

THE INCOMPLETENESS PHENOMENON

MARTIN GOLDSTERN
HAIM JUDAH

Über formal unentscheidbare Mathematica und verwandte

Von Kurt Gödel

Die Entwicklung der Mathematik hat bekanntlich dazu geführt, daß die formalisierten Systeme, in der Art, wie sie mechanischen Regeln von den verschiedensten derzeit aufgestellten formalen Systemen (von Principia Mathematica (PM)¹⁾ bis hin zu den neueren (von J. v. Neumann weiter angelehnt) andererseits. Diese beiden Systeme sind heute in der Mathematik allgemein formalisiert, d. h. auf einige wenige Axiome und Schlussregeln zurückgeführt. Es liegt daher die Vermutung nahe, daß diese Axiome und Schlussregeln dazu ausreichen, überhaupt jeden denkbaren Beweis zu führen. Im folgenden wird gezeigt, daß dies nicht der Fall ist, sondern daß es in den beiden angeführten Systemen sogar relativ einfache Probleme aus der Theorie der gewöhnlichen ganzen Zahlen gibt²⁾, die sich aus den Axiomen nicht entscheiden lassen. Dieser Umstand liegt nicht etwa an der speziellen Natur der aufgestellten Systeme, sondern gilt für eine sehr weite Klasse formaler Systeme, zu denen insbesondere alle gehören, die aus den beiden angeführten durch Hinzufügung endlich vieler Axiome entstehen³⁾, vorausgesetzt, daß durch die hinzugefügten Axiome keine falschen Sätze von der in Fußnote⁴⁾ angegebenen Art beweisbar werden.

¹⁾ Vgl. die in $\equiv$ erscheinende Zusammenfassung der Resultate dieser Arbeit.

²⁾ Zu den Axiomen des Systems PM rechnen wir insbesondere auch: Das Unendlichkeitsaxiom (in der Form: es gibt genau abzählbar viele Individuen), das Reduktions- und das Auswahlaxiom (für alle Typen).

³⁾ Vgl. A. Fraenkel, Zahl Vorlesungen über die Grundlegung der Mengenlehre, Wissensch. u. Hyp. Bd. XXXI, J. v. Neumann, Die Axiomatisierung der Mengenlehre, Math. Zentrbl. 27, 1928.

⁴⁾ D. h. genauer, es gibt unentscheidbare Sätze, in denen außer den logischen Konstanten — (nicht), $\forall$ (oder), $\exists$ (für alle), $\neg$ (keine anderen Begriffe vorkommen als $+$ (Addition), $\cdot$ (Multiplikation), beide bezogen auf natürliche Zahlen, wobei auch die Präfixe (λ) auch nur auf natürliche Zahlen beziehen dürfen. In solchen Sätzen könnte also nur Zahlenvariable, niemals Funktionsvariable vorkommen.

⁵⁾ Dabei werden in PM nur solche Axiome als verschieden gezählt, die aus einander nicht bloß durch Typenwechsel entstehen.

jede mathematische Aussage, die in den beiden Systemen überprüfbar ist, auch im Widerspruch mit

(Widerspruch mit)

† Journ. f. reine u. angew. Math. 154 (1925), 162 (1929) — Die nachfolgenden Überlegungen gelten auch für die im letzten Jahrgang von D. Hilbert u. seinen Mitarbeitern aufgestellten formalen Systeme, soweit diese bisher vorliegen. Vgl. D. Hilbert Math. Ann. 88, 45; Abh. math. math. Sem. d. Univ. Hamburg I (1922) VI (1928) P. Bernays Math. Ann. 90; J. v. Neumann Math. Zentrbl. 26, 1927 W. Ackermann Math. Ann. 93.

The Incompleteness Phenomenon



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

The Incompleteness Phenomenon

A New Course in Mathematical Logic

Martin Goldstern
Haim Judah

*Bar Ilan University
Ramat Gan, Israel*



CRC Press

Taylor & Francis Group

Boca Raton London New York

CRC Press is an imprint of the
Taylor & Francis Group, an **informa** business

AN AK PETERS BOOK

Editorial, Sales, and Customer Service Office

A K Peters, Ltd.
289 Linden Street
Wellesley, MA 02181

Copyright © 1995 by A K Peters, Ltd.

All rights reserved. No part of the material protected by this copyright notice may be reproduced or utilized in any form, electronic or mechanical, including photocopying, recording, or by any information storage and retrieval system, without written permission from the copyright owner.

