fewer logical connectives than before. The continuation of this procedure obviously leads to the definitive form of the sequent in finitely many steps.

That, conversely, every sequent without variables for which a reduction rule is available is true follows from the fact that a false sequent, as is easily verified, would be transformed into another false sequent by every permissible reduction step, or that in the case of a reduction step according to 13.22, the choice of $\mathcal{U}$ or $\mathcal{B}$ could be made in such a way that this is the case.

13.82. These considerations can be extended without difficulty to the case of sequents containing $\forall$ -symbols ranging over only finitely many numbers. The reduction of the $\forall$ here corresponds to that of the $\&$.

13.82. If we proceed to the infinite domain of objects of all natural numbers the statement of a reduction rule for an arbitrary derivable sequent is in general no longer as simple. Since it is here no longer true that all formulae are decidable we may, for example, be forced at times to make use of the permission to retain the transformed antecedent formula in reduction steps according to 13.51 to 13.53, whereas this formula could always be omitted in the case of a finite domain (13.81, 13.82).

As an example I shall give a reduction rule for the proposition mentioned at 10.6: "Fermat's last theorem is either true or not true" which, according to its finitist interpretation at that point, is not a true proposition: after the replacement of the and written as a sequent, this proposition has the form:

$$\rightarrow \neg \{ [\neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u)] \\ \& [\neg \neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u)] \}$$

This is reduced as follows: First we obtain (13.23):

$$[\neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u)] \\ \& [\neg \neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u)] \rightarrow 1 = 2$$

By two reductions according to 13.52 we obtain

$$\neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u), \\ \neg \neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u) \rightarrow 1 = 2;$$

Further (13.53):

$$\neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u) \\ \rightarrow \neg \neg \forall x \forall y \forall z \forall u \neg (u > 2 \ \& \ x^u + y^u = z^u)$$

The reduction of this logical basic sequent must now be completed along the lines described in general at 13.92.

13.90. In the following I shall prove that reduction rules can be

given for all sequents occurring in an arbitrarily given derivation once the derivation has been transformed according to Paragraph 12.

From this the consistency will then follow at once:

For if a sequent of the form $\rightarrow \mathcal{U} \& \neg \mathcal{U}$ were derivable, then $\rightarrow 1=2$, for example, would also be derivable. This is so since $\rightarrow \mathcal{U}$ as well as $\rightarrow \neg \mathcal{U}$ follow from $\rightarrow \mathcal{U} \& \neg \mathcal{U}$ by $\&$ -elimination, hence also (5.243) $\neg 1=2 \rightarrow \mathcal{U}$ and $\neg 1=2 \rightarrow \neg \mathcal{U}$; by "reductio" we obtain $\rightarrow \neg \neg 1=2$, and by "elimination of the double negation" $\rightarrow 1=2$. (In the same way any arbitrary proposition can be derived from a contradiction.) Yet no reduction rule can be stated for the sequent $\rightarrow 1=2$, since there is no reduction step that might possibly be applied to it, nor is it in definitive form (13.4) since $1=2$ is false.

13.91. Of the mathematical basic sequents I am requiring that reduction rules have been given for them and that these rules do not make use of the permission, which exists for reduction steps carried out according to 13.5, to retain the transformed antecedent formula.

For all customary number-theoretical axioms such rules are easily stated. Let us look at the examples mentioned at 6.2 in particular; these must first be written as sequents and the $\supset$ replaced by $\&$ and $\neg$; the resulting sequent can then be reduced by first eliminating the $\forall$ -symbols according to 13.1 and by replacing

their associated variables by arbitrary numerals and by then proceeding as described at 13.81. After all, the formulae which result are indeed "true".

13.92. Logical basic sequents are to be reduced according to the following simple rule:

Suppose that a sequent of the form $\mathcal{U} \rightarrow \mathcal{U}$ is given. We first replace the free variables by arbitrary numerals (13.11), then the minimal terms by those numerals that represent their values (13.12). The latter procedure must be repeated until no further minimal terms occur - for it can certainly happen that new minimal terms arise during the computation. The sequent finally has the form $\mathcal{U}^* \rightarrow \mathcal{U}^*$.

The succedent formula $\mathcal{U}^*$ is then reduced by means of reduction steps according to 13.21, 13.22 and, if necessary, 13.12 until it has the form $\neg \mathcal{C}$ or is a minimal formula. In the case of reductions according to 13.21 or 13.22 the replacement numerals or formulae may be chosen arbitrarily.

If the succedent formula has now become a true minimal formula then the reduction procedure is already at an end (13.4).

If the succedent formula has become a false minimal formula then further reduction steps must be carried out according to 13.51, 13.52

and 13.12 in such a way that the antecedent formula $\mathcal{U}^*$ undergoes precisely the same transformations, in the same order, as the succeedent formula $\mathcal{U}^*$ did earlier. If the antecedent formula has taken on the form $\forall x F(x)$, for example, it must be replaced by a formula $F(n)$ and for the replacement numeral n the same numeral must be taken that was chosen in the corresponding reduction of the succeedent formula. Reduction steps according to 13.52 are dealt with correspondingly. The antecedent formula thus eventually becomes equal to the succeedent formula and the procedure is once again at an end since the definitive form (13.4) has been reached.

If the succeedent formula has taken on the form $\neg \mathcal{C}$, a reduction according to 13.23 must first be carried out. The sequent then runs: $\mathcal{U}^*, \mathcal{C} \rightarrow 1=2$. As in the previous case, this sequent is reduced, in such a way that the antecedent formula $\mathcal{U}^*$ is transformed in the same way as was the succeedent formula $\mathcal{U}^*$, so that finally $\neg \mathcal{C}$ appears in its place. Then the sequent runs $\neg \mathcal{C}, \mathcal{C} \rightarrow 1=2$. By means of 13.53 it is reduced to $\mathcal{C} \rightarrow \mathcal{C}$. This is another logical basic sequent; the formula $\mathcal{C}$ contains at least one logical connective less than $\mathcal{U}^*$, and this procedure will consequently end after finitely many steps. A reduction rule has thus been given for arbitrary logical basic sequents.

13.93. In a similar way arbitrary sequents of the form $\mathcal{U} \& \mathcal{B} \rightarrow \mathcal{U}$ or $\mathcal{U} \& \mathcal{B} \rightarrow \mathcal{B}$ or $\mathcal{U}, \mathcal{B} \rightarrow \mathcal{U} \& \mathcal{B}$ or $\forall x F(x) \rightarrow F(t)$ or $\mathcal{U}, \neg \mathcal{U} \rightarrow 1=2$ or $\neg \neg \mathcal{U} \rightarrow \mathcal{U}$ may be reduced, a fact which will be used later.

Here, too, the free variables and minimal terms are first replaced according to 13.11 and 13.12. The sequent $\mathcal{U}^* \& \mathcal{B}^* \rightarrow \mathcal{U}^*$ then has a form which also occurred in the reduction of the logical basic sequent $\mathcal{U}^* \& \mathcal{B}^* \rightarrow \mathcal{U}^* \& \mathcal{B}^*$ according to 13.92; hence the reduction of the sequent in question can be completed in the same way as that of the latter. The same holds true for $\mathcal{U}^* \& \mathcal{B}^* \rightarrow \mathcal{B}^*$ and, correspondingly, for $(\forall x F(x))^* \rightarrow (F(t))^*$; here the basic sequent $(\forall x F(x))^* \rightarrow (\forall x F(x))^*$ must be used. In the case of $\mathcal{U}^*, \mathcal{B}^* \rightarrow \mathcal{U}^* \& \mathcal{B}^*$

a reduction step according to 13.22 must be carried out; from it either $\mathcal{U}^*, \mathcal{B}^* \rightarrow \mathcal{U}^*$ or $\mathcal{U}^*, \mathcal{B}^* \rightarrow \mathcal{B}^*$ follows, whichever we wish. The reduction is then continued in exactly the same way as that of the basic sequent $\mathcal{U}^* \rightarrow \mathcal{U}^*$ or $\mathcal{B}^* \rightarrow \mathcal{B}^*$; the additional antecedent formula is disregarded and presents no problem. In the case of $\mathcal{U}^*, \neg \mathcal{U}^* \rightarrow 1=2$ a reduction step according to 13.53 yields $\mathcal{U}^* \rightarrow \mathcal{U}^*$, hence another basic sequent.

In the case of $\neg \neg \mathcal{U}^* \rightarrow \mathcal{U}^*$ reduction steps are carried out on the succedent formula according to 13.21, 13.22 and 13.12 until it has the form $\neg \mathcal{C}$ or is a minimal formula. If it has become a true minimal formula, then the reduction is at an end. If it has assumed the form $\neg \mathcal{C}$ then it is reduced according to 13.23 to $\neg \neg \mathcal{U}^*, \mathcal{C} \rightarrow 1=2$, further (13.53) to $\mathcal{C} \rightarrow \neg \mathcal{U}^*$ then (13.23) to $\mathcal{C}, \mathcal{U}^* \rightarrow 1=2$. The same procedure is followed in the case where the succedent formula has become a false minimal formula; in the case $\rightarrow \mathcal{U}^*$ is

obtained first and then $\mathcal{U}^* \rightarrow 1=2$.

In both cases we have obtained a sequent which has also occurred in the reduction of the logical basic sequent $\mathcal{U}^* \rightarrow \mathcal{U}^*$ according to the procedure stated at 13.92 (or by a procedure that is not essentially different). Once again we need only follow the procedure stated there in order to complete the reduction of the sequent in hand.

It should be noted that in any reduction steps according to 13.5 in the reduction procedures at 13.92 and 13.93 the antecedent formula involved was never retained.

Paragraph 14

REDUCTION STEPS ON DERIVATIONS⁽²⁰⁾

In order to reduce arbitrary derived sequents we shall state a procedure by which certain reduction steps are carried out on the entire derivation of the sequent concerned. For this purpose I shall modify somewhat the notion of a derivation used so far (14.1) and shall then explain how an individual reduction step is to be carried out on such a derivation (14.2).

14.1. Modification of the Notion of a Derivation.

The new notion of a derivation results from the old one (5.2) as follows:

5.22 continues to apply even though the "end-sequent" of the derivation may now also contain antecedent formulae (so that we can speak of a "derivation for a sequent"). The symbols $\forall$, $\exists$ and $\supset$ must not occur in the derivation. No sequent of the derivation may be used to obtain more than one further sequent (by the application of a rule of inference).

It is easily seen that a derivation in the old sense can be transformed into a derivation with the same end-sequent which also satisfies this condition. We need merely work backwards from the end-sequent and write down correspondingly often those sequents which have been used more than once together with the sequents used for their derivation.

Mathematical basic sequents must fulfil the requirement 13.91; together with these all their "reduction instances", i.e., all sequents which may occur in the course of a given reduction procedure, are also admitted as mathematical basic sequents.

As logical basic sequents we may take arbitrary sequents of the form

$$\begin{aligned} & \mathcal{U} \rightarrow \mathcal{U} \quad \text{or} \quad \mathcal{U} \& \mathcal{B} \longrightarrow \mathcal{U} \quad \text{or} \quad \mathcal{U} \& \mathcal{B} \longrightarrow \mathcal{B} \\ \text{or } & \mathcal{U}, \mathcal{B} \longrightarrow \mathcal{U} \& \mathcal{B} \quad \text{or} \quad \forall x \mathcal{F}(x) \longrightarrow \mathcal{F}(t) \quad \text{or} \\ & \mathcal{U}, \neg \mathcal{U} \longrightarrow 1 = 2 \quad \text{or} \quad \neg \neg \mathcal{U} \longrightarrow \mathcal{U}, \end{aligned}$$

as well as all sequents which may occur in the reduction of one of these sequents according to 13.92, 13.93.

Structural transformations in their old form are no longer permissible.

Among the rules of inference we retain the rule of the $\forall$ -introduction and of "complete induction" with the following modification: a

$\forall$ -introduction or a "complete induction" in whose sequents no free variables other than a occur, remains permissible if in all associated sequents not containing the variable a the minimal terms which occur are replaced by their "numerical values" until all minimal terms have been eliminated (for motivation cf. 14.22).

The following new "rule of the $\neg$ -introduction" is added: From

$$\mathcal{T}, \mathcal{U} \rightarrow 1=2 \quad \text{results} \quad \mathcal{T} \rightarrow \neg \mathcal{U}.$$

One further rule of inference is still added - the "chain rule" -:

From a sequence of sequents (at least one) of arbitrary form a

sequent of the following kind results: for its succedent formula

we take the succedent formula of any one of the sequents of the

sequence. If this formula is a false minimal formula, any other

false minimal formula may be taken. For its antecedent formulae we

write down, in arbitrary order, all antecedent formulae of the sequent

concerned, together with all antecedent formulae of earlier sequents

in the sequence. In carrying out this inference we may omit formulae

for which the following holds: the same formula occurs already among

the formulae written down (i.e. those not omitted); or: the formula

is the same as the succedent formula of a sequent occurring earlier

in the sequence than the sequent from whose antecedent formulae it is

taken. Other antecedent formulae may be inserted among the formulae

already written down. Finally, the completed sequent may be transformed

further by replacing any one of its bound variables one or more times by another variable according to 5.244.

The "chain rule" has thus been formulated flexibly enough to allow for the transformation of a derivation in the old sense, which we assume to be already freed of the symbols $\forall$, $\exists$ and $\Rightarrow$ by the method described in Paragraph 12 (and which we also suppose to fulfil the conditions for functions, predicates and axioms in 13.12, 13.3, 13.91), into a derivation in the new sense without any change in its end-sequent.

Reason: All structural transformations are special cases of the "chain rule". The omitted rules of inference may be replaced by the new basic sequents that have taken their place, together with the "chain rule", e.g., the $\&$ -introduction: $\mathcal{T} \rightarrow \mathcal{U}$ and $\Delta \rightarrow \mathcal{B}$ and $\mathcal{U}, \mathcal{B} \rightarrow \mathcal{U} \& \mathcal{B}$ by the "chain rule" yields $\mathcal{T}, \Delta \rightarrow \mathcal{U} \& \mathcal{B}$. The $\forall$ -elimination: $\mathcal{T} \rightarrow \forall x \mathcal{F}(x)$ and $\forall x \mathcal{F}(x) \rightarrow \mathcal{F}(t)$ by the "chain rule" yields $\mathcal{T} \rightarrow \mathcal{F}(t)$. The

$\&$ -elimination and the "elimination of the double negation" are replaced correspondingly. Finally the "reductio": from $\mathcal{U}, \mathcal{T} \rightarrow \mathcal{B}$ and $\mathcal{U}, \Delta \rightarrow \neg \mathcal{B}$ and $\mathcal{B}, \neg \mathcal{B} \rightarrow 1=2$ we obtain by the "chain rule" $\mathcal{T}, \Delta, \mathcal{U} \rightarrow 1=2$, and by $\neg$ -introduction finally $\mathcal{T}, \Delta \rightarrow \neg \mathcal{U}$.

The new notion of a derivation is thus not narrower than the old one and for the purpose of stating a reduction rule for any one of

the sequents occurring in a derivation we can without loss of generality assume as given a derivation in the new sense for the sequent concerned.

In applying a rule of inference below I shall label as "premisses" those sequents from which a new sequent, the "conclusion", results.

That the "chain rule" in its intuitive meaning constitutes a "correct" inference is fairly obvious. It can after all be shown that this rule is replaceable by the old rules of inference and the structural transformations.

In formulating the "chain rule" it was permitted that no actual use was made of some premisses. This proves to be of practical value for the reduction procedure. The extensive replacement of the rules of inference by combinations of basic sequents and the "chain rule" is also motivated by convenience; it has the virtue of changing the original vertical arrangement of inferences into a horizontal arrangement.

Finally, I shall also pre-suppose that it has been stated for each sequent of a given derivation whether it is a basic sequent and of what kind or from what preceding sequents and by what rules of inference it has been obtained; I assume in general that it has been stated how the individual sequents, formulae, etc., involved in an application of a rule of inference, correspond to the designations used in the associated general schema: in this way the need for resolving possible ambiguities does not arise.

14.2. Reduction steps on derivations.

I shall now define the notion of a reduction step on a derivation (14.1) and at the same time prove the following: in such a step the derivation concerned is transformed into another derivation and its end-sequent is hereby modified in the following way:

The possible occurrences of free variables are replaced by arbitrarily chosen numerals; then any minimal terms that may be present are replaced by their "numerical values" until all minimal terms have been eliminated; and, furthermore, at most one reduction step according to 13.2 or 13.5 is carried out on the sequent. (It may thus happen that an end-sequent without free variables or terms remains entirely unchanged.)

The reduction step on derivations is unambiguous except in the cases in which the end-sequent undergoes one or more transformations according to a reduction step on sequents involving a choice (13.11, 13.21, 13.22); here the choice may be made arbitrarily; if this has been done the reduction step is then also unambiguous.

If the end-sequent of the derivation is in definitive form according to 13.4, then no reduction step is defined for this derivation. In other cases we carry out a reduction step whose definition now follows (recursively). In the following we therefore assume that the end-sequent is not in definitive form.

14.21. If the end-sequent of the derivation is a basic sequent then the

reduction step on it is carried out according to the reduction rules 13.91 - 13.93, which clearly also cover all basic sequents in their present sense: a replacement of all possible occurrences of free variables and terms must here take place, followed merely by precisely one step according to 13.2 or 13.5 (or none at all if the definitive form has already been reached). The claims made above concerning the reduction step on derivations are then obviously fulfilled.

14.22. We now consider the case where the end-sequent is the result of the application of a rule of inference and we presuppose that for the derivations of the premisses the notion of a reduction step is already defined and the validity of the associated assertions demonstrated.

The reduction step on the entire derivation begins with the following preliminary (replacement of free variables and minimal terms):

We begin by replacing all occurrences of free variables in the end-sequent by arbitrarily chosen numerals. Then we replace the same variables (i.e. the variables that were replaced in the end-sequent) in the entire derivation by the same numerals and replace the remaining free variables by 1, with one important exception: the free variable occurring in a $\forall$ -introduction or "complete induction" and designated by a at 5.25 must not be replaced in the premisses $\mathcal{P} \rightarrow \mathcal{F}(a)$ or $\mathcal{F}(a), \Delta \longrightarrow \mathcal{F}(a+1)$, nor in any sequent belonging to the derivation of that sequent.

Next we replace all minimal terms occurring in the derivation one by one by their "numerical values", with one important exception: no replacement takes place in the premises of a $\forall$ -introduction or "complete induction" containing a , nor in any sequent belonging to the derivation of that sequent.

Both of these replacement procedures obviously leave the derivation correct. Essential to this in the replacement of free variables is, first, the special condition for the variable a in the case of a $\forall$ -introduction and a "complete induction" as formulated at 5.25, further the requirement (14.1) that every derivational sequent serves as a premise for at most one application of a rule of inference. These two facts make it actually possible to separate completely from the rest of the variables the variables to be replaced so that by this distinction no error is introduced into any application of a rule of inference.

In the case of a term replacement the special requirement formulated at 14.1 for the $\forall$ -introduction and the "complete induction" is important (which is why it was introduced); for the original normal form of these rules of inference (5.25) may be destroyed by the replacement.

After this "preliminary" comes the actual reduction step according to the following rules. Yet if the end-sequent is now already in definitive form the reduction step terminates at this point.

14.23. Suppose that the end-sequent is the result of a $\forall$ -introduction or a $\neg$ -introduction. It is then eliminated and its premise taken for the new end-sequent, where, in the case of a $\forall$ -introduction, every occurrence of the free variable a must be replaced throughout the derivation of this premise by an arbitrarily chosen numeral and every minimal term by its "numerical value", subject to the same restrictions as at 14.22; not to be replaced, however, are terms in which the variable a occurred earlier.

The derivation has obviously remained correct and the end-sequent has become a reduced end-sequent in the sense of 13.21 or 13.23.

14.24. Suppose that the end-sequent is the result of a "complete induction". The numerical value of the term z will be denoted by the numeral n ; m shall be the numeral for the number smaller by 1 (if n is not equal to 1). The free variable a in the derivation of the premise $F(a), \Delta \rightarrow F(a+1)$ is replaced successively by the numerals 1, 2, 3, etc. up to m , subject to the same restriction as at 14.22, and all minimal terms that may have resulted are then replaced by their "numerical values", also subject to the same restriction as at 14.22. The derivation as a whole is then completed by the application of the "chain rule" which makes it possible to derive the end-sequent $\mathcal{T}, \Delta \rightarrow (F(n))^*$ once again from $\mathcal{T} \rightarrow (F(1))^*$ and the newly derives sequents $(F(1))^*, \Delta \rightarrow (F(2))^*$ and $(F(2))^*, \Delta \rightarrow (F(3))^*$ etc. up to $(F(m))^*, \Delta \rightarrow (F(n))^*$.

The asterisk denotes in each case the changes that have occurred through the replacement of minimal terms. By virtue of the preparatory replacement of terms (14.22) and the further replacement of terms here carried out all occurrences of minimal terms have finally been eliminated so that the related $\mathcal{F}^*$ -expressions have indeed become equal to one another, even if they had not been equal before. If n equals 1, then we merely put 1 for a and by the "chain rule" the end-sequent $\mathcal{T}, \Delta \rightarrow (\mathcal{F}(1))^*$ from $\mathcal{T} \rightarrow (\mathcal{F}(1))^*$ and $(\mathcal{F}(1))^*, \Delta \rightarrow (\mathcal{F}(1))^*$.

14.25. The last case to be considered is that where the end-sequent is the conclusion of a "chain rule" inference. This is the most difficult reduction since the chain rule in some sense amasses the difficulties of all inferences.

That premise whose succedent formula provides the succedent formula of the end-sequent I shall call the "major premise". If the succedent formula of the end-sequent is a false minimal formula we choose as major premise the first premise (in the given order) whose succedent formula is also a false minimal formula. This does not change the correctness of the "chain rule", even if a later premise was the major premise before; it may merely happen that certain antecedent formulae of the end-sequent can no longer be regarded as taken from the premises but rather as newly adjoined.

From these preliminaries it follows that the major premise can in no case be in definitive form (13.4), for otherwise the end-sequent would obviously also have to be in definitive form and this was assumed not

to be the case. Hence a reduction step can be carried out on the derivation of the major premise. In this respect I shall distinguish four cases which will be dealt with in ~~them~~ (14.251 - 14.254).

14.251. Suppose that the major premise undergoes a change according to 13.2 in the reduction step on its derivation. In that case the end-sequent is subjected to the appropriate reduction step for sequents according to 13.2; any choice that arises is to be made arbitrarily. The reduction step for derivations is then carried out on the derivation of the major premise and wherever a choice exists the same choice is to be made as before. The succedent formulae of both sequents are now the same once again (up to possible re-designations of bound variables) and the "chain rule" is once again correct. In this case the reduction step for the whole derivation is thus completed.

14.252. Suppose that the major premise undergoes a change according to 13.5 in the reduction step on its derivation and that the affected antecedent formula is one of the formulae that has been included among the antecedent formulae of the end-sequent (when the latter was formed by the "chain rule") or that it was omitted because an equal formula had already occurred among the antecedent formulae. In that case the reduction step is carried out on the derivation of the major premise and, so that the "chain rule" becomes again correct, the end-sequent is modified according to the corresponding reduction step on sequents (13.5). I.e., if the affected antecedent formula was itself absorbed into the end-sequent then the same reduction step is here carried out

on that formula; but if it was omitted because it agreed with an already existing formula then the reduction step is carried out on the latter formula and it is retained regardless of whether the corresponding formula in the reduction of the premise is omitted or retained.

14.253. (Principal Case.) Suppose that the major premise, say

$\Delta \rightarrow C$, undergoes a change according to 13.5 in the reduction step on its derivation and that the affected antecedent formula (V) is a formula that was not included among the antecedent formulae of the end-sequent because it agreed with the succedent formula of an earlier premise; suppose further that this premise, call it $T \rightarrow V$, undergoes a change during the reduction step on its derivation which, in that case, must necessarily be a change according to 13.2. (Since V cannot be a minimal formula.) - Suppose that the end-sequent of the whole derivation has the form $C \rightarrow D$. I shall distinguish three individual cases depending on whether V has the form $\forall x F(x)$, $\exists x F(x)$ or $\neg A$. The treatment of the three cases is not essentially different.

