

Pure and Applied Mathematics
A Wiley-Interscience Series of Texts, Monographs, and Surveys

Convexity and Optimization in $\mathbb{R}^n$

Leonard D. Berkovitz



CONVEXITY AND OPTIMIZATION IN $\mathbb{R}^n$

LEONARDO B. ABSCOVITE
Purdue University



A Wiley-Interscience Publication
JOHN WILEY & SONS, INC.

CONVEXITY AND OPTIMIZATION IN $\mathbb{R}^n$

CONVEXITY AND OPTIMIZATION IN $\mathbb{R}^n$

LEONARDO B. ABREU
Purdue University



A Wiley-Interscience Publication
JOHN WILEY & SONS, INC.

This book is printed on acid-free paper. 

Copyright © 2002 by John Wiley & Sons, Inc., New York, NY. All rights reserved.

Published simultaneously in Canada.

No part of this publication may be reproduced, stored in a retrieval system or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, scanning or otherwise, except as permitted under Sections 107 or 108 of the 1976 United States Copyright Act, without either the prior written permission of the Publisher, or authorization through payment of the appropriate per-copy fee to the Copyright Clearance Center, 222 Rosewood Drive, Danvers, MA 01923, (978) 7508400 for (978) 7508744. Requests to the Publisher for permission should be addressed to the Permissions Department, John Wiley & Sons, Inc., 605 Third Avenue, New York, NY 10158-0001; (212) 8506011, fax (212) 8506050, E-mail: PERMREQ @ WILEY.COM.

For ordering and customer service, call 1-800-CALL-WILEY.

Library of Congress Cataloging-in-Publication Data

Berkson, Leonard David. 1926-

Convexity and optimization in R / [supervised by] Leonard D. Berkson.

p. cm.—(Pure and applied mathematics)

1. A Wiley-Interscience publication.

Includes bibliographical references and index.

ISBN 0-471-55201-0 (pbk.: alk. paper)

1. Convex sets. 2. Mathematical optimization. I. Title. II. Pure and applied mathematics (Wiley & Sons Association).

QA646.B46 2002

517.94—dc22

2001040391

Printed in the United States of America.

10 9 8 7 6 5 4 3 2 1

To my wife, Anne

CONTENTS

Preface	ix
I Topics in Real Analysis	1
1. Introduction	1
2. Vectors in $\mathbb{R}^n$	1
3. Algebra of Sets	4
4. Metric Topology of $\mathbb{R}^n$	5
5. Limits and Continuity	6
6. Basic Property of Real Numbers	11
7. Compactness	14
8. Equivalent Norms and Cartesian Products	16
9. Fundamental Existence Theorem	18
10. Linear Transformations	21
11. Differentiation in $\mathbb{R}^n$	22
II Convex Sets in $\mathbb{R}^n$	26
1. Lines and Hyperplanes in $\mathbb{R}^n$	26
2. Properties of Convex Sets	33
3. Separation Theorems	42
4. Supporting Hyperplanes: Extreme Points	56
5. Systems of Linear Inequalities: Theorems of the Alternative	61
6. Affine Geometry	68
7. More on Separation and Support	80
III Convex Functions	87
1. Definition and Elementary Properties	87
2. Subgradients	102
3. Differentiable Convex Functions	106
4. Alternative Theorems for Convex Functions	112
5. Application to Game Theory	117

IV Optimization Problems	128
1. Introduction	128
2. Differentiable Unconstrained Problems	129
3. Optimization of Convex Functions	137
4. Linear Programming Problems	138
5. First-Order Conditions for Differentiable Nonlinear Programming Problems	143
6. Second-Order Conditions	161
V Convex Programming and Duality	178
1. Problem Statement	178
2. Necessary Conditions and Sufficient Conditions	181
3. Perturbation Theory	188
4. Lagrangian Duality	200
5. Geometric Interpretation	207
6. Quadratic Programming	213
7. Duality in Linear Programming	215
VI Simplex Method	221
1. Introduction	221
2. Extreme Points of Feasible Set	223
3. Preliminaries to Simplex Method	230
4. Phase II of Simplex Method	234
5. Termination and Cycling	249
6. Phase I of Simplex Method	251
7. Revised Simplex Method	255
Bibliography	261
Index	263

PREFACE

This book presents the mathematics of finite-dimensional optimization, focusing those aspects of convexity that are useful in this context. It provides a basis for the further study of convexity, of more general optimization problems, and of numerical algorithms for the solution of finite-dimensional optimization problems. The intended audience consists of beginning graduate students in engineering, economics, operations research, and mathematics and strong undergraduates in mathematics. This was the audience in a one-semester course at Purdue, MA 521, from which this book evolved.

Ideally, the prerequisites for reading this book are good introductory courses in real analysis and linear algebra. In teaching MA 521, I found that while the mathematics students had the real analysis prerequisites, many of the other students who took the course because of their interest in optimization did not have this prerequisite. Chapter I is for those students and readers who do not have the real analysis prerequisite; in it I present those concepts and results from real analysis that are needed. Except for the Weierstrass theorem on the existence of a minimum, the “heavy” or “deep” theorems are stated without proof. Students without the real variables prerequisite found the material difficult at first, but most managed to assimilate it at a satisfactory level. The advice to readers for whom this is the first encounter with the material in Chapter I is to make a serious effort to master it and to return to it as it is used in the sequel.

To address as wide an audience as possible, I have not always presented the most general result or argument. Thus, in Chapter II I chose the “cleanest point” approach to separation theorems, rather than more generally valid arguments, because I believe it to be more intuitive and straightforward for the intended audience. Readers who wish to get the best possible separation theorems in finite dimensions should read Sections 6 and 7 of Chapter II. In proving the Farkas Lemma Theorem, I used a penalty function argument due to Motzkin rather than more technical arguments involving linearizations. I limited the discussion of duality to Lagrangian duality and did not consider Fenchel duality, since the latter would require the development of more mathematical machinery.

The numbering system and reference system for theorems, lemmas, remarks, and corollaries is the following. Within a given chapter, theorems, lemmas, and remarks are numbered consecutively in each section, preceded by the section number. Thus, the first theorem of Section 1 is Theorem 1.1, the second, Theorem 1.2, and so on. The same applies to lemmas and remarks. Corollaries are numbered consecutively within each section without a reference to the section number. Reference to a theorem in the same chapter is given by the theorem number. Reference to a theorem in a chapter different from the current one is given by the theorem number preceded by the chapter number in Roman numerals. Thus, a reference in Chapter IV to Theorem 4.2 in Chapter II would be Theorem II.4.2. References to lemmas and remarks are similar. References to corollaries within the same section are given by the number of the corollary. References to corollaries in a different section of the same chapter are given by prefixing the section number to the corollary number; references in a different chapter are given by prefixing the chapter number in Roman numerals to the preceding.

I thank Rita Sturges and John Gregory for reading the current notes version of this book and for their corrections and suggestions for improvement. I thank Terry Conboe for preparing the figures. I also thank Betty Glick for typing, seemingly innumerable versions and revisions of the notes for MA 321. Her skill and cooperation contributed greatly to the success of this project.

LEONARD D. BERKOVITZ

West Lafayette, Indiana

I

TOPICS IN REAL ANALYSIS

1. INTRODUCTION

The serious study of convexity and optimization problems in $\mathbb{R}^n$ requires some background in real analysis and in linear algebra. In teaching a course based on notes from which this text evolved, the author and his colleagues assumed that the students had an undergraduate course in linear algebra but did not necessarily have a background in real analysis. The purpose of this chapter is to provide the reader with most of the necessary background in analysis. Not all statements will be proved. The reader, however, is expected to understand the definitions and the theorems and is expected to follow the proofs of statements whenever the proofs are given. The bad news is that many readers will find this chapter to be the most difficult one in the text. The good news is that careful study of this material will provide background for many other courses and that subsequent chapters should be easier. If necessary, the reader should return to this chapter when encountering this material later on.

2. VECTORS IN $\mathbb{R}^n$

By *column n -space*, or $\mathbb{R}^n$, we mean the set of all *column* $\mathbf{x} = (x_1, \dots, x_n)$, where the x_i , $i = 1, \dots, n$ are real numbers. Thus, $\mathbb{R}^n$ is a generalization of the familiar two- and three-dimensional spaces $\mathbb{R}^2$ and $\mathbb{R}^3$. The elements $\mathbf{x}$ of $\mathbb{R}^n$ are called *vectors* or *points*. We will often identify the vector $\mathbf{x} = (x_1, \dots, x_n)$ with the $n \times 1$ matrix

$$\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}$$

and abuse the use of the equal sign to write

$$\mathbf{x} = \begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix}.$$

In this case we shall call $\mathbf{x}$ a column vector. We shall also identify $\mathbf{x}$ with the $1 \times n$ matrix $(x_1, \dots, x_n)$ and write $\mathbf{x} = (x_1, \dots, x_n)$. In this case we shall call $\mathbf{x}$ a row vector. It will usually be clear from the context whether we consider $\mathbf{x}$ to be a row vector or a column vector. When there is danger of confusion, we will identify $\mathbf{x}$ with the column vector and use the transpose symbol, which is a superscript t , to denote the row vector. Thus $\mathbf{x}^t = (x_1, \dots, x_n)$.

We define two operations, vector addition and multiplication by a scalar. If $\mathbf{x} = (x_1, \dots, x_n)$ and $\mathbf{y} = (y_1, \dots, y_n)$ are two vectors, we define their sum $\mathbf{x} + \mathbf{y}$ to be

$$\mathbf{x} + \mathbf{y} = (x_1 + y_1, \dots, x_n + y_n).$$

For any scalar, or real number, α we define $\alpha\mathbf{x}$ to be

$$\alpha\mathbf{x} = (\alpha x_1, \dots, \alpha x_n).$$

We assume that the reader is familiar with the properties of these operations and knows that under these operations $\mathbb{R}^n$ is a vector space over the reals.

Another important operation is the inner product, or dot product, of two vectors, denoted by $\langle \mathbf{x}, \mathbf{y} \rangle$ or $\mathbf{x} \cdot \mathbf{y}$ and defined by

$$\langle \mathbf{x}, \mathbf{y} \rangle = \sum_{i=1}^n x_i y_i.$$

Again, we assume that the reader is familiar with the properties of the inner product. We use the inner product to define the norm $\|\cdot\|$ or length of a vector $\mathbf{x}$ as follows:

$$\|\mathbf{x}\| = \langle \mathbf{x}, \mathbf{x} \rangle^{1/2} = \left[\sum_{i=1}^n x_i^2 \right]^{1/2}.$$

This norm is called the euclidean norm. In $\mathbb{R}^2$ and $\mathbb{R}^3$ the euclidean norm reduces to the familiar length. It is straightforward to show that the norm has the following properties:

$$\|\mathbf{x}\| \geq 0 \quad \text{for all } \mathbf{x} \text{ in } \mathbb{R}^n, \quad (1)$$

$$\|\mathbf{x}\| = 0 \quad \text{if and only if } \mathbf{x} \text{ is the zero vector } \mathbf{0} \text{ in } \mathbb{R}^n, \quad (2)$$

$$\|\alpha\mathbf{x}\| = |\alpha| \|\mathbf{x}\| \quad \text{for all real numbers } \alpha \text{ and vectors } \mathbf{x} \text{ in } \mathbb{R}^n. \quad (3)$$

The norm has two additional properties, which we will prove:

For all vectors $\mathbf{x}$ and $\mathbf{y}$ in $\mathbb{R}^n$

$$|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\| \|\mathbf{y}\| \quad (4)$$

with equality holding if and only if one of the vectors is a scalar multiple of the other.