Library of Congress Cataloging-in-Publication Data

Goldstern, M. (Martin)

The incompleteness phenomenon : a new course in mathematical logic
/ Martin Goldstern, Haim Judah.

p. cm.

Includes index.

ISBN 1-56881-029-6

1. Incompleteness theorems. I. Judah, Haim. II. Title.

QA9.54.J83 1995

511.3--dc20

95-39361
CIP

Cover illustration: Gödel's *Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme*, page one. Courtesy of the Archives, Institute for Advanced Study.

This book is dedicated to
Rafael Guendelman



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Contents

Foreword	ix
Introduction	xi
1 The Framework of Logic	1
1.1. Induction	1
1.2. Sentential Logic	24
1.3. First Order Logic	41
1.4. Proof Systems	76
2 Completeness	95
2.1. Enumerability	95
2.2. The Completeness Theorem	110
2.3. Nonstandard Models of Arithmetic	120
3 Model Theory	131
3.0. Introduction	131
3.1. Elementary Substructures and Chains	135
3.2. Ultraproducts and Compactness	148
3.3. Types and Countable Models	158
4 The Incompleteness Theorem	187
4.0. Introduction	187
4.1. The Language of Peano Arithmetic	189
4.2. The Axioms of Peano Arithmetic	191
4.3. Basic Theorems of Number Theory in PA	195
4.4. Encoding Finite Sequences of Numbers	204
4.5. Gödel Numbers	209
4.6. Substitution	214
4.7. The Incompleteness Theorem	216
4.8. Other Axiom Systems	219
4.9. Bounded Formulas	221
4.10. A Finer Analysis of 4.4 and 4.5	230
4.11. More on Recursive Sets and Functions	234
Index	243



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Foreword

You may have wondered: does not the shoemaker go barefoot? Mathematics boasts of being the epitome of exactness, but what is the exact meaning of proof? Construction? Computation?

Or you may be very ambitious and wonder whether we can prove theorems concerning the collection of all possible mathematical theories.

Or you may have resigned yourself to having no exact answer, as you cannot “pull yourself out of the mud,” at most you can philosophize about it.

Or you may wonder: is mathematics one body or is it fragmented into many branches; i.e. can we put it all in one framework?

Or you may be philosophically inclined and wonder whether having a proof and being true are the same.

However there is a branch of mathematics dealing exactly with those problems: LOGIC. It is one of the oldest intellectual disciplines (see Aristotle) yet also one which has developed enormously in this century.

Yes! Mathematics can deal with these problems and give exact answers with proof; i.e. we can define relevant notions and give answers.

Yes! We can define what a proof is, and show in a sense that being true and having a proof are the same (Gödel’s completeness theorem).

Yes! We cannot raise ourselves out of the mud: we cannot prove in our system that it does not have a contradiction (Gödel’s incompleteness theorem).

Yes! We can have a general theory of mathematical theories (model theory).

Yes! We can define what it means to be computable i.e. having an algorithm (for this purpose mathematical machines were invented, and you probably have met their offspring, the computers).

Saharon Shelah



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Introduction

This book is a course in Mathematical Logic. It is divided into four chapters which can be taught in two semesters. The first two chapters provide a basic background in mathematical logic. All details are explained for students not so familiar with the abstract method used in mathematical logic. The last two chapters are more sophisticated, and here we assume that the reader will be able to fill in more details; in fact, this ability is an essential step for this sphere of mathematical thinking.

Mathematical logic is the most abstract branch of mathematical thought, the most abstract human discipline. The main objective in this area is to understand the logic implicit in all mathematical thought. The difference between logic (considered as a branch of philosophy) and mathematical logic is that in mathematical logic we use and develop mathematical methods. That is, we use mathematical theorems to investigate and explain the logic implicit in mathematics. It should be clear that some of the results can also shed light on more general questions in epistemology and philosophy of science, but this is not the subject of this book. The main result in basic mathematical logic is that every “reasonable” mathematical system is intrinsically incomplete. This means that axiom systems cannot capture all semantical truths. This can also be expressed in the following way: If we assume that the human mind works in a way similar to an ideal computer, then there are mathematical problems which can never be solved by mathematicians. This is one aspect of Gödel’s famous incompleteness theorem, and the study of this phenomenon of incompleteness will be the main focus of this book. We think that the material of this book should be part of the basic background of every student in any discipline which employs deductive and formal reasoning as a part of its methodology. This definitely includes a large part of the social sciences.