Suppose first that V has the form $\forall x F(x)$. In that case an antecedent formula $F(x)$ is adjoined in the reduction step according to 13.51 on $\Delta \rightarrow C$, and $\forall x F(x)$ is either retained or omitted; in the reduction step on $T \rightarrow \forall x F(x)$, which must be carried out according to 13.21, the same symbol x may be chosen for the numeral to be substituted so that $T \rightarrow F(x)$ results. We now form three "chain rule"

inferences: the first one contains for its premisses those of the original "chain rule" inference, but with $T \rightarrow F(n)$ in place of $T \rightarrow \forall x F(x)$; its conclusion: $\textcircled{D} \rightarrow F(n)$. A correct result. The second "chain rule" inference contains for its premisses those of the original "chain rule" inference, but with the sequent that was reduced according to 13.51 in place of $\Delta \rightarrow \mathcal{C}$; its conclusion: $\textcircled{D}, F(n) \rightarrow \mathcal{D}$. This is also a correct "chain rule" inference. The third "chain rule" inference again yields the end-sequent ~~from~~ $\textcircled{D} \rightarrow \mathcal{D}$ from $\textcircled{D} \rightarrow F(n)$ and $\textcircled{D}, F(n) \rightarrow \mathcal{D}$ - Together with each one of the sequents used we must of course write down the complete derivation of each one of them so that now an altogether correct derivation again results.

If $\mathcal{V}$ has the form $\mathcal{U} \wedge \mathcal{B}$, then an antecedent formula $\mathcal{U}$ or $\mathcal{B}$ is adjoined in carrying out a reduction step on $\Delta \rightarrow \mathcal{C}$ according to 13.52. $T \rightarrow \mathcal{U} \wedge \mathcal{B}$ becomes either $T \rightarrow \mathcal{U}$ or $T \rightarrow \mathcal{B}$, as desired; the choice should be made so that the same formula occurs as in $\Delta \rightarrow \mathcal{C}$. The procedure is the continued exactly as in the previous case.

If $\mathcal{V}$ has the form $\neg \mathcal{U}$, then $\Delta \rightarrow \mathcal{C}$ is reduced to $\Delta^{(w)} \rightarrow \mathcal{U}$ and $T \rightarrow \neg \mathcal{U}$ to $T, \mathcal{U} \rightarrow 1=2$. We now form, as before, two "chain rule" inferences with the conclusions $\textcircled{D}, \mathcal{U} \rightarrow 1=2$ and $\textcircled{D} \rightarrow \mathcal{U}$. With their order interchanged, these two inferences again yield $\textcircled{D} \rightarrow \mathcal{D}$ by a third "chain rule" inference. This is so since $\mathcal{D}$, like $\mathcal{C}$ and $1=2$, is a false

minimal formula.

14.254. We are still left with the following possibilities: the major premise remains unchanged in the reduction step on its derivation; or: its change is of the kind assumed at 14.253 and the premise $\mathcal{T} \rightarrow \mathcal{V}$ remains unchanged in the reduction step on its derivation. -In both cases we carry out the reduction step on the derivation of the premisses that have remained unchanged and this completes the reduction. Yet in the particular case of a reduction step on the derivation of the premisses according to 14.253 (where the end-sequent, i.e., the premise, remains unchanged) we proceed somewhat differently, viz.: this reduction step is to be carried out yet without forming the "third chain rule inference" presented for this purpose; in its place we put rather the two premisses of that "chain rule" inference in place of the conclusion of that inference in the sequence of the premisses of that "chain rule" inference which terminates the derivation as a whole. This obviously leaves the "chain rule inference" correct. The end-sequent is not changed. The definition of a reduction step on a derivation is thus completed.

Paragraph 15

ORDINAL NUMBERS AND PROOF OF FINITENESS

It remains to show that a successive application of a reduction step on a given derivation always leads to the definitive form (of the end-sequent) in finitely many steps regardless of the choices made in

those cases in which a choice exists. In doing so, we shall at the same time have given a reduction rule (13.6) for arbitrary derived sequents, since the reduction of the derivation of the sequent (according to Paragraph 14) automatically involves the reduction of the sequent (according to Paragraph 13).

In order to prove the finiteness of the procedure we shall have to show that each reduction step in a definite sense "simplifies" a derivation. For this purpose I shall correlate with each derivation an "ordinal number" representing a measure for the "complexity" of the derivation (15.1, 15.2). It can then indeed be shown that with every reduction step on a derivation the ordinal number of that derivation (in general) diminishes (15.3). However, the finiteness of the reduction procedure is hereby not immediately guaranteed; for the ordering of the derivations (corresponding to the well-ordering of their ordinal numbers) is of a special kind since it may happen that in terms of its complexity a derivation ranks above infinitely many other derivations. E.g., a derivation whose end-sequent has taken on the form $\rightarrow \forall x F(x)$, as a result of a "complete induction" and a $\forall$ -introduction, must be regarded as more complex than any one of the infinitely many special instances obtained by substituting individual numerals for x and resolving the "complete induction" (14.23, 14.24). The situation may be complicated still further by a multiple resting of such instances. Thus the "ordinal numbers" here have the nature of "transfinite ordinal numbers" (cf. footnote 21) and the inductive comprehension of their totality is not possible by ordinary complete induction but only by

"transfinite induction" whose validity requires a special verification (15.4).

15.1. Definition of ordinal numbers (recursive).

As "ordinal numbers" I shall use certain positive finite decimal fractions formed according to the following rule:

Ordinal numbers with the characteristic 0 are precisely the following numbers: 0.1, 0.11, 0.111, 0.1111, ... i.e. in general: every number with the characteristic 0 whose mantissa consists of finitely many 1's; also the number 0.2.

Zeros may not be appended to these expressions, neither here nor below; this achieves uniqueness of notation. - I shall call one mantissa smaller than another mantissa if this relationship holds between the numbers that result from the prefixing of these mantissae by "0".

The mantissa of an ordinal number with the characteristic $\xi + 1$ ($\xi \geq 0$) is obtained by taking several mutually distinct ordinal numbers (at least one) with the characteristic ξ , ordering their mantissae according to size, so that the largest occurs first, the smallest last, and by then writing them down in that order from left to right, separating any two successive mantissae by $\xi + 1$ zeros. All numbers obtainable in this way from ordinal numbers with the characteristic ξ , and no others, are ordinal numbers with the characteristic $\xi + 1$.

Examples of ordinal numbers:

0.111, 1.1101, 1.2, 2.111, 2.2010011010011, 3.2010020001

It can be determined uniquely of a given number with the characteristic $\mathfrak{g} + 1$ from what numbers with the characteristic $\mathfrak{g}$ it has been generated by the above rule. For a number with the characteristic $\mathfrak{g}$ can obviously have no more than $\mathfrak{g}$ consecutive zeros in any one place.

Further details about the ordering of the ordinal numbers follow at 15.4.

15.2. The correlation of ordinal numbers with derivations.

With every given derivation (in the sense of 14.1) we can correlate a unique appropriate ordinal number calculated according to the following recursive rule:

The following observation must here be kept in mind: the maximum number ($\mathfrak{v}$) of consecutive zeros in the mantissa is larger than 1 and all of its sections that are separated by successions of $\mathfrak{v}$ zeros, except for the last one, begin with the numeral 2, the last section consists only of 1's.

If the end-sequent of the derivation is a basic sequent, the derivation receives an ordinal number of the form 2.2001 ... 1, where the number of 1's must be chosen to be larger by one than the total number of logical connectives occurring in the sequent.

Now suppose that the end-sequent is the conclusion of the application of a rule of inference and that for the derivation of the premisses their associated ordinal numbers are already known. From these the ordinal

number of the whole derivation is calculated as follows:

If the end-sequent is the conclusion of a $\vee$ - or $\rightarrow$ -introduction then the numeral 1 is adjoined to the ordinal number for the derivation of the premise. By virtue of the stated properties of arbitrary ordinal numbers for derivations we have obviously another correct ordinal number in accordance with 15.1.

If the end-sequent is the conclusion of a "chain rule" inference, we focus our attention on the mantissae of the ordinal numbers of the derivations for the premisses; suppose that ν is the maximum number of consecutive zeros in all of these mantissae. Should there be equal mantissae among them, we distinguish these by adjoining to one of them $\nu + 1$ zeros and one 1, to another $\nu + 1$ zeros and two ones, etc.; this principle is to be applied to every occurrence of equal mantissae. The mantissae thus obtained are mutually distinct; they are then written down from left to right according to size (the largest one first) and two successive mantissae are in each case to be separated by $\nu + 2$ zeros; finally $\nu + 2$ zeros and one 1 are adjoined at the end. The result is the mantissa of the ordinal numbers for the whole derivation. For its characteristic we take the smallest natural number which exceeds the maximum number of consecutive zeros in the mantissa by 0 or more and which, first, exceeds by at least two the maximum number of consecutive zeros in any one of the ordinal numbers for the derivations of the premisses and which, second, is no smaller than twice the total number of logical connectives in the succedent formula of any one of the

premisses preceding the major premise (14.25).

If the end-sequent is the conclusion of a "complete induction", then the ordinal number of the whole derivation receives a mantissa of the form $201..10..01$; where the number of consecutive 1's is to be chosen greater by one than the total number of consecutive 1's in the corresponding place in the larger one of the mantissae of the ordinal numbers for the derivations of both premisses (or either one of them, if both are equal); i.e., if the latter mantissa begins with 200, one 1 is to be chosen; in every other case it must begin with 201 .. 10, in which case one more 1 than here is to be chosen. The total number of consecutive zeros must be $\nu + 2$, where ν is the maximum number of consecutive zeros in the two mantissae mentioned. As characteristic we take the smallest natural number that exceeds the maximum number of consecutive zeros in the mantissa by zero or more and which first is not smaller by two than the corresponding maximum number of zeros in any one of the two ordinal numbers used and which second is not smaller than twice the total number of logical connectives in the formula $\mathcal{F}(1)$.

It is easily seen that this newly formed number is another correct ordinal number (15.1) and possesses moreover the special properties stated above.

15.3. A reduction step diminishes the ordinal number of a derivation.

We must now prove that with every reduction step on a derivation

according to 14.2 the ordinal number of the newly resulting derivation becomes in general smaller than that of the old derivation. I shall show: the characteristic does not increase; the mantissa decreases in all cases in which the end-sequent is not already in definition form after the replacement of the free variables and terms (14.21, 14.22); the maximum number of consecutive zeros in the mantissa furthermore remains unchanged except in the case of a reduction according to 14.253 where it increases by exactly two.

I shall again proceed recursively, i.e., I shall prove the assertion by complete induction.

For derivations whose end-sequent is a basic sequent the result follows from the method of correlating ordinal numbers with such derivations together with the fact that in the reduction step the sequent undergoes a change according to 13.2 or 13.5, and here the total number of occurring logical connectives is diminished. (If the definition form of the derivation is achieved earlier then the ordinal number remains unchanged.) What is important here is that in changes according to 13.5, the altered antecedent formula is always omitted, c.f. 13.91 - 13.93.

Suppose now that the end-sequent is the result of the application of a rule of inference and that the assertion has already been proved for the derivations of the premisses.

The preliminary step (14.22) has obviously no influence on the ordinal

number of the derivation. If the definition form of the end-sequent results already with this step then the ordinal number therefore remains unchanged. If this is not the case then the following holds:

If the end-sequent is the conclusion of a $\forall$ - or $\neg$ -introduction then the assertion follows at once from the method of correlating ordinal numbers with such a derivation.

Even if the end-sequent is the conclusion of a "complete induction" the truth of the assertion follows easily. The "complete induction" is, after all, transformed into a "chain rule" inference; this does not lead to an increase in the characteristic of the ordinal number; although the mantissa may become much longer, it nevertheless diminishes since the mantissa of the ordinal number of one of the two original derivations of the premisses must always occur at the beginning of that mantissa. The maximum number of consecutive zeros ($\neg + 2$) remains unchanged.

Suppose finally that the end-sequent is the conclusion of a "chain rule" inference. The selection of the premise of an earlier sequent as major premise (14.25) does not alter the mantissa of the ordinal number; the characteristic may on the other hand diminish because certain succedent formulae of the premisses no longer contribute to its calculation.

The reduction step now takes the form of either 14.251 or 14.252. Here one of the mantissae of the ordinal numbers for the derivations of premisses is diminished without a change in the maximum number of

consecutive zeros occurring in it. This has obviously a simultaneous diminishing effect of the mantissa for the ordinal number of the total derivation. The number of zeros is after all still $\nu + 2$; the diminished mantissa may conceivably occur in a later place of the sequence, which is ordered by size; if the mantissa was one of several equal mantissae then one less 1 is adjoined to the remaining mantissae; yet in all cases the first mantissa in the sequence of mantissae separated by $\nu + 2$ zeros which has not remained the same must be smaller than before; consequently the total mantissa has certainly also been diminished. The characteristic does not increase.

In a reduction step according to 14.253 the ordinal number of the derivation is altered as follows: let us first consider the ordinal numbers for the two derivations which conclude with the newly formed first or second "chain rule" inference. For these two derivations the situation is the same as that in the previous case, i.e.: the two mantissae are smaller than the mantissa of the ordinal number of the original derivation; the maximum number of consecutive zeros ($\nu + 2$) has remained the same; the characteristics have not increased. We now introduce the third "chain rule" inference and form the ordinal number of the new total derivation: Its mantissa begins with one of the two earlier mantissae followed by $\nu + 3$ zeros (usually $\nu + 4$); it is consequently smaller than the mantissa of the original ordinal number; the maximum number of consecutive zeros is $\nu + 4$, hence larger by two than before; the characteristic of the total derivation, finally, cannot have increased, for the total number of logical connectives in the

succedent formula $F(\kappa)$, or $\mathcal{U}$ or $\mathcal{B}$, or $\mathcal{U}$, is smaller than that in the formula $\mathcal{V}$, i.e., in $\mathcal{V} \wedge F(\kappa)$, or $\mathcal{U} \wedge \mathcal{B}$, or $\neg \mathcal{U}$; hence the sum of twice the number of logical connectives in the former formulae with $\mathcal{V} + 4$ zeros, which determines the new characteristic, is not larger than the sum of twice the number of logical connectives in the latter formulae with $\mathcal{V} + 2$ zeros; nor could the characteristic of the original derivation be smaller than the latter sum since $\mathcal{V}$ was one of the succedent formulae which contributed to its calculation.

In a reduction step according to 14.254 the situation is the same as in the case of 14.251 and 14.252 unless we are dealing with an exceptional case. Yet even a special case can be dealt with without difficulty on the basis of our previous considerations; here one of the mantissae of the ordinal numbers for the derivations of the premisses is no longer replaced by one smaller mantissa, as was done above, but by two; yet the effect is the same in every desired respect. The characteristic is not increased; its maximum number of consecutive zeros before the reduction was not smaller by two than twice the total number of logical connectives in so that the contributions of $F(\kappa)$, or $\mathcal{U}$ or $\mathcal{B}$, or $\mathcal{U}$, after the reduction, cannot lead to an increase.

It has thus been proved that in a reduction step the ordinal number (usually) diminishes. The most important point was our consideration concerning the characteristic of the ordinal number in discussing the reduction steps 14.253 and 14.254; this is the idea which enables us

to recognize a simplification of the derivation in such a reduction step in spite of the apparent increase in complexity. The simplification consists precisely in the fact that the premisses of the "third chain rule" inference are "interwoven" to a lesser degree (viz., to a degree corresponding to the total number of logical connectives in the succedent formula of the first premise which is also the antecedent formula of the second premise) than the premisses of the first and second and the premisses of the original "chain rule" inference. The method of correlating an ordinal number with a "chain rule" inference (15.2) is formulated from the above point of view; all other details follow more or less automatically.

15.4 Demonstration of the finiteness of the reduction procedure.

Some facts - needed below - about the ordering according to size of the ordinal numbers:

With every number α with the characteristic $\wp$ ($\wp \geq 0$) I correlate the system $G(\alpha)$ of those ordinal numbers with the characteristic $\wp + 1$ in whose formation according to 15.1 the number α was the largest of the ordinal numbers with the characteristic $\wp$ that were used. Every ordinal number with the characteristic $\wp + 1$ belongs uniquely to one such system $G(\alpha)$. If α_1 is smaller than α_2 , then every number of $G(\alpha_1)$ is also smaller than every number of $G(\alpha_2)$. The ordering of the systems $G(\alpha)$ corresponds therefore to the ordering of the numbers α . The following holds for the ordering of the numbers (with the characteristic $\wp + 1$)

within a system $G(\alpha)$: the smallest number within $G(\alpha)$ is the number $\alpha + 1$. The remaining numbers of $G(\alpha)$ correspond order-isomorphically to the totality of those numbers with the characteristic $\beta + 1$ which are smaller than $\alpha + 1$ in the following way: Every number of $G(\alpha)$, except for $\alpha + 1$, results from $\alpha + 1$ through the adjunction of $\beta + 1$ zeros followed by the mantissa of any one of the numbers with the characteristic $\beta + 1$ which is smaller than $\alpha + 1$. The ordering of this mantissa is hereby carried over.

The correctness of all these assertions follows easily from the definition of the ordinal numbers. The reader may find it beneficial to examine the ordering of the ordinal numbers with the characteristics 1, as well as 2 and 3, for example, using this definition. (21)

I now assert (theorem of "transfinite induction"):

All ordinal numbers (15.1) are "accessible" in the following sense by running through them in the order of increasing magnitude: the first number, 0.1, is considered as "accessible"; if all numbers smaller than a number β have furthermore been recognized as "accessible" then β is also considered as "accessible".

Proof. 0.1 is accessible by hypothesis, hence also 0.11, hence also 0.111, etc., and it follows in general by complete induction that every number smaller than 0.2 is accessible. Hence 0.2 is also accessible, and thus all numbers with the characteristic 0.

I now apply a complete induction, i.e., I assume that the accessibility of all numbers up to and including those with the characteristic $\mathfrak{s}$ ($\mathfrak{s} \geq 0$) has already been proven and that it is now to be proved for numbers with the characteristic $\mathfrak{s} + 1$. The first of these numbers, i.e., the number with the mantissa 1, is provable. Now note that we have already run through the numbers with the characteristic $\mathfrak{s}$. To every such number α corresponds a system $G(\alpha)$ of numbers with the characteristic $\mathfrak{s} + 1$; this system consists of the number $\alpha + 1$ and a system order-isomorphic with those numbers with the characteristic $\mathfrak{s} + 1$ that are smaller than $\alpha + 1$. To run through the numbers with the characteristic $\mathfrak{s} + 1$ now amounts merely to a running through of the systems $G(\alpha)$ in the same way in which we ran through the numbers α with the characteristic $\mathfrak{s}$; for if a number $\alpha + 1$ has been recognized as accessible then all remaining numbers of the system $G(\alpha)$ obviously become accessible at the same time; we need merely run through this system in exactly the same way in which we have already run through the isomorphic system of the numbers (with the characteristic $\mathfrak{s} + 1$) smaller than $\alpha + 1$. In this way we can run through all numbers with the characteristic $\mathfrak{s} + 1$ by virtue of having run through the numbers with the characteristic $\mathfrak{s}$. To the totality of numbers α (with the characteristic $\mathfrak{s}$) smaller than a number α_0 corresponds, in the case of the number $\alpha_0 + 1$ (with the characteristic $\mathfrak{s} + 1$), the totality of numbers belonging to the systems $G(\alpha)$ (where $\alpha < \alpha_0$).

Conclusion. By means of the "theorem of transfinite induction" the finiteness of the reduction procedure for arbitrary derivations now follows at once. If the finiteness of the reduction procedure has already been proven for all derivations whose ordinal number is smaller than a number β then this also holds for every derivation with the ordinal number β ; for by a single reduction step the latter derivation is transformed into a derivation with a smaller ordinal number or a derivation in definitive form. If the derivation was already in definitive form then there was nothing more to prove.) Thus the property of the finiteness of the reduction procedure carries over from the totality of the derivations with the ordinal numbers smaller than β to the derivations with the ordinal number β ; by the theorem of transfinite induction it therefore holds for all derivations with arbitrary ordinal numbers. This concludes the consistency proof.

SECTION V.

REFLECTIONS ON THE CONSISTENCY PROOF

Paragraph 16

THE FORMS OF INFERENCE USED IN THE CONSISTENCY PROOF

I shall review in the following the inferences and specific concepts used in the consistency proof from two aspects: First I shall examine to what extent they can be considered as indisputable (16.1) second, in connection with the theorem of Gödel (2.32), to what extent they correspond to the methods of proof contained in formalized elementary number theory and in what way they go beyond these methods (16.2).

16.1 In terms of the indisputability of the methods of proof used, the critical point is the proof of the finiteness of the reduction procedure (15.4). We shall come back to this point later. All other techniques of proof used in the consistency proof can certainly be considered as "finitist" in the sense outlined in detail in Section III. This cannot be "proved" if for no other reason than the fact that the notion of "finitist" is not unequivocally formally defined and cannot in fact be delineated in this way. All we can do is to examine every individual inference from this point of view and try to assess whether that inference is in harmony with the finitist sense of the concepts that occur and make sure that it does not rest on a possibly inadmissible "actualist" interpretation of these concepts. I shall discuss briefly the here most relevant passages of the consistency proof:

The objects of the consistency proof, as of proof theory in general, are certain symbols and expressions, such as terms, formulae, sequents, derivations, ordinal numbers, not to forget the natural numbers. All these objects are defined (3.2, 5.2, 14.1, 15.1) by construction rules analogously to the definition of the natural numbers (8.11); in each case such a rule indicates how more and more such objects can be constructed step by step. - It must here be presupposed that in formalized elementary number theory certain specific "functions", "predicates" and "axioms" have been stipulated which satisfy the conditions laid down for these objects (13.12, 13.3, 13.91). Strictly speaking, this presupposition introduces a transfinitely used "if - then" into the consistency proof; yet this "if - then" is obviously harmless since the proof need not be regarded as meaningful at all until that presupposition has actually been made and its conditions have been shown to be satisfied.

A number of functions and predicates were furthermore applied to these objects and they were decidably defined in the sense of 8.12. E.g., the function "the end-formula of a derivation", the predicate "containing at least one $\forall$ - or $\exists$ -symbol" and many others. The following functions, in particular, were also decidably defined, as is easily verified: "the derivation resulting from a derivation by a transformation according to Paragraph 12", "the derivation resulting from a derivation by a reduction step in which the conditions of a possible choice were unequivocally specified" (14.2), "the ordinal number of a derivation" (15.2).

Furthermore, propositions of the following kind were proved by complete induction: "for all sequents", "for all derivations" etc., whose validity for each individual sequent or derivation was decidable. E.g.: "The figure resulting from a derivation by a reduction step is another derivation and the transformation of the end-sequent fulfils certain conditions" (14.2); "in carrying out a reduction step we diminish the ordinal number" (15.3).

In applying the concept "all" in the consistency proof, I have not used the clumsy finitist expression given in 10.11 for it; here the distinction between the actualist finitist interpretations has no bearing on our reasoning in any case.

The negation of a transfinite proposition occurs only once in the entire proof (at 13.90) and only in a harmless form in which the proposition concerned leads to a quite elementary contradiction. The negation can actually be avoided altogether if for "consistency" the following positive expression is used: "every derivation has an end-formula which does not have the form $\mathcal{A} \rightarrow \mathcal{A}$." Here the "not" is no longer transfinite.

16.11. What can be said, finally, about the proof of the finiteness of the reduction procedure (15.4)?

The notion of "accessibility" in the "theorem of transfinite induction" is of a very special kind. It is certainly not decidable in advance whether it is going to apply to an arbitrary given number; from the point of view explained in Paragraph 9, this concept therefore has no immediate sense since an "actualist sense" has after all been rejected.

It gains a sense merely by being predicated of an individual number for which its validity is simultaneously proved. It is quite permissible to introduce concepts in this way; the same situation arises, after all, in the case of all transfinite propositions if a finitist sense is to be ascribed to them, c.f. Paragraph 10. With the statement that "if all numbers smaller than β have already been recognized as accessible then β is also accessible" the definition of the notion of "accessibility" is already formulated in conformity with this interpretation. No circularity is of course involved in this formulation; the definition is, on the contrary, entirely constructive; for β is counted as accessible only when all numbers smaller than β have previously been recognized as accessible. The "all" occurring here is of course to be interpreted finitistically (10.11); in each case we are after all dealing with a totality with a constructive rule for generating all elements.