For all vectors $\mathbf{x}$ and $\mathbf{y}$ in $\mathbb{R}^n$

$$\|\mathbf{x} + \mathbf{y}\| \leq \|\mathbf{x}\| + \|\mathbf{y}\|, \quad (5)$$

with equality holding if and only if one of the vectors is a nonnegative scalar multiple of the other.

The inequality (4) is known as the Cauchy–Schwarz inequality. The reader should not be discouraged by the slickness of the proof that we will give or by the slickness of other proofs. Usually, the first time that a mathematical result is discovered, the proof is rather complicated. Often, many very smart people then look at it and constantly improve the proof until a relatively simple and clever argument is found. Therefore, do not be intimidated by clever arguments. New to the proof.

Let $\mathbf{x}$ and $\mathbf{y}$ be any pair of fixed vectors in $\mathbb{R}^n$. Let λ be any real scalar. Then

$$\begin{aligned} 0 &\leq \|\mathbf{x} + \lambda\mathbf{y}\|^2 = \langle \mathbf{x} + \lambda\mathbf{y}, \mathbf{x} + \lambda\mathbf{y} \rangle \\ &= \|\mathbf{x}\|^2 + 2\lambda\langle \mathbf{x}, \mathbf{y} \rangle + \|\mathbf{y}\|^2\lambda^2 \quad \text{for all } \lambda. \end{aligned} \quad (6)$$

This says that the quadratic in λ on the right either has a double real root or has no real roots. Therefore, by the quadratic formula, its discriminant must satisfy

$$4\langle \mathbf{x}, \mathbf{y} \rangle^2 - 4\|\mathbf{x}\|^2\|\mathbf{y}\|^2 \leq 0.$$

Transposing and taking square roots give

$$|\langle \mathbf{x}, \mathbf{y} \rangle| \leq \|\mathbf{x}\| \|\mathbf{y}\|.$$

Equality holds if and only if the quadratic has a double real root, say $\lambda = \alpha$. But then, from (6), $\|\mathbf{x} + \alpha\mathbf{y}\| = 0$. Hence by (2) $\mathbf{x} + \alpha\mathbf{y} = \mathbf{0}$, and $\mathbf{x} = -\alpha\mathbf{y}$.

To prove (5), we write

$$\begin{aligned} \|\mathbf{x} + \mathbf{y}\|^2 &= \langle \mathbf{x} + \mathbf{y}, \mathbf{x} + \mathbf{y} \rangle = \|\mathbf{x}\|^2 + 2\langle \mathbf{x}, \mathbf{y} \rangle + \|\mathbf{y}\|^2 \\ &\leq \|\mathbf{x}\|^2 + 2\|\mathbf{x}\| \|\mathbf{y}\| + \|\mathbf{y}\|^2 = (\|\mathbf{x}\| + \|\mathbf{y}\|)^2, \end{aligned} \quad (7)$$

where the inequality follows from the Cauchy–Schwarz inequality. If we now take square roots, we get the triangle inequality. We obtain equality in (7), and hence in the triangle inequality, if and only if $\langle \mathbf{x}, \mathbf{y} \rangle$ is nonnegative and either $\mathbf{y}$ is a scalar multiple of $\mathbf{x}$ or $\mathbf{x}$ is a scalar multiple of $\mathbf{y}$. But under these circumstances, for $\langle \mathbf{x}, \mathbf{y} \rangle$ to be nonnegative, the multiple must be nonnegative. The inequality (5) is called the triangle inequality because when it is applied to vectors in $\mathbb{R}^2$ or $\mathbb{R}^3$ it says that the length of a third side of a triangle is less than or equal to the sum of the lengths of the other two sides.

Exercise 1.1. For any vectors x and y in $\mathbb{R}^2$, show that $|x + y|^2 + |x - y|^2 = 2|x|^2 + 2|y|^2$. Interpret this relation as a statement about parallelograms in $\mathbb{R}^2$ and $\mathbb{R}^3$.

In elementary courses it is shown that in $\mathbb{R}^2$ and $\mathbb{R}^3$ the cosine of the angle θ , $0 \leq \theta \leq \pi$, between two nonzero vectors x and y is given by the formula

$$\cos \theta = \frac{\langle x, y \rangle}{|x| |y|}.$$

Thus, x and y are orthogonal if and only if $\langle x, y \rangle = 0$. In $\mathbb{R}^n$, the Cauchy-Schwarz inequality allows us to give meaning to the notion of angle between nonzero vectors as follows. Since $|\langle x, y \rangle| \leq |x| |y|$, the absolute value of the quotient

$$\frac{\langle x, y \rangle}{|x| |y|} \quad (8)$$

is less than or equal to 1, and thus (8) is the cosine of an angle between zero and π . We define this angle to be the angle between x and y . We say that x and y are orthogonal if $\langle x, y \rangle = 0$.

3. ALGEBRA OF SETS

Given two sets A and B , the set A is said to be a subset of B , or to be contained in B , and written $A \subset B$, if every element of A is also an element of B . The set A is said to be a proper subset of B if A is a subset of B and there exists at least one element of B that is not in A . This is written as $A \subsetneq B$. Two sets A and B are said to be equal, written $A = B$, if $A \subset B$ and $B \subset A$. The empty set, or null set, is the set that has no members. The empty set will be denoted by $\emptyset$. The notation $x \in S$ will mean " x belongs to the set S " or " x is an element of the set S ."

Let $\{S_\alpha\}_{\alpha \in A}$ be a collection of sets indexed by an index set A . By the union of the sets $\{S_\alpha\}_{\alpha \in A}$, denoted by $\bigcup_{\alpha \in A} S_\alpha$, we mean the set consisting of all elements that belong to at least one of the S_α . Note that if not all of the S_α are empty, then $\bigcup_{\alpha \in A} S_\alpha$ is not empty. By the intersection of the sets $\{S_\alpha\}_{\alpha \in A}$, denoted by $\bigcap_{\alpha \in A} S_\alpha$, we mean the set of points x with the property that x belongs to every set S_α . Note that for a collection $\{S_\alpha\}_{\alpha \in A}$ of nonempty sets, the intersection can be empty.

Let X be a set, say $\mathbb{R}^2$ for example, and let S be a subset of X . By the complement of S (relative to X), written cS , we mean the set consisting of those points in X that do not belong to S . Note that $c\emptyset$ is the empty set. By convention, we take the complement of the empty set to be X . We thus always have $c(\emptyset) = X$. Also, if $A \subset B$, then $cA \supset cB$. The reader should convince himself or herself of the truth of the following lemma, known as de Morgan's

law, either by writing out a proof or drawing Venn diagrams for some simple cases.

Lemma 3.1. Let $\{S_\alpha\}_{\alpha \in A}$ be a collection of subsets of a set X . Then

$$\begin{aligned}\bigcup_{\alpha \in A} S_\alpha &= \left[\bigcap_{\alpha \in A} (X \setminus S_\alpha) \right]^c, \\ \bigcap_{\alpha \in A} S_\alpha &= \left[\bigcup_{\alpha \in A} (X \setminus S_\alpha) \right]^c.\end{aligned}$$

4. METRIC TOPOLOGY OF $\mathbb{R}^n$

Let X be a set. A function d that assigns a real number to each pair of points (x, y) with $x \in X$ and $y \in X$ is said to be a *metric*, or *distance function*, on X if it satisfies the following:

$$d(x, y) \geq 0, \quad \text{with equality holding if and only if } x = y, \quad (1)$$

$$d(x, y) = d(y, x) \quad (2)$$

$$d(x, y) \leq d(x, z) + d(z, y) \quad \text{for all } x, y, z \in X. \quad (3)$$

The last inequality is called the *triangle inequality*.

In $\mathbb{R}^n$ we define a function d on pairs of points x, y by the formula

$$d(x, y) = \|x - y\| = \left(\sum_{i=1}^n (x_i - y_i)^2 \right)^{1/2}. \quad (4)$$

Example 4b. Use the properties of the norm to show that the function d defined by (4) is a metric, or distance function, on $\mathbb{R}^n$.

A set X with metric d is called a *metric space*. Although metric spaces occur in many areas of mathematics and applications, we shall confine ourselves to $\mathbb{R}^n$. In $\mathbb{R}^n$ and $\mathbb{R}^2$ the function defined in (4) is the ordinary euclidean distance.

We now present some important definitions to the reader, who should master them and try to improve his or her understanding by drawing pictures in $\mathbb{R}^2$.

The (open) *ball* centered at x with radius $r > 0$, denoted by $B(x, r)$, is defined to be the set of points y whose distance from x is less than r . We write this in symbols as

$$B(x, r) = \{y : \|y - x\| < r\}.$$

The *closed ball* $\overline{B}(x, r)$ with center at x and radius $r > 0$ is defined by

$$\overline{B}(x, r) = \{y : \|y - x\| \leq r\}.$$

Let $S \subseteq \mathbb{R}^n$. A point x is an *interior point* of S if there exists an $r > 0$ such that $B(x, r) \subseteq S$. A set S need not have any interior points. For example, consider the set of points in the plane of the form $(x_1, 0)$ with $0 < x_1 < 1$. This is the interval $0 < x_1 < 1$ on the x_1 -axis. It has no interior points, considered as a set in $\mathbb{R}^2$. On the other hand, if we consider this as a set in $\mathbb{R}^1$, every point is an interior point. This leads to the important observation that for a given set, whether or not it has interior points may depend on in which Euclidean space $\mathbb{R}^n$ we take the set to be lying, or embedded.

If the set of interior points of a set S is not empty, then we call the set of interior points the *interior* of S and denote it by $\text{int} S$.

A set X is said to be *open* if all points of X are interior points. Thus an equivalent definition of an open set is that it is equal to its interior. If we go back to the definition of interior point, we can restate the definition of an open set as follows. A set S is open if for every point x in S there exists a positive number $r > 0$, which may depend on x , such that the ball $B(x, r)$ is contained in S . The last definition, while being wordy, is the one that the reader should picture mentally.

A point x is said to be a *limit point* of a set S if for every $\epsilon > 0$ there exists a point $x_\epsilon \neq x$ such that x_ϵ belongs to S and $x_\epsilon \in B(x, \epsilon)$. The point x_ϵ will in general depend on x . A set need not have any limit points, and a limit point of a set need not belong to the set. An example of a set without a limit point is the set $\mathbb{N}$ of positive integers, considered as a set in $\mathbb{R}^1$. For an example of a limit point that does not belong to a set, consider the set $S = \{n\pi : n = 1, 2, 3, \dots\}$ in $\mathbb{R}^1$. Zero is a limit point of the set, yet zero does not belong to S . We shall denote the set of limit points of a set S by S' .

Exercise 4.1. (a) Sketch the graph of $y = \sin(1/x)$, $x > 0$.

(b) Consider the graph as a set in $\mathbb{R}^2$ and find the limit points of this set.

The *closure* of a set S , denoted by $\bar{S}$, is defined to be the set $S \cup S'$; that is, $\bar{S} = S \cup S'$. A set S is said to be *closed* if S contains all of its limit points; that is, $S' \subseteq S$. A set S is closed if and only if $\bar{S} = S$. To see this, note that $\bar{S} = S$ and the definitions $\bar{S} = S \cup S'$ imply that $S' \subseteq S \cup S'$. Hence $S' \subseteq S$. On the other hand, if $S' \subseteq S$, then $S \cup S' = S$, and so $\bar{S} = S \cup S' = S$.