Chapter 1 contains the basic material a student has to know about mathematical languages and logical systems. The most fundamental tool in this book is the concept of mathematical induction. Section 1.1 introduces this concept in a very general way through the notion of an “inductive structure.” This notion is essential for everything in the rest of the book. In Section 1.2 we study propositional logic and tautologies. In Section 1.3 we present “first order logic” as a typical example of a mathematical language, and the notion of a “model,” i.e., a mathematical structure which allows us to interpret symbols of the language. For example, a model for the language of groups (in which we can speak about “multiplication” and “inverses”) will be a group. In this section we also study the concept of validity, i.e.,

semantical truth. In Section 1.4 we will study the concept of a “formal proof.” We will define an axiom system and a deductive tool called *modus ponens*. Together they try to capture the notion of logical truth.

The main objective of Chapter 2 is to show that the syntactical concept of “provability” and the semantical concept of “validity” coincide. This theorem is called the completeness theorem (since it shows that the logical system presented in 1.4 is “complete”). An equivalent version of this theorem says that every axiom system which does not contain a contradiction will be realized in some model. This theorem was discovered by Gödel; the proof we present is due to Henkin. The main idea is as follows: If we want a sentence such as “Reagan is Batman” to be true in some model, we can consider a model with an element x , which has two names, “Reagan” and “Batman.” More generally, we will build a model from names (i.e., syntactical objects), and two names will denote the same object only if our axioms tell us that they have to. (A philosophical question for the reader: Is Reality no more than a Henkin model? We doubt it.) We start this chapter with a study of the concept of “enumerability,” which plays an important role in the proof of the completeness theorem. In Section 2.2 we present Henkin’s proof. We close the chapter by applying the completeness theorem to show that there are nonstandard models of number theory.

In Chapter 3 we deal with model theory. Model theory investigates the relation a logical theory has to its models. We exhibit several tools used in model theory, and we show how to apply them to classical problems of model theory such as finding the number of nonisomorphic countable models of a first order theory.

The last chapter is mainly devoted to Gödel’s famous incompleteness theorem, which says that any reasonable axiom system for the natural numbers is intrinsically incomplete. We start by proving that the Peano axioms for number theory are incomplete and then give a more general version of this theorem. We prove the incompleteness theorem in a simplified form to make the proof more accessible to beginners who are not especially interested in mathematical logic. The proof of the incompleteness theorem is conceptually based on the “liar’s paradox,” which was already known to the ancient Greeks. Gödel’s novel idea was to encode the language into the formal system itself. We present the details of this main tool of Gödel’s proof. The main change from the traditional proof is that we prove everything semantically, working in the model of natural numbers rather than talking about derivations from the Peano axioms. In this way we can simplify the proofs, while keeping near to the main idea underlying the incompleteness phenomenon. We close Chapter 4 with an introduction to recursion theory, a branch of mathematical logic which is closely related to the methods used in the proof of the incompleteness theorem.

We have included exercises at the end of each section. They are an intrinsic part of this book. We believe that it is impossible to *understand* mathematics without actually *doing* mathematics.

This book is based on Judah's logic lectures, given in Berkeley and Bar Ilan. We want to thank all our friends who have read the manuscript at various stages and have made important contributions. We especially thank Boaz Tsaban for preparing the index and Tzvi Scarr and Andy Lewis, who wrote early parts of the first two chapters, for their dedication to the project.

Martin Goldstern
Haim Judah



Taylor & Francis

Taylor & Francis Group

<http://taylorandfrancis.com>

Chapter 1

The Framework of Logic

1.1. Induction

Induction is the main tool used to prove theorems in mathematical logic. The best way to develop an intuitive feel for *inductive proofs* is to look at some examples. We start this section with examples that are close to ordinary experience, followed by mathematical examples. In the middle of the section we establish our *induction principle* by defining *inductive structures*. We conclude this section with a proof that the usual language for sentential logic is an inductive structure.

1.1.1. Example. Everyone has a name

We begin with the fact that everyone has parents.
Let's assume that parents with names always give names to their children.
Adam and Eve had names.
So we conclude that every person who ever lived had a name.
For if not, let Person be the first person with no name.
Person was not Adam or Eve, who had names.
So Person had parents.
Person's parents had names, since Person is the first person without a name.
But by assumption, the parents must have given Person a name.
So it cannot be that Person had no name.
Thus there cannot have been a first person without a name.
So everyone had a name.