About the proof of the theorem of transfinite induction the following must be said: From the way the notion of "accessibility" was defined it follows that in proving this theorem a "running through" of all ordinal numbers in ascending magnitude must take place. In dealing with the numbers with the characteristic 0 the following is to be observed: the infinite totality of the numbers smaller than 0.2 is overcome by one single idea: the proof can be carried out arbitrarily far into this totality; hence it may be considered as completed for the entire totality. This "potential" interpretation of the "running through" of an infinite totality must be applied throughout the entire proof.

The occurrence of a transfinite induction hypothesis in the complete

induction on $\mathfrak{f}$ is to be interpreted in the sense of 10.5 and is therefore unquestionable. In the inference: "if the number $\alpha + 1$ has been recognized as accessible then all remaining numbers of the system $G(\alpha)$ are accessible" a transfinite "if - then" occurs. Objections were raised against this concept at 11.1; yet these do not apply to the present case for the very reason that the hypothesis is here not to be interpreted hypothetically but rather as follows: only after having reached $\alpha + 1$ can we successfully run through the numbers of $G(\alpha)$ (viz., in exact correspondence with the way in which we ran through the numbers up to $\alpha + 1$).

Now let us consider the induction step as a whole, i.e., the re-interpretation of the running through of the $\mathfrak{f} + 1$ -systems in terms of the running through of the $\mathfrak{f}$ -systems. This is undoubtedly the most critical point of the argument. Yet I believe that if we think about it deeply enough we cannot dispute the remarkable plausibility of the argument here used. We might for example visualize the initial cases with the characteristics 1, 2, 3 in detail. After all, as the characteristic grows nothing new is basically added; the method of progression always remains the same. It must of course be admitted that the complexity of the multiply nested infinities which must be "run through" grows considerably; this running through must always be regarded as "potential", as was done in the case of the characteristic 0. The difficulty lies in the fact that although the precise finitist sense of the "running through" of the $\mathfrak{f}$ -numbers is reasonably perspicuous in the initial cases it becomes of such great complexity in the general case that it is only remotely visualizable; yet this must be regarded as

sufficient for an acceptable basis upon which the possibility of the running through of the $\omega + 1$ -numbers can be convincingly justified.

The "conclusion", finally, adds nothing essentially new. The proposition that the reduction procedure for a derivation is finite regardless of how possible choices may be made, contains a transfinite "there is", viz., with respect to the total number of reduction steps. This proposition is of the same kind as the proposition conceiving the "accessibility"; in each special case it receives its definite sense only through the proof of its validity for this case; this corresponds to the finitist interpretation (10.3). For the purpose of the consistency proof alone, incidentally, the notion of a "choice" is dispensable since we are here dealing only with the reduction of a derivation with the end-sequent $\rightarrow 1 = 2$ and since here all reduction steps are unequivocal and do not depend on choices. The total number of steps is not specified in advance; we can merely make certain statements about it and these become more and more indefinite as the ordinal number of the derivation increases. (The place of a direct statement of such a number is taken by its ^{"stability"}~~"stability"~~). This can undoubtedly still be regarded as being in harmony with the finitist view.

Altogether I am inclined to believe that in terms of the fundamental distinction between disputable and indisputable methods of proof (Paragraph 9), the proof of the finiteness of the reduction procedure (15.4) can still be considered as indisputable so that the consistency proof represents a real vindication of the disputable parts of elementary number theory.

16.2. In order to examine the extent to which the consistency proof coincides with the theorem of Gödel (2.32) we would first have to correlate natural numbers with the objects of proof theory (formulae, derivations, etc.) corresponding to the way in which it was done in Gödel's paper cited in footnote 3, and would also have to introduce the required functions and predicates for these objects as functions and predicates for the corresponding natural numbers. Then the consistency proof becomes a proof with the natural numbers as objects. In order to obtain a formally delineated formalism we would have to limit the possibilities of definition provided for above to definite schemata which can easily be chosen general enough to allow for the definition of all functions and predicates required in proof theory; cf., for example, Gödel's version.

The forms of inference in the consistency proof are then none others than those presented in our formalization of number theory; only the proof of finiteness (15.4) occupies again a special position. It is impossible to see how the latter proof could be carried out with the techniques of elementary number theory. For this reason the consistency proof is in harmony with Gödel's theorem.

In this connection the following two facts are of interest whose proof will not be given since it would lead too far afield:

1. If the inference of complete induction is omitted from formalized elementary number theory then the consistency proof can be formulated without essential change in such a way that - after having carried out the mentioned re-interpretation into a proof about natural numbers - the techniques of elementary number theory (including complete induction)

suffice completely.

2. The consistency proof for the whole of elementary number theory, reinterpreted with the natural numbers as objects, can be carried out with techniques from analysis.⁽²²⁾

The special position of the inference of complete induction is due to the following fact: if this inference is omitted then a definite upper bound can be given for the total number of reduction steps required for the reduction of a given sequent. Yet if the inference of complete induction is included then this number in its dependence on choices, can become arbitrarily large. This is so since in the re-interpretation of this rule of inference (14.24) the total number of required reduction steps for the sequent $\mathcal{T}, \Delta \longrightarrow \mathcal{F}(z)$ obviously depended on the number n (the value of z) and this number may depend on a choice, as is the case if z is a free variable, and must therefore first be replaced by an arbitrarily chosen numeral n . In this case it may happen that there exists no general bound for the total number of reduction steps required in the reduction of the sequent $\mathcal{T}, \Delta \longrightarrow \mathcal{F}(z)$.

This fact must be the reason why in the earlier consistency proofs the rule of complete induction could not be included (2.4).

Paragraph 17

CONSEQUENCES OF THE CONSISTENCY PROOF

First I shall discuss the question to what extent the consistency proof remains applicable if the "elementary number theory" formulated in Section II is extended by the addition of new concepts and methods (17.1), then I shall point out its transferability to other branches of mathematics

(17.2), and shall finally examine certain objections by the "intuitionists" against the significance of consistency proofs as such (17.3).

17.1. For the value of a consistency proof it is very significant whether the stipulated formalism for the particular mathematical theory involved, in our case elementary number theory, really fully contains that theory (cf. 3.3, 5.3). Yet in practice elementary number theory is not subject to any formal restrictions; it can always be extended further by new kinds of specific concepts, possibly also by the application of new kinds of forms of inference. How does this affect the consistency proof? Well, whenever the present framework is exceeded an extension of the consistency proof to the newly incorporated techniques is required. The consistency proof is already designed in such a way that this is possible to a very large degree without difficulties.

If new functions or predicates for natural numbers are introduced, for example, then a decision rule in accordance with 8.12 must be given for them; if additional mathematical axioms are introduced then a reduction rule must be given for them in accordance with 13.91 (cf. Paragraph 6 and 10.14). Specific non-decidable concepts in the sense of 6.3 present no difficulties either since they can be eliminated by the method described at that point. All these requirements are easily fulfilled as long as the introductions are, in some customary sense, "correct", and the axioms "true".

Even new kinds of inferences may be carried out which are not representable in the present formalism. In fact every formally defined system contained

in elementary number theory is necessarily incomplete in the sense that there are number-theoretical theorems of an elementary character whose truth can be proven by plausible finitist inferences yet not by means of methods of proof of the system proper.⁽²³⁾ This fact was advanced as an argument against the value of consistency proofs.⁽²⁴⁾ Yet my consistency proof remains unaffected by it; quite generally it can here be said: if an elementary number-theoretical theorem can be proved by means of inferences not belonging to my formalism, then the statement of a reduction rule for this theorem according to 13.91 will include the theorem in the consistency proof. The theorem given as an example by Godel has the quite elementary form $\forall x \mathcal{P}(x)$, where $\mathcal{P}$ represents a decidable predicate about the natural numbers; the fact that the finitist truth of this theorem has been recognized means that $\mathcal{P}(n)$ is true for each individual n , and from this the reducibility of the sequent $\longrightarrow \forall x \mathcal{P}(x)$ according to 13.21, 13.4 follows at once.

The concept of the reduction rule has in fact been kept general enough so that it is not tied to any definite logical formalism but corresponds rather to the general concept of "truth", certainly to the extent to which that concept has any clear meaning at all (cf. 13.8).

If a new form of inference is to be included as such in elementary number theory as formulated so far, it must be suitably included in the reduction procedure. (An example might be a "~~t~~ransfinite induction" up to a fixed "number of the second number class".)

Yet if specific concepts and forms of inference from analysis, which are

after all also used in proofs of number-theoretical theorems, are to be included in elementary number theory then the consistency proof can in general not be extended to these additions in a straightforward way; here difficulties arise whose resolutions are still outstanding.

17.2. The consistency proof for elementary number theory can nevertheless be transferred without difficulty to a number of other branches of mathematics. This can be done quite generally in the case of such mathematical theories whose objects are given by a construction rule corresponding to that for the natural numbers (8.11). A particularly simple and in all cases applicable kind of such a rule is this: first a definite number of primitive symbols is given and it is then stated that each one of these symbols designates an object; if a primitive symbol is adjoined to the designation of an object then this results in another designation of an object. (In short: "Every finite sequence of primitive symbols designates an object of the theory".)

In such theories functions and predicates are then introduced by decidable definitions (8.12) and the same logical forms of inference are used as those in elementary number theory. The consistency proof carries over at once with the only difference that the place of the "numerals" is taken by the "object symbols" of the theory and this changes nothing in essence.

Such branches of mathematics are, for example, extensive parts of algebra (polynomials as objects are indeed finite combinations of symbols); from the realm of geometry, e.g., combinatorial topology; even large parts of analysis may be represented in this way if the concept of a real number is not used in its most general form. Finally important parts

of proof theory also belong here (cf. 16.1).

The addition of negative numbers, fractional numbers, diophantine equations, etc., to the natural numbers as objects in elementary number theory proper (3.31) can be incorporated in the consistency proof in the same way. All propositions about these objects can of course also be re-interpreted as propositions about the natural numbers, as mentioned at 3.31, by correlating these new objects in an appropriate way with the natural numbers. The same is also true of all other theories of the kind mentioned; a one-to-one correspondence can after all always be established between "finite combinations of symbols" and the natural numbers ("denumerability"). Yet this is unnecessarily cumbersome and unnatural for the requirements of the consistency proof.

17.3. (Cf. Paragraph 9). On the part of the intuitionists the following objection is raised against the significance of consistency proofs:⁽²⁵⁾ even if it had been demonstrated that the disputable forms of inference cannot lead to mutually contradictory results, these results would nevertheless be propositions without sense and their investigation therefore an idle pastime; real knowledge could be gained only by means of indisputable intuitionist (or finitist, as the case may be) forms of inference.

Let us, for example, consider the existential proposition cited at 10.6, for which the statement of a number whose existence is asserted is not possible. According to the intuitionist view this proposition is therefore without sense; an existential proposition can after all be sensibly asserted only if an numerical example is available.

What can we say to this?

Does such a proposition have any cognitive value? To be sure, a certain practical value of propositions of this kind lies first of all in the following possibility of application advanced by opponents of the intuitionist interpretation:

They might possibly serve as a source for the derivation of simple propositions, possibly representable by minimal formulae (3.24), which are themselves finitist and intuitionistically meaningful and which must be true by virtue of the consistency proof.

Furthermore, an existential proposition $\exists x F(x)$, e.g., for which no example is given, nevertheless serves the purpose of making a search for a proof for the proposition $\forall x \neg F(x)$ unnecessary; for there can be no such proof since a contradiction would otherwise result.

These are certainly reasons which make proofs of theorems by means of "actualist" forms of inference seem not entirely useless, apart from the "aesthetic value" of mathematical research as such.

Thus propositions of actualist mathematics seem to have a certain utility yet still no sense. The major part of my consistency proof, however, consists precisely in ascribing a finitist sense to actualist propositions, viz.: for every arbitrary proposition, as long as it is provable, a reduction rule according to 13.6 can be stated and this fact represents the finitist sense of the proposition concerned and this is gained precisely through the consistency proof.

This "finitist sense" can admittedly be rather complicated for even simply formed propositions and has in general a looser connection with the (actualistically determined) form of the proposition than is the case in the realm of finitist reasoning.

In this way the above mentioned existential proposition, e.g., also receives a finitist sense, yet this sense is weaker than that of a finitistically proven existential proposition, since it does not assert that an example can be given.

A quite different question is what significance can still be attached to the actualist sense of the propositions. The proof certainly reveals that it is possible to reason consistently "as though" everything in the infinite domain of objects were as actualistically determined as in finite domains (cf. Paragraph 9). Yet whether and in how far anything "real" corresponds to the actualist sense of a transfinite proposition - apart from what its restricted finitist sense expresses - is a question which the consistency proof does not answer.

NEW VERSION OF THE CONSISTENCY PROOF FOR ELEMENTARY NUMBER THEORY

In the following I shall present a new version of the consistency proof contained in Section IV of an earlier paper;⁽²⁶⁾ only this time the main emphasis will be placed on developing the fundamental ideas and on making every single step of the proof as lucid as possible. For this purpose I shall in places dispense with the explicit exposition of all details; viz., in those places where this is unimportant for the understanding of the context as a whole and where it can furthermore be supplied by the

reader himself without much difficulty.

Sections I and III of the earlier paper contained deliberations the knowledge of which need not be presupposed for an appreciation of the logic of the consistency proof, even though they are indispensable for the understanding of its purpose. In Section II, I had developed quite a detailed formalization of elementary number theory which preserved a close affinity with mathematical practice. This formalization is of great value now as ever; although a complete formal system could have been written down from the start it seems to me that by doing so an essential part of the context as a whole would have fallen by the wayside.

Added to this must be the fact that the formal representation of the forms of inference (Paragraph 5 of the earlier paper), which was directly assimilated to mathematical practice, with the characteristic notion of the "sequent", proves already quite suitable for meta-mathematical investigations, in fact, judging by my own experience, it is better suited to most purposes than the methods of representation generally customary to date.

Nevertheless it can not be said that the "most natural" logical calculus, simply because it corresponds most closely to real reasoning, is also the most suitable calculus for proof theoretical investigations. For the consistency proof, in particular, a somewhat different version has proved to be even more suitable and will therefore be adopted in this paper. I am referring to that formalization of the logical forms of inference which I had already developed in my dissertation⁽²⁷⁾ as the "LK-calculus". A knowledge of that paper is not however presupposed. I shall furthermore adopt only few basic concepts from Section II of the earlier consistency

proof and shall refer to them as such.

The constructivist proof of the "theorem of transfinite induction" (up to $\aleph_0$), article 15.4 of the earlier paper, is retained unchanged as the conclusion of the consistency proof and will not be revised for the time being; cf. the concluding remarks at the end of the present paper.

Paragraph 1

NEW FORMALIZATION OF NUMBER-THEORETICAL PROOFS.

I shall formulate the concepts involved and in each case add some explanatory remarks.

1.1. "Formula".

The definition of a formula is adopted from the earlier paper (Article 3.2), yet with the following simplification:

Only 1 is used as a numeral. Functions are not admitted (cf. however the concluding remarks) with the exception of a single one, the successor function, which is denoted by a prime: a' has the same intuitive meaning as $a + 1$. By means of this function symbol the natural numbers can now be represented by 1, $1'$, $1''$, $1'''$ etc. Terms are therefore now always of the form 1 or $1'$ or $1''$ etc. or a or a' or a'' etc., where a stands for an arbitrary free variable. The former we call numerical terms and they therefore correspond to the earlier numerals; the latter variable terms. Predicate symbols are admitted as before according to need; it is required only (Article 13.3 of the earlier paper) that they are decidably defined, i.e., that it can be decided of every individual

natural number whether the predicate does or does not hold. On the basis of these concepts of terms and predicates the old definition of a formula (Article 3.23) is now preserved, yet the logical connective $\supset$ will no longer be used. This represents no significant restriction since it is well known that $\supset$ can be replaced by $\wedge$ and $\neg$ or $\vee$ and $\neg$. In addition we could still eliminate $\vee$ and $\exists$, as was done in the earlier paper (Paragraph 12); yet this is unnecessary since by being in exact correspondence with $\wedge$ and $\vee$ these connectives cause no difficulties whatever in the "LK-calculus".

Example of a formula:

$$\forall x (x > 1 \wedge \exists y (y''' = x))$$

where x is a free variable and y and y are bound variables.

Three simple auxiliary concepts will still be needed below:

A prime formula is a formula containing no logical connectives.

Example: $x''' = 1''$.

The terminal connective of a formula which is not a prime formula is that logical connective which is added last in the construction of that formula (according to Article 3.23 of the earlier paper).

The degree of a formula is the total number of logical connectives occurring in it.

Examples: A prime formula has degree 0. The formula

$\forall x (x > 1 \wedge \exists y (y''' = x))$ has degree 3 and its terminal connective is the $\vee$.

1.2. "Sequent".

A sequent is an expression of the form

$$\mathcal{U}_1, \mathcal{U}_2, \dots, \mathcal{U}_\mu \longrightarrow \mathcal{B}_1, \mathcal{B}_2, \dots, \mathcal{B}_\nu$$

where arbitrary formulae may take the place of

$$\mathcal{U}_1, \mathcal{U}_2, \dots, \mathcal{U}_\mu, \mathcal{B}_1, \mathcal{B}_2, \dots, \mathcal{B}_\nu$$

The $\mathcal{U}$'s are called the antecedent formulae, the $\mathcal{B}$'s the succedent formulae of the sequent. Both the antecedent and the succedent parts of the sequent may be empty.

Suppose that it is known of each antecedent and succedent formula of a sequent without free variables whether it is "true" or "false". Then the sequent is "false" if all of its antecedent formulae are true and all of its succedent formulae are false. (Moreover, a sequent which has neither antecedent nor succedent formulae is also false.) In every case the sequent is "true".

Explanatory Remarks. We shall make use of the definition of "true" and "false" only in connection with the concept of the "basic sequent" and here the $\mathcal{U}$ and $\mathcal{B}$ will be prime formulae without free variables and therefore immediately decidable. In general the concept of the "truth" of a formula is of course not formally defined at all. The definition can nevertheless serve quite generally to explain the intuitive meaning of a sequent, but it should still be added that a sequent with free variables is considered to be true if and only if every arbitrary substitution of numerals for free variables yields a true sequent. The intuitive meaning of a sequent without free variables can be expressed briefly as follows:

"If the assumptions $\mathcal{A}_1, \dots, \mathcal{A}_n$ hold then at least one of the propositions $\mathcal{B}_1, \dots, \mathcal{B}_m$ holds."

In the earlier paper I had introduced the concept of a sequent with only one succedent formula for the immediate purpose of providing a natural representation of mathematical proofs (Paragraph 5). Considerations of this kind may in fact also lead to the new symmetric concept of a sequent in situations where the aim is a particularly natural representation of case distinctions (cf. Paragraph 4 of the earlier paper and 5.26 in particular). For a $\vee$ -elimination can now be represented simply as follows: From $\rightarrow \mathcal{A} \vee \mathcal{B}$ we infer $\rightarrow \mathcal{A}, \mathcal{B}$, which reads: "There exists the two possibilities $\mathcal{A}$ as well as $\mathcal{B}$." Yet it must be admitted that this new concept of a sequent in general already constitutes a departure from the "natural" and that its introduction is primarily justified by the considerable formal advantages exhibited by the representation of the forms of inference following below which this concept makes possible.

It should still be pointed out that the intuitive sense of a sequent is to be considered to coincide with the given definition in those cases in which the sequent possesses no antecedent formulae or no succedent formulae: if there are no antecedent formulae, the sequent expresses the fact that any one of the propositions $\mathcal{B}_1, \dots, \mathcal{B}_m$ holds, this time independently of any assumptions. If there are no succedent formulae, the sequent expresses the fact that on the basis of the assumptions $\mathcal{A}_1, \dots, \mathcal{A}_n$ no possibility remains open, i.e.: the assumptions are incompatible, they lead to a contradiction.

A sequent without antecedent and succedent formulae, the "empty sequent", therefore indicates that without any assumptions at all a contradiction results, i.e.: if this sequent is derivable in a system then that system itself is inconsistent.

Example of a sequent.

$$\forall x (x' > 1) \longrightarrow a > 1 \vee a = 1, 1' > 1, 1'' = 1$$

where a and b are free variables and x a bound variable.

1.3. "Inference Figure."

An inference figure (the formal counterpart of an inference) consists of a line of inference, a lower sequent, written below the line, and upper sequents (one or more), written above the line. The lower sequent here stands for the conclusion of the inference which has been drawn from the premises represented by the upper sequents.

The only inference figures admitted into our formalism are those obtainable from one of the following twenty inference figure schemata by a substitution of the following kind:

Any formulae may be put in place of $\mathcal{U}$, $\mathcal{B}$, $\mathcal{D}$, $\mathcal{E}$;

$\forall x F(x)$ or $\exists x F(x)$ may be replaced by any formula of this form; furthermore $F(a)$ or $F(t)$ may be replaced by that formula which results from $F(x)$ by the substitution of an arbitrary free variable a or an arbitrary term t for the bound variable x .

Γ , Δ , $\mathbb{H}$ and Λ may be replaced by arbitrary, possibly empty sequences of formulae, separated by ~~comma~~ commas.

The following restriction on variables is to be observed: the free variable designated by a - which we call the eigen-variable of the inference figure concerned - may not occur in the lower sequent of this inference figure.

The Inference Figure Schemata:

1.31. Schemata for Structural Inference Figures:

$$\begin{array}{l}
 \text{Thinning:} \quad \frac{\Gamma \rightarrow \mathcal{O}}{\mathcal{J}, \Gamma \rightarrow \mathcal{O}} \qquad \frac{\Gamma \rightarrow \mathcal{O}}{\Gamma \rightarrow \mathcal{O}, \mathcal{D}} \\
 \\
 \text{Contraction:} \quad \frac{\mathcal{J}, \mathcal{J}, \Gamma \rightarrow \mathcal{O}}{\mathcal{J}, \Gamma \rightarrow \mathcal{O}} \qquad \frac{\Gamma \rightarrow \mathcal{O}, \mathcal{D}, \mathcal{D}}{\Gamma \rightarrow \mathcal{O}, \mathcal{D}} \\
 \\
 \text{Interchange:} \quad \frac{\mathcal{A}, \mathcal{D}, \mathcal{E}, \Gamma \rightarrow \mathcal{O}}{\mathcal{A}, \mathcal{E}, \mathcal{D}, \Gamma \rightarrow \mathcal{O}} \qquad \frac{\Gamma \rightarrow \mathcal{O}, \mathcal{D}, \mathcal{E}, \mathcal{A}}{\Gamma \rightarrow \mathcal{O}, \mathcal{E}, \mathcal{D}, \mathcal{A}} \\
 \\
 \text{Cut:} \quad \frac{\Gamma \rightarrow \mathcal{O}, \mathcal{D} \quad \mathcal{D}, \mathcal{A} \rightarrow \wedge}{\Gamma, \mathcal{A} \rightarrow \mathcal{O}, \wedge}
 \end{array}$$

The two formulae in the last schema designated by $\mathcal{D}$ are called cut formulae, their degree the degree of the cut.

1.32. Schemata for Operational Inference Figures:

$$\mathcal{K}: \left\{ \begin{array}{l} \frac{\Gamma \rightarrow \mathcal{O}, \mathcal{U} \quad \Gamma \rightarrow \mathcal{O}, \mathcal{V}}{\Gamma \rightarrow \mathcal{O}, \mathcal{U} \wedge \mathcal{V}} \qquad \frac{\mathcal{U}, \Gamma \rightarrow \mathcal{O}}{\mathcal{U} \wedge \mathcal{V}, \Gamma \rightarrow \mathcal{O}} \qquad \frac{\mathcal{V}, \Gamma \rightarrow \mathcal{O}}{\mathcal{U} \wedge \mathcal{V}, \Gamma \rightarrow \mathcal{O}} \end{array} \right.$$

$$\begin{array}{l}
 V: \left\{ \frac{\mathcal{U}, T \rightarrow \textcircled{0} \quad \mathcal{V}, T \rightarrow \textcircled{0}}{\mathcal{U} \vee \mathcal{V}, T \rightarrow \textcircled{0}} \quad \frac{T \rightarrow \textcircled{0}, \mathcal{U}}{T \rightarrow \textcircled{0}, \mathcal{U} \vee \mathcal{V}} \quad \frac{T \rightarrow \textcircled{0}, \mathcal{V}}{T \rightarrow \textcircled{0}, \mathcal{U} \vee \mathcal{V}} \right. \\
 \\
 \forall: \left\{ \frac{T \rightarrow \textcircled{0}, F(a)}{T \rightarrow \textcircled{0}, \forall x F(x)} \quad \frac{F(z), T \rightarrow \textcircled{0}}{\forall x F(x), T \rightarrow \textcircled{0}} \right. \\
 \\
 \exists: \left\{ \frac{F(a), T \rightarrow \textcircled{0}}{\exists x F(x), T \rightarrow \textcircled{0}} \quad \frac{T \rightarrow \textcircled{0}, F(z)}{T \rightarrow \textcircled{0}, \exists x F(x)} \right. \\
 \\
 \neg: \left\{ \frac{\mathcal{U}, T \rightarrow \textcircled{0}}{T \rightarrow \textcircled{0}, \neg \mathcal{U}} \quad \frac{T \rightarrow \textcircled{0}, \mathcal{U}}{\neg \mathcal{U}, T \rightarrow \textcircled{0}} \right.
 \end{array}$$

That formula in the schema which contains the logical connective is called the principal formula of the operational inference figure concerned.