A set can be neither open nor closed. In $\mathbb{R}^2$ consider $B(0, 1)$, the ball with center at the origin and radius 1. It should be intuitively clear that all points x in $\mathbb{R}^2$ with $|x| = 1$ are limit points of $B(0, 1)$. Now consider the set

$$S = B(0, 1) \cup \{x = (x_1, x_2) : |x| = 1, x_1 \geq 0\}.$$

(The reader should sketch this set.) Points $x = (x_1, x_2)$ with $|x| = 1$ and $x_1 > 0$ are not interior points of S since for such an x , no matter how small we choose $r > 0$, the ball $B(x, r)$ will not belong to S . Hence S is not open. Also, S is not closed since points $x = (x_1, x_2)$ with $|x| = 1$ and $x_1 < 0$ are limit points of S , yet they do not belong to S .

It follows from our definitions that $\mathbb{R}^n$ itself is both open and closed. By convention, we will also take the empty set to be open and closed.

THEOREM 4.1. *Complements of closed sets are open. Complements of open sets are closed.*

Proof. If $S = \emptyset$, where $\emptyset$ denotes the empty set, or if $S = \mathbb{R}^n$, then there is nothing to prove. Let $S \neq \emptyset$, $S \neq \mathbb{R}^n$ and let S be closed. Let x be any point in cS . Then $x \notin S$ and x is not a limit point of S . If we recall the definition of " x is a limit point of S ," we see that the statement " x is not a limit point of S " means that there is an $\epsilon_0 > 0$ such that all points $y \neq x$ with $y \in B(x, \epsilon_0)$ are in cS . Since $x \notin S$, this means that $B(x, \epsilon_0) \subset cS$, and so cS is open. Now let $S = \emptyset$, $S \neq \mathbb{R}^n$ and let S be open. Let x be a limit point of cS . Then for every $\epsilon > 0$ there exists a point $x_\epsilon \neq x$ in cS such that $x_\epsilon \in B(x, \epsilon)$. Since $x_\epsilon \in cS$, the ball $B(x_\epsilon, \epsilon)$ does not belong to S . Thus x cannot belong to the open set S , for then we would be able to find an ϵ_0 such that $B(x, \epsilon_0)$ did belong to S . Hence $x \in cS$, so cS is closed.

Exercise 4.3. Show that for $x \in \mathbb{R}^n$ and $r > 0$ the set $B(x, r)$ is open; that is, show that an open ball is open.

Exercise 4.4. Show that for $x \in \mathbb{R}^n$ and $r > 0$ the closed ball $\overline{B(x, r)}$ is closed.

Exercise 4.5. Show that any finite set of points $x_1, \dots, x_k$ in $\mathbb{R}^n$ is closed.

Exercise 4.6. Show that in $\mathbb{R}^n$ no point x with $\|x\| = 1$ is an interior point of $B(0, 1)$.

Exercise 4.7. Show that for any set S in $\mathbb{R}^n$ the set $\bar{S}$ is closed.

Exercise 4.8. Show that for any set S the closure $\bar{S}$ is equal to the intersection of all closed sets containing S .

THEOREM 4.2. (i) Let $\{O_\alpha\}_{\alpha \in A}$ be a collection of open sets. Then $\bigcup_{\alpha \in A} O_\alpha$ is open.
 (ii) Let $O_1, \dots, O_n$ be a finite collection of open sets. Then $\bigcap_{i=1}^n O_i$ is open.
 (iii) Let $\{F_\alpha\}_{\alpha \in A}$ be a collection of closed sets. Then $\bigcap_{\alpha \in A} F_\alpha$ is closed.
 (iv) Let $F_1, \dots, F_n$ be a finite collection of closed sets. Then $\bigcup_{i=1}^n F_i$ is closed.

Note that an infinite collection of open sets need not have an intersection that is open. For example, in $\mathbb{R}^1$, for each $n = 1, 2, 3, \dots$ let

$$O_n = (1 - 1/n, 1 + 1/n) = \{x \in \mathbb{R}^1 : 1 - 1/n < x < 1 + 1/n\}.$$

In $\mathbb{R}^1$ each O_n is open, and $\bigcap_{n=1}^\infty O_n = \{1\}$, a point. By Exercise 4.5, the set consisting of the point $x = 1$ is closed. Similarly, the union of an infinite number of closed sets need not be closed. To see this, for each $n = 1, 2, 3, \dots$ let $F_n = [0, 1 - 1/n(1 + 1/n)] = \{x \in \mathbb{R}^1 : 0 \leq x \leq 1 - 1/n(1 + 1/n)\}$. Each F_n is closed, and

$\bigcup_{i=1}^n F_i = [0, 1] = \{x \in \mathbb{R} : 0 \leq x < 1\}$, which is not closed since it does not contain the limit point $x = 1$.

We now prove the theorem. To establish (i), we take an arbitrary element x in $\bigcup_{i=1}^n O_i$. Since x is in the union, it must belong to at least one of the sets $\{O_i\}_{i=1, \dots, n}$. Say O_{i_0} . Since O_{i_0} is open, there exists an $r_0 > 0$ such that $B(x, r_0) \subseteq O_{i_0}$. Hence $B(x, r_0) \subseteq \bigcup_{i=1}^n O_i$ and thus the union is open. To establish (ii), let x be an arbitrary point in $\bigcap_{i=1}^n O_i$. Then $x \in O_i$ for each $i = 1, \dots, n$. Since each O_i is open, for each $i = 1, \dots, n$ there exists an $r_i > 0$ such that $B(x, r_i) \subseteq O_i$. Let $r = \min\{r_i : i = 1, \dots, n\}$. Then $B(x, r) \subseteq B(x, r_i) \subseteq O_i$ for each $i = 1, \dots, n$. Hence $B(x, r) \subseteq \bigcap_{i=1}^n O_i$ and thus the intersection is open.

The most efficient way to establish (iii) is to write

$$\bigcap_{i=1}^n F_i = \bigcap_{i=1}^n c(F_i) = c\left[\bigcup_{i=1}^n c(F_i)\right],$$

where the second equality follows from Lemma 3.1. Since F_i is closed, cF_i is open and so is $\bigcup_{i=1}^n (cF_i)$. Hence the complement of this set is closed, and therefore so is $\bigcap_{i=1}^n F_i$. To establish (iv), write

$$\bigcup_{i=1}^n F_i = \bigcup_{i=1}^n c(cF_i) = c\left[\bigcap_{i=1}^n c(F_i)\right].$$

Again, cF_i is open for each $i = 1, \dots, n$. Hence the intersection $\bigcap_{i=1}^n (cF_i)$ is open and the complement of the intersection is closed.

The reader who wishes to test his or her mastery of the relevant definitions should prove statements (iii) and (iv) of the theorem directly.

6. LIMITS AND CONTINUITY

A *sequence* in $\mathbb{R}^n$ is a function that assigns to each positive integer k a vector or point $\mathbf{a}_k$ in $\mathbb{R}^n$. We usually write the sequence as $\{\mathbf{a}_k\}_{k=1}^\infty$ or $\{\mathbf{a}_k\}$, rather than in the usual function notation. Examples of sequences in $\mathbb{R}^2$ are

- (i) $\{\mathbf{a}_k\} = \left\{k, k^2\right\}$,
- (ii) $\{\mathbf{a}_k\} = \left\{\left(\cos \frac{1}{k}, \sin \frac{1}{k}\right)\right\}$,
- (iii) $\{\mathbf{a}_k\} = \left\{\left(\frac{1-1^k}{2}, \frac{1}{2}\right)\right\}$,
- (iv) $\{\mathbf{a}_k\} = \left\{\left(1-1^k + \frac{1}{k}, 1-1^k - \frac{1}{k}\right)\right\}$.

The reader should plot the points in these sequences.

A sequence $\{x_k\}$ is said to have a limit y or to converge to y if for every $\epsilon > 0$ there is a positive integer $K(\epsilon)$ such that whenever $k > K(\epsilon)$, $x_k \in B(y, \epsilon)$. We write

$$\lim_{k \rightarrow \infty} x_k = y \quad \text{or} \quad x_k \rightarrow y.$$

The sequences (i), (ii), and (iv) in (1) do not converge, while the sequence (iii) converges to $0 = (0, 0)$. Sequences that converge are said to be convergent; sequences that do not converge are said to be divergent. A sequence $\{x_k\}$ is said to be bounded if there exists an $M > 0$ such that $|x_k| < M$ for all k . A sequence that is not bounded is said to be unbounded. The sequence (i) is unbounded; the sequences (ii), (iii), and (iv) are bounded.

The following theorem lists some basic properties of convergent sequences.

THEOREM 5.1. (i) The limit of a convergent sequence is unique.

(ii) Let $\lim_{k \rightarrow \infty} x_k = x_0$, let $\lim_{k \rightarrow \infty} y_k = y_0$, and let $\lim_{k \rightarrow \infty} r_k = r$, where $\{r_k\}$ is a sequence of scalars. Then $\lim_{k \rightarrow \infty} (x_k + y_k)$ exists and equals $x_0 + y_0$, and $\lim_{k \rightarrow \infty} r_k x_k$ exists and equals rx_0 .

(iii) A convergent sequence is bounded.

Exercise 5.1. Prove Theorem 5.1.

Let $\{x_n\}$ be a sequence in $\mathbb{R}^2$ and let $r_1 < r_2 < r_3 < \cdots < r_n < \cdots$ be a strictly increasing sequence of positive integers. Then the sequence $\{x_{r_n}\}_{n=1}^{\infty}$ is called a subsequence of $\{x_n\}$. Thus, $\{x_{2k}\}_{k=1}^{\infty} = \{(2k, 2k)\}$ is a subsequence of (i). The sequence $\{x_{n^2}\}_{n=1}^{\infty} = \{(n^2, n^2)\}$ is a subsequence of (ii). Each of the sequences $\{x_{2n}\} = \{(1-1)^2, (1-1)^2, (1-1)^2, \dots, (1-1)^2, \dots\} = \{(1-1)^2, (1-1)^2, \dots\}$ and $\{x_{2n+1}\} = \{(1-1)^2, (1-1)^2, (1-1)^2, \dots, (1-1)^2, \dots\}$ is a subsequence of (iv). The reader should plot the various subsequences.

We can now formulate a very useful criterion for a point to be a limit point of a set. (The reader should not confuse the notions of limit and limit point. They are different.)

Lemma 5.1. A point x is a limit point of a set S if and only if there exists a sequence $\{x_k\}$ of points in S such that, for each k , $x_k \neq x$ and $x_k \rightarrow x$.

Proof. Let x be a limit point of S . Then for every positive integer k there is a point x_k in S such that $x_k \neq x$ and $x_k \in B(x, 1/k)$. For every $\epsilon > 0$, there is a positive integer $K(\epsilon)$ satisfying $1/K(\epsilon) = \epsilon$. Hence for $k > K(\epsilon)$, we have $1/k < \epsilon$ and $B(x, 1/k) \subset B(x, \epsilon)$. Hence the sequence $\{x_k\}$ converges to x . Conversely, let there exist a sequence $\{x_k\}$ of points in S with $x_k \neq x$ and $x_k \rightarrow x$. Let $\epsilon > 0$ be arbitrary. Since $x_k \rightarrow x$, there exists a positive integer $K(\epsilon)$ such that, for $k > K(\epsilon)$, $x_k \in B(x, \epsilon)$. Since by hypothesis $x_k \neq x$ for all k , it follows that we have satisfied the requirements of the definition that x is a limit point of S .

Remark 5.1. Let $\{x_k\}$ be a sequence of points belonging to a set S and converging to a point x . Then x must belong to $\bar{S}$. If x is in S , there is nothing to prove. If x were not in S , then since $x_k \in S$ for each k , it follows that $x_k \neq x$. Lemma 5.1 then implies that x is a limit point of S . Thus $x \in \bar{S}$. If S is a closed set, then since $\bar{S} = S$, the point x belongs to S .