In the previous example we were using a strong assumption about reality, namely that the initial conditions determine the future of the system forever. Clearly, systems like this do not exist in reality, but in mathematics we deal with ideal objects that are not subject to any external influence or the influence of time.

The first mathematical objects were the numbers:

1, 2, 3, 4, 5, 6, 7, 8, 9, 10

What are the natural numbers? This is a good question. The first thing we can say is that it is **not a mathematical question**, but rather, a philosophical question about mathematics. There is controversy, as always in philosophy, about the nature of the natural numbers, and the various opinions are strongly influenced by the positions the philosophers have on the existence of objects in reality. We mention a few of these positions without further remarks:

- (1) The numbers are abstract objects in a world of ideas.
- (2) The number 5, say, is what all objects with 5 components (or all sets with 5 members) have in common
- (3) The number 5 is a human category used to communicate.

Are there infinitely many numbers? Again this is a philosophical question. The following argument is usually given to “prove” that there are infinitely many numbers.

Assume that there are not infinitely many numbers. Then there must be the biggest one, call it n . Then $n + 1$ is bigger than n , a contradiction.

There are several problems with this argument. One objection is: If n is a number, why should it follow that $n + 1$ is a number? Let us replace the concept of “number” by “conceivable number,” where we call a number n “conceivable” if we can imagine somebody owning n dollars.

Thus, 5 is a conceivable number, 10000 is a conceivable number, and even 10^9 , a thousand millions (also called billion or milliard), is a conceivable number. What about 10^{12} ? What about $n + 1$, where n is the total value of all property anybody on earth owns? And if this is still conceivable, is 10^{100} conceivable? $10^{10^{100}}$?

We can also consider the following (related) argument: The universe, as perceived by us, is finite. So how can there be infinitely many numbers?

A second problem is the following: Even if we agree that for every number n there is a number $n + 1$, does that mean that the *infinite* totality of all numbers exists? Assume we want to make a list of all numbers. Even if we know that whenever we write down the number n , there will be room

for the number $n + 1$ (and time to write it down, too), does that mean we will ever have a complete list?

It is possible to avoid all this discussion by using the axiomatic method and stipulating that the natural numbers are any universe of objects and operations that satisfies our list of axioms.

For our purpose we will assume that the reader has a good feel for the natural numbers, the operations of addition, multiplication, and exponentiation, and the $<$ -relation.

We will write $\mathbb{N}$ for the set of natural numbers (including 0). Thus, $\mathbb{N} = \{0, 1, 2, \dots\}$.

1.1.2. Example. Show that for all $n > 0$ the following holds:

$$(1 + 2 + \dots + n) = \frac{n \cdot (n + 1)}{2} \quad (*)$$

Proof: by induction on n .

First Stage: $n = 1$. The sum on the left side consists of the single term 1, and the expression on the right side is $1 \cdot (1 + 1)/2 = 1$.

Second Stage:

$n = k + 1$. Assume that we already know that $1 + \dots + k = \frac{k(k+1)}{2}$. We need to show that:

$$(1 + \dots + k + (k + 1)) = \frac{(k + 1)((k + 1) + 1)}{2}.$$

So we start with:

$$(1 + \dots + k + (k + 1)) = (1 + \dots + k) + (k + 1).$$

We already know that the first sum is equal to $\frac{k(k+1)}{2}$, so we get

$$= \frac{k(k + 1)}{2} + (k + 1) = \frac{k(k + 1)}{2} + \frac{(k + 1)2}{2} = \frac{(k + 1)(k + 2)}{2}$$

which concludes the proof.

Why is the proof complete? If $(*)$ does not hold for all n , then there must be a first number n for which $(*)$ does not hold. But we just showed that there can be no such “first number”: n cannot be 1, by the “first

stage," and if $n > 1$, then $(*)$ must hold for $n - 1$. In the second stage we showed that $(*)$ must then hold for n also.

1.1.3. Example. Every convex polygon with $n \geq 3$ vertices has exactly $\frac{n(n-3)}{2}$ diagonals.

Proof: Here the first case is for $n = 3$. The polygon in this case is a triangle and it has no diagonals. This is also what the formula says.