1.33. Schema for CJ - Inference Figures (the formal counterpart of complete inductions):

$$\frac{F(a), T \rightarrow \textcircled{0}, F(a')}{F(1), T \rightarrow \textcircled{0}, F(z)}$$

The degree of the CJ - Inference Figure is the degree of that formula in the schema which is designated by $F(1)$ - and which is, of course, the same as that of $F(a)$, $F(a')$ and $F(z)$.

Example of an Inference Figure:

$$\frac{\rightarrow \underline{a}' = 1', \quad 1 < 1'' \text{ \& } \underline{a} = 1''}{\rightarrow \underline{a}' = 1', \quad \exists \underline{z} (1 < \underline{z} \text{ \& } \underline{a} = 1'')}$$

where $\underline{a}$ is a free, and $\underline{z}$ a bound variable.

Explanatory remarks about the inference figure schemata will follow below in connection with the concept of a derivation.

1.4. "Basic Sequents."

We shall distinguish "logical" and "mathematical" basic sequents.

A logical basic sequent is a sequent of the form $\mathcal{D} \rightarrow \mathcal{D}$,
where an arbitrary formula may stand for $\mathcal{D}$.

A mathematical basic sequent is a sequent consisting entirely of prime formulae and becoming a "true" sequent with every arbitrary substitution of numerical terms for possible occurrences of free variables.

The "truth" of a prime formula without free variables is, according to our assumption of the decidability of all predicates, always verifiable. Whether or not a sequent with free variables is a mathematical basic sequent is of course not generally decidable; nor is this actually essential.

Examples of Basic Sequents:

$$\begin{aligned}
 \forall \underline{x} \exists \underline{y} (\underline{x}'' = \underline{a} \wedge \underline{y} > \underline{x}) &\longrightarrow \forall \underline{x} \exists \underline{y} (\underline{x}'' = \underline{a} \wedge \underline{y} > \underline{x}) \\
 \underline{a} = \underline{b}, \underline{b} = \underline{c} &\longrightarrow \underline{a} = \underline{c} \\
 \underline{a} < 1 &\longrightarrow \\
 &\longrightarrow 1' > 1 \\
 \underline{a} = \underline{b} &\longrightarrow \underline{a} = \underline{b} \\
 &\longrightarrow \underline{a}' > \underline{a} \\
 &\longrightarrow 1''' = 1 \pmod{1''}
 \end{aligned}$$

1.5. "Derivation."

A derivation is a figure in tree form consisting of a number of sequents (at least one) with one lowest sequent, the end-sequent, and certain uppermost sequents which must be basic sequents; the connection between the uppermost sequents and the end-sequent is established by inference figures.

It should be intuitively obvious how this is meant; yet let me paraphrase it again as follows: suppose that an end-sequent is given. This sequent is either already an uppermost sequent - in which case it alone constitutes at once the entire derivation - or it is the lower sequent of an inference figure. Every upper sequent of this inference figure is in turn either an uppermost sequent of the derivation or the lower sequent of a further inference figure, etc.

The reader should always visualize a derivation quite intuitively as a tree-like structure, then the transformations on a derivation to be performed in Paragraph 3 become most easily intelligible.

Example of a Derivation:

$$\begin{array}{c}
 \begin{array}{c} \rightarrow I = I \\ \hline \rightarrow \underline{I} = \underline{I} \\ \hline \rightarrow \forall \underline{x} (\underline{x} = \underline{x}) \end{array} \quad \begin{array}{c} \underline{a} = \underline{a} \rightarrow \underline{a}' = \underline{a}' \\ \hline I = I \rightarrow \underline{a} = \underline{a}' \\ \hline \text{Cut} \end{array} \quad \begin{array}{c} \underline{I}'' = \underline{I}''' \rightarrow \underline{I}''' = \underline{I}''' \\ \hline \forall \underline{x} (\underline{x} = \underline{x}) \rightarrow \underline{I}''' = \underline{I}''' \\ \hline \text{Cut} \end{array} \\
 \hline
 \rightarrow \underline{I}''' = \underline{I}'''
 \end{array}$$

CT-inference figure

CT

CT-inference figure

CT

For a further example, cf. 1.6.

Another auxiliary concept which will be needed later:

A path in a derivation is, briefly speaking, a sequence of sequents which must be followed in descending from an individual uppermost sequent to the end-sequent. At each step the path leads via one of the upper sequents of an inference figure to the lower sequent of that inference figure.

It is furthermore immediately obvious what is meant by the following:
a sequent in the derivation stands above or below another sequent occurring in the same path (i.e. not only immediately above or below it, but any number of steps apart). It is understood that wherever the notion of "above" or "below" is used, the sequents concerned belong to a common path; otherwise the concept is meaningless.

1.6. Explanatory Remarks about the new Formalization of Number-Theoretical Proofs.

As a result of the revised concept of a derivation a formalization of number-theoretical proofs is given which distinguishes itself from my earlier "natural" version mainly in two points. First: the rules of inference belonging to the logical connectives, i.e. the "introduction" and "elimination" of a logical connective, have now been re-formulated throughout in such a way that the lower sequent always contains the "principal formula" whereas the upper sequents contain the associated side formulae. To the earlier "introduction" of a logical connective now corresponds the occurrence of that connective in a succedent formula of the lower sequent, to the "elimination" of a logical connective corresponds the occurrence of that connective in an antecedent formula of the lower sequent. The reader should convince himself of the equivalence of the old and new versions by examining, for example, the $\vee$ -rules of inference (disregarding, for the time being, the occurrence of several succedent formulae). The "cut" and the logical basic sequents must be used in the proof of equivalence. Cf. the derivation with the $\vee$ -introduction on the left and the subsequent " $\vee$ -elimination", given as an example in 1.5.

This part of the conversion from the former to the new rules of inference amounts to an abandoning of the natural succession of the propositions in number-theoretical proofs and to the introduction in its place of an artificial arrangement of the propositions along special lines with the result that in operational inferences the simpler proposition now always comes first and is followed by the more complex proposition, viz., the proposition with the additional connective.

This re-arrangement proves of practical value for the consistency proof because of the essential role which the concept of the complexity of a derivation and, with it, the complexity of a particular formula (which increases as the degree of the formula increases) plays in the consistency proof.

The second important distinction vis-à-vis the old concept of a derivation consists in the symmetrization of the sequents by the admission of arbitrarily many succedent formulae. This makes it unfortunately more difficult to grasp the intuitive meaning of the various inference schemata and to persuade oneself of their "correctness". To overcome this difficulty the reader should first conceive of the presence of only one succedent formula and should then convince himself that the inference remains correct even if several succedent formulae occur and also if no succedent formula occurs. As the reader becomes more familiar with this concept of a derivation he should be able to realize that transformations of derivations and other proof-theoretical investigations can be carried out particularly simply and elegantly with this concept. The decisive advantages are these:

There exists a complete symmetry between $\&$ and $\vee$, $\vee$ and $\exists$. All of the connectives $\&$, $\vee$, $\vee$, $\exists$ and $\neg$ have, to a great extent, equal status in the system; no connective ranks notably above any other connective. The special position of the negation, in particular, which constituted a troublesome exception in the natural calculus (cf. Articles 4.56 and 5.26 of the earlier paper), has been completely removed in a seemingly magical way. The manner in which this observation

is expressed is undoubtedly justified since I myself was completely surprised by this property of the "LK-calculus" when first formulating that calculus. The "law of the excluded middle" and the "elimination of the double negation" are implicit in the new inference schemata - the reader may convince himself of this by deriving both of them from the new calculus - yet they have become completely harmless and no longer play the least special role in the consistency proof that follows.

If we think of the $\mathcal{T}$, Δ , $\oplus$, $\wedge$ as removed from the inference figure schemata we see that the schemata are of the greatest simplicity and likeness in the sense that none of them any longer contains anything that is not absolutely essential; the $\mathcal{T}$, Δ , $\oplus$, $\wedge$ constitute an appendage which signifies merely that additional antecedent and succedent formulae are carried forward unchanged from the upper sequent to the lower sequent.

The new formulation of the concept of the "mathematical basic sequent" still requires an explanation. In the earlier paper this concept was defined differently (Articles 5.23 and 10.14). It turns out however that the former basic sequents are derivable in the new system. An example which typifies the general aspects of the situation may make this clear:

The following "mathematical basic sequent" in the old sense

$$\longrightarrow \forall \underline{x} \forall \underline{y} \neg (\underline{x} = \underline{y} \& \neg \underline{y} = \underline{x})$$

is derivable thus:

$$\begin{array}{c}
 \underline{a = b} \longrightarrow \underline{b = a} \\
 \hline
 \underline{a = b} \wedge \neg \underline{b = a} \longrightarrow \underline{b = a} \\
 \hline
 \neg \underline{b = a}, \underline{a = b} \wedge \neg \underline{b = a} \longrightarrow \\
 \underline{a = b} \wedge \neg \underline{b = a}, \underline{a = b} \wedge \neg \underline{b = a} \longrightarrow \\
 \hline
 \underline{a = b} \wedge \neg \underline{b = a} \longrightarrow \\
 \hline
 \longrightarrow \neg (\underline{a = b} \wedge \neg \underline{b = a}) \\
 \hline
 \longrightarrow \forall y \neg (\underline{a = y} \wedge \neg \underline{y = a}) \\
 \hline
 \longrightarrow \forall x \forall y \neg (\underline{x = y} \wedge \neg \underline{y = x})
 \end{array}$$

All usual "mathematical basic sequents" in the old sense are derivable in the same way from intuitively synonymous mathematical basic sequents in the new sense.⁽²⁸⁾ The fact that the new concept of a derivation is actually equivalent with that of the earlier paper - apart from the restriction which results from the initially limited admission of functions in the new system - can be verified without great difficulty from the observations made above and I shall discuss it no further.⁽²⁹⁾

Paragraph 2

SURVEY OF THE CONSISTENCY PROOF

It is to be shown that every derivation is consistent;⁽³⁰⁾ this may be paraphrased by saying that no derivation has an empty end-sequent.

For from a contradiction, $\longrightarrow \mathcal{U}$ and $\longrightarrow \neg \mathcal{U}$ we can first of all derive the sequents $\longrightarrow \neg \mathcal{U}$ and $\neg \mathcal{U} \longrightarrow$ and from them, by means of a cut, the empty sequent. (Conversely, from the empty sequent every arbitrary sequent can be derived by "thinnings".)

It makes sense that we should begin by proving the consistency of simple derivations, then of more complex ones, using the consistency of the simpler derivations, and so forth. We thus proceed "inductively". It is furthermore not implausible that this procedure repeatedly requires the examination of an already infinite sequence of derivations before a more complex class can be tackled; for example, first all derivations consisting of only one sequent, then all derivations consisting of two sequents, etc. Yet this means actually that we are applying a "transfinite induction". The pattern of this analysis is in practice of course considerably more involved than in the case of the given example.

The proof is carried out in three stages:

1. The consistency of an arbitrary derivation is reduced to the consistency of all "simpler" derivations. This is done by defining an - unequivocal - reduction step for arbitrary "inconsistent derivations", i.e. derivations with the empty sequent as end-sequent; this step transforms such a derivation into a "simpler" derivation with the same end-sequent. The definition of this reduction step forms the contents of Paragraph 3.
2. Then a transfinite ordinal number is correlated with every derivation and it is shown that in a reduction step the inconsistent derivation concerned is turned into a derivation with a smaller ordinal number. In this way the so far only loosely determined concept of "simplicity" receives its precise sense: the larger the ordinal number of a derivation the greater is its "complexity" in the context of this consistency proof. This is the contents of Paragraph 4.

3. From this the consistency of all derivations then obviously follows by "transfinite induction". The inference of transfinite induction, which is still a rather "disputable" inference up to this point, may not be presupposed in the consistency proof nor proved as in set theory. This inference requires rather a separate justification by means of indisputable "constructive" forms of inference. At the end of Paragraph 4 the reader is at this time referred to the earlier paper in this connection.

Paragraph 3

A REDUCTION STEP ON AN INCONSISTENT DERIVATION

3.1. Underlying Ideas.

Suppose that a derivation is given whose end-sequent is the empty sequent. This derivation is to be transformed into a (in some sense) simpler derivation with the same end-sequent. What is here meant by "simpler" can at present only be stated roughly and will be made precise later through the ordinal numbers.

What are the considerations that make us suspect at all that, given a proof for a contradiction, there already exists an even simpler way of proving such a contradiction? By a contradiction is meant a proposition of a quite simple structure, for example " $1=2$ ". If such a simple proposition can be proved by means of a complex proof it is reasonable to suspect that the proof can be simplified. The following argument might conceivably be used: Somewhere in the proof there must after all occur a proposition of maximal complexity. In that case it must be

) assumed that this "complexity extremum" (in the formalization of the proof this might be a formula with the highest degree occurring in the derivation) must somehow be "reducible". The occurrence of this proposition is in the general case conceivable only in this way that the proposition is introduced into the proof by the inference of the "introduction" of its terminal connective and is then used again through the inference of the "elimination" of precisely this connective. Yet if a connective is first introduced and then again eliminated it can be left out altogether by direct passage from the preceding sub-propositions to the corresponding succeeding sub-propositions. (31)

This is the basic idea underlying the "operational reduction" to be outlined below. In actual fact, however, the situation will turn out not to be quite as simple as assumed in the argument just sketched. One of the difficulties that may arise is the occurrence of a complete induction in the proof; viz., in the case where the proposition with the maximal number of connectives in question is not directly proved by an "introduction" inference but rather by a complete induction. This requires a further special kind of reduction step which will be called a "CJ-reduction". The form of this reduction step is extremely simple and precisely what we would expect: if the term in the schema of the CJ-inference figure is a numerical term, thus denoting an individual number, the complete induction can naturally be replaced by a number of ordinary inferences - in our formalization a number of "cuts". This constitutes the "CJ-reduction".

If a CJ-inference figure occurs in the derivation whose t is a variable

term - and this is in fact normally the case - then this figure cannot of course be reduced immediately in this way. Yet the reduction procedure may be arranged in such a way that with successive reduction steps more and more variable terms are gradually replaced by numerical terms so that eventually even initially irreducible CJ-inference figures become in turn reducible. This remark is incidental. Here we are actually concerned only with the definition of one single reduction step so that regardless of the nature of the given inconsistent derivation at least one place can be found in it to which a reduction can be applied.

Let us suppose therefore that there is no place in the derivation in which a CJ-reduction can be carried out. Then, as will be shown in detail below, a "operational reduction" is always feasible. On the other hand, it cannot be expected that a formula of highest degree in the entire derivation is always amenable to reduction. As mentioned before, this formula may have been introduced by a CJ-inference figure and this figure can contain a variable t . It is nonetheless possible in each case to locate a formula in the derivation which represents a "relative extremum", viz., a formula which is introduced by the introduction of its terminal connective and whose further use in the derivation then consists in the elimination of that connective, and which is therefore reducible. Why such a formula must always exist is best seen within the context of the proof following below (3.43).

The following phenomenon should still be pointed out: it may for example happen that the formula which is intended to form the starting point of the operational reduction is used again in the derivation not

only once but several times. (An example: suppose that the formula has the form $\forall x F(x)$, and from it are inferred $F(1)$ and $F(1'')$ or in another place perhaps even $\forall x F(x) \vee \mathcal{U}$.)

In the general case all that can be achieved is that in one place of application the formula is used in the form of an elimination of its terminal connective. About the remaining places nothing can be said. In this general case the formula can therefore not be reduced away completely; we can merely bring about a simplification of this one place of application which at this point makes the passage via the formula redundant. The occurrence of this formula in the remaining places must nevertheless remain unaffected. It turns out that this is sufficient.

These preliminaries have been carried out against the background of the "natural proof" with the natural succession of the individual propositions. For their application to our formalism developed in Paragraph 1, a corresponding translation must be made: To the "introduction" of a connective here corresponds its occurrence in a succedent formula of the lower sequent, to the "elimination" of that connective its occurrence in an antecedent formula of the lower sequent of the operational inference figure. All other details will follow from the precise formal development now to be carried out; the preliminaries ought not and cannot of course do more than indicate to the reader in a superficial way the main ideas of the procedure and in doing so facilitate the understanding of the actual presentation.

3.2. Elimination of redundant free variables in preparation of the reduction step. - The "ending".

We begin with the definition of the "reduction step on an inconsistent derivation" by stipulating: before the reduction step proper the following simple transformation must be carried out:

All free variables in the derivation are replaced by the numeral 1; excepted from this is however every eigen-variable (1.3.) of an inference figure in all derivation sequents occurring above the lower sequent of the inference figure concerned.

What is the effect of this preliminary step? Actually, a free variable normally serves as eigen-variable of an inference figure and may here occur only above the lower sequent of this inference figure; its occurrence in the lower sequent itself is of course also expressly forbidden by the restriction on variables (1.3). Wherever else free variables may thus still occur they are completely redundant and can equally well be replaced by 1. It is fairly obvious that this leaves the derivation correct. The empty end-sequent remains of course unchanged.

We furthermore require a simple auxiliary concept - the ending of a derivation - which is defined thus: the ending consists of all those derivation sequents that are encountered if we ascend each individual path (1.5) from the end-sequent and stop as soon as we arrive at the line of inference of an operational inference figure. Thus the lower sequent of this inference figure in each case still belongs to the ending but its upper sequents do so no longer. If a path crosses no line of inference of an operational inference figure at all then it is of course completely included in the ending.

$$\frac{\mathcal{F}(\alpha), \mathcal{T} \longrightarrow \mathcal{M}, \mathcal{F}(\alpha')}{\mathcal{F}(1), \mathcal{T} \longrightarrow \mathcal{M}, \mathcal{F}(\kappa)}$$

Among the inference figures, the ending obviously contains only structural and CJ-inference figures.

We now distinguish two cases:

1. The ending of our inconsistent derivation contains at least one CJ-inference figure. In that case a CJ-reduction is carried out, Cf. 3.3.
2. The ending contains no CJ-inference figure. In that case an operational reduction is carried out (3.5) after a further preparatory step (3.4).

3.3. The CJ-reduction.

If the ending of the given inconsistent derivation contains at least one CJ-inference figure after the stated preparatory step, then the reduction step proper consists in the transformation of the derivation described next.

We select a CJ-inference figure in the ending which is such that it occurs above no other CJ-inference figure. (i.e.: the derivational path which goes through the lower sequent of the selected CJ-inference figure must not cross the line of inference of any CJ-inference figure between that sequent and the end-sequent.) In order to make the reduction step unambiguous an appropriate procedure for the unique determination of the CJ-inference figure to be selected must still be given; there is a simple way in which this can be done.

The CJ-inference figure has the form:

$$\frac{F(\alpha), T' \longrightarrow \textcircled{A}, F(\alpha')}{F(1), T \longrightarrow \textcircled{A}, F(\kappa)}$$

where n designates a numerical term. For by virtue of the preparations made no variable term could here possibly occur; in fact the lower sequent cannot contain a single free variable: after the preparatory step free variables can occur only above inference figures with one eigen-variable and no such figure occurs below our CJ-inference figure. Indeed, the section of the derivational path between the lower sequent and the end-sequent of this figure from here on crosses only lines of inference of structural inference figures.

This CJ-inference figure is now replaced by a system of structural inference figures of the following kind:

$$\begin{array}{c}
 \frac{\frac{\frac{F(1), \pi \rightarrow \odot, F(1)}{\quad} \quad \frac{F(1'), \pi' \rightarrow \odot, F(1'')}{\quad}}{\frac{F(1), \pi, \pi' \rightarrow \odot, \odot, F(1'')}{\quad}}}{\frac{F(1), \pi \rightarrow \odot, F(1'') \quad F(1'), \pi' \rightarrow \odot, F(1''')}{\quad}}}{\frac{F(1), \pi, \pi' \rightarrow \odot, \odot, F(1'')}{\quad}}}{\frac{F(1), \pi \rightarrow \odot, F(1'')}{\quad}} \\
 \vdots \\
 \frac{F(1), \pi \rightarrow \odot, F(n)}{\quad}
 \end{array}$$

Above the sequents $F(1), \pi \rightarrow \odot, F(1')$ and $F(1'), \pi' \rightarrow \odot, F(1'')$ etc. we write in each case that

section of the derivation which precedes $F(a), \pi \rightarrow \odot, F(a')$,

where we replace the free variable a in the entire section - except in the case where it at the same time happens to be the eigen-variable

of an inference figure occurring in that section, in all sequents occurring above the lower sequent of that inference figure - by the numerical terms l or l' or l'' etc. From the sequent

$\mathcal{F}(1), \mathcal{T} \rightarrow \mathcal{F}(n)$ downwards the ending is finally continued by

adjoining the unchanged remainder of the old derivation. To put it precisely: All derivational paths which did not go through this sequent have been preserved unchanged and those which did go through it remain unchanged from the end-sequent up to this point.

If n is equal to 1, then the reduction proceeds somewhat differently: in that case the lower sequent of the CJ-inference figure runs

$\mathcal{F}(1), \mathcal{T} \rightarrow \textcircled{0}, \mathcal{F}(1)$. This sequent is derived from the logical basic sequent $\mathcal{F}(1) \longrightarrow \mathcal{F}(1)$ by thinnings and interchanges, as required.

Whatever preceded this lower sequent in the derivation is omitted; everything else is retained unchanged, as in the general case.

It is easily seen that in the CJ-reduction step the given inconsistent derivation is in all parts transformed into another correct inconsistent derivation. All we need to realize here in essence is that the replacement of α by a numerical term turns every inference figure into another correct inference figure.

Comments about the nature of the reduction step should no longer be required; as stated at 3.1, its intuitive meaning is exceedingly simple: a complete induction up to a definite number is replaced by a corresponding number of ordinary inferences.

3.4. Preliminaries and preparatory step for an operational reduction.

We must now deal with the case where the inconsistent derivation contains no CJ-inference figure in its ending after the preparatory step 3.2.

The "operational reduction" to be carried out in this case is initiated by a further preparatory step (3.42) whose purpose it is to eliminate all possible occurrences of thinnings and logical basic sequents from the ending since these would otherwise give rise to bothersome exceptions in the actual operational reduction.

For this purpose and also for the sake of its further use we must first examine the structure of the ending more closely.

3.41. The ending of our derivation contains only structural inference figures. Its uppermost sequents are the uppermost sequents of the entire derivation or the lower sequents of operational inference figures. The ending contains no free variables (since it contains no inference figures with eigen-variables). This is all quite obvious.

We now introduce two simple auxiliary concepts:

Equal sequent formulae in the upper sequents and the lower sequent of a structural inference figure corresponding to one another according to the inference figure schema will be called clustered.

Clustered are, for example, the three formulae designated by $\mathcal{D}$ in the schema of a contraction, likewise the first of the formulae designated by π in the upper sequent and the first of the formulae designated by π in the lower sequent, the second formula and the second formula, etc.; the two cut formulae of a cut are clustered; etc.

The totality of all formulae in the ending of the derivation obtained by starting with a particular formula and collecting all of its clustered formulae, then all formulae clustered with these etc. is called a formula cluster; we can also say: the cluster associated with the relevant first formula.