Remark 5.2. Let S be a set in $\mathbb{R}^n$. If x belongs to $\bar{S}$, then there is a sequence of points $\{x_k\}$ in S such that $x_k \rightarrow x$. To see this, note that if x is in S , then x belongs either to S or to S' . If $x \in S'$, the assertion follows from Lemma 5.1. If $x \notin S$ and $x \notin S'$, we may take $x_k = x$ for all positive integers k .

Let S be a subset (proper or improper) of $\mathbb{R}^n$. Let f be a function with domain S and range contained in $\mathbb{R}^m$. We denote this by $f: S \rightarrow \mathbb{R}^m$. Thus, $f = (f_1, \dots, f_m)$, where each f_i is a real-valued function.

Let a be a limit point of S . We say that "the limit of $f(x)$ as x approaches a exists and equals L " if for every $\epsilon > 0$ there exists a $\delta > 0$, which may depend on ϵ , such that for all $x \in B_\delta(a)$, $x \neq a$, we have $f(x) \in B_\epsilon(L)$. We write

$$\lim_{x \rightarrow a} f(x) = L.$$

It is not hard to show that, if $L = (L_1, \dots, L_m)$, $\lim_{x \rightarrow a} f(x) = L$ if and only if, for each $i = 1, \dots, m$, $\lim_{x \rightarrow a} f_i(x) = L_i$.

It is also not hard to show that if a limit exists it is unique and that the algebra of limits, as stated in the next theorem, holds.

Theorem 5.2. Let $\lim_{x \rightarrow a} f(x) = L$ and let $\lim_{x \rightarrow a} g(x) = M$. Then

$$\lim_{x \rightarrow a} [f(x) + g(x)]$$

exists and equals $L + M$. Also $\lim_{x \rightarrow a} [cf(x)]$ exists and equals cL .

A useful sequential criterion for the existence of a limit is now given.

Theorem 5.3. $\lim_{x \rightarrow a} f(x) = L$ if and only if for every sequence $\{x_k\}$ of points in S such that $x_k \neq a$ for all k and $x_k \rightarrow a$ it is true that $\lim_{k \rightarrow \infty} f(x_k) = L$.

Proof. Let $\lim_{x \rightarrow a} f(x) = L$ and let $\{x_k\}$ be an arbitrary sequence of points in S with $x_k \neq a$ for all k and with $x_k \rightarrow a$. Let $\epsilon > 0$ be given. Since $\lim_{x \rightarrow a} f(x) = L$, there exists a $\delta(\epsilon) > 0$ such that, for all $x \neq a$ and in $B_\delta(a) \cap S$, $f(x) \in B_\epsilon(L)$. Since $x_k \rightarrow a$, there exists a positive integer $K(\delta)$ such that whenever $k > K(\delta)$, $x_k \in B_\delta(a)$. Thus, for $k > K(\delta)$, $f(x_k) \in B_\epsilon(L)$, and so $f(x_k) \rightarrow L$.

We now prove the implication in the other direction. Suppose that $\lim_{x \rightarrow a} f(x)$ does not equal L . Then there exists an $\epsilon_0 > 0$ such that for every $\delta > 0$ there exists a point x_δ in S satisfying the conditions $x_\delta \neq a$, $x_\delta \in B_\delta(a)$

and $f(x_k) \notin B(\delta, v_k)$. Take δ through the sequence of values $\delta = 1/k$, where k runs through the positive integers. We then get a sequence $\{x_k\}$, with $x_k \in E$, $x_k \neq a$, and $x_k \rightarrow a$, with the further property that $f(x_k) \notin B(1/k, v_k)$. Thus $f(x_k)$ does not converge to L , and the theorem is proved.

Let x_0 be a point of S that is also a limit point of S ; that is, $x_0 \in S \cap E$. We say that f is *continuous at x_0* if

$$\lim_{x \rightarrow x_0} f(x) = f(x_0).$$

If $x_0 \in S$ and x_0 is not a limit point of S , we will take f to be continuous at x_0 by convention.

From the corresponding property of limits it follows that f is continuous at x_0 if and only if each of the real-valued functions $f_1, \dots, f_n$ is continuous at x_0 . It follows from Theorem 5.2 that if f and g are continuous at x_0 , then so are $f + g$ and af .

The function f is said to be *continuous on S* if it is continuous at every point of S .

Exercise 5.1. Show that if $\lim_{x \rightarrow a} f(x) = L$, then L is unique.

Exercise 5.2. Prove Theorem 5.2.

Exercise 5.3. Let f be a real-valued function defined on a set S in $\mathbb{R}^n$. Show that if f is continuous at a point x_0 in S and if $f(x_0) < 0$, then there exists a $\delta > 0$ such that $f(x) < 0$ for all x in $B(x_0, \delta) \cap S$.

Exercise 5.4. Show that the real-valued function $f: \mathbb{R}^n \rightarrow \mathbb{R}^1$ defined by $f(x) = \|x\|$ is continuous on $\mathbb{R}^n$ (i.e., show that the norm is a continuous function). Hint: Use the triangle inequality to show that, for any pair of vectors x and y , $||x| - |y|| \leq |x - y|$.

6. BASIC PROPERTY OF REAL NUMBERS

A set E in $\mathbb{R}^1$ is said to be *bounded* if there exists a positive number M such that $|x| \leq M$ for all $x \in E$. Another way of stating this is $E \subset B(0, M)$. The number M is said to be a *bound* for the set. Note that if a set is bounded, then there are infinitely many bounds.

We now shall deal with sets in $\mathbb{R}^1$ only. In $\mathbb{R}^1$ a set E being bounded means that, for every x in E , $|x| \leq M$, or $-M \leq x \leq M$. A set E in $\mathbb{R}^1$ is said to be *bounded above* if there exists a real number A such that $x \leq A$ for all x in E . A set E is *bounded below* if there exists a real number B such that $x \geq B$ for all x in E . The number A is said to be an *upper bound* of E , the number B a *lower bound*.

A number U is said to be a *least upper bound* (l.u.b.), or *supremum* (sup.), of a set S

- (i) if U is an upper bound of S and
- (ii) if U' is another upper bound of S and then $U' \geq U$.

If a set has a least upper bound, then the least upper bound is unique. To see this, let U_1 and U_2 be two least upper bounds of a set S . Then since U_2 is an upper bound and U_1 is a least upper bound, $U_2 \geq U_1$. Similarly, $U_1 \geq U_2$ and so $U_1 = U_2$.

Condition (ii) states that no number $A < U$ can be an upper bound of S . Therefore condition (ii) can be replaced by the equivalent statement:

- (ii') For every $\epsilon > 0$ there is a number x_ϵ in S with $U - \epsilon < x_\epsilon \leq U$.

Let S be the set of numbers $1 - 1/n$, $n = 1, 2, 3, \dots$. Then one is the l.u.b. of this set. Note that one does not belong to S . Now consider the set $S_0 = S \cup \{1\}$. The number 1 is the l.u.b. of this set and belongs to S_0 . The reader should be clear about the distinction between an upper bound and a supremum, or l.u.b.

A number L is said to be a *greatest lower bound* (g.l.b.), or *infimum* (inf.), of a set S

- (i) if L is a lower bound of S and
- (ii) if L' is another lower bound of S and then $L' \leq L$.

Observations analogous to those made following the definition of l.u.b. hold for the definition of g.l.b. We leave their formulation to the reader.

Exercise 6.1. Let $L = \text{g.l.b.}$ of a set S . Show that there exists a sequence of points $\{x_n\}$ in S such that $x_n \rightarrow L$. Show that the sequence $\{x_n\}$ can be taken to be nonincreasing, that is, $x_{n+1} \leq x_n$ for every n . Does L have to be a limit point of S ?

Exercise 6.2. Let S be a set in $\mathbb{R}^1$. We define $-S$ to be $\{-x : x \in S\}$. Thus $-S$ is the set that we obtain by replacing each element x in S by the element $-x$. Show that S is bounded below if and only if $-S$ is bounded above. Show that α is the g.l.b. of S if and only if $-\alpha$ is the l.u.b. of $-S$.

We can now state the basic property of the real numbers $\mathbb{R}^1$, which is sometimes called the *completeness property*.

Basic Property of the Reals

Every set $S \subset \mathbb{R}^1$ that is bounded $\left(\begin{smallmatrix} \text{above} \\ \text{below} \end{smallmatrix}\right)$ has a $\left(\begin{smallmatrix} \text{l.u.b.} \\ \text{g.l.b.} \end{smallmatrix}\right)$.

It follows from Exercise 8.2 that every set in $\mathbb{R}^1$ that is bounded above has a l.u.b. if and only if every set that is bounded below has a g.l.b. Thus, one statement of the completeness property can be interpreted as two equivalent statements.

Not every number system has this property. Consider the rational numbers $\mathbb{Q}$. They possess all of the algebraic and order properties of the reals but do not possess the completeness property. We will outline the argument showing that the rationals are not complete. For details see Rudin [1976]. Recall that a rational number is a number that can be expressed as a quotient of integers and recall that $\sqrt{2}$ is not rational. (The ancient Greeks knew that $\sqrt{2}$ is not rational.) Consider the set S of rational numbers x defined by $S = \{x: x \text{ rational, } x^2 < 2\}$. Clearly, 2 is an upper bound. It can be shown that for any x in S there is an x' in S such that $x' > x$. Thus no element of S can be an upper bound for S , let alone a l.u.b. It can also be shown that for any rational number y not in S that is an upper bound of S there is another rational number y' that is not in S , that is, also an upper bound of S and satisfies $y' < y$. Thus no rational number not in S can be a l.u.b. of S . Hence, there is no l.u.b. of S in the rationals. The number $\sqrt{2}$ in the reals is the candidate for the title of l.u.b., but $\sqrt{2}$ is not in the system of rational numbers.

Exercise 8.3. Let A and B be two bounded sets of real numbers with $A \subset B$. Show that

$$\begin{aligned}\sup\{a: a \in A\} &\leq \sup\{b: b \in B\}, \\ \inf\{a: a \in A\} &\geq \inf\{b: b \in B\}.\end{aligned}$$

Exercise 8.4. Let A and B be two sets of real numbers.

(i) Show that if A and B are bounded above, then

$$\sup\{a + b: a \in A, b \in B\} = \sup\{a: a \in A\} + \sup\{b: b \in B\}.$$

(ii) Show that if A and B are bounded below, then

$$\inf\{a + b: a \in A, b \in B\} = \inf\{a: a \in A\} + \inf\{b: b \in B\}.$$

Exercise 8.5. Let a and b be real numbers and let $J = \{x: a \leq x \leq b\}$. A real-valued function f is said to have a right-hand limit ρ at a point a_0 where $a \leq a_0 < b$, if for every $\epsilon > 0$ there exists a $\delta(\epsilon)$ such that whenever $0 < x - a_0$

$< \delta_0$) the inequality $|f(x) - \rho| < \epsilon$ holds. We shall write

$$f(x_0+) = \lim_{x \rightarrow x_0+} f(x) = \rho.$$

- (ii) Formulate a definition for a left-hand limit at a point x_0 where $a < x_0 < b$. If l is the left-hand limit, we shall write

$$f(x_0-) = \lim_{x \rightarrow x_0-} f(x) = l.$$

- (iii) Show that f has a limit at a point x_0 where $a < x_0 < b$, if and only if f has right- and left-hand limits at x_0 and these limits are equal.