Now, assume that the formula is correct for a polygon with n vertices. Given a convex polygon with $n + 1$ vertices, we use our "induction hypothesis" in the following way: take a subset of n vertices and connect it (using one diagonal of the original polygon) to get a closed polygon. This polygon has $\frac{n(n-3)}{2}$ diagonal all of which are also diagonals of the original polygon. In addition, one edge on this polygon is a diagonal of the original polygon. Finally we need to add $n - 2$ diagonals going from the extra vertex to each one of the other vertices except its two neighbors. The total number is therefore:

$$\frac{n(n-3)}{2} + 1 + (n-2) = \frac{(n+1)(n-2)}{2}$$

The above examples are typical inductive proofs in mathematics. This argument can in general be used in the following way to show that all the natural numbers have a certain property P :

- (a) Show that the number 0 has the property P .
- (b) Use the assumption that k has the property P in order to show that $k + 1$ has the property P .
- (c) Conclude from (a) and (b) that all natural numbers have the property P .

Why do we accept (c)? We do so because mathematical intuition says that the numbers can be generated as follows:

First we have the number 0 (by (a), 0 has property P).

From 0 we can go to the number $0 + 1 (= 1)$ (by (b) and the above, 1 also has the property P).

From 1 we can go to the number $1 + 1 (= 2)$ (by (b) and the above, 2 also has the property P).

From 2 we can go to the number $2 + 1 (= 3)$ (by (b) and the above, 3 also has the property P).

$\vdots$

From k we can go to the number $k + 1$ (using (b) and the above).

$\vdots$

Looking at this process of “generating numbers,” we can see that an inductive proof is a **schematization** of an infinite proof starting from 0 and continuing through all the numbers.

The point here is that the above description of the natural numbers may be schematized as follows: We start by assuming the existence of 0 and the existence of the operator $+1$ (this means to add 1). Then we say that n is a number if:

- (a) n is 0, or
- (b) There is a number m such that $n = m + 1$.

We can see that in this process two different concepts are involved, namely: “0” is an object given *a priori*, and it is the basis from which we build the rest of the numbers, which are formed from “0” by repeatedly applying the process of “adding 1.” This process of “adding 1” is the second concept implicit in this presentation of the natural numbers. This process is also given *a priori*, and we can think of it as the “method” of getting a new object from a previously given object. Such “methods” are usually described by functions.

An abstract view of this situation gives us the following important concepts that will be the main ingredient of all our mathematical constructions:

- (1) **Blocks:** are the objects, given *a priori* (like the object “0” in the above example).
- (2) **Operators:** are the methods, given *a priori* and used to create new objects from the previously created objects (like “adding 1” in the above example).

Now we may assume that we have a set B of blocks and a set K of operators.

What can we do with B and K ? We want to form a collection of objects using the elements of B and the operations in K . The first objects will be the objects in B . Then we can apply the operations of K to the elements of B to get new objects. Then we can again apply the operations of K to these new objects to get more objects, and we can continue this process for ever, to get a collection of objects which is denoted by $C(B, K)$.

To exemplify this construction, define the following operation:

$$\begin{aligned} s : \mathbb{N} &\rightarrow \mathbb{N} \\ n &\mapsto s(n) = n + 1. \end{aligned}$$

This function s is called the successor function. Now if $B = \{0\}$ and $K = \{s\}$, then $\mathbb{N} = C(B, K)$.

We will give two more examples of inductive structures:

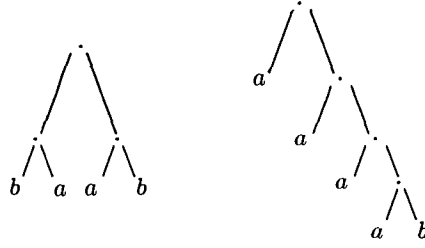


Figure 1.

1.1.4. Example. Let us consider a two element set $B = \{a, b\}$, and let D be the collection of all finite sequences of members of B . Let $K = \{f, g\}$, where

$$f : D \rightarrow D$$

$$x \mapsto f(x) = axa$$

and

$$g : D \rightarrow D$$

$$x \mapsto g(x) = bxb.$$

Then the $C(B, K)$ is the set of all sequences of odd length which are their own mirror image.

1.1.5. Example. Our set of blocks will again be the two-element set $\{a, b\}$. We will consider finite sequences that contain a , b , $.$, and the square brackets $[,]$. Our only operator will be the 2-place function f defined by

$$f(x, y) = [x.y].$$

The following are examples of elements of $C(\{a, b\}, \{f\})$:

$$a$$

$$b$$

$$[a.b]$$

$$[a.a]$$

$$[[b.a].[a.b]]$$

$$[a.[a.[a.[a.b]]]].$$

When we deal with an inductive structure, it is sometimes convenient to associate to each element of $C(B, K)$ its “syntax tree.” We will not formalize this concept, but only give a few examples. The syntax trees for the last two objects from example 1.1.5 are given in figure 1.