About the form of this cluster we can say the following:

With every cluster is associated a cut in the sense that its cut formulae belong to the cluster. This is so since every formula which occurs somewhere in the ending, as is evident from the structural inference figure schemata, is always clustered with a formula in the sequent standing immediately below it, except when it is a cut formula. Since the end-sequent of our derivation is empty, we must at some point reach such a cut in tracing a cluster downwards towards the end-sequent.

We now start with this cut and trace the location of the cluster upwards from the two cut formulae belonging to the cluster. With the following result: That portion of the cluster which is obtained by starting with the left cut formula - we call it the left side of the cluster - is in tree-form; a branching takes place if in coming from below we reach a contraction whose $\circ$ belongs to the cluster; a branch may terminate at some point if the $\circ$ of a thinning or the uppermost sequent of the ending is reached; in that case we speak of an uppermost formula of the cluster. All formulae of the left side of the cluster are succedent formulae of the sequents concerned. Exactly analogous remarks apply to the right side of the cluster obtained by starting with the right cut formula; it too is in tree-form, etc., all its formulae are antecedent

formulae. It follows further that no cut formulae other than the two formulae from which we started belong to the cluster; hence the cut associated with a cluster is uniquely determined and so are therefore the concepts of the left side and the right side of the cluster. No formulae of the cut other than the cut formulae belong to the cluster. All formulae belonging to the cluster occur above the lower sequent of the cut. (i.e.: all sequents containing cluster formulae occur above that sequent.) The left and right sides together therefore constitute the whole cluster.

The correctness of all these assertions is easily seen by tracing the cluster mentally from the cut formulae upwards and by visualizing with the help of the schemata of the structural inference figures the kinds of procedure which alone lead to new clustered formulae.

3.42. We can now turn to the preparatory step for the operational reduction which, as said earlier, is intended to accomplish the elimination of all thinnings and logical basic sequents from the ending. This can clearly be done. After all, a "thinning" represents only a weakening of the intuitive sense of a sequent; if a contradiction can be derived from the weakened sequent the same can obviously also be derived from the stronger upper sequent alone; and a logical basic sequent, being a pure tautology, is also dispensable in the context of mere structural transformations.

The procedure almost suggests itself. Let us begin with the thinnings: We select a thinning above which - in the ending - no other thinning

occurs. We then simply cancel its lower sequent and then use the upper sequent in its place. In order to leave the derivation correct we continue downwards and in the next lower sequent cancel the formula clustered with the formula $\mathcal{D}$ in the thinning as well as the formula clustered with the latter in the subsequent lower sequent, etc. Can this procedure lead to new difficulties? Actually, a contraction may arise in which a $\mathcal{D}$ of the upper sequent is to be cancelled. All the better, the upper sequent becomes then equal with the lower sequent; the contraction becomes redundant and we have finished. There may be other occasions in the procedure in which the upper and lower sequents of an inference figure become equal; in that case we simply omit the inference figure and write the sequent down only once. If we encounter a cut in which the formula to be cancelled is a cut formula we cancel the other upper sequent of the cut together with whatever stands above it and derive the lower sequent from the remaining upper sequent alone by thinnings and interchanges (as far as necessary).

The new thinnings which arise are again eliminated by the same procedure. That this procedure terminates, thus ridding the ending of thinnings completely, follows from the fact that with each reduction step we find ourselves lower down in the derivation (measured in terms of the total number of cuts up to the end-sequent, for example).

We leave it to the reader to give an exact demonstration of the feasibility of the indicated procedure as well as to formulate it unambiguously; this presents no essential difficulties.

Next we eliminate the logical basic sequents: In the ending such a sequent can now occur only as the upper sequent of a cut since no contractions and interchanges are applicable to it; the lower sequent of the cut is therefore, as is easily seen, equal to the other upper sequent. We therefore simply omit the cut and have thus finished.

As a result we finally obtain an inconsistent derivation whose ending has the same properties as those stated above with the additional property of containing no thinnings and no logical basic sequents (as uppermost derivation sequents).

3.43. Further Preliminaries to the Operational Reduction.

I now assert: There exists at least one formula cluster in the ending of our derivation, with at least one uppermost formula both on its left side and on its right side, which is the principal formula of an operational inference figure.

At this point the connection between our formal procedure and the fundamental ideas sketched in 3.1 becomes apparent: the concept of the formula cluster makes it possible for us to grasp in its entirety the collection of all occurrences of a "proposition" in the "proof" (i.e.: formula in the derivation). A principal formula as the uppermost formula on the left side corresponds to an introduction instance of the terminal connective of the proposition concerned; a principal formula on the right side - which is, after all, an antecedent formula - corresponds to a subsequent elimination instance of that connective. The cut associated with the cluster represents nothing more than the formal establishment

of the connection between the two instances made necessary by the particular structure of our formalism. The fact that branchings of a cluster occur corresponds to the difficulty discussed at the end of 3.1; branchings on the right side, for example, represent a multiple application of the proposition. That branchings can appear both on the left and the right is due to the general symmetry of our formalism and renders more difficult a transfer of the fundamental ideas to each individual detail of the reduction. Yet it suffices if we have a reasonable conception of the fundamental ideas and continue to let ourselves be guided simply by formal analogies; this is precisely what I have done in formulating the consistency proof.

We must now prove the above assertion which can be interpreted as asserting the existence of a suitable place for an operational reduction in our derivation.

In this connection we first observe that our derivation must contain at least one operational inference figure. If this were not the case the ending would represent the entire derivation. This would mean that a "false" sequent has been derived from mathematical basic sequents which contain no free variables, and are therefore "true" sequents, by means of the application of structural inference figures alone and without thinnings. At the same time the only formulae occurring in the whole derivation are prime formulae without free variables, thus decidable formulae, so that it can be decided of each sequent whether it is true or false. (A formula with logical connectives cannot occur because no such connective occurs in the basic sequents and because none could

have been introduced by the possible inference figures.) This would mean that at least one inference figure occurs whose lower sequent is "false" whereas its upper sequents are "true". This is easily seen to be impossible.

In order to prove the above assertion we now examine all those paths of the ending whose uppermost sequent is the lower sequent of an operational inference figure. We follow these paths from the top down and record whether in the sequents which we encounter a formula occurs which belongs to the same cluster as one of the principal formulae standing immediately above it (or whether it itself is a principal formula). This is usually so in the case of the uppermost sequents of our paths and as we continue downwards in a path this property is generally inherited. It is preserved trivially in passing through contractions and interchanges (by the definition of cluster). If we reach a cut in which two paths of the considered type meet it may however happen that this property is not transferred to the lower sequent; yet this can arise at most in the case where the cluster belonging to the cut formulae contains a principal formula on both sides. This is precisely the case specified in the assertion. Since the empty end-sequent does not possess the mentioned property in any case, the assertion is proved as long as this case really is the only possible one in which the property under discussion may fail to be passed on in tracing out the paths under examination. To this needs to be added only one more case, viz., the case in which, coming from above, a cut is encountered whose other upper sequent belongs to none of the paths examined and can therefore occur only in paths of the ending

that are bordered above by mathematical basic sequents. This upper sequent can then contain only prime formulae and the cut formulae are therefore also prime; indeed, a formula occurring in the traced upper sequent and belonging to the same cluster as a principal formula cannot be a cut formula since its degree is greater than 0 and it is therefore clustered with a formula with the same property in the lower sequent.

This concludes the proof of the existence of a formula cluster suitable for an operational reduction.

Now one last auxiliary concept that will be of central importance for the definition of the "measure of complexity" of a derivation:

By the level of a derivational sequent we mean the highest degree of any cut or of a CJ-inference figure whose lower sequent stands below the sequent concerned. If there is no such inference figure then the level is equal to 0.

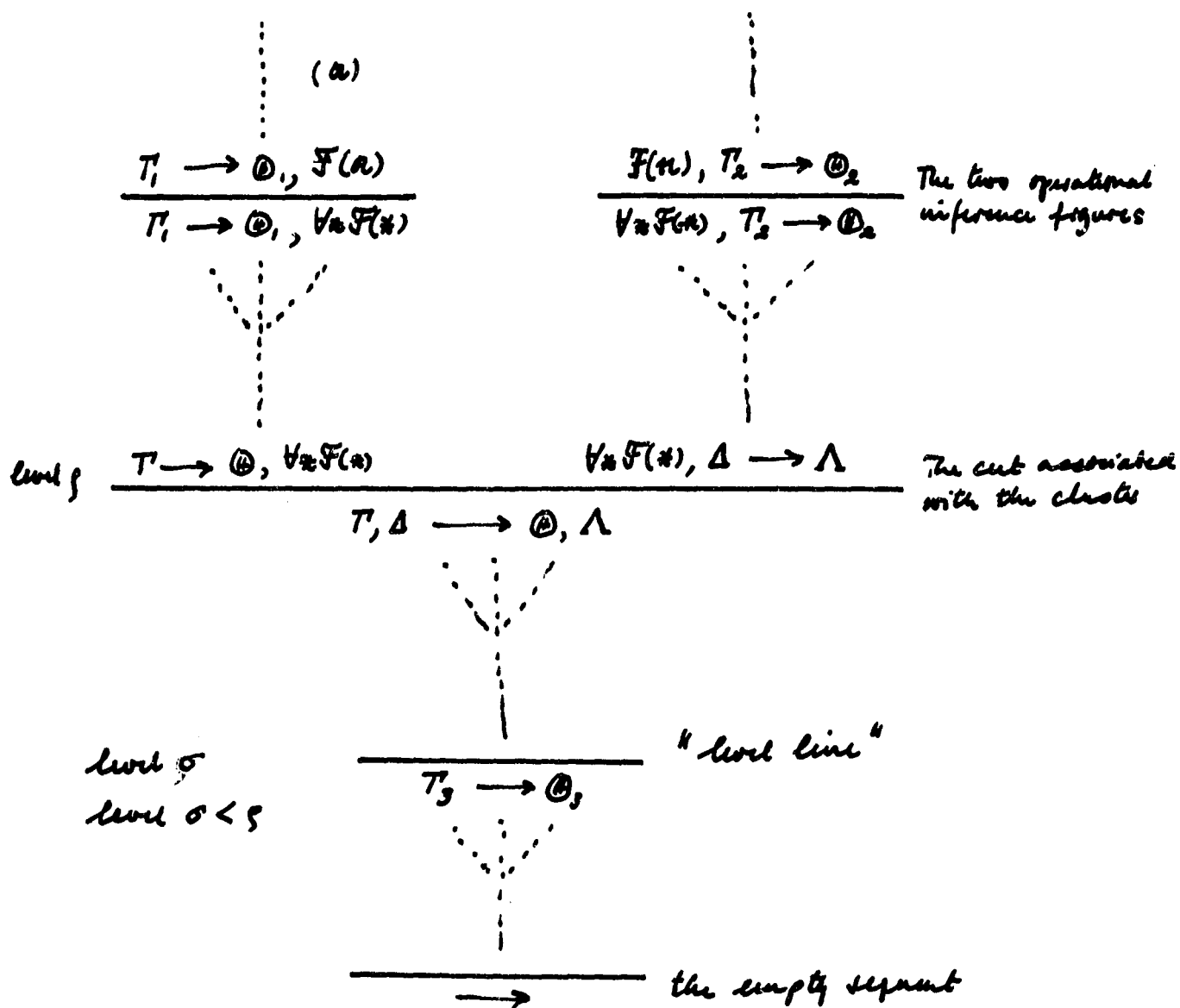
Comments about the importance of this concept will follow further below.

3.5. The Operational Reduction.

Now the operational reduction proper can be defined. Given is an inconsistent derivation whose ending includes at least one formula cluster containing on each side at least one principal formula of an operational inference figure. We select such a formula cluster and from each of its sides one uppermost formula of the kind mentioned. In order to make this step unambiguous a certain procedure concerning the type of choice to be made must be specified; this is not difficult.

We shall first deal with the case in which the terminal connective of the clustered formulae is a $\forall$. The remaining cases are dealt with almost in the same way and can be disposed of later in a few words.

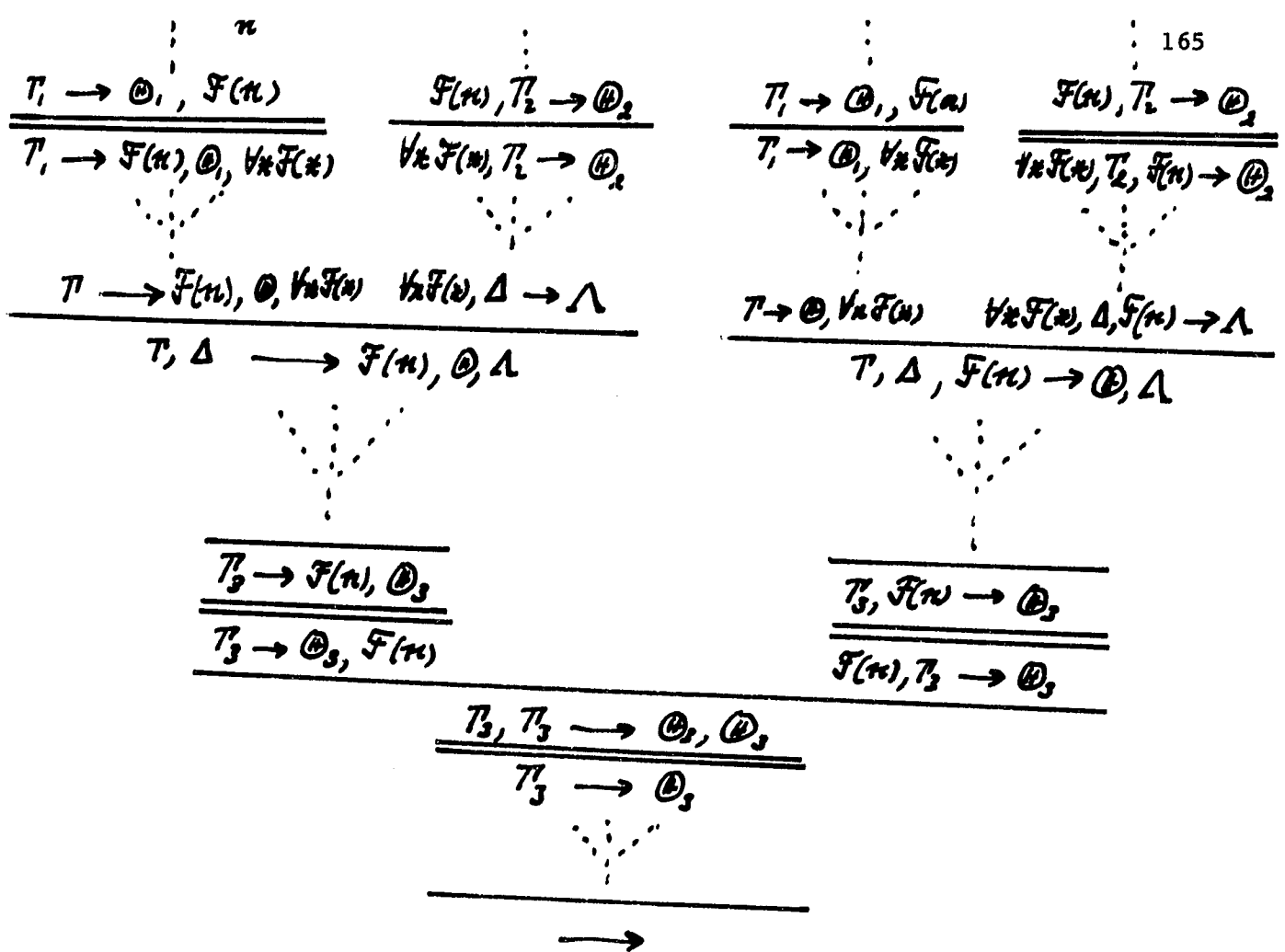
The derivation therefore looks like this:



Explanatory remarks:

The dots are intended to indicate that further paths may enter from both sides in arbitrary fashion into the traced paths. In addition entire derivational sections of any form whatever may stand above the operational inference figures. The term μ can only be a numerical term since no inference figure with an eigen-variable can occur below it (3.2, 3.41). Suppose that $\mathcal{T}_3 \rightarrow \mathcal{Q}_3$ is the first sequent encountered in tracing the path from $\mathcal{T}, \Delta \rightarrow \mathcal{Q}, \Lambda$ to the end-sequent which is of a lower level than the upper sequent of the cut belonging to the cluster. (Such a sequent must always exist since the level of the end-sequent equals 0, yet that of the upper sequent of the cut in question at least 1 since the degree of the cut itself is at least equal to 1.) It may happen that the sequent $\mathcal{T}, \Delta \rightarrow \mathcal{Q}, \Lambda$ is already the desired sequent; the above diagram must then be interpreted correspondingly. It may of course equally well happen that an upper sequent of the cut is itself already the lower sequent of the operational inference figure; and, finally, the sequent $\mathcal{T}_3 \rightarrow \mathcal{Q}_3$ may be identical with the end-sequent; all of this makes no difference to the reduction.

The reduction step consists now in the transformation of the derivation into the form indicated by the following diagram:



How the diagram is intended should basically be obvious. The old derivation for $\pi_3 \rightarrow \mathcal{Q}_3$ is written down twice side by side and the first instance is modified in such a way that the left operational inference figure vanishes; in the section of the derivation standing immediately above it every occurrence of the free variable x is here replaced by the numerical term n - except again where it happens to be used simultaneously as the eigen-variable of an occurring inference figure in the sequents standing above the lower sequent of that inference figure -; the formula $\forall x F(x)$ is then nevertheless re-introduced, but this time by a thinning; everything else is left exactly as it was before

with the single exception that in the path going through

$\pi_1 \rightarrow F(\kappa), \Theta_1, \forall x F(x)$ the formula $F(\kappa)$ is carried

along as an additional succedent formula. It can be seen at once by reference to the inference figure schemata that this leaves all inference figures correct; the same is true of the replacement of κ by κ .

In the second instance of the old derivation of $\pi_3 \rightarrow \Theta_3$ the procedure is analogous. Here the right operational inference figure vanishes without necessitating the replacement of a variable; and the formula $F(\kappa)$ is carried along down as an additional antecedent formula.

From the two sequents $\pi_3 \longrightarrow F(\kappa), \Theta_3$ and $\pi_3, F(\kappa) \longrightarrow \Theta_3$ the old sequent $\pi_3 \longrightarrow \Theta_3$

is then obtained by a new cut together with the applications of interchanges and contractions and the rest of the old derivation is taken over unchanged.

The reader can convince himself without difficulty that the reduction step here defined turns the given derivation into another entirely correct derivation in the sense of our formalism.

Remarks about the significance of this reduction step.

Let us recall the fundamental ideas of the operational reduction (3.1) and compare with it the formal presentation just given. The two operational inference figures represent an introduction and elimination of the $\forall$ in $\forall x F(x)$. According to the original fundamental idea the two inference figures should have been omitted and the $\forall x F(x)$ replaced by the "simpler" $F(\kappa)$ -whose degree is smaller by 1 -; the place of the

cut with the cut formulae $\forall x F(x)$ should have been taken by a new cut with the cut formulae $F(x)$. Yet there exists the already mentioned difficulty that the formula $\forall x F(x)$ may have several application instances, even several introduction instances - i.e., the formula cluster may have branchings on both sides and contain several uppermost formulae. It is therefore actually necessary, both in connection with the cancellation of the left operational inference figure and that of the right operational inference figure, to retain also the old cut with $\forall x F(x)$; a "simplification" has nevertheless been achieved in each case by the omission of an operational inference figure above this cut. (Although interchanges and a thinning have taken the place of this figure, these "do not count" in the determination of the "complexity" of the derivation.- The fact that $\forall x F(x)$ is reintroduced by a thinning is motivated only by convenience since its reappearance further down in the derivation must be expected in any case and since this is the most convenient way of obtaining the new form of the derivation from the old one.)

Further down in the new derivation then follows the "new cut" with the cut formulae $F(x)$. Precisely why has this cut been placed below the "level line"? (Basically it could have been introduced at any stage below the two $\forall x F(x)$ - cuts up to the end of the derivation; we would merely also have had to write down twice the section of the derivation from these cuts up to the new cut with $F(x)$ as an additional antecedent or succedent formula and to leave unchanged the section below the new cut.)

This leads us to the purpose of the notion of a level in general. What actually matters here is that in the reduction a "simplification" of the derivation is achieved in a sense to be made precise in the next paragraph through the ordinal numbers. Yet at first glance the new form of the derivation looks more complex than the old form: one and the same section of the derivation now occurs twice, although in each case somewhat simpler than before because of the omission of an operational inference figure. In defining a measure of complexity for derivations it will therefore be easy to achieve that each individual section standing above the new cut is valued somewhat lower than the corresponding section of the old derivation. Yet how is it to be accomplished, once the new cut has been added, that the entire section of the derivation up to $\mathcal{T}_3 \rightarrow \mathcal{B}_3$

is valued lower than the old derivation up to the same sequent? The new cut has a lower degree than the old cut; it is this feature to which we must cling. The new cut is thus placed below the collection of all, cuts whose degree is equal to that of the old cut so that after the reduction the collection of cuts above any one of these cuts of high degree is no larger than before, but is at most the same or a "simplified" collection. On the other hand, the new cut and everything below it now extend over a larger collection than before. This is compensated for, however, by the fact that all of these cuts are of lower degree than the old cut. Our success in achieving a lowering of the ordinal number of the derivation through the reduction will depend merely on our exploiting these facts properly when assigning ordinal numbers below.

An exceeding importance will thus have to be attached to the degree of a

cut in this connection.

In this discussion it was tacitly assumed as normal that the cut of larger degree generally occur above the cuts of smaller degree in the derivation. Since in reality this may of course not be the case the "degree" is replaced by the concept of the "level", and all this means is that cuts of a lower degree above cuts of a higher degree are treated as though they also possessed the higher degree; once this is done the main ideas stated above carry over without difficulty.

In determining the level of arbitrary derivation sequents, the CJ-inference figures are furthermore treated like cuts since in the course of their reduction they would be resolved into cuts of the same degree in any case.

The form of the reduction step for other connectives.

We must still specify how the reduction step is to be modified if the terminal connective of the cluster formulae is not a $\vee$, as in the explicitly presented case, but a $\&$, $\exists$, $\forall$ or $\rightarrow$. The differences are only minor:

If the cluster formulae have the form $\mathcal{U} \& \mathcal{B}$, we imagine the above diagrams suitably modified; in place of $V \# F(\pi)$ stands $\mathcal{U} \& \mathcal{B}$, and the operational inference figures run thus:

$$\frac{\mathcal{T}_1 \rightarrow \mathcal{C}_1, \mathcal{U} \quad \mathcal{T}_1 \rightarrow \mathcal{C}_1, \mathcal{B}}{\mathcal{T}_1 \rightarrow \mathcal{C}_1, \mathcal{U} \& \mathcal{B}} \quad \text{and} \quad \frac{\mathcal{U}, \mathcal{T}_2 \rightarrow \mathcal{C}_2}{\mathcal{U} \& \mathcal{B}, \mathcal{T}_2 \rightarrow \mathcal{C}_2} \quad \text{or} \quad \frac{\mathcal{B}, \mathcal{T}_2 \rightarrow \mathcal{C}_2}{\mathcal{U} \& \mathcal{B}, \mathcal{T}_2 \rightarrow \mathcal{C}_2}$$

In the new derivation $\mathcal{U} \& \mathcal{B}$ now also takes the place of $\forall x F(x)$, in place of $F(x)$ occurs $\mathcal{U}$ or $\mathcal{B}$, depending on which of the two possible forms the right operational inference figure (the " $\&$ -elimination") has had. Above the place from which the left operational inference figure is omitted we retain only the derivation of $\mathcal{T}_i \rightarrow \mathcal{Q}, \mathcal{U}$ or of $\mathcal{T}_i \rightarrow \mathcal{Q}, \mathcal{B}$, the other derivation is omitted. (This corresponds to the replacement of a by x in the $\forall$ -case.) The rest of the procedure is exactly the same as above; even the indicated differences completely suggest themselves.

If the terminal connective of the cluster formulae is a $\exists$ or $\vee$ the reduction proceeds completely symmetrically to the cases $\vee$ and $\exists$. Right and left are here interchanged.

If the cluster formulae finally have the form $\rightarrow \mathcal{U}$, nothing changes essentially: the formula $F(x)$ in the new derivation then corresponds to the formula $\mathcal{U}$, except that, as a consequence of the omission of the left operational inference figure, the latter formula occurs as an additional antecedent formula and, correspondingly, as a consequence of the omission of the right operational inference figure, as a succedent formula. In both cases the formula is carried forward up to the sequent $\mathcal{T}_i \rightarrow \mathcal{Q}$, as usual; the only difference is the fact that now the left and the right upper sequent of the "new cut", i.e., the complete derivational sections standing above it, must be interchanged with one another.