Exercise 6.8. Let f be as in Exercise 6.5. A real-valued function f defined on I is said to be nondecreasing if for any pair of points x_1 and x_2 in I with $x_2 > x_1$ we have $f(x_2) \geq f(x_1)$.

- Formulate the definition of nonincreasing function.
- Show that if f is a nondecreasing function on I , then f has a right-hand limit at every point x_1 where $a < x_1 < b$, and left-hand limit at every point x_2 where $a < x_2 < b$.
- Show that if I is an open interval $[a, b) = \{x : a < x < b\}$ and if f is nondecreasing and bounded below on I , then $\lim_{x \rightarrow a+} f(x)$ exists. If f is bounded above, then $\lim_{x \rightarrow a+} f(x)$ exists.

7. COMPACTNESS

A property that is of great importance in many contexts is that of compactness. The definition that we shall give is one that is applicable in very general contexts as well as in $\mathbb{R}^n$. We do this so that the reader who takes more advanced courses in analysis will not have to unlearn anything. After presenting the definition of compactness, we will state, without proof, two theorems that give necessary and sufficient conditions for a set X in $\mathbb{R}^n$ to be compact. We frequently shall use the properties stated in these theorems when working with compact sets.

Let X be a set in $\mathbb{R}^n$. A collection of open sets $\{O_\alpha\}_{\alpha \in A}$ is said to be an open cover of X if every $x \in X$ is contained in some O_α . A set X in $\mathbb{R}^n$ is said to be compact if for every open cover $\{O_\alpha\}_{\alpha \in A}$ of X there is a finite collection of sets $O_1, \dots, O_k$ from the original collection $\{O_\alpha\}_{\alpha \in A}$ such that the finite collection $O_1, \dots, O_k$ is also an open cover of X . In mathematical jargon this property is stated as "Every open cover has a finite subcover."

To illustrate, consider the set X in $\mathbb{R}^1$ defined by $X = [0, 1) = \{x : 0 \leq x < 1\}$. For each positive integer k , let $O_k = [0, 1 - 1/k) = \{x : 0 \leq x < 1 - 1/k\}$. Then $\{O_k\}_{k=1}^\infty$ is an open cover of X , but no finite collection of sets in the cover will

cover S . Thus, $[0, 1)$ is not compact. It is difficult to give a meaningful example of a compact set based on the definitions, since we must show that every open cover has a finite subcover.

THEOREM 7.1. In $\mathbb{R}^n$ a set S is compact if and only if every sequence $\{x_k\}$ of points in S has a subsequence $\{x_{k_j}\}$ that converges to a point in S (Bolzano–Weierstrass property).

THEOREM 7.2. In $\mathbb{R}^n$ a set S is compact if and only if S is closed and bounded.

We refer the reader to Bartle and Sherbert [1999] or Rudin [1976] for proofs of these theorems.

We emphasize that the criteria for compactness given in Theorems 7.1 and 7.2 are not necessarily valid in contexts more general than $\mathbb{R}^n$.

Theorem 7.2 provides an easily applied criterion for compactness. Using Theorem 7.2, we see immediately that $S = [0, 1) = \{x : 0 \leq x < 1\}$ is not compact in $\mathbb{R}^1$. (S is not closed.) We also see that $S_1 = [0, 1] = \{x : 0 \leq x \leq 1\}$ is compact, since it is closed and bounded.

We conclude with another characterization of compactness that is valid in more general contexts than $\mathbb{R}^n$. A set S is said to have the *finite-intersection property* if for every collection of closed sets $\{F_\alpha\}$ such that $\bigcap F_\alpha \cap S \neq \emptyset$ there is a finite subcollection $F_1, \dots, F_n$ such that $\bigcap_{i=1}^n F_i \cap S \neq \emptyset$. We shall illustrate the definition by exhibiting a set that fails to have the finite-intersection property. The set $(0, 1) = \{x : 0 < x < 1\}$ in $\mathbb{R}^1$ does not have the finite-intersection property. Take $F_n = [-1/(n+1), 1/n]$, $n = 1, 2, 3, 4, \dots$. Then $\bigcap_{n=1}^\infty F_n = [1, \infty) \cap (0, 1)$, yet the intersection of any finite subcollection of the F_n has nonempty intersection with $(0, 1)$ and so is not contained in $\emptyset \cap (0, 1)$.

THEOREM 7.3. A set S is compact if and only if it has the finite-intersection property.

The proof is an exercise in the use of Lemma 3.1. Let S have the finite-intersection property. Let $(\mathcal{O}_\beta)$ be an open cover of S . Then $S \subset \bigcup \mathcal{O}_\beta = \complement \bigcap (\complement \mathcal{O}_\beta)$, and so $\complement S \supset \complement (\bigcap \complement \mathcal{O}_\beta) = \bigcap (\complement \mathcal{O}_\beta)$. Each set $\complement \mathcal{O}_\beta$ is closed, so there exist a finite number of sets $\mathcal{O}_{\beta_1}, \dots, \mathcal{O}_{\beta_n}$ such that $\complement S = \bigcap_{i=1}^n \complement \mathcal{O}_{\beta_i}$. Hence $S \subset \complement (\bigcap_{i=1}^n \complement \mathcal{O}_{\beta_i}) = \bigcup_{i=1}^n \mathcal{O}_{\beta_i}$, and so $\{(\mathcal{O}_\beta)\}$ has a finite subcover $\mathcal{O}_{\beta_1}, \dots, \mathcal{O}_{\beta_n}$. Thus we have shown that S is compact. Now let S be compact. Let $\{F_\alpha\}$ be a collection of closed sets such that $\bigcap F_\alpha \cap S \neq \emptyset$. Then $\complement (\bigcap F_\alpha) \cap \complement S = \complement S$, and so $\bigcup (\complement F_\alpha) \supset \complement S$. Thus $\{(\complement F_\alpha)\}$ is an open cover of S . Since S is compact, there exist a finite subcollection of sets $F_1, \dots, F_n$ such that $\bigcup_{i=1}^n \complement F_i \supset \complement S$. Therefore $\bigcap_{i=1}^n F_i \cap S = \complement (\bigcup_{i=1}^n \complement F_i) \cap S \subset \complement S$, and thus S has the finite-intersection property.

Remark 7.1. Note that the finite-intersection property can be stated in the following equivalent form. Every collection of closed sets $\{F_\alpha\}$ such that $\bigcap_{i=1}^\infty F_i \cap S = \emptyset$ has a finite subcollection $F_1, \dots, F_n$ such that $\bigcap_{i=1}^n F_i \cap S = \emptyset$.

8. EQUIVALENT NORMS AND CARTESIAN PRODUCTS

A norm on a vector space $\mathbb{R}^n$ is a function $\|\cdot\|$ that assigns a real number $\|\mathbf{x}\|$ to every vector $\mathbf{x}$ in the space and has the properties (1)–(3) and (5) listed in Section 2, with (2) replaced by $\|\cdot\|$. Once a norm has been defined, one can define a function ρ by the formula $\rho(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|$. Properties (1)–(3) and (5) of the norm imply that the function ρ satisfies the properties (1)–(3) of a metric, and therefore is a metric, or distance function. Thus a distance has been defined, open sets are defined in terms of the distance, as before. Convergence is defined in terms of open sets.

To illustrate these ideas, we give an example of another norm in $\mathbb{R}^n$, defined by

$$\|\mathbf{x}\|_\infty = \|\mathbf{x}\|_\infty = \max\{|x_1|, \dots, |x_n|\}.$$

The corresponding distance is defined by

$$\rho(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|_\infty = \max\{|x_1 - y_1|, \dots, |x_n - y_n|\}.$$

Since a norm determines the family of open sets, and since the family of open sets determines convergence, the question arises, “Given two norms, ν_1 and ν_2 , do they determine the same family of open sets?” The answer, which is not difficult to show, is, “If there exist positive constants m and M such that for all $\mathbf{x}$ in $\mathbb{R}^n$,

$$m\nu_2(\mathbf{x}) \leq \nu_1(\mathbf{x}) \leq M\nu_2(\mathbf{x}), \quad (1)$$

then a set is open under the ν_1 norm if and only if it is open under the norm ν_2 .” Two norms ν_1 and ν_2 that satisfy (1) are said to be *equivalent norms*.

An extremely important fact is that in $\mathbb{R}^n$ all norms are equivalent. For a proof see Fleming [1977]. In optimization theory, and in other areas, it may happen that the analysis of a problem or the development of a computational algorithm is more convenient in one norm than in another. The equivalence of the norms guarantees that the convergence of an algorithm is intrinsic, that is, is independent of the norm.

Let $A_1, \dots, A_m$ be sets with $A_i \subseteq \mathbb{R}^{n_i}$, $i = 1, \dots, m$. By the *cartesian product* of these sets, denoted by

$$\prod_{i=1}^m A_i = \{\mathbf{x} \mid \mathbf{x}_1 \in A_1 \text{ and } \dots \text{ and } \mathbf{x}_m \in A_m\}$$

we mean the set A in $\mathbb{R}^{n_1 + \dots + n_m}$ consisting of all possible points in $\mathbb{R}^{n_1 + \dots + n_m}$ obtained by taking the first n_1 components from a point in A_1 , the next n_2 components from a point in A_2 , the next n_3 components from a point in A_3 ,

and so forth, until the last n_m components, which are obtained by taking a point in A_m . For example, $\mathbb{R}^n$ can be considered to be the cartesian product of $\mathbb{R}^1$ taken n times with itself:

$$\mathbb{R}^n = \underbrace{\mathbb{R}^1 \times \mathbb{R}^1 \times \cdots \times \mathbb{R}^1}_{n \text{ times}}.$$

If A_1 is the set in $\mathbb{R}^2$ consisting of points on the circumference of the circle with center at the origin and radius 1 and if $A_2 = [0, 1]$ is in $\mathbb{R}^1$, then $A_1 \times A_2$ is the set in $\mathbb{R}^3$ that is the lateral surface of a cylinder of height 1 whose base is the closed disk in the (x_1, x_2) plane with center at the origin and radius equal to 1.

Let $X = \mathbb{R}^{n_1} \times \cdots \times \mathbb{R}^{n_m}$. Let x be a vector in X . Then $x = (x_1, \dots, x_m)$ where $x_i = (x_{i1}, \dots, x_{in_i})$, $i = 1, \dots, m$. If for vectors x and y in X and scalars λ in $\mathbb{R}^1$ we define

$$(x + y) = (x_1 + y_1, \dots, x_m + y_m) \quad \lambda x = (\lambda x_1, \dots, \lambda x_m)$$

and let $n = n_1 + n_2 + \cdots + n_m$, then X can be identified with the usual vector space $\mathbb{R}^n$. The euclidean norm in X would then be given by

$$\|x\| = \left[\|x_1\|^2 + \cdots + \|x_m\|^2 \right]^{1/2} = \left[\sum_{i=1}^m \sum_{j=1}^{n_i} (x_{ij})^2 \right]^{1/2}. \quad (2)$$

On the other hand, in keeping with the definitions used in more general contexts, the norm on X is defined by

$$\|x\|_1 = \sum_{i=1}^m \|x_i\|, \quad (3)$$

where $\|x_i\|$ is the euclidean norm of x_i in $\mathbb{R}^{n_i}$ defined by $\|x_i\| = (x_{i1}x_{i1} + \cdots + x_{in_i}x_{in_i})^{1/2} = [\sum_{j=1}^{n_i} (x_{ij})^2]^{1/2}$. Since all norms on $\mathbb{R}^n$ are equivalent, the norm defined by (3) is equivalent to the euclidean norm defined by (2).

The following lemma is an immediate consequence of the definition of $\|\cdot\|_1$.