Now we will give a formal definition of $C(B, K)$:

1.1.6. Definition.

- (a) Every block is in $C(B, K)$ (that is, every element in B is also an element of $C(B, K)$).
- (b) If F is an n -place operator in K , and $c_1, \dots, c_n$ are elements of $C(B, K)$ then $F(c_1, \dots, c_n)$ is an element of $C(B, K)$.
- (c) Every element of $C(B, K)$ is obtained by (a) or (b).

We call $C(B, K)$ the set **generated** from B by K . If $C = C(B, K)$, then (B, K) is called an **inductive structure** on C .

Note that a given set C can be generated by different inductive structures (see exercise 7).

For example, every set can be viewed as an inductive structure by simply taking $B = C$ (i.e. taking each element of the set as a block). We also do not exclude the possibility that the result of applying an operator is again a “block.” For example, we could view the natural numbers as an inductive system with the set of blocks $= \{0, 2\}$ and one operator (the successor operation). However, this is very unnatural: Why do we need 2 as a block, if it can already be obtained from the other block, 0, by applying the successor operation twice?

In general we want an inductive structure to be a **simple description of a set** — the simpler the description, the easier it is to prove properties by induction. Often there is a unique “natural” way to define an inductive structure on a set C .

1.1.7. Example. The natural inductive structure on the set of natural numbers is given by choosing $\{0\}$ as the set of blocks, and the successor operation as the only operator.

1.1.8. Definition. Let $C = C(B, K)$ be an inductive structure, and let P be a property that elements of C may or may not have. Let F be an n -place operator in K .

We say that “ F preserves P ,” if:

Whenever $a_1, \dots, a_n$ satisfy the property P , then
 $F(a_1, \dots, a_n)$ also satisfies the property P .

1.1.9. Induction law. Let $C = C(B, K)$ be an inductive structure with B as the set of blocks and K as the set of operators such that the following is true:

- (a) Every block satisfies the property P , and
- (b) Every operator preserves the property P .

Then:

- (*) Every element of C satisfies the property P .

1.1.10. Notation. When we prove that a property P holds for all elements x of an inductive structure $C(B, K)$, we usually start by saying:

We will prove P by induction on $C(B, K)$

or

We will prove $P(x)$ by induction on x .

Such a proof consists of two parts: In the first part we deal with the blocks, i.e., we show that all blocks have the property P . (This is called the *induction basis*.) In the second part we deal with the operators, i.e., we show that all operators preserve P . This is called the *inductive step*.

Usually the essence of these proofs is in the second part. Sometimes we can give a general argument that works for all operators, sometimes we have to deal with each operator separately.

To show that an n -ary operator F preserves the property P , we assume that $a_1, \dots, a_n$ are arbitrary elements in our inductive structure satisfying P , and we have to show that $F(a_1, \dots, a_n)$ satisfies P .

The assumption “ $a_1, \dots, a_n$ satisfy P ” is often called the “induction hypothesis” or “inductive assumption.”

Induction is one of the most natural ways to deal with infinite objects. There are many examples of mathematical objects which can be seen as inductive structures. One example that we have in mind is the language for “sentential logic.” We will study the mathematical properties of this language below. We will then generalize these properties to other inductive structures. The language of sentential logic will be defined starting from a set of blocks and a set of operators. The blocks will be basic sentences like:

New York is a city.

2 is odd.

The operators will be the logical connectives used to build more complex sentences from the blocks and other sentences. For instance, “and” and “if and only if” are logical connectives.

Before we define the language of sentential logic, we will make a few remarks about formal languages. We consider an “alphabet” or “set of symbols” S . The set of all “words” in our formal language will be all

elements of S^+ , the set of all finite sequences of elements from S . (We allow only sequences of length ≥ 1 , i.e., we will not consider the empty sequence as a word.) It does not matter what the true nature of these “symbols” is, we only demand that

No symbol is also a finite sequence of symbols. (*)

We do not strictly distinguish between a symbol x and the one element sequence containing only the symbol x . This causes no ambiguities because of (*).