This completes the definition of a reduction step on an inconsistent derivation.

Paragraph 4

THE ORDINAL NUMBERS

CONCLUDING REMARKS

4.1. The Transfinite Ordinal Numbers below ε_0 .

I shall now define the ordinal numbers to be used. These will not be written as decimal fractions, as in the earlier paper; this time I shall adopt the notation customary in set theory. (In spite of this all definitions and proofs given in the following paragraphs are entirely "finitist" and are of an especially elementary nature in this respect as were the corresponding sections in the earlier proof. Here we are not really concerned with a study of transfinite induction, cf. below.)

Recursive definition of the ordinal numbers, also of equality and the order relation ($<$) among these:

The system $\mathcal{G}_0$ consists of the number 0. We define: $0 = 0$ and not $0 < 0$.

Suppose that the numbers of the system $\mathcal{G}_s$ (where s is a natural number or 0) are already defined, as well as $=$ and the $<$ -relation among these. An arbitrary number of the system $\mathcal{G}_{s+1}$ then has the form

$$\omega^{\alpha_1} + \omega^{\alpha_2} + \dots + \omega^{\alpha_n}$$

where the α_i 's are the numbers of the system G_s , with
 $\alpha_1 \geq \alpha_2 \geq \dots \geq \alpha_n$; n designates
 a natural number. The number 0 also once again belongs to the system
 G_{s+1} .

A G_{s+1} -number β is equal or smaller or larger than a
 G_{s+1} -number r if their representations coincide or if the
 first non-coinciding "exponent" α in the representation of β
 is smaller or larger than the corresponding exponent in the representation
 of r . If $\beta = r + \dots$, then $\beta > r$. 0 is considered
 to be smaller than any other number. $\beta > r$ means of course the same as
 $r < \beta$.

This completes the definition. It is easily seen that each system
 includes all preceding systems and that the relations of "smaller than"
 and "equal" between two numbers are independent of the system to which
 these numbers are considered to belong. It also follows quite clearly
 that of a given expression it can always be decided whether it is an
 ordinal number or not and that of two given ordinal numbers it can be
 decided (in a simple way) whether they are equal or which is the smaller
 one. (These concepts are therefore indeed "finitist".)

For present purposes the symbols '0', '+', and ' ω ', as
 well as the "exponentiation" occurring in the representation of numbers
 are to be interpreted quite formally and no particular sense needs to
 be associated with them such as regarding ω as "an infinite number"
 and the '+'-symbol as corresponding to "addition". Such

visualizations are of use merely for the understanding of the context as a whole. Solely for the purpose of comparing the size of the individual systems the following might still be said by using concepts and results from set theory:

The system $\mathcal{G}_1$ consists of the numbers: $0, \omega^0, \omega^0 + \omega^0, \dots$
 i.e., in the usual notation: $0, 1, 2, \dots$,
hence of the 0 and the natural numbers.

The limit number of the system is ω .

$\mathcal{G}_2$ contains already all numbers below ω^ω , viz:

$$0, \omega^0, \omega^0 + \omega^0, \dots, \omega^{\omega^0}, \omega^{\omega^0} + \omega^0, \omega^{\omega^0} + \omega^{\omega^0}, \\ \omega^{\omega^0} + \omega^{\omega^0} + \omega^0, \dots, \omega^{\omega^0 + \omega^0}, \dots, \omega^{\omega^0 + \omega^0 + \omega^0}, \dots,$$

thus: $0, 1, 2, \dots, \omega, \omega+1, \dots, \omega \cdot 2, \omega \cdot 2 + 1, \dots, \omega^2, \dots, \omega^3, \dots$

in general all polynomials $\omega^{\nu_1} \cdot \mu_1 + \dots + \omega^{\nu_0} \cdot \mu_0$;
 the ν and μ designate natural numbers or 0; $\nu_1 > \nu_2 > \dots > \nu_0$.

$\mathcal{G}_3$ contains all numbers below ω^ω (i.e., $\omega^{(\omega)}$); in the

following, multiple exponentiations are to be interpreted correspondingly).

G_4 contains all numbers below $\omega^{\omega^{\omega^{\omega}}}$, etc.

The limit number of all systems taken together is the number ε_0 , the "first ε -number".

We shall use the symbol 1 as an abbreviation for ω^0 . We also need the concept of the "natural sum" of two (non-zero) ordinal numbers which is defined as follows:⁽³²⁾

Suppose that $\alpha = \omega^{\gamma_1} + \omega^{\gamma_2} + \dots + \omega^{\gamma_\mu}$ and
 $\beta = \omega^{\delta_1} + \omega^{\delta_2} + \dots + \omega^{\delta_\nu}$ ($\mu \geq 1, \nu \geq 1$).

The "natural sum" $\alpha \# \beta$ is then obtained by arranging the $\mu + \nu$ constituents ω^{γ} and ω^{δ} by size and joining them back together again by '+'-symbols, the largest constituent first, the smallest last, with equal constituents of course side by side. In this way another correct ordinal number obviously results.

An example: If

$$\alpha = \omega^{\omega'+1} + \text{ and } \beta = \omega^{\omega^{\omega'+1+1}+1} + \omega^{\omega'+1} + \omega'$$

$$\text{then } \alpha \# \beta = \omega^{\omega^{\omega'+1+1}+1} + \omega^{\omega'+1} + \omega^{\omega'+1} + \omega'+1$$

In all cases $\alpha \# \beta = \beta \# \alpha$. The natural sum of even arbitrarily many ordinal numbers is independent of the order of the individual summations. $\alpha \# \beta > \alpha$.

If $\alpha^* < \alpha$ then $\alpha^* \# \beta < \alpha \# \beta$. These facts are easily proven.

4.2. The Correlation of Ordinal Numbers with Derivations.

Suppose that an arbitrary derivation is given. Its ordinal number is calculated by passing downward from the uppermost sequents and assigning to each individual derivational sequent as well as to each line of inference an ordinal number (> 0) on the basis of the following stipulations:

Each uppermost sequent receives the ordinal number 1 (i.e., ω^0).

Suppose that the ordinal numbers of the upper sequents of an inference figure have already been determined. The ordinal number of the line of inference is then obtained as follows:

If the inference figure is structural then the ordinal number of the upper sequent is adopted unchanged, or, in the case of a cut, the natural sum of the ordinal numbers of the two upper sequents is formed.

If the inference figure is operational then $+1$ is adjoined to the ordinal number of the upper sequent; yet if the figure has two upper sequents, the larger of the two ordinal numbers is selected and $+1$ is adjoined to it.

If a CJ-inference figure is finally encountered - whose upper sequent has the ordinal number $\omega^{\alpha_1} + \dots + \omega^{\alpha_n}$ ($n \geq 1$) - then $\omega^{\alpha_1} + 1$ is taken as the ordinal number of the line of

inference. If $\alpha_1 = 0$ then this number is of course ω' .

From the ordinal number of a line of inference - call it α - the ordinal number of the lower sequent of the inference figure concerned is obtained in the following way:

If the level of the lower sequent is the same as that of the upper sequent then the ordinal number of the lower sequent is equal to α . If its level is lower by 1, then the ordinal number of the lower sequent is ω^α . If lower by 2, the ordinal number is ω^{ω^α} , lower by 3: $\omega^{\omega^{\omega^\alpha}}$, etc.

The ordinal number which is finally obtained for the end-sequent of the derivation is the ordinal number of the derivation.

The reader can easily convince himself that the mentioned operations always yield new genuine ordinal numbers in accordance with their definition. -For the time being I shall not comment on this method of correlating ordinal numbers; it is really quite simple; of special interest is only the evaluation of the CJ-inference figures and that of the different levels; in both cases this valuation will be most easily understood through its effect later on.

4.3. The Decrease of the Ordinal Number in the Course of a Reduction Step on an Inconsistent Derivation.

It still remains to show that in the course of a reduction step according to Paragraph 3 the ordinal number of an inconsistent derivation decreases.

This no longer presents any special difficulties; all we need to do is to examine the correctness of this assertion carefully for each individual case.

The preparatory step 3.2. obviously leaves the ordinal number entirely unaffected. What is the situation in the case of a CJ-reduction (3.3.)?

Suppose that the ordinal number of the upper sequent of the CJ-inference figure is $\omega^{\alpha_1} + \dots + \omega^{\alpha_v}$ ($v \geq 1$), that of the line of inference therefore $\omega^{\alpha_1 + 1}$. This is also

at once the ordinal number of the lower sequent whose level cannot be lower than that of the upper sequent since the cluster cuts associated with $\mathcal{F}(1)$ and $\mathcal{F}(n)$ and which have the same degree as the CJ-inference figure, must still occur further down in the derivation. Let

us now examine the figure which has replaced the CJ-inference figure in the reduction (first for n not equal to 1). In the new derivation each one of its uppermost sequents obviously receives the same ordinal number $\omega^{\alpha_1} + \dots + \omega^{\alpha_v}$. Furthermore, all sequents of the replacement figure have the same level, viz., that level which the two sequents of the CJ-inference figures had before.

(The newly occurring cuts have of course the same degree as that of the CJ-inference figure). The ordinal number of the lowest sequent of this figure is therefore obviously equal to the natural sum of all numbers $\omega^{\alpha_1} + \dots + \omega^{\alpha_v}$.

Consequently it begins: $\omega^{\alpha_1 + 1}$.

It is therefore smaller than $\omega^{\alpha_1 + 1}$, according to the definition of "smaller than" for the ordinal numbers.

From this it now follows easily that the ordinal number of the entire derivation has also been decreased. After all, from the CJ-inference figure downwards nothing has changed in the derivation, in fact, all levels have here also remained the same. The decrease which has occurred at one place is preserved in calculating the ordinal number further up to the end-sequent; what is essential is that in proceeding downwards only structural inference figures are encountered and that the following holds: If $\alpha^* < \alpha$, then $\omega^{\alpha^*} < \omega^\alpha$ and $\alpha^* \# \beta < \alpha \# \beta$. (Suppose that α , α^* and $\beta \neq 0$.) Both requirements are satisfied at once by definition.

Now the purpose of the ω^{α_1+1} in the evaluation of a CJ-inference figure also becomes clear: in the reduction the figure breaks up into a number of cuts; and in some sense the n-fold multiple of one and the same derivational section occurs. In order to achieve a decrease in the ordinal number we must therefore choose as the ordinal number of the original derivational section up to the CJ-inference figure the "limit number" of all "n-fold multiples" of the ordinal number of the upper sequent, i.e., $\omega^{\alpha_1+1} = "\omega^{\alpha_1} \cdot \omega"$. (The expressions in " " serve of course only as illustrations; after all they are not even defined in this context.)

Now there remains only the case where n equals 1: in the new derivation the sequent $F(1), \pi \longrightarrow \odot, F(1)$ receives the ordinal number 1. In the old derivation its ordinal number was at least equal to ω^1 . Here we have an obvious decrease which is at

the same time inherited by the ordinal number of the entire derivation.

This proves the decrease of the ordinal number of an inconsistent derivation in a CJ-reduction. There still remains the case of the operational reduction. Here it must first be observed that through the further preparatory step (3.42) no increase in the ordinal number can occur. The proof of this fact presents certain difficulties in spite of the obvious external simplification of the derivation in this step. I shall sketch only briefly what kind of reasoning is here required - the reader interested only in bare essentials may skip this paragraph -:

The omission or adjunction of formulae and other transformations within structural inference figures except cuts have no influence whatever on the ordinal number. This is different in the case of the cancellation of a cut through the omission of an upper sequent together with everything standing above it. If we disregard for the time being the change in level which this cancellation entails then a decrease in the ordinal number results from the replacement of the natural sum of two numbers by only one of these two numbers. Added to this must be the fact that through the omission of a cut the level of a whole collection of sequents above this cut may be reduced to a greater or lesser extent (not only in the ending but in the entire derivation). In order to recognize that this rather entangled transformation cannot affect an increase in the ordinal number of the entire derivation we argue thus: we imagine that we can fix the level quite arbitrarily. We begin with the old derivation, omit the cut and,

at first, leave all levels untouched. Then we gradually adjust these levels to the values which the transformed derivation really should have according to the definition of a level, by carrying out a succession of single steps of the following kind: the level of the upper sequents of one inference figure whose lower sequent has a lower level than the upper sequent is in each case diminished by 1. It is easily seen that the entire adjustment of levels can in fact be made up of such operations. (We begin from below.) What exactly happens to the ordinal numbers in such a single adjustment of level? Suppose that before the adjustment the ordinal numbers of the upper sequents are α and β (if there is only one such number we simply think of the second number below as not being present). After the adjustment they then take the form ω^α and ω^β . (Except if one upper sequent is an uppermost sequent of the derivation, in which case its ordinal number was and remains equal to 1 and this simplifies the following discussion further.) Before the adjustment, the ordinal number of the line of inference was thus either α or $\alpha \# \beta$ or $\alpha + 1$, or $\beta + 1$ or ω^{α_1+1} (where $\alpha = \omega^{\alpha_1} + \dots$ in the case of a CJ-inference figure), depending on the kind of inference figure involved. After the adjustment, the ordinal number takes the form of either ω^α or $\omega^\alpha \# \omega^\beta$ or $\omega^\alpha + 1$, or $\omega^\beta + 1$, or ω^{α_1+1} . Now to the lower sequent: If the difference in level between it and the upper sequent was equal to 1 before and is therefore now equal to 0, then the adjustment has brought about a change in the ordinal number of this sequent from ω^α to ω^α , or from $\omega^{\alpha \# \beta}$ to $\omega^\alpha \# \omega^\beta$ or from

$\omega^{\alpha+1}$ to $\omega^{\alpha} + 1$, or from $\omega^{\beta+1}$ to $\omega^{\beta} + 1$,
 or finally from $\omega^{\omega^{\kappa_1}+1}$ to $\omega^{\omega^{\kappa_1} + \dots + 1}$. In each case the
ordinal number has either remained the same or has become smaller
 and this should be verified by the reader from case to case by means
 of the definition of "smaller than". If the difference in level
 between the upper and lower sequents was greater than 1 nothing has
 essentially changed: in each case the mentioned numbers are augmented
 by an equal number of exponentiations with ω . This property of
 non-increasing transfers to the ordinal number of the entire derivation
 and this number can therefore rise neither in a single step of the
 described adjustment of level nor, quite generally, in the preparatory
 step for the operational reduction as a whole.

We now come to the operational reduction proper (3.5) in which we must
 demonstrate a decrease of the ordinal number. We shall again base our
 discussion upon the case presented in detail above (with $\mathcal{A}$ as the
 connective to be reduced). The ordinal number of each of the two lines
of inference standing immediately above the sequents $\mathcal{T}_3 \rightarrow \mathcal{F}(\kappa), \mathcal{A}_3$
 and $\mathcal{T}_3, \mathcal{F}(\kappa) \rightarrow \mathcal{A}_3$
 in the new derivation - which we denote by α_1 and α_2 - is smaller
 than the ordinal number α of the "level line" in the old
 derivation. This is so since the derivational sections standing above
 the lines of inference essentially correspond to one another; all
 levels in particular are the same as those in the old derivation -
 the levels of the sequents standing immediately above the mentioned
 lines of inference are equal to $\mathcal{A}$ throughout -; in each

case only one operational inference figure has disappeared and been replaced by structural inference figures which have no influence on the ordinal number. At this point a decrease in the ordinal number has therefore taken place which is preserved throughout the subsequent structural inference figures up to the mentioned lines of inference.

Also: the sequent $\mathcal{T}_3 \longrightarrow \mathcal{Q}_3$ has of course the same level σ in the new derivation as in the old one; $\sigma < \wp$. The sequent $\mathcal{T}_3, \mathcal{T}'_3 \longrightarrow \mathcal{Q}_3, \mathcal{Q}'_3$ has of course the level σ . The level τ of the upper sequents of the "new cut" satisfies $\wp > \tau \geq \sigma$. The inequality on the right is trivial; and that $\wp > \tau$ is recognized thus: by the definition of a level the τ is equal to the larger of the two numbers σ and the "degree of $\mathcal{F}(n)$ ". If $\tau = \sigma$ then $\tau < \wp$, since $\sigma < \wp$. If τ equals the degree of $\mathcal{F}(n)$, then $\tau < \wp$ since the degree of $\mathcal{F}(n)$ is smaller than the degree of $\forall x \mathcal{F}(x)$ and since $\wp$ is at least equal to the latter.

Let us first suppose that the differences between the levels

$\wp$, τ and σ are minimal, i.e., that $\wp = \tau + 1$ and $\tau = \sigma$. In this case our demonstration is completed as follows:

In the old derivation the level line had the ordinal number α , the sequent $\mathcal{T}_3 \longrightarrow \mathcal{Q}_3$ therefore the ordinal number ω^α .

In the new derivation the lines of inference corresponding to this level line have the ordinal numbers α_1 and α_2 , both are smaller than α and the upper sequents of the new cut therefore have the ordinal numbers ω^{α_1} and ω^{α_2} ; the sequent $\mathcal{T}_3 \longrightarrow \mathcal{Q}_3$ receives the ordinal

number $\omega^{\alpha_1} + \omega^{\alpha_2}$. (Without loss of generality we may assume that $\alpha_1 \geq \alpha_2$.) The latter number is obviously smaller than ω^α ; and we have thus finished. For on the basis of an already repeatedly applied argument this decrease transfers to the ordinal number of the end-sequent and therefore to the derivation as a whole. (Below $\pi_s \rightarrow \textcircled{H}_s$ nothing has of course changed.)

If the distances between the levels $\wp$, τ and σ are greater, our argument is not essentially changed. The place of the inequality

$$\omega^\alpha > \omega^{\alpha_1} + \omega^{\alpha_2} \quad (\text{where } \alpha > \alpha_1 \geq \alpha_2)$$

is then simply taken by the inequality

$$\omega \cdots \omega \cdots \omega^\alpha > \omega \cdots \omega \cdots \omega^{\alpha_1} + \omega \cdots \omega \cdots \omega^{\alpha_2}$$

and the latter inequality is also easily seen to be valid.

It now becomes apparent how through the method of definition of the ordinal numbers in connection with the notion of level the difficulties associated with the apparent increase in complexity of a derivation as a result of the operational reduction have been overcome. The main idea is: in the reduction the same derivational section occurs twice, although both times somewhat simplified. In the general case, however, $\alpha < \alpha_1 + \alpha_2$, where we suppose α_1 and α_2 to be smaller than α . Yet for the exponential expression holds: $\omega^\alpha > \omega^{\alpha_1} + \omega^{\alpha_2}$. (precisely as in the case of the natural numbers, for ω we can put any number ≥ 3 .)

The "simplification" of the figure as a whole has thus been achieved as long as it is always possible to insert an exponentiation; and this is made possible by the fact that the degree of the new cut is smaller than the degree of the old $\forall x F(x)$ -cut. It was for the purpose of exploiting this fact that the general concept of a level was introduced and applied in the correlation of ordinal numbers.

The cases where the connective to be reduced is a $\wedge$, $\exists$, $\vee$, or $\neg$ are so similar that a special discussion of them becomes superfluous.

The decrease of the ordinal number of an inconsistent derivation in the reduction step has thus been proved.

4.4. Concluding Remarks.

If we had not admitted CJ-inference figures into our formalism it would be possible to make do with the natural numbers as ordinal numbers. In order to realize this the reader should omit 4.1 and in 4.2 replace ω by 3 throughout and "natural sum" simply by "sum". Sums and powers are to be understood in the way customary for the natural numbers. 4.3 then remains valid throughout, as is easily verified; the CJ-reduction would here of course have to be left out. The consistency proof could then be concluded by an ordinary complete induction instead of a transfinite induction.

As a result of the admission of the CJ-inference figures, and therefore

for our formalism, the following remarkable connection between the magnitude of the ordinal number of a derivation and the highest degree of the formulae occurring in the derivation holds: the ordinal number of a derivation in which only formulae of degree 0 occur is smaller than ω^ω (i.e., ω^ω in our notation). If the highest degree of a formula equals 1 then its ordinal number is smaller than ω^{ω^ω} , if the degree equals 2, then the ordinal number is smaller than $\omega^{\omega^{\omega^\omega}}$, etc. This is not difficult to prove.

These theorems are of course meaningful only relative to our special correlation of ordinal numbers. Yet it is reasonable to assume that by and large this correlation is already fairly optimal, i.e., that we could not make do with essentially lower ordinal numbers. In particular the totality of all our derivations cannot be handled by means of ordinal numbers of which all lie below a number which is smaller than ϵ_0 . For transfinite induction up to such a number is itself provable in our formalism; a consistency proof carried out by means of this induction would therefore contradict Gödel's theorem (given, of course, that the other techniques of proof used, especially the correlation of ordinal numbers, have not assumed forms that are non-representable in our number-theoretical formalism). By the same round-about argument we can presumably also show that certain sub-classes of derivations cannot be handled by ordinal numbers below certain numbers of the form $\omega^{\dots\omega}$. It is quite likely that one day a direct approach to the proof of such impossibility theorems will be found.

If we include arbitrary functions in our formalism then the consistency proof remains valid with minor modifications: all that needs to be shown is that at some point in the reduction, following the first preparatory step, for example, all terms without free variables can be evaluated and replaced by their numerical values. It is presupposed that all functions can be effectively calculated for all given numerical values. There still arise certain formal difficulties from the fact that although term may be calculable, a corresponding term in another place in the same inference figure may still contain a variable (cf. Article 14.22 of the earlier paper); yet these difficulties do not affect the main ideas here involved.

In principle the contents of Section V of the earlier paper also remain valid for the new version of the consistency proof. I have not given a new proof of the "reducibility" of arbitrary derivable sequents; nor do I attach any special importance to this. (I had previously advanced it as an argument against radical intuitionism - article 17.3 - , but it is not particularly essential for this purpose.)

Transfinite Induction.

I have not given a new proof of the transfinite induction which concludes the consistency proof since I intend to discuss the questions involved at this point separately at some later date. For the conclusion of the present proof the earlier proof of the "theorem of transfinite induction" (Articles 15.4 and 15.1) is therefore to be adopted for the time being. For this purpose the new ordinal numbers

must be made to correspond with the decimal fractions used in the earlier paper; this presents no special difficulties. (Both systems are after all of the same "order type ϵ_0 ".)

The transfinite induction occupies quite a special position within the consistency proof. Whereas all other forms of inference used are of a rather quite elementary kind from the point of view of being "finitist" - this applies to the new proof as much as it does to old one - this cannot be maintained of the transfinite induction. Here we therefore have a task of a different kind: we are not merely required to prove transfinite induction - this is not particularly difficult and possible in various ways - but rather to prove it on a finitist basis, i.e., to establish clearly that it is a form of inference which is in harmony with the principle of the constructivist interpretation of infinity; an undertaking which is no longer purely mathematical but which nevertheless forms part of a consistency proof.

We might be inclined to doubt the finitist character of the "transfinite" induction, even if only because of its suspect name. In its defense it should here merely be pointed out that most somehow constructivist orientated authors place special emphasis on building up constructively (up to ω^ω , for example) an initial segment of the transfinite number sequence (within the "second number class"). And in the consistency proof and in possible future extensions of it we are certainly dealing only with an initial segment, a "section" of the second number class, even though this is an already comparatively

) extensive segment, and which must probably be extended still considerably further for a consistency proof for analysis. Yet I fail to see at what "place" constructive certainty is here supposed to end and where a further extension of transfinite induction is therefore thought to become disputable. I think rather that the reliability of the transfinite numbers required for the consistency proof compares with that of the early segment, up to ω^2 , for example, in the same way as the reliability of a numerical calculation extending over a hundred pages with that of a calculation of a few lines: it is merely a considerably vaster undertaking to convince oneself of this certainty from beginning to end! A detailed discussion of these matters (whose exposition in the earlier paper - Article 16.11 - seems to me now to be somewhat too sketchy) will, as said before, follow at a later date.