LEMMA 3.1. A sequence $x_n = (x_{n1}, \dots, x_{nm})$ of points in X converges to a point $x_0 = (x_{01}, \dots, x_{0m})$ of X in the metric determined by (3) if and only if for each $j = 1, \dots, m$ the sequence $\{x_{nj}\}$ converges to x_{0j} in the euclidean norm on $\mathbb{R}^{n_j}$.

The next theorem can be proved by using Lemma 3.1 and Theorem 7.1 and by selecting appropriate subsequences.

THEOREM 3.1. Let $A_1, \dots, A_m$ be metrics with $A_1 = \mathbb{R}^n$ and for each A_j be compact, $j = 1, \dots, m$. Then $\bigcap_{n=1}^{\infty} A_n$ is compact in $\mathbb{R}^n$ where $n = n_1 + n_2 + \cdots + n_m$.

Exercise 8.1. (a) In $\mathbb{R}^2$ sketch the set $B_{\infty}(\mathbf{0}, 1) = \{\mathbf{x} : |\mathbf{x}|_{\infty} < 1\}$.

(b) Show that $|\cdot|_{\infty}$ is a norm (i.e., that it satisfies (1)–(3) and (5) of Section 2).

(c) Can you draw a sketch in $\mathbb{R}^2$ that would indicate that a set is open under the metric $d(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|$ if and only if it is open under the metric $\rho(\mathbf{x}, \mathbf{y}) = \|\mathbf{x} - \mathbf{y}\|_{\infty}$?

(d) Is it true that $\|\mathbf{x} + \mathbf{y}\|_{\infty} = \|\mathbf{x}\|_{\infty} + \|\mathbf{y}\|_{\infty}$ if and only if $\mathbf{y} = c\mathbf{x}$, $c > 0$?

Exercise 8.2. (a) Show that the function $|\cdot|_1$, defined by (3) is a norm on $\mathbb{R}$; that is, $|\cdot|_1$ satisfies (1)–(3) and (5) of Section 2.

(b) If $\mathbb{X} = \mathbb{R}^2 = \mathbb{R}^2 \times \mathbb{R}^2$, sketch the set $B_1(\mathbf{0}, 1) = \{\mathbf{x} : |\mathbf{x}|_1 < 1\}$.

Exercise 8.3. Without using the fact that all norms in $\mathbb{R}^n$ are equivalent, show that $|\cdot|$ and $|\cdot|_1$, defined by (2) and (3), respectively, are equivalent norms on $\mathbb{R}$.

Exercise 8.4. Show that the following functions are continuous:

(a) $f: \mathbb{R}^2 \rightarrow \mathbb{R}^2$, where $f(\mathbf{x}, \mathbf{y}) = (\mathbf{x} + \mathbf{y})$ (addition is continuous);

(b) $f: \mathbb{R}^2 \times \mathbb{R}^2 \rightarrow \mathbb{R}^2$, where $f(\lambda, \mathbf{x}) = \lambda\mathbf{x}$ (scalar multiplication is continuous); and

(c) $f: \mathbb{R}^2 \times \mathbb{R}^2 \rightarrow \mathbb{R}$, where $f(\mathbf{x}, \mathbf{y}) = \langle \mathbf{x}, \mathbf{y} \rangle$ (inner product is continuous).

Exercise 8.5. Prove Theorem 8.1.

Exercise 8.6. Show that if $\mathbf{Q}$ is a positive-definite real symmetric matrix, then $\phi(\mathbf{x}) = (\mathbf{x}, \mathbf{Q}\mathbf{x})^{1/2}$ defines a norm on $\mathbb{R}^n$. *Hint:* Look at Section 1.

9. FUNDAMENTAL EXISTENCE THEOREM

The basic problem in optimization theory is as follows: Given a set S (not necessarily in $\mathbb{R}^n$) and a real-valued function f defined on S , does there exist an element x_0 in S such that $f(x_0) \leq f(x)$ for all x in S , and if so find it. The problem as stated is a minimization problem. The point x_0 is said to be a *minimizer* for f , or to *minimize* f . We also say that f attains a *minimum* on S at x_0 . The problems of “maximizing f over S ” or finding an x^0 in S such that $f(x^0) \geq f(x)$ for all x , can be reduced to the problem of minimizing the function $-f$ over S . To see this, observe that if $x^0 = f(x^0) \leq -f(x)$ for all x in S , then $f(x^0) \geq f(x)$ for all x in S , and conversely.

The first question that we address is, “Does a minimizer exist?” To point out some of the difficulties here, we look at some very simple examples in $\mathbb{R}^1$.

Let $S = \{x \mid x \geq 1\}$ and let $f(x) = 1/x$. It is clear that $\inf\{f(x) \mid x \geq 1\} = 0$, yet there is no $x_0 \in S$ such that $f(x_0) = 0$. Thus a minimizer does not exist, or f does not attain its minimum on S . As a second example, let $S = [1, 2]$ and let $f(x) = 1/x$. Now, $\inf\{f(x) \mid 1 \leq x \leq 2\} = \frac{1}{2}$, yet for no x_0 in S do we have $f(x_0) = \frac{1}{2}$. Of course, if we modify the last example by taking $S = [1, 2] = \{x \mid 1 \leq x \leq 2\}$, then $x_0 = 2$ is the minimizer. For future reference, note that if $S = (1, 2] = \{x \mid 1 < x \leq 2\}$, then $x_0 = 2$ is still the minimizer.

We now state and prove the fundamental existence theorem for minimization of functions with domains in $\mathbb{R}^n$.

THEOREM 8.1. *Let S be a compact set in $\mathbb{R}^n$ and let f be a real-valued continuous function defined on S . The f attains a maximum and a minimum on S .*

Before taking up the proof, we point out that the theorem gives a sufficient condition for the existence of a maximizer and minimizer. It is not a necessary condition, as the example in $\mathbb{R}^1$, $S = [1, 2]$, $f(x) = 1/x$, shows. In most advanced courses the reader may encounter the notions of upper semicontinuity and lower semicontinuity and learn that continuity can be replaced by upper semicontinuity to ensure the existence of maximizers and that continuity can be replaced by lower semicontinuity to ensure the existence of minimizers. Finally, the reader should study the proof of this theorem very carefully and be sure to understand it. The proof of the theorem in $\mathbb{R}^n$ is the model for the proofs of the existence of minimizers in more general contexts than $\mathbb{R}^n$.

Proof. If f is continuous on S , then so is $-f$. If $-f$ attains a minimum on S , then f attains a maximum on S . Thus we need only prove that f attains a minimum.

We first prove that the set of values that f assumes on S is bounded below, that is, that the set $\mathcal{B} = \{f(x) = f(x_0) \mid x_0 \in S\}$ is bounded below. To show this, we suppose that $\mathcal{B}$ were not bounded below. Then for each positive integer k there would exist a point x_k in S such that $f(x_k) < -k$. Since S is compact, the sequence $\{x_k\}$ of points in S has a subsequence $\{x_{k_j}\}$ that converges to a point x_0 in S . The continuity of f implies that $\lim_{j \rightarrow \infty} f(x_{k_j}) = f(x_0)$. Thus, by the definition of limit (taking $\epsilon = 1$) we have that there exists a positive integer K such that, for $j > K$, $f(x_{k_j}) > f(x_0) - 1$. This contradicts the defining property of $\{x_k\}$, namely that, for each k , $f(x_k) < -k$.

Since the set $\mathcal{B}$ is bounded below, it has an infimum, say L . Therefore, by definition of infimum, for each positive integer k there exists a point x_k in S such that

$$L \leq f(x_k) < L + \frac{1}{k}. \quad (1)$$

Since S is compact, the sequence $\{x_k\}$ has a subsequence $\{x_{k_j}\}$ that converges

to a point x_0 in S . The continuity of f implies that $\lim_{j \rightarrow \infty} f(x_{k_j}) = f(x_0)$. Therefore, if we replace δ by δ_j in (1) and let $j \rightarrow \infty$, we get that $L \leq f(x_0) \leq L$. Hence

$$f(x_0) = L \leq f(x) \quad \text{for all } x \in S.$$

This proves the theorem.

If the reader understands the proof of Theorem 9.1, and understands the use of compactness, he or she should have no trouble in proving the following theorem. Actually, Theorem 9.1 is a corollary of Theorem 9.2.

Theorem 9.2. Let S be a compact subset of $\mathbb{R}^n$ and let $f = (f_1, \dots, f_m)$ be a continuous function defined on S with range in $\mathbb{R}^m$. Then the set

$$f(S) = \{y \in \mathbb{R}^m : y = f(x), x \in S\}$$

is compact.

In mathematical jargon this result is stated as follows: The continuous image of a compact set is compact.

Exercise 9.1. Prove Theorem 9.2.

Exercise 9.2. Let A be a real $n \times n$ symmetric matrix with entries a_{ij} and let $Q(x)$ be the quadratic form

$$Q(x) = x^T A x = \langle x, Ax \rangle = \sum_{i,j=1}^n a_{ij} x_i x_j.$$

Show that there exists a constant $M > 0$ such that, for all x , $Q(x) \leq M \|x\|^2$. Show that if $Q(x) > 0$ for all $x \neq 0$, then there exists a constant $m > 0$ such that, for all x , $Q(x) \geq m \|x\|^2$. *Hint:* First show that the conclusion is true for $n(SO, 1) = \{x \in \mathbb{R}^n : \|x\| = 1\}$.

Exercise 9.3. Let f be a real-valued function defined on $\mathbb{R}^n$ such that

$$\lim_{\|x\| \rightarrow \infty} f(x) = +\infty.$$

By this is meant that for each $M > 0$ there exists an $R > 0$ such that if $\|x\| \geq R$, then $f(x) > M$. Such a function f is said to be *coercive*. Show that if f is continuous on $\mathbb{R}^n$ and is coercive, then f attains a minimum on $\mathbb{R}^n$.

18. LINEAR TRANSFORMATIONS

A linear transformation from $\mathbb{R}^n$ to $\mathbb{R}^m$ is a function T with domain $\mathbb{R}^n$ and range contained in $\mathbb{R}^m$ such that for every x, y in $\mathbb{R}^n$ and every pair of scalars α, β

$$T(\alpha x + \beta y) = \alpha T(x) + \beta T(y).$$

Example 18.1. $T: \mathbb{R}^2 \rightarrow \mathbb{R}^2$: Rotate each vector in $\mathbb{R}^2$ through an angle θ .

Example 18.2. $T: \mathbb{R}^3 \rightarrow \mathbb{R}^3$ defined by

$$T(x_1, x_2, x_3) = (x_1, x_2, 0).$$

This is the projection of $\mathbb{R}^3$ onto the x_1x_2 plane.

Example 18.3. Let A be an $m \times n$ matrix. Define

$$Tx = Ax. \quad (1)$$

In the event that the range space is $\mathbb{R}^1$, that is, we have a linear transformation from $\mathbb{R}^n$ to $\mathbb{R}^1$, we call the transformation a linear functional. The following is an example of a linear functional L . Let u be a fixed vector in $\mathbb{R}^n$. Then

$$L(x) = \langle u, x \rangle \quad (2)$$

defines a linear functional.

We now show that (1) and (2) are essentially the only examples of a linear transformation and linear functional, respectively. Recall that the standard basis in $\mathbb{R}^n$ is the basis consisting of the vectors $e_1, \dots, e_n$, where $e_i = (0, 0, \dots, 0, 1, 0, \dots, 0)$ and 1 is the i th entry.