If x and y denote symbols, we write xy to denote the sequence with first element x and second element y , i.e., (x, y) . (Here, x may be equal to y .) We call the sequence (x, y) a “pair” and say that the sequence (x, y) has “length” 2. The set of all pairs of elements of S is called S^2 .

Similarly, we call a sequence (x, y, z) of length 3 a “triple” or “3-tuple.” A sequence $(x_1, \dots, x_n)$ of length n will be called “an n -tuple” or “word of length n .” The set of all n -tuples is written as S^n . So S^1 is the set of all words of length 1, which is essentially the same as S .

For two sequences x and y we write xy for the concatenation of these two sequences. Similarly for xyz , etc.

We will not explain what a “finite sequence” is. We assume that the reader is familiar with basic facts such as $(xy)z = x(yz)$ whenever x, y, z are finite sequences, and

$$\text{If } xr = ys, \text{ then } x = y \text{ and } r = s$$

whenever x and y are symbols and r and s are sequences.

We will often use letters such as “ x ” to denote symbols of our language. It is important to keep in mind that such a letter “ x ” is not the symbol itself, but only a name for the actual symbol. Thus it is possible that the letter “ x ” on one page denotes a different symbol than the letter “ x ” on some other page. Also, different letters may denote the same symbol (or sequence of symbols).

Before giving the explicit definitions of the language for sentential logic, let us introduce some sets.

Let $B = \{A_1, A_2, A_3, \dots\}$ be a set of distinct symbols (i.e., whenever i and j are distinct natural numbers, then “ A_i ” and “ A_j ” stand for distinct symbols) and let F be the following set with 6 elements: $F = \{\neg, \wedge, \vee, \rightarrow, \leftrightarrow, \mid\}$.

The elements of B are called **sentential symbols**. They will serve as blocks when we build our language as an inductive structure. The elements of F are called **connectives**. $\neg$ is called “not”, $\wedge$ is called “and”, $\vee$ is

called “or”, $\rightarrow$ is called “implies”, $\leftrightarrow$ is called “iff” (= if and only if), and $|$ is called “nand”.

1.1.11. Definition. The alphabet for the sentential language will be the set $S = B \cup F \cup \{ (,) \}$. S^+ is the collection of finite sequences from S .

The following are examples of members of S^+ :

$$\begin{aligned} &A_1 \\ &A_3 \\ &(\leftrightarrow \wedge A_1 (|)) \\ &()A_1 \rightarrow A_2. \end{aligned}$$

1.1.12. Definition. We define some operations over S^+ . For any two elements α, β of S^+ we define $F_{\neg}(\alpha)$, $F_{\wedge}(\alpha, \beta)$, $F_{\vee}(\alpha, \beta)$, $F_{\rightarrow}(\alpha, \beta)$, $F_{\leftrightarrow}(\alpha, \beta)$ and $F_{|}(\alpha, \beta)$ by

$$\begin{aligned} F_{\neg}(\alpha) &= (\neg\alpha) \\ F_{\wedge}(\alpha, \beta) &= (\alpha \wedge \beta) \\ F_{\rightarrow}(\alpha, \beta) &= (\alpha \rightarrow \beta) \\ F_{\vee}(\alpha, \beta) &= (\alpha \vee \beta) \\ F_{\leftrightarrow}(\alpha, \beta) &= (\alpha \leftrightarrow \beta) \\ F_{|}(\alpha, \beta) &= (\alpha | \beta) \end{aligned}$$

Each of these operators has domain S^+ and range in S^+ .

1.1.13. Definition. Let $K = \{F_{\neg}, F_{\wedge}, F_{\vee}, F_{\rightarrow}, F_{\leftrightarrow}, F_{|}\}$. Then we define $\mathcal{L}$ to be $C(B, K)$.

We will call $\mathcal{L}$ the *sentential language*, or the language of sentential logic. The elements of the sentential language will be called *sentential formulas*.

The following are examples of sentential formulas:

$$\begin{aligned} &(A_1 \vee (A_2 \vee A_3)) \\ &(((\neg A_1) \wedge A_2) \vee (A_1 \wedge (\neg A_2))) \\ &((A_1 \vee A_2) \leftrightarrow (A_2 \vee A_1)) \\ &(A_1 | A_1). \end{aligned}$$

For example $(A_1 \vee (A_2 \vee A_3))$ is a sentential formula, because

$$(A_1 \vee (A_2 \vee A_3)) = F_{\vee}(A_1, F_{\vee}(A_2, A_3)).$$

Again we will give examples of “syntax trees” of elements of $C(B, K)$.