GALLEY PROOF

14.3. If a reduction rule is known for a sequent then a reduction rule can also be stated for every sequent which has resulted from the former by a structural transformation. Viz.: An interchange of antecedent formulae (5.241) does not affect the reduction procedure. If an antecedent formula was omitted which was equal to another antecedent formula (5.242) we reduce the newly arisen sequent in the same way as the old one, yet if the omitted formula would have been subject to a reduction step according to 13.5, we apply this reduction step to the formula equal to the omitted formula and then retain the latter formula; this is after all permissible.

If a formula was adjoined to the antecedent formulae (5.243), we first carry out the required reductions on it according to 13.11 and 13.12 and continue the rest of the reduction up to the definitive form as if this formula were not even present.

A re-designating of a bound variable (5.244) does not necessitate a change in the reduction rule.

In the following I shall repeatedly make tacit use of the fact that a reduction rule for a sequent which results from another sequent by a structural transformation can be obtained from the reduction rule or the former sequent.

14.4. Now it still remains to show that a reduction rule can always be given for a sequent which results from those sequents by the application of a rule of inference for which reduction rules are already known. After the transformation according to Paragraph 12 the following rules of inference can still be applied in the derivation: $\forall$ -introduction, $\forall$ -elimination, $\&$ -introduction, $\&$ -elimination, "reductio", "elimination of the double negation" and "complete induction". I shall deal with them in that order.

14.41. Suppose that we are given a $\forall$ -introduction: "From $\mathcal{P} \rightarrow \mathcal{F}(a)$ follows $\mathcal{P} \rightarrow \forall x \mathcal{F}(x)$ ". Assume a reduction rule to be known for the sequent $\mathcal{P} \rightarrow \mathcal{F}(a)$. The reduction of $\mathcal{P} \rightarrow \forall x \mathcal{F}(x)$ must begin with the replacement of (possible) occurrences of free variables by arbitrarily chosen numerals (13.11). Suppose that

$\mathcal{T}^* \rightarrow \forall x \mathcal{F}^*(x)$ results. If no free variables occurred then $\mathcal{T}^* \rightarrow \forall x \mathcal{F}^*(x)$ stands again for $\mathcal{T} \rightarrow \forall x \mathcal{F}(x)$.

(A corresponding argument is to apply below.) Then all (possible) minimal terms must be replaced by their numerical values (13.12) resulting in $\mathcal{T}^{**} \rightarrow \forall x \mathcal{F}^{**}(x)$. This sequent is reduced according to 13.21 to $\mathcal{T}^{**} \rightarrow \mathcal{F}^{**}(n)$, where n is to be chosen arbitrarily. Now any minimal terms that may have newly arisen must still be replaced by their numerical values in accordance with 13.12, resulting in $\mathcal{T}^{**} \rightarrow \mathcal{F}^{***}(n)$.

The reduction of the sequent $\mathcal{T} \rightarrow \mathcal{F}(a)$ must also begin with the replacement of the free variables. For this replacement we may in particular use the same numerals that were chosen in the reduction of $\mathcal{T} \rightarrow \forall x \mathcal{F}(x)$, as well as the symbol $\overset{n}{\lambda}$ for the replacement of a so that the sequent $\mathcal{T}^* \rightarrow \mathcal{F}^*(n)$ results. Now must follow the replacement of possible minimal terms and from this $\mathcal{T}^{**} \rightarrow \mathcal{F}^{***}(n)$ obviously results, i.e., the same sequent as above. By virtue of the reduction rule for $\mathcal{T} \rightarrow \mathcal{F}(a)$ a reduction rule must now be statable for this sequent; hence a reduction rule has also been obtained for $\mathcal{T} \rightarrow \forall x \mathcal{F}(x)$.

14.42. Suppose we are given a $\forall$ -elimination: "From $\mathcal{T} \rightarrow \forall x \mathcal{F}(x)$ results $\mathcal{T} \rightarrow \mathcal{F}(t)$ ". $\mathcal{T} \rightarrow \mathcal{F}(t)$ is again first subjected to (possibly) necessary reduction steps according to 13.11 and 13.12; suppose that $\mathcal{T}^* \rightarrow \mathcal{F}^*(n)$ results. In the reduction of $\mathcal{T} \rightarrow \forall x \mathcal{F}(x)$, which must begin with reduction steps according to 13.11 (if necessary), 13.12 (if necessary), and

then a step according to 13.21, possibly followed by further steps according to 13.12, the numerals to be substituted may obviously be chosen so that these steps also yield the sequent $\mathcal{T}^* \rightarrow \mathcal{F}^*(n)$.

We therefore have a reduction rule for that sequent and hence also for $\mathcal{T} \rightarrow \mathcal{F}(t)$.

14.43. The $\exists$ -introduction and the $\exists$ -elimination are dealt with quite analogously to the $\forall$ -introduction and $\forall$ -elimination. Here the reduction step according to 13.21 is replaced by a step according to 13.22.

14.44. In dealing with the three rules of inference still remaining I make use of the following lemma: "If reduction rules are known for two sequents of the form $\mathcal{T} \rightarrow \mathcal{D}$ and $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ in which no free variables and no minimal terms occur then a reduction rule can also be given for the sequent $\mathcal{T}, \Delta \rightarrow \mathcal{C}$." (The meaning of the symbols $\mathcal{T}$, Δ , $\mathcal{C}$ and $\mathcal{D}$ is the same as that defined at 5.250, $\mathcal{D}$ also stands for an arbitrary formula.) The proof of this lemma which represents the major part of the consistency proof follows at 14.6. Here I shall first show how the lemma is to be applied to the rules of the "reductio", the "elimination of the double negation" and the "complete induction".

14.441. Suppose that a "reductio" is given: "From $\mathcal{U}, \mathcal{T} \rightarrow \mathcal{B}$ and $\mathcal{U}, \Delta \rightarrow \neg \mathcal{B}$ follows $\mathcal{T}, \Delta \rightarrow \neg \mathcal{U}$." We first reduce $\mathcal{T}, \Delta \rightarrow \neg \mathcal{U}$ (if required) according to 13.11 and 13.12; suppose that the result is $\mathcal{T}^*, \Delta^* \rightarrow \neg \mathcal{U}^*$. This we reduce according to 13.23 to $\mathcal{T}^*, \Delta^*, \mathcal{U}^* \rightarrow 1=2$. In the

reduction of $\mathcal{U}, \mathcal{T} \rightarrow \mathcal{B}$, on the other hand, we can choose the numerals to be substituted in the reduction steps according to 13.11 and 13.12, which are carried out first, so that from these steps a sequent of the form $\mathcal{U}^*, \mathcal{T}^* \rightarrow \mathcal{B}^*$ results. In the same way it can be achieved that $\mathcal{U}, \Delta \rightarrow \neg \mathcal{B}$ assumes the form $\mathcal{U}^*, \Delta^* \rightarrow \neg \mathcal{B}^*$ after the appropriate reduction steps. This then yields the sequent $\mathcal{U}^*, \Delta^*, \mathcal{B}^* \rightarrow 1=2$ by 13.23. Reduction rules are therefore known for the sequents $\mathcal{U}^*, \mathcal{T}^* \rightarrow \mathcal{B}^*$ and $\mathcal{U}^*, \Delta^*, \mathcal{B}^* \rightarrow 1=2$ and, by the lemma, therefore also for $\mathcal{U}^*, \mathcal{T}^*, \mathcal{U}^*, \Delta^* \rightarrow 1=2$, i.e., (14.3) also for $\mathcal{T}^*, \Delta^*, \mathcal{U}^* \rightarrow 1=2$. We have thus a reduction rule for $\mathcal{T}, \Delta \rightarrow \neg \mathcal{U}$.

14.442. Suppose that we are given an "elimination of the double negation": "From $\mathcal{T} \rightarrow \neg \neg \mathcal{U}$ we obtain $\mathcal{T} \rightarrow \mathcal{U}$." The reductions of $\mathcal{T} \rightarrow \mathcal{U}$ according to 13.11 and 13.12 which may first be necessary can be carried out analogously on $\mathcal{T} \rightarrow \neg \neg \mathcal{U}$. We must therefore still reduce a sequent $\mathcal{T}^* \rightarrow \mathcal{U}^*$ in which free variables and minimal terms no longer occur and this will simultaneously yield a reduction rule for $\mathcal{T}^* \rightarrow \neg \neg \mathcal{U}^*$.

It is sufficient to state a reduction rule for the sequent $\neg \neg \mathcal{U}^* \rightarrow \mathcal{U}^*$. For we can then apply the lemma and from the availability of reduction rules for $\mathcal{T}^* \rightarrow \neg \neg \mathcal{U}^*$ and $\neg \neg \mathcal{U}^* \rightarrow \mathcal{U}^*$ conclude the stability of a reduction rule for $\mathcal{T}^* \rightarrow \mathcal{U}^*$. $\neg \neg \mathcal{U}^* \rightarrow \mathcal{U}^*$ can be reduced easily according to the

following rule (cf. 14.1): We reduce the succedent formula according to 13.21, 13.22, and 13.12 until it has the form $\neg \mathcal{C}$ or is a minimal formula. If it has become a correct minimal formula the reduction is finished. If it has assumed the form $\neg \mathcal{C}$ we continue the reduction according to 13.23 and obtain $\neg \neg \mathcal{U}^*, \mathcal{C} \rightarrow 1=2$, further (by 13.53) we obtain $\mathcal{C} \rightarrow \neg \mathcal{U}^*$, then (by 13.23) we obtain $\mathcal{C}, \mathcal{U}^* \rightarrow 1=2$.

In the case where the succedent formula has become a false minimal formula we proceed in the same way; in the latter case we first obtain $\rightarrow \neg \mathcal{U}^*$, and then $\mathcal{U}^* \rightarrow 1=2$.

In both cases we have now obtained a sequent which also occurs in the reduction of the logical basic sequent $\mathcal{U}^* \rightarrow \mathcal{U}^*$ according to the procedure stated at 14.1. We need therefore merely follow the procedure stated at that point in order to complete the reduction of the sequent.

14.443. Suppose that a "complete induction" is given: "From $\mathcal{T} \rightarrow \mathcal{F}(1)$ and $\mathcal{F}(\alpha), \Delta \rightarrow \mathcal{F}(\alpha+1)$ we obtain $\mathcal{T}, \Delta \rightarrow \mathcal{F}(t)$."

In $\mathcal{T}, \Delta \rightarrow \mathcal{F}(t)$ we first replace all (possible) occurrences of free variables by arbitrarily chosen numerals (13.11) and obtain

$\mathcal{T}^*, \Delta^* \rightarrow \mathcal{F}^*(t^*)$. Then (if necessary) we carry

out reduction steps according to 13.12 and achieve in this way that finally every occurrence of t^* has been replaced by the numeral

κ which represents the value of the term t^* (which no longer contains a variable). The sequent has thus become $\mathcal{T}^*, \Delta^* \rightarrow \mathcal{F}^*(\kappa)$.

Now we carry out further reduction steps according to 13.12 (if necessary)

until all minimal terms have been eliminated. The sequent then has the form $\mathcal{T}^{**}, \Delta^{**} \longrightarrow (\mathcal{F}^*(n))^*$.

In the reduction of $\mathcal{T} \rightarrow \mathcal{F}(1)$ and $\mathcal{F}(a), \Delta \rightarrow \mathcal{F}(a+1)$, which must begin with the replacement of possible occurrences of free variables, we can actually choose the numerals to be substituted so that they agree with the numerals chosen previously and can replace the variable a , which after all did not occur in $\mathcal{T}, \Delta \rightarrow \mathcal{F}(t)$, by any one of the numerals from 1 to m , where m denotes the number 1 smaller than n . It then follows that for each one of the sequents $\mathcal{T}^* \rightarrow \mathcal{F}^*(1)$ and $\mathcal{F}^*(1), \Delta^* \rightarrow \mathcal{F}^*(1+1)$ and $\mathcal{F}^*(2), \Delta^* \rightarrow \mathcal{F}^*(2+1)$ etc. up to $\mathcal{F}^*(m), \Delta^* \rightarrow \mathcal{F}^*(m+1)$ reduction rules are statable. If these sequents are then reduced by the reduction steps prescribed in 13.12, there result obviously sequents of the following form for which reduction rules are therefore also statable:

$\mathcal{T}^{**} \longrightarrow (\mathcal{F}^*(1))^*$ and $(\mathcal{F}^*(1))^*, \Delta^{**} \longrightarrow (\mathcal{F}^*(2))^*$ and $(\mathcal{F}^*(2))^*, \Delta^{**} \longrightarrow (\mathcal{F}^*(3))^*$ etc. up to $(\mathcal{F}^*(m))^*, \Delta^{**} \longrightarrow (\mathcal{F}^*(n))^*$.

We now apply the lemma: From the reduction rules for $\mathcal{T}^{**}$

$\mathcal{T}^{**} \longrightarrow (\mathcal{F}^*(1))^*$ and $(\mathcal{F}^*(1))^*, \Delta^{**} \longrightarrow (\mathcal{F}^*(2))^*$

we obtain a reduction rule for the sequent $\mathcal{T}^{**}, \Delta^{**} \longrightarrow (\mathcal{F}^*(2))^*$,

from it and from the reduction rule for $(\mathcal{F}^*(2))^*, \Delta^{**} \longrightarrow (\mathcal{F}^*(3))^*$

we obtain a reduction rule for $\mathcal{T}^{**}, \Delta^{**} \longrightarrow (\mathcal{F}^*(3))^*$ etc.;

finally it follows that a reduction rule is statable for the sequent

$\mathcal{T}^{**}, \Delta^{**} \longrightarrow (\mathcal{F}^*(n))^*$, hence also for $\mathcal{T}, \Delta \rightarrow \mathcal{F}(t)$,

since this sequent had actually already been reduced above to the

form of the former sequent.

14.5. Now the proof of the "lemma" is still outstanding. At this point I should like to add a few remarks which may contribute to an easier understanding of the proof.

What is the reason for the special position of the "lemma"? Let us examine the kind of finitist interpretation that takes the place of the "actualist truth" through the reduction concept: the concepts $\forall$ and $\exists$ are interpreted in a quite natural way ("reduced", 13.21 and 13.22), and the associated rules of inference (14.41 - 14.43) are dealt with in a correspondingly effortless manner. Not so for $\neg$; $\neg \mathcal{U}$ is interpreted as $\mathcal{U} \rightarrow 1=2$ (13.23) and in order to reduce this form further the reduction steps on antecedent formulae (13.5) are necessary. To the intuitive sense of the $\rightarrow$ there therefore corresponds a comparatively artificial and less immediately comparable reduction procedure. The difficulties which the $\supset$ and $\neg$ present to a finitist interpretation (Paragraph 11) make it indeed impossible to state a "natural" procedure.

A typical form of inference exhausting the intuitive meaning of the $\rightarrow$ is in fact the following: "From the assumptions $\mathcal{T}$ follows $\mathcal{D}$. From the assumption $\mathcal{D}$ and further assumptions Δ follows $\mathcal{C}$. Then $\mathcal{C}$ also follows from the assumptions $\mathcal{T}$, Δ ." This form of inference is implicit both in the "reductio" and in the "complete induction". Hence the reliance on the lemma (14.44) in dealing with these two rules of inference.

In the proof of the lemma the difficulty now consists in bridging the gap between the actualist meaning of the $\rightarrow$ according to which the mentioned form of inference is trivially "true" and the dissimilar finitist interpretation given by the reduction concept. The fundamental idea of the proof is this: in reducing $\mathcal{T} \rightarrow \mathcal{D}$ the $\mathcal{D}$ is referred back to "something simpler" (13.21 - 13.23). The same is done with the antecedent formula $\mathcal{D}$ in the reduction of $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ (13.51 - 13.53). From this we generally obtain two new sequents $\mathcal{T} \rightarrow \mathcal{D}^*$ and $\mathcal{D}^*, \Delta \rightarrow \mathcal{C}$; this method can be continued (complete induction on the number of logical connectives in $\mathcal{D}$) until a minimal formula takes the place of $\mathcal{D}$, and we have thus a trivial case. Yet this method does not suffice if in the reduction of the antecedent formula $\mathcal{D}$ that formula is retained. The consideration of this possibility requires a further reduction argument of a special kind (14.63).

14.6. Proof of the Lemma.

The lemma runs: "If reduction rules are known for two sequents of the form $\mathcal{T} \rightarrow \mathcal{D}$ and $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ in which no free variables and no minimal terms occur then a reduction rule can also be given for the sequent $\mathcal{T}, \Delta \rightarrow \mathcal{C}$."

The latter sequent will be called the mix sequent of the two other sequents; the formula $\mathcal{D}$ its mix formula.

In order to prove the lemma I apply a complete induction on the number of logical connectives occurring in the mix formula. I

therefore assume that the total number of these connectives is equal to a definite number s and that the lemma has already been proven for smaller s or that s is equal to 0.

14.60. Suppose therefore that two definite sequents $\mathcal{T} \rightarrow \mathcal{D}$ and $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ without free variables and minimal terms, with s logical connectives in the formula $\mathcal{D}$, are given and that for each sequent a reduction rule is known. It must then be shown that a reduction rule can also be given for the mix sequent $\mathcal{T}, \Delta \rightarrow \mathcal{C}$.

14.61. I shall first deal with the case where the sequent $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ is already in definitive form. If $\mathcal{C}$ is a true minimal formula then $\mathcal{T}, \Delta \rightarrow \mathcal{C}$ is also already in definitive form. The same holds if $\mathcal{C}$ is a false minimal formula and if in Δ a false minimal formula occurs. The case remains where $\mathcal{C}$ and $\mathcal{D}$ are false minimal formulae. In that case $\mathcal{T}, \Delta \rightarrow \mathcal{C}$ is reduced according to precisely the same rule as that provided for $\mathcal{T} \rightarrow \mathcal{D}$. Since $\mathcal{C}$ and $\mathcal{D}$ are both false minimal formulae their difference is here immaterial; and the formulae designated by Δ may be ignored altogether in the reduction (cf. 14.3).

14.62. Suppose that the sequent $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ is not yet in definitive form. In relation to the first reduction step to be carried out on the sequent I then distinguish three cases:

1. Suppose that $\mathcal{C}$ is no minimal formula.
2. Suppose that $\mathcal{C}$ is a false minimal formula and that the first prescribed reduction step for the sequent

$\mathcal{D}, \Delta \rightarrow \mathcal{C}$ (according to 13.5) does not affect the antecedent formula $\mathcal{D}$.

3. Suppose that $\mathcal{C}$ is a false minimal formula and that the mentioned reduction step (according to 13.5) affects the antecedent formula $\mathcal{D}$.

I shall now deal with each of the three cases separately.

14.621. Suppose that the first case arises. The first reduction step to be carried out on the mix sequent $\mathcal{P}, \Delta \rightarrow \mathcal{C}$ is then the step which alone is applicable according to 13.21, 13.22, 13.23, where the choice of n or $\mathcal{U}$ or $\mathcal{B}$ is free if $\mathcal{C}$ has the form $\forall x F(x)$ or $\mathcal{U} \& \mathcal{B}$. Suppose that after this reduction step (and, if necessary, successive steps according to 13.12 until no further minimal terms occur) the sequent runs $\mathcal{P}, \Delta'' \rightarrow \mathcal{C}''$. The first reduction step on the sequent $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ must necessarily be of the same kind, and in the case of a choice, the same choice may be made as above so that after the first reduction step (and possibly further necessary steps according to 13.12) this sequent assumes the form $\mathcal{D}, \Delta'' \rightarrow \mathcal{C}''$. Now the following assertion still remains to be proved which is again a special case of the lemma: "On the basis of the known reduction rules for the sequents $\mathcal{P} \rightarrow \mathcal{D}$ and $\mathcal{D}, \Delta'' \rightarrow \mathcal{C}''$ a reduction rule for their mix sequent $\mathcal{P}, \Delta'' \rightarrow \mathcal{C}''$ is also statable." I shall postpone the proof of this assertion for the time being.

14.622. Suppose that the second case arises. After its first reduction step (according to 13.5) (and possibly successive steps according to 13.12 until no further minimal terms occur) the sequent

$\mathcal{D}, \Delta \rightarrow \mathcal{C}$ runs $\mathcal{D}, \Delta'' \rightarrow \mathcal{C}''$. The reduction procedure of $\mathcal{T}, \Delta \rightarrow \mathcal{C}$ must then begin with the steps required to yield $\mathcal{T}, \Delta' \rightarrow \mathcal{C}''$ from the former sequent. In that case the

following assertion still remains to be proved, which once again is a special case of the lemma: "On the basis of the known reduction rules for the sequents $\mathcal{T} \rightarrow \mathcal{D}$ and $\mathcal{D}, \Delta'' \rightarrow \mathcal{C}''$

a reduction rule for their mix sequent

$\mathcal{T}, \Delta'' \rightarrow \mathcal{C}''$ is also statable."

14.623. Suppose that the third case arises. I distinguish three sub-cases depending on whether $\mathcal{D}$ has the form $\forall x F(x)$, $\mathcal{U} \& \mathcal{B}$ or $\neg \mathcal{U}$, i.e., depending on whether the first prescribed reduction step on $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ takes the form of 13.51, 13.52 or 13.53. The treatment of these three cases is not essentially different.

14.623. 1. Suppose that $\mathcal{D}$ has the form $\forall x F(x)$. In that case the first reduction step turns the sequent $\mathcal{D}, \Delta \rightarrow \mathcal{C}$, i.e., $\forall x F(x), \Delta \rightarrow \mathcal{C}$, into $F(n), \forall x F(x), \Delta \rightarrow \mathcal{C}$ or $F(n), \Delta \rightarrow \mathcal{C}$. The sequent $\mathcal{T} \rightarrow \mathcal{D}$ is equal to $\mathcal{T} \rightarrow \forall x F(x)$ and its first reduction step must therefore yield $\mathcal{T} \rightarrow F(m)$ (according to 13.21), with arbitrarily chosen m . In particular, we can choose the numeral n for m and obtain $\mathcal{T} \rightarrow F(n)$.

If $\mathcal{F}(\kappa)$ contains minimal terms we subject it, and the sequent dealt with before, to further reductions according to 13.12, as prescribed, until no further minimal terms occur. The two sequents then run $\pi \rightarrow (\mathcal{F}(\kappa))^*$ and $(\mathcal{F}(\kappa))^*, \forall x \mathcal{F}(x), \Delta \rightarrow \mathcal{C}$ or $(\mathcal{F}(\kappa))^*, \Delta \rightarrow \mathcal{C}$. If no minimal terms had occurred $(\mathcal{F}(\kappa))^*$ shall stand for the formula $\mathcal{F}(\kappa)$. Now I first consider the case where $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ has assumed the second form, viz., $(\mathcal{F}(\kappa))^*, \Delta \rightarrow \mathcal{C}$. Here reduction rules for the sequents $\pi \rightarrow (\mathcal{F}(\kappa))^*$ and $(\mathcal{F}(\kappa))^*, \Delta \rightarrow \mathcal{C}$ are known; I now apply the induction hypothesis, according to which the lemma is assumed to be already proven for mix formulae with fewer logical connectives than those contained in $\mathcal{D}$; from this it follows that a reduction rule is also statable for the mix sequent of the two given sequents, i.e., for the sequent $\pi, \Delta \rightarrow \mathcal{C}$. For the mix formula $(\mathcal{F}(\kappa))^*$ obviously contains one logical connective less, viz., the $\forall$, than the formula $\mathcal{D}$ which, as we know, equals $\forall x \mathcal{F}(x)$. This completes the present case.

If $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ should have assumed the more complex form $(\mathcal{F}(\kappa))^*, \forall x \mathcal{F}(x), \Delta \rightarrow \mathcal{C}$, however, the following assertion still remains to be proved: "On the basis of the known reduction rules for the sequents $\pi \rightarrow \forall x \mathcal{F}(x)$ and $\forall x \mathcal{F}(x), (\mathcal{F}(\kappa))^*, \Delta \rightarrow \mathcal{C}$ a reduction rule is also statable for their mix sequent $\pi, (\mathcal{F}(\kappa))^*, \Delta \rightarrow \mathcal{C}$."

The proof for this will be postponed for the time being; once it has been carried out the induction hypothesis can be applied as before and from the fact that reduction rules are known for the sequents

$\mathcal{T} \rightarrow (\mathcal{F}(\kappa))^*$ and $(\mathcal{F}(\kappa))^*, \mathcal{T}, \Delta \rightarrow \mathcal{C}$ it

can be inferred that a reduction rule is also statable for their mix sequent $\mathcal{T}, \mathcal{T}, \Delta \rightarrow \mathcal{C}$ and hence for $\mathcal{T}, \Delta \rightarrow \mathcal{C}$.