THEOREM 18.1. Let T be a linear transformation from $\mathbb{R}^n$ to $\mathbb{R}^m$. Then relative to the standard basis in $\mathbb{R}^n$ and $\mathbb{R}^m$ there exists an $m \times n$ matrix A such that $Tx = Ax$.

Let $x \in \mathbb{R}^n$; then $x = \sum_{j=1}^n x_j e_j$. Hence

$$Tx = T\left(\sum_{j=1}^n x_j e_j\right) = \sum_{j=1}^n x_j T(e_j). \quad (3)$$

For each $j = 1, \dots, n$, $T e_j$ is a vector in $\mathbb{R}^m$, so if $a_1^j, \dots, a_m^j$ is the standard basis in $\mathbb{R}^m$,

$$T e_j = \sum_{i=1}^m a_{ij} a_i^j. \quad (4)$$

Substituting (4) into (3) and changing the order of summation give

$$Tx = \sum_{j=1}^n x_j \left(\sum_{i=1}^n a_{ij} e_i^T \right) = \sum_{i=1}^n \left(\sum_{j=1}^n a_{ij} x_j \right) e_i^T.$$

Therefore if A is the $n \times n$ matrix whose j th column is $(a_{1j}, \dots, a_{nj})^T$, then Tx expressed as an n -vector $(x_1, \dots, x_n)^T$ is given by

$$\begin{pmatrix} x_1 \\ \vdots \\ x_n \end{pmatrix} = Tx = Ax.$$

From (4) we get that the j th column of A gives the coordinates of Te_j relative to $e_1^T, \dots, e_n^T$.

In the case of a linear functional L , the matrix A reduces to a row vector $\mathbf{a}$ in $\mathbb{R}_n$ so that

$$L(\mathbf{x}) = (\mathbf{a}, \mathbf{x}).$$

Exercise 10.1. Show that a linear functional L is a continuous mapping from $\mathbb{R}^n$ to $\mathbb{R}$.

11. DIFFERENTIATION IN $\mathbb{R}^n$

Let D be an open interval in $\mathbb{R}^1$ and let f be real-valued function defined on D . The function f is said to have a derivative or to be differentiable at a point x_0 in D if

$$\lim_{h \rightarrow 0} \frac{f(x_0 + h) - f(x_0)}{h} \quad (1)$$

exists. The limit is called the derivative of f at x_0 and is denoted by $f'(x_0)$.

The question now arises as to how one generalizes this concept to real-valued functions defined on open sets D in $\mathbb{R}^n$ and to functions defined on open sets in $\mathbb{R}^n$ with range in $\mathbb{R}^m$. In elementary calculus the notion of partial derivative for real-valued functions defined on open sets D in $\mathbb{R}^n$ is introduced. For each $i = 1, \dots, n$ the partial derivative with respect to x_i at a point $\mathbf{x}_0$ is defined by

$$\frac{\partial f}{\partial x_i}(\mathbf{x}_0) = \lim_{h \rightarrow 0} \frac{f(x_{01}, \dots, x_{0,i-1}, x_{0i} + h, x_{0,i+1}, \dots, x_{0n}) - f(x_{01}, \dots, x_{0n})}{h} \quad (2)$$

provided the limit on the right exists. It turns out that the notion of partial derivative is not the correct generalization of the notion of derivative.

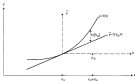


Figure 1.1.

To motivate the natural generalization of the notion of derivative to functions with domain and range in spaces of dimension greater than 1, we examine the notion of derivative for functions f defined on an open interval D in $\mathbb{R}^1$ with range in $\mathbb{R}^1$.

We rewrite (1) as

$$f(x_0 + h) = \frac{f(x_0 + h) - f(x_0)}{h} + q(h),$$

where $q(h) \rightarrow 0$ as $h \rightarrow 0$. This in turn can be rewritten as

$$f(x_0 + h) - f(x_0) = f'(x_0)h + q(h), \quad (2)$$

where $q(h)/h \rightarrow 0$ as $h \rightarrow 0$. The term $f'(x_0)h$ is a linear approximation to f near x_0 and h is a "good" approximation for small h in the sense that $q(h)/h \rightarrow 0$ as $h \rightarrow 0$. See Figure 1.1.

If we consider h to be an arbitrary real number, then $f'(x_0)h$ defines a linear functional L on $\mathbb{R}^1$ by the formula $L(h) = f'(x_0)h$. Thus if f is differentiable at x_0 , there exists a linear functional (or linear transformation) on $\mathbb{R}^1$ such that

$$f(x_0 + h) - f(x_0) = L(h) + q(h) \quad (3)$$

where $q(h)/h \rightarrow 0$ as $h \rightarrow 0$. Conversely, let there exist a linear functional L on $\mathbb{R}^1$ such that (3) holds. Then $L(h) = ah$ for some real number a , and we may write

$$f(x_0 + h) - f(x_0) = ah + q(h).$$

If we divide by $h \neq 0$ and then let $h \rightarrow 0$, we get that $f'(x_0)$ exists and equals α . Thus we could have used (4) to define the notion of derivative and could have defined the derivative to be the linear functional L , which in this case is determined by the number α . The reader might feel that this is a very convoluted way of defining the notion of derivative. It has the merit, however, of being the definition that generalizes correctly to functions with domain in $\mathbb{R}^n$ and range in $\mathbb{R}^m$ for any $n \geq 1$ and any $m \geq 1$.

Let D be an open set in $\mathbb{R}^n$, $n \geq 1$, and let $F = (f_1, \dots, f_m)$ be a function defined on D with range in $\mathbb{R}^m$, $m \geq 1$. The function F is said to be differentiable at a point x_0 in D , or to be Fréchet differentiable at x_0 , if there exists a linear transformation $T(x_0)$ from $\mathbb{R}^n$ to $\mathbb{R}^m$ such that for all h in $\mathbb{R}^n$ with h sufficiently small

$$F(x_0 + h) - F(x_0) = T(x_0)h + o(h), \quad (5)$$

where $\|o(h)\|/\|h\| \rightarrow 0$ as $\|h\| \rightarrow 0$.

We now represent $T(x_0)$ as a matrix relative to the standard bases $e_1, \dots, e_n$ in $\mathbb{R}^n$ and $e_1^*, \dots, e_m^*$ in $\mathbb{R}^m$. Let α be a real number and let $h = \alpha e_j$ for a fixed j in the set $1, \dots, n$. Then $\alpha e_j = (\alpha, \dots, \alpha, 0, \dots, 0)$, where the α occurs in the j th component. From (5) we have

$$F(x_0 + \alpha e_j) - F(x_0) = T(x_0)(\alpha e_j) + o(\alpha e_j).$$

Hence for $\alpha \neq 0$

$$\frac{F(x_0 + \alpha e_j) - F(x_0)}{\alpha} = T(x_0)e_j + \frac{o(\alpha e_j)}{\alpha}.$$

Since $\|e_j\| = 1$ and for $j = 1, \dots, n$,

$$\begin{aligned} f_j(x_0 + \alpha e_j) - f_j(x_0) &= f_j(x_{01}, \dots, x_{0j-1}, x_{0j} + \alpha, x_{0j+1}, \dots, x_{0n}) \\ &\quad - f_j(x_{01}, \dots, x_{0n}). \end{aligned}$$

It follows on letting $\alpha \rightarrow 0$ that for $j = 1, \dots, n$ the partial derivatives $f_j'(x_0)/\partial x_j$ exist and

$$T(x_0)e_j = \begin{pmatrix} \frac{\partial f_1}{\partial x_j}(x_0) \\ \vdots \\ \frac{\partial f_m}{\partial x_j}(x_0) \end{pmatrix}.$$

Since the coordinates of $T(\mathbf{x}_0)\mathbf{e}_j$ relative to the standard basis in $\mathbb{R}^n$ are given by the j th column of the matrix representing $T(\mathbf{x}_0)$, the matrix representing $T(\mathbf{x}_0)$ is the matrix

$$\left(\frac{\partial \xi_i}{\partial x_j}(\mathbf{x}_0) \right).$$

If we define $\nabla \xi_i$ to be the row vector

$$\nabla \xi_i = \left(\frac{\partial \xi_i}{\partial x_1}, \dots, \frac{\partial \xi_i}{\partial x_n} \right), \quad i = 1, \dots, n,$$

and $\nabla \mathbf{f}$ to be the matrix

$$\begin{pmatrix} \nabla \xi_1 \\ \vdots \\ \nabla \xi_n \end{pmatrix},$$

then we have shown that

$$T(\mathbf{x}_0) = \nabla \mathbf{f}(\mathbf{x}_0).$$

The matrix $\nabla \mathbf{f}(\mathbf{x}_0)$ is called the *Jacobian matrix* of $\mathbf{f}$ at $\mathbf{x}_0$. It is sometimes also written as $J(\mathbf{f})(\mathbf{x}_0)$. Note that we have shown that if $\mathbf{f}$ is differentiable at $\mathbf{x}_0$, then each component function ξ_i has all first partial derivatives existing at $\mathbf{x}_0$.

THEOREM 11.1. *Let B be an open set in $\mathbb{R}^n$ and let $\mathbf{f}$ be a function defined on B with range in $\mathbb{R}^n$. If $\mathbf{f}$ is differentiable at a point $\mathbf{x}_0$ in B , then $\mathbf{f}$ is continuous at $\mathbf{x}_0$.*

This follows immediately from (5) on letting $\mathbf{h} \rightarrow \mathbf{0}$.

We will now show that the existence of partial derivatives at a point does not imply differentiability. We shall consider real-valued functions f defined on $\mathbb{R}^2$. For typographical convenience we shall represent points in the plane as (x, y) instead of (x_1, x_2) .

Example 11.1. Let

$$f(x, y) = \begin{cases} \frac{xy^2}{x^2 + y^2}, & (x, y) \neq (0, 0), \\ 0, & (x, y) = (0, 0). \end{cases}$$

Since $f(x, 0) = 0$ for all x , $\partial f / \partial y$ exists at $(0, 0)$ and equals 0. Similarly $\partial f / \partial x(0, 0) = 0$. Along the line $y = kx$, the function f becomes

$$f(x, kx) = \frac{kx^3}{x^2 + k^2x^2} = \frac{k}{1 + k^2}, \quad x \neq 0,$$

and so $\lim_{(x,y) \rightarrow (0,0)} f(x,y) = 0/(1+1^2)$. Thus $\lim_{(x,y) \rightarrow (0,0)} f(x,y)$ does not exist. Therefore f is not continuous at the origin and so by Theorem 11.3 cannot be differentiable at the origin.

One might now ask, what if the function is continuous at a point and has partial derivatives at the point? Must the function be differentiable at the point? The answer is no, as the next example shows.

Example 11.3. Let

$$f(x,y) = \begin{cases} \frac{xy}{\sqrt{x^2+y^2}} & (x,y) \neq (0,0), \\ 0 & (x,y) = 0. \end{cases}$$

As in Example 11.1, $\partial f/\partial x$ and $\partial f/\partial y$ exist at the origin and are both equal to zero. Note that

$$0 \leq (x) = (x)^2 = (x)^2 + 2x(0) + (0)^2$$

implies that $x^2 + y^2 \geq 2x(0)$. Hence for $(x,y) \neq 0$

$$|f(x,y)| = \frac{|xy|}{\sqrt{x^2+y^2}} \leq \frac{1}{2} \frac{(x^2+y^2)}{\sqrt{x^2+y^2}} = \frac{1}{2} \sqrt{x^2+y^2}.$$

Therefore $\lim_{(x,y) \rightarrow (0,0)} f(x,y) = 0$, and so f is continuous at $(0,0)$.