Figure 2.

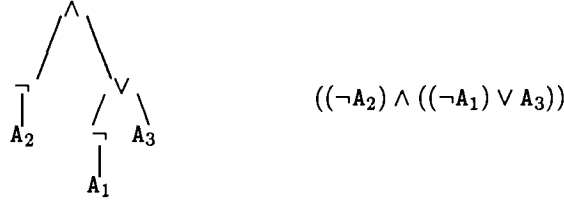


Figure 3.

For example, the syntax tree of $(A_1 \wedge A_3)$ is given in figure 2, and the syntax tree of $((\neg A_2) \wedge ((\neg A_1) \vee A_3))$ is given in figure 3.

It is not immediately obvious which elements of S^+ are in $C(B, K)$. The next proposition will help us in deciding which elements of S^+ are sentential formulas.

1.1.14. Lemma. Let α be in $\mathcal{L}$. Then the number of right parentheses in α is equal to the number of left parentheses.

Proof: by induction. Let P be the property of having an equal number of right and left parentheses.

- (a) If α is a block then α has the property P
(since the number of left parentheses = the number of right parentheses = 0).
- (b) Assume that α and β have the property P , then

$$\begin{array}{ll}
 F_{\neg}(\alpha) &= (\neg\alpha) \text{ has the property } P \\
 F_{\wedge}(\alpha, \beta) &= (\alpha \wedge \beta) \text{ has the property } P \\
 &\dots \\
 F_{|}(\alpha, \beta) &= (\alpha|\beta) \text{ has the property } P.
 \end{array}$$

So by the induction law every element of $\mathcal{L}$ has the property P (i.e. has an equal number of right and left parentheses). This ends the proof of 1.1.14.

If α and β are two strings of symbols from $\mathcal{L}$, say $\alpha = \alpha_1 \cdots \alpha_n$ and $\beta = \beta_1 \cdots \beta_m$, then by $\alpha\beta$ we mean the concatenation of the two strings, so in this case, $\alpha\beta = \alpha_1 \cdots \alpha_n \beta_1 \cdots \beta_m$.

1.1.15. Definition. We say that α in S^+ is an *initial segment* of β in S^+ if there exists γ in S^+ such that $\alpha\gamma = \beta$.

1.1.16. Example. $\wedge A_1$ is an initial segment of $\wedge A_1)A_2$.
 $(A_1\wedge$ is a initial segment of $(A_1 \wedge A_2)$.
 No string is an initial segment of itself.

1.1.17. Lemma. If α is in $\mathcal{L}$ and α' is an initial segment of α , then α' has more left parentheses than right parentheses.

Proof: by induction.

Induction Basis:

If α is in B , then there are no initial segments, so α has the property.

Induction Step, Case 1: $\alpha = (\neg\beta)$. α' can be:

- (1.a) "("
- (1.b) "(\neg"
- (1.c) "(\neg\beta'" (where β' is an initial segment of β)
- (1.d) "(\neg\beta"

In cases (1.a) and (1.b), α' has one left parenthesis and no right parentheses, so it has more left parentheses than right parentheses.

Case (1.c): If α' is $(\neg\beta'$ then by the induction hypothesis β' has more left parentheses than right parentheses. α' has the same number of right parentheses as β' , but one more left parenthesis. Hence also α' has more left than right parentheses.

Case (1.d): If $\alpha' = (\neg\beta$, then as β has the same number of left and right parentheses (by lemma 1.1.14), so α' has more left than right parentheses, since it has one left parenthesis more than β .

Induction Step, Case 2: $\alpha = (\beta \wedge \gamma)$. So α' is an initial segment of $(\beta \wedge \gamma)$. α' must be one of the following:

- (2.a) "("
- (2.b) "(\beta'", where β' is an initial segment of β
- (2.c) "(\beta"
- (2.d) "(\beta\wedge"
- (2.e) "(\beta \wedge \gamma'", where γ' is an initial segment of γ
- (2.f) "(\beta \wedge \gamma"

In cases (2.a), (2.c), (2.d), (2.f), the number of left parentheses of α' is exactly one more than the number of right parentheses, because β (and γ) have the same number of right and left parentheses.

Case (2.b): β' has more left parentheses than right parentheses, so α' , having an additional left parentheses (but the same number of right parentheses as β'), has more left parentheses than right parentheses.