14.623. 2. Suppose that $\mathcal{D}$ has the form $\mathcal{U} \& \mathcal{B}$. Then the first reduction step on the sequent $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ yields $\mathcal{U}, \mathcal{U} \& \mathcal{B}, \Delta \rightarrow \mathcal{C}$ (or $\mathcal{B}, \mathcal{U} \& \mathcal{B}, \Delta \rightarrow \mathcal{C}$) or $\mathcal{U}, \Delta \rightarrow \mathcal{C}$ (or $\mathcal{B}, \Delta \rightarrow \mathcal{C}$). In the first reduction step on the sequent $\mathcal{T} \rightarrow \mathcal{U} \& \mathcal{B}$ a choice can be made in such a way that $\mathcal{T} \rightarrow \mathcal{U}$ (or $\mathcal{T} \rightarrow \mathcal{B}$) results (according to 13.22).

If $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ has assumed the form without $\mathcal{U} \& \mathcal{B}$ we apply the induction hypothesis at once: Since reduction rules are known for the sequents $\mathcal{T} \rightarrow \mathcal{U}$ (or $\mathcal{T} \rightarrow \mathcal{B}$) and $\mathcal{U}, \Delta \rightarrow \mathcal{C}$ (or $\mathcal{B}, \Delta \rightarrow \mathcal{C}$) and since the mix formula $\mathcal{U}$ (or $\mathcal{B}$) contains fewer logical connectives than $\mathcal{U} \& \mathcal{B}$, a reduction rule is also statable for the mix sequent $\mathcal{T}, \Delta \rightarrow \mathcal{C}$.

In the other case the following assertion is still to be proved:

"On the basis of the known reduction rules for the sequents

$\mathcal{T} \rightarrow \mathcal{U} \& \mathcal{B}$ and $\mathcal{U} \& \mathcal{B}, \mathcal{U}, \Delta \rightarrow \mathcal{C}$

(or $\mathcal{U} \& \mathcal{B}, \mathcal{B}, \Delta \rightarrow \mathcal{C}$) a reduction rule is also statable for their mix sequent $\mathcal{T}, \mathcal{U}, \Delta \rightarrow \mathcal{C}$ (or $\mathcal{T}, \mathcal{B}, \Delta \rightarrow \mathcal{C}$)."

For if this has been proven it follows once again through the application of the induction hypothesis that given reduction rules for $\mathcal{T} \rightarrow \mathcal{U}$ (or $\mathcal{T} \rightarrow \mathcal{B}$) and $\mathcal{U}, \mathcal{T}, \Delta \rightarrow \mathcal{C}$ (or $\mathcal{B}, \mathcal{T}, \Delta \rightarrow \mathcal{C}$) a reduction rule is also statable for the

mix sequent $\pi, \pi, \Delta \rightarrow \mathcal{C}$ and hence for $\pi, \Delta \rightarrow \mathcal{C}$.

14.623. 3. Suppose that $\mathcal{D}$ has the form $\neg \mathcal{U}$. The first reduction step then turns the sequent $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ into $\neg \mathcal{U}, \Delta \rightarrow \mathcal{U}$ or $\Delta \rightarrow \mathcal{U}$. In its first reduction step (according to 13.23) the sequent $\pi \rightarrow \neg \mathcal{U}$ then becomes $\mathcal{U}, \pi \rightarrow 1 = 2$.

If $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ has assumed the form $\Delta \rightarrow \mathcal{U}$ then we apply the induction hypothesis at once: since reduction rules are known for the sequents $\Delta \rightarrow \mathcal{U}$ and $\mathcal{U}, \pi \rightarrow 1 = 2$ and since the mix formula $\mathcal{U}$ contains fewer logical connectives than

$\neg \mathcal{U}$, a reduction rule is also statable for the mix sequent $\Delta, \pi \rightarrow 1 = 2$. The same therefore also holds for $\pi, \Delta \rightarrow \mathcal{C}$; for $\mathcal{C}$, like $1 = 2$, is a false minimal formula.

In the other case the following assertion is still to be proved:

"On the basis of the reduction rules known for the sequents

$\pi \rightarrow \neg \mathcal{U}$ and $\mathcal{U}, \Delta \rightarrow \mathcal{U}$ a reduction rule is also statable for their mix sequent $\pi, \Delta \rightarrow \mathcal{U}$." If this has been proven it follows again by the use of the induction hypothesis that given the reduction rules for $\pi, \Delta \rightarrow \mathcal{U}$ and $\mathcal{U}, \pi \rightarrow 1 = 2$ a reduction rule is also statable for the mix sequent $\pi, \Delta, \pi \rightarrow 1 = 2$ and hence also for $\pi, \Delta \rightarrow \mathcal{C}$.

14.63. Conclusion of the Proof. In several of the cases discussed an assertion was made whose proof had been postponed. In each case this assertion had the following form: "On the basis of the known reduction rules for the sequent $\pi \rightarrow \mathcal{D}$ and a sequent of the form

$\mathcal{D}, \Delta^* \rightarrow \mathcal{C}^*$ which has resulted from $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ by one or several reduction steps carried out according to the appropriate reduction rule, a reduction rule is also statable for their mix sequent $\pi, \Delta^* \rightarrow \mathcal{C}^*$. Here the sequents $\pi \rightarrow \mathcal{D}$ and $\mathcal{D}, \Delta^* \rightarrow \mathcal{C}^*$ contained no free variables and no minimal terms.

This assertion is quite obviously of the same kind as that made at 14.60 and for which the entire proof was intended. The mix formula $\mathcal{D}$ is the same as that in the earlier assertion; the sequent $\pi \rightarrow \mathcal{D}$ plays the same role; in place of $\mathcal{D}, \Delta \rightarrow \mathcal{C}$, however, there now occurs a sequent obtained from the latter by one or several reduction steps.

In order to prove the new assertion we now apply precisely the same inferences as before (14.61 to 14.623.3.); hence there (possibly) remains to be proved another assertion of the same kind, where the second sequent once again results from $\mathcal{D}, \Delta^* \rightarrow \mathcal{C}^*$ by at least one reduction step.

Continuing in this way we must reach the end in finitely many steps, i.e., the completion of the proof. This is so since the continual reduction of the sequent $\mathcal{D}, \Delta \rightarrow \mathcal{C}$ which after all, proceeds according to the reduction rule stated for that sequent, must (13.6) lead to definitive form in finitely many steps so that here no further re-interpretation is required (14.61) (as long as the case in which no re-interpretation in terms of a new assertion is required

did not arise even earlier).

In the transformation of the derivation in Paragraph 12 only quite harmless, entirely finitist concepts and inferences were required.

Of a special kind is the concept of the "reduction rule" which is central to the consistency proof. The proposition: "for a certain sequent a reduction rule is known" contains the concepts "all" and "there exists" to the extent to which it asserts that the reduction rule concerned exists and that the reduction procedure to be carried out according to the rule is defined for all possible choices of numerals to be substituted in the case where a choice arises in the reduction (13.6), and that the procedure terminates in finitely many steps, i.e., that once again there exists a natural number in each case which indicates the total number of steps. (This number generally depends on the choices made.)

The two "there exist" -concepts in the reducibility proof have actually always been used finitistically in the sense of 10.3. Hence the expressions: "a rule is known, given that a rule is statable". At 14.2, e.g., the reduction rule for logical basic sequents was stated precisely and the total number of required reduction steps can be inferred at once. In 14.3 - 14.44 it was stated in each case how an already existing reduction rule must be modified in order to obtain from it a reduction rule for a further sequent. In the remaining proof the transfinite "there exists", in connection with "there exists a reduction rule" has always been used in the finitist

sense of such a rule being given or (in the case of "introduction" -inferences) a new rule being stated.

Corresponding remarks hold for the "there exist" in relation to the total number of reduction steps; with the formulation of a reduction rule on the basis of already known reduction rules is always connected the possibility of determining the total number of newly arising (or disappearing) reduction steps.

In the lemma an essentially novel element is added by the transfinite use of the concept "follows" in expressions of the form "if a certain proposition holds then a certain other proposition also holds". Here we must recall the objections which were raised in 11.1 against the quite general use of this concept. It turns out however that in the consistency proof the "follows" occurs only in one connection: "If reduction rules are known for two particular sequents then a reduction rule is also statable for a certain third sequent formed from the former sequent." From the finitist standpoint this use of "follows" is unobjectionable; after all, no nesting whatever of "follows" -concepts occurs; here the "follows" is to be understood simply as an expression for the fact that by means of finitistically correct inferences the validity of a proposition (free from "follows" -concepts) is derivable from the validity of another proposition (also free from "follows" -concepts). (The "follows" is interpreted "meta-theoretically", as it were.) The forms of inference of the "follows" -introduction and "follows" -elimination are in harmony with this interpretation (cf. 11.1), and these are

precisely the inferences occurring in the proof of the lemma
(14.6) and in its applications (at 14.441, 14.442 and 14.443).

Complete inductions occurred repeatedly in the consistency proof
(at 14.6, 14.63, 14.443, and in still other places). These are to
be interpreted according to 10.5 and in this sense they are quite
unobjectionable even in the case where the induction hypothesis
is a transfinite proposition.

I hope that these reflections have helped to make the finitist
character of the methods of proof used in the consistency proof
sufficiently credible.

15.11. I consider it not impossible that the inferences used in
the consistency proof can be re-interpreted in terms of still more
elementary ones so that the methods of proof that have to be
presupposed as correct and which are no longer justified can be
further diminished.

FOOTNOTES

- 1) A detailed and very readable discussion of these questions is contained in D. Hilbert's paper: Über das Unendliche, Math. Annalen 95 (1926), pp. 161-190.
- 2) In this connexion cf. also:
H. Weyl, Über die neue Grundlagenkrise der Mathematik, Math. Zeitschrift 10 (1921), pp. 39-79; and
A. Fraenkel, Zehn Vorlesungen über die Grundlegung der Mengenlehre (or the relevant sections in Fraenkel's textbook on set theory).
- 3) K. Gödel, Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I, Monatsh. f. Math. u. Phys. 38 (1931), pp. 173-198.
- 4) W. Ackermann, Begründung des "tertium non datur" mittels der Hilbertschen Theorie der Widerspruchsfreiheit, Math. Annalen 93 (1925), pp. 1-36;
J. von Neumann, Zur Hilbertschen Beweistheorie, Math. Zeitschrift 26 (1927), pp. 1-46;
J. Herbrand, Sur la non-contradiction de l'Arithmétique, Journ. f. d. reine u. angew. Math. 166 (1932), pp. 1-8;
G. Gentzen, Untersuchungen über das logische Schliessen, Math. Zeitschrift 39 (1935), pp. 176-210, 405-431 (or the English translation by M.E. Szabo, Investigations into Logical Deduction, American Philosophical Quarterly, Vol. I, Number 4, Oct. 1964 and Vol. II, Number 3, July 1965).
- 5) There already exist several such formalizations and the present one follows more or less the established lines.
- 6) Since the notion of a "formula" is used quite generally for formalized propositions, the special case defined here should really be called a "number-theoretical formula". However, since no other "formulae" occur in this paper, this modifier may be omitted. Corresponding remarks apply to the notions of "term", "function symbol" etc.

- 7) I shall not interpret such a formula as "valid for arbitrary substitutions of numbers", as is usually customary in formal logic, since free variables are used in a more general sense in mathematical proofs; for example, cf. 4.53. Here, as in the case of bound $\mathfrak{f}$ -variables, we should more appropriately speak of "indeterminates" instead of "variables", yet, for better or worse, "variable" has become the generally accepted expression.
- 8) For this we could write a single formula of the form

$$(\dots ((\mathfrak{u}_1, \mathfrak{f} \mathfrak{u}_2) \mathfrak{f} \dots) \mathfrak{f} \mathfrak{u}_\mu) \supset \mathfrak{B}$$
 However, this would obscure the original structure of the mathematical proof; after all, in the proof the proposition "if $\mathfrak{u}_1$ and $\mathfrak{u}_2 \dots$ and $\mathfrak{u}_\mu$ hold, then $\mathfrak{B}$ holds" never occurred explicitly, the various propositions $\mathfrak{u}_1, \mathfrak{u}_2, \dots, \mathfrak{u}_\mu$ occurred rather as assumptions and the proposition $\mathfrak{B}$ as a consequence of these assumptions.
- 9) In my "Investigations into Logical Deduction" I used the word "sequent" in a more general sense which is here not necessary. For the benefit of the readers of that paper it should be pointed out that the logical formalism developed here corresponds essentially to the "NK-calculus" of the "Investigations". The "LK-calculus" is also suitable for the consistency proof. In fact, the proof then becomes even simpler in parts, although less "natural".
- 10) For "propositional logic" ($\&, \vee, \supset, \neg$) cf.:
 Hilbert-Ackermann, Grundzüge der theoretischen Logik, p. 33;
 for "predicate logic" ($\forall, \exists$ added) Cf.:
 K. Gödel, Die Vollständigkeit der Axiome des logischen Funktionenkalküls, Monatsh. f. Math. u. Phys. 37 (1930), pp. 349-360.
 The formalizations of the forms of inference used there can easily be shown to be equivalent with the formalization which I have chosen. (Cf. the proofs of equivalence in Section V of my "Investigations into Logical Deduction".)
- 11) Cf. W. Ackermann, Zum Hilbertschen Aufbau der reellen Zahlen, Math. Annalen 99 (1928), pp. 118-133.

- 12) The "Peano Axioms" for the natural numbers are the result of such efforts. (For example, cf. E. Landau, Grundlagen der Analysis, 1930) These axioms also contain complete induction, which I have included in the forms of inference. There is no fundamental difference between forms of inference and axioms, since logical forms of inference can also be formulated as "logical axioms" such as $\mathcal{U} \& \mathcal{B} \longrightarrow \mathcal{U}$ for the &-elimination, etc.
- 13) A proof for the "redundancy of the $\mathcal{L}$ " can be found in the book:
Hilbert-Bernays, Grundlagen der Mathematik, I (1934), pp. 422-457.
- 14) Cf. the papers by Hilbert and Weyl cited in footnotes 1) and 2).
- 15) Cf. D. Hilbert, Über das Unendliche, Math. Annalen 95 (1926), pp. 161-190.
- 16) Cf. A. Heyting, Die formalen Regeln der intuitionistischen Logik, Sitzungsberichte d. Preuss. Akad. d. Wiss., phys.-math. Kl. (1930), pp. 42-56.
- 17) K. Gödel, Zur intuitionistischen Arithmetik und Zahlentheorie, Ergebnisse eines math. Koll., Heft 4 (1933), pp. 34-38. - The result mentioned above was also discovered somewhat later by P. Bernays and myself independently of Gödel. Gödel also replaces $\mathcal{U} \supset \mathcal{B}$ by $\neg(\mathcal{U} \& \neg \mathcal{B})$, this is unnecessary in my system of rules of inference since I do not use propositional variables.
- 18) For example, cf. P. Bachmann, Die Elemente der Zahlentheorie, III, 10.
- 19) I could here also use any other false minimal formula.
- 20) Footnote added during the correction of the galley proof:
Articles 14.1 to 16.11 have been inserted in February 1936 in place of an earlier text.

- 21) Readers acquainted with set theory should note: The system of "ordinal numbers" here used is well-ordered by the $<$ -relation, and the numbers with the characteristics 0, 1, 2, 3, 4, 5, etc. correspond, in that order, to the transfinite ordinal numbers $\omega + 1$; $2^{\omega+1} = \omega + \omega$; $2^{\omega+\omega} = \omega \cdot \omega$; $2^{\omega \cdot \omega} = \omega^2$; $2^{\omega^2} = \omega(\omega)$; $2^{\omega(\omega)} = \omega[\omega(\omega)]$; etc.; the entire system corresponds to the "first ϵ -number". (In order to prove this the reader need merely consider the fact that the transition from the numbers with the characteristic ξ to the numbers with the characteristic $\xi + 1$ described above correspond to the definition rule of the power of 2, and then apply the rules of transfinite arithmetic.) The "theorem of transfinite induction" asserts nothing but the validity of transfinite induction for this segment of the second number class. The disputable aspects of general set theory do not, of course, enter into the consistency proof, since the corresponding concepts and theorems are here developed quite independently in a more elementary form than in set theory, where they are used in a much greater generality. - Similar connections between mathematical proofs or theorems and the theory of well-ordering, especially of the numbers of the second number class, are established in a paper by A. Church, A proof of freedom from contradiction, Proc. Nat. Acad. of Sc. (1935), pp. 275-281; and: E. Zermelo, Grundlagen einer allgemeinen Theorie der mathematischen Satzsysteme I, Fund. Math. 25 (1935), pp. 136-146.
- 22) Also cf. K. Gödel, Über Vollständigkeit und Widerspruchsfreiheit, Ergebnisse eines math. Koll., Heft 3 (1932), pp. 12-13.
- 23) Cf. P. Finsler, Formale Beweise und die Entscheidbarkeit, Math. Zeitschr. 25 (1926), pp. 676-682, and the paper by K. Gödel cited in footnote 3).
- 24) Cf. the paper by P. Finsler cited in footnote 23).
- 25) For example, cf.: L.E.J. Brouwer, Intuitionistische Betrachtungen über den Formalismus, Sitzungsber. d. Preuss. Akad. d. Wiss., phys.-math. Kl. (1928), pp. 48-52; and A. Heyting, Mathematische Grundlagenforschung - Intuitionismus - Bewistheorie, Ergebnisse d. Math. und ihrer Grenzgebiete 3 (1935), Heft 4.
- 26) G. Gentzen, Die Widerspruchsfreiheit der reinen Zahlentheorie, Math. Ann. 112 (1936), pp. 493-565.
- 27) G. Gentzen, Untersuchungen über das logische Schliessen, Math. Z. 39 (1935), pp. 176-210 and 405-431. In the paper cited in footnote 26), a formalism was introduced in Section IV that differs somewhat from the formalism developed in Section II. It was specifically designed for the proof in question and has no general significance.

- 28) It should be mentioned, incidentally, that all logical basic sequents are also derivable in the new system and I therefore do not really have to admit such sequents any longer. Their retention has of course certain formal advantages.
- 29) The proof of equivalence is to a large extent already given by the proof for the equivalence of the calculi NK and LK carried out in Section V of my dissertation.
- 30) In the earlier paper I have proved more generally the "reducibility" of the end-sequent of arbitrary derivations. Here I shall confine myself to consistency; this makes certain simplifications possible.
- 31) The same reasoning, incidentally, underlies the proof of the "Hauptsatz" of my dissertation.
- 32) Cf. G. Hessenberg, Grundbegriffe der Mengenlehre, Sonderdruck a. d. Abh. d. Friesschen Schule, N.F., I. Bd., Heft 4; pp. 479-706, Göttingen 1906.

BIBLIOGRAPHY

Benacerraf, P. and H. Putnam

- (1) Philosophy of Mathematics
Prentice-Hall, Inc., 1964

Beth, E.W.

- (1) The Foundations of Mathematics
North-Holland Publishing Company, Amsterdam, 1965

Heyting, A.

- (1) Intuitionism
North-Holland Publishing Company, Amsterdam, 1956

Hilbert, D. and W. Ackermann

- (1) Grundzüge der theoretischen Logik
Springer-Verlag, Berlin, 1928

Hilbert, D. and P. Bernays

- (1) Grundlagen der Mathematik, Volume I
Springer-Verlag, Berlin, 1934
- (2) Grundlagen der Mathematik, Volume II
Springer-Verlag, Berlin, 1939

Gentzen, G.

- (1) Über die Existenz unabhängiger Axiomsysteme zu unendlichen Satzsystemen
Mathematische Annalen, CVII, 1932
- (2) Über das Verhältnis zwischen intuitionistischer und klassischer Arithmetik
Unpublished galley proof, Math. Ann., 1933
- (3) Untersuchungen über das logische Schliessen, I & II
Mathematische Zeitschrift, 39, 1935

(translated by M.E. Szabo under the title:
Investigations into Logical Deduction, I & II
American Philosophical Quarterly,
Vol. I, No. 4, Oct. 1964, and Vol. II, No. 3, July 1965)

- (4) Anfängliche Fassung des Widerspruchsfreiheitsbeweises
Unpublished galley proof, Math. Ann., 1935
- (5) Die Widerspruchsfreiheit der reinen Zahlentheorie
Mathematische Annalen, CXII, 1936
- (6) Die Widerspruchsfreiheit der Stufenlogik
Mathematische Zeitschrift, XLI, Vol. 3, 1936
- (7) Der Unendlichkeitsbegriff in der Mathematik
Semester-Berichte, Münster in/W., 9. Semester,
Winter 1936-37
- (8) Unendlichkeitsbegriff und Widerspruchsfreiheit der Mathematik
IXe Congrès Int. d. Phil. VI, Logique et Mathématique,
Paris, 1937
- (9) Die gegenwärtige Lage in der math. Grundlagenforschung
Forschungen zur Logik und zur Grundlegung der
exakten Wissenschaften, New Series, No. 4,
Leipzig, Hirzel, 1938
- (10) Neue Fassung des Widerspruchsfreiheitsbeweises für die reine Zahlentheorie
Forschungen zur Logik und zur Grundlegung der
exakten Wissenschaften, New Series, No. 4, 1938,
Leipzig, Hirzel.
- (11) Beweisbarkeit und Unbeweisbarkeit von Anfangsfällen der transfiniten Induktion in der reinen Zahlentheorie
- (12) Zusammenfassung von mehreren vollständigen Induktionen zu einer einzigen
Unpublished and undated

Kleene, S.C.

- (1) Introduction to Metamathematics
North-Holland Publishing Company, Amsterdam, 1952

Kneebone, G.T.

- (1) Mathematical Logic and the Foundations of Mathematics
D. Van Nostrand Co. Ltd., Toronto, 1963.

GLOSSARY

Abbild	counterpart
Abgrenzung	demarcation delineation delimitation
Aneinanderreihung	enumeration
Anfangsstück	initial segment
Anordnung	ordering
An-sich	actual actualist
Auffassung	interpretation view
Aussagenverknüpfung	logical composition of propositions
Aussagenverknüpfungszeichen	logical connective
Äusserstes Verknüpfungszeichen	terminal connective
Bedenklich	disputable
Bedeutung	meaning significance
Begriff	concept notion
Begriffsbildung	specific concept
Beilegen	ascribe
Bestimmt	individual definite determinate
Beweismittel	method technique
Bund	cluster

Durchlaufung	running through
Eigenvariable	eigen-variable
Endform	definitive form
Endformel	end-formula
Endlichkeit	finiteness
Endsequenz	end-sequent
Endstück	ending
Entscheidbar	decidable
Ergebnis	conclusion
Erkenntniswert	cognitive value
Erreichbar	accessible
Faden	path
Finit	finitist
Formelbund	formula cluster
Gipfelpunkt	extremum
Grenzziehung	demarcation delimitation delineation
Grundsequenz	basic sequent
Herleitung	derivation
Hilfsmittel	technique of proof
Hilfsbegriff	auxiliary concept
Hinterformel	succedent formula
Höhe	level
Inhaltlicher Sinn	intuitive sense intuitive meaning

Kettenschluss	chain rule (inference)
Korrekt	correct true valid well-formed
Mischsequenz	mix-sequent
Mitteilungszeichen	syntactic variable
Nacheinander	vertical
Nebeneinander	horizontal
Numerus	characteristic
Obersequenz	upper sequent
Oberste Sequenz	uppermost sequent
Reine Zahlentheorie	elementary number theory
Richtig	true correct valid well-formed well defined
Schluss	inference
Schlussstrich	line of inference
Schlussweise	form of inference
Schnitt	cut
Sequenz	sequent
Sinn	sense meaning
Stammbaumförmig	in tree form
Strukturänderung	structural transformation
Struktur-Schlussfigur	structural inference figure

Transfinit	transfinite
übereinandergeschachtelt	nested
Verbundene Formeln	clustered formulae
Verdünnung	thinning
Verknüpfungsschlussfigur	operational inference figure
Vertauschung	interchange
Verzweigung	branching
Vorderformel	antecedent formula
Wahlfreiheit	choice
Widerspruch	contradiction
Widerspruchsfreiheit	consistency
Widerspruchsherleitung	inconsistent derivation
Zahlzeichen	numeral
zugehörig	associated relevant appropriate
Zuordnung	correlation
Zusammenziehung	contraction