If f were differentiable at $(0,0)$, then for any $\mathbf{h} = (h_1, h_2)$

$$f(\mathbf{0} + \mathbf{h}) = f(\mathbf{0}) + (Df)(\mathbf{0}) \cdot \mathbf{h} + \eta(\mathbf{h}),$$

where $\eta(\mathbf{h})/|\mathbf{h}| \rightarrow 0$ as $\mathbf{h} \rightarrow 0$. Since $(Df)(\mathbf{0}) = (0\partial f/\partial x)(\mathbf{0}), (0\partial f/\partial y)(\mathbf{0}) = (0,0)$ and since $f(\mathbf{0}) = 0$, we would have

$$f(\mathbf{h}) = \eta(\mathbf{h}), \quad \mathbf{h} \neq 0.$$

From the definition of f we get, writing (h_1, h_2) in place of (x,y) ,

$$\eta(\mathbf{h}) = \frac{h_1 h_2}{\sqrt{(h_1)^2 + (h_2)^2}}.$$

We require that $\eta(\mathbf{h})/|\mathbf{h}| \rightarrow 0$ as $\mathbf{h} \rightarrow 0$. But this is not true since

$$\frac{\eta(\mathbf{h})}{|\mathbf{h}|} = \frac{h_1 h_2}{(h_1)^2 + (h_2)^2},$$

and by Example 11.1, the right-hand side does not have a limit as $\mathbf{h} \rightarrow 0$.

A real-valued function f defined on an open set D in $\mathbb{R}^n$ is said to be of class C^k on D if all of the partial derivatives up to and including those of order k exist and are continuous on D . A function $\mathbf{F}$ defined on D with range in $\mathbb{R}^p$, $m > 1$, is said to be of class C^k on D if each of its component functions f_i , $i = 1, \dots, m$, is of class C^k on D .

We now show that if ϕ is a real-valued function defined on an open interval $D = \{x \mid a < x < b\}$ in $\mathbb{R}$, then $\phi(x_0) + h - \phi(x_0)$ is a "good" approximation to $\phi(x_0 + h) - \phi(x_0)$ for small h at any point x_0 in D at which ϕ is differentiable. If ϕ is of class C^2 in D , then a better approximation can be obtained using Taylor's theorem from elementary calculus, one form of which states that for x_0 in D and h such that $x_0 + h$ is in D

$$\phi(x_0 + h) - \phi(x_0) = \sum_{j=1}^{k-1} \frac{\phi^{(j)}(x_0)}{j!} h^j + \frac{\phi^{(k)}(\xi) h^k}{k!}, \quad (5)$$

where $\phi^{(j)}$ denotes the j th derivative of ϕ and ξ is a point lying between x_0 and $x_0 + h$. If $k = 1$, then the summation term in (5) is absent and we have a statement of the mean-value theorem. Since $\phi^{(k)}$ is continuous,

$$\phi^{(k)}(\xi) = \phi^{(k)}(x_0) + o(h),$$

where $o(h) \rightarrow 0$ as $h \rightarrow 0$. Substituting this relation into (5) gives

$$\phi(x_0 + h) - \phi(x_0) = \sum_{j=1}^k \frac{\phi^{(j)}(x_0) h^j}{j!} + \eta(h), \quad (7)$$

where $\eta(h) = o(h)h^k/k!$. Thus $\eta(h)h^k \rightarrow 0$ as $h \rightarrow 0$. For small h , the polynomial in (7) is thus a "good" approximation to $\phi(x_0 + h) - \phi(x_0)$ in the sense that the error committed in using the approximation tends to zero faster than h^k .

We now generalize (7) to the case in which f is a real-valued function of class C^{k+1} or C^{k+2} on an open set D in $\mathbb{R}^n$. We restrict our attention to functions of class C^{k+1} or C^{k+2} because for functions of class C^{k+1} with $k > 2$ the statement of the result is very cumbersome, and in this text we shall only need $k = 1, 2$.

THEOREM 11.2. Let D be an open set in $\mathbb{R}^n$, let x_0 be a point in D , and let f be a real-valued function defined on D . Let h be a vector such that $x_0 + h$ is in D .

(a) If f is of class C^{k+1} on D , then

$$f(x_0 + h) - f(x_0) = \nabla f(x_0) \cdot h + \eta_1(h), \quad (8)$$

where $\eta_1(h)/\|h\| \rightarrow 0$ as $\|h\| \rightarrow 0$.

(b) If f is of class C^{k+2} on D , then

$$f(x_0 + h) - f(x_0) = \nabla f(x_0) \cdot h + \frac{1}{2} h \cdot H(x_0) h + \eta_2(h), \quad (9)$$

where $\alpha_k(\mathbf{h})/\|\mathbf{h}\|^2 \rightarrow 0$ as $\mathbf{h} \rightarrow \mathbf{0}$ and

$$M(\mathbf{x}_0) = \left(\frac{\partial^2 f}{\partial x_j \partial x_k}(\mathbf{x}_0) \right). \quad (10)$$

Proof. Let $\phi(t) = f(\mathbf{x}_0 + t\mathbf{h})$. Then

$$\phi(0) = f(\mathbf{x}_0) \quad \text{and} \quad \phi(1) = f(\mathbf{x}_0 + \mathbf{h}). \quad (11)$$

It follows from the chain rule that ϕ is a function of class C^2 on an open interval containing the closed interval $[0, 1] = \{t \in \mathbb{R} : t \leq 1\}$ in its interior whenever f is of class C^2 . If f is of class C^2 , then using the chain rule, we get

$$\phi'(t) = \sum_{j=1}^n \frac{\partial f}{\partial x_j}(\mathbf{x}_0 + t\mathbf{h})h_j = \langle \nabla f(\mathbf{x}_0 + t\mathbf{h}), \mathbf{h} \rangle. \quad (12)$$

If f is of class C^3 , we get

$$\phi''(t) = \sum_{j=1}^n \sum_{k=1}^n \frac{\partial^2 f}{\partial x_j \partial x_k}(\mathbf{x}_0 + t\mathbf{h})h_j h_k = \langle \mathbf{h}, M(\mathbf{x}_0 + t\mathbf{h})\mathbf{h} \rangle, \quad (13)$$

where M is the matrix defined in (10).

Let f be of class C^2 . If we take $k = 1$, $\alpha_0 = 0$, and $h = 1$ in (8), recall that the summation term in (8) is absent when $d = 1$, and use (11) and (12), we get

$$f(\mathbf{x}_0 + \mathbf{h}) - f(\mathbf{x}_0) = \phi(1) - \phi(0) = \langle \nabla f(\mathbf{x}_0 + \theta\mathbf{h}), \mathbf{h} \rangle, \quad (14)$$

where $0 < \theta < 1$. Since f is of class C^2 , the function ∇f is continuous and thus

$$\nabla f(\mathbf{x}_0 + \theta\mathbf{h}) = \nabla f(\mathbf{x}_0) + \alpha_1(\mathbf{h}),$$

where $\alpha_1(\mathbf{h}) \rightarrow \mathbf{0}$ as $\mathbf{h} \rightarrow \mathbf{0}$. Substituting this relation into (14) and using the Cauchy–Schwarz inequality

$$|\langle \alpha_1(\mathbf{h}), \mathbf{h} \rangle| \leq \|\alpha_1(\mathbf{h})\| \|\mathbf{h}\|$$

give (8), with $\alpha_0(\mathbf{h}) = \langle \alpha_1(\mathbf{h}), \mathbf{h} \rangle$.

If f is of class C^3 , we proceed in similar fashion. We take $k = 2$, $\alpha_0 = 0$, and $h = 1$ in (8) and use (11), (12), and (13) to get

$$f(\mathbf{x}_0 + \mathbf{h}) - f(\mathbf{x}_0) = \langle \nabla f(\mathbf{x}_0), \mathbf{h} \rangle + \frac{1}{2} \langle \mathbf{h}, M(\mathbf{x}_0 + \theta\mathbf{h})\mathbf{h} \rangle, \quad (15)$$

where $\delta < \bar{\epsilon} < 1$. Since f is of class $C^{(2)}$, each entry of H is continuous. Hence

$$f(\mathbf{x}_0 + \mathbf{h}) = f(\mathbf{x}_0) + H(\mathbf{h}),$$

where $H(\mathbf{h})$ is a matrix with entries $m_{ij}(\mathbf{h})$ such that $m_{ij}(\mathbf{h}) \rightarrow 0$ as $\|\mathbf{h}\| \rightarrow 0$. Substituting this relation into (1.7) and using the relations

$$|\langle \mathbf{h}, M(\mathbf{h})\mathbf{h} \rangle| \leq \|\mathbf{h}\| \|M(\mathbf{h})\mathbf{h}\| \leq m(\mathbf{h}) \|\mathbf{h}\|^2,$$

where $m(\mathbf{h}) \rightarrow 0$ as $\|\mathbf{h}\| \rightarrow 0$, gives (9) with $g_{ij}(\mathbf{h}) = \langle \mathbf{h}, M(\mathbf{h})\mathbf{h} \rangle$.

COROLLARY 1 (DARBOUX'S THEOREM). Let D be an open set in $\mathbb{R}^n$ and let f be a real-valued function defined in D . If f is of class $C^{(2)}$ in D , then for each $\mathbf{x}_0$ in D and $\mathbf{h}$ such that $\mathbf{x}_0 + \mathbf{h}$ is in D

$$f(\mathbf{x}_0 + \mathbf{h}) - f(\mathbf{x}_0) = \langle \nabla f(\mathbf{x}_0), \mathbf{h} \rangle, \quad (10)$$

where $\delta < \bar{\epsilon} < 1$. If f is of class $C^{(2)}$, then

$$f(\mathbf{x}_0 + \mathbf{h}) - f(\mathbf{x}_0) = \langle \nabla f(\mathbf{x}_0), \mathbf{h} \rangle + \frac{1}{2} \langle \mathbf{h}, M(\mathbf{x}_0) \mathbf{h} \rangle, \quad (11)$$

where $\delta < \bar{\epsilon} < 1$.

Proof. Equation (10) follows from (10) and equation (1.7) from (11).

We now give a condition under which the existence of partial derivatives does imply differentiability.

THEOREM 11.3. Let D be an open set in $\mathbb{R}^n$, let f be a function defined on D with range in $\mathbb{R}^r$, and let $\mathbf{x}_0$ be a point in D . If f is of class $C^{(1)}$ on some open ball $B(\mathbf{x}_0, r)$ centered at $\mathbf{x}_0$, then f is differentiable at $\mathbf{x}_0$.

Proof. Since each component function f_i , $i = 1, \dots, r$, is of class $C^{(1)}$, from (8) of Theorem 11.2 we have that

$$f(\mathbf{x}_0 + \mathbf{h}) - f(\mathbf{x}_0) = \langle \nabla f(\mathbf{x}_0), \mathbf{h} \rangle + \eta(\mathbf{h}),$$

where $\eta(\mathbf{h})/\|\mathbf{h}\| \rightarrow 0$ as $\|\mathbf{h}\| \rightarrow 0$.

Let $\eta(\mathbf{h}) = (\eta_1(\mathbf{h}), \dots, \eta_r(\mathbf{h}))$. Then

$$f(\mathbf{x}_0 + \mathbf{h}) - f(\mathbf{x}_0) = \nabla f(\mathbf{x}_0)\mathbf{h} + \eta(\mathbf{h}),$$

where $\eta(\mathbf{h})/\|\mathbf{h}\| \rightarrow 0$ as $\|\mathbf{h}\| \rightarrow 0$. Therefore f is differentiable at $\mathbf{x}_0$.

II

CONVEX SETS IN $\mathbb{R}^n$

1. LINES AND HYPERPLANES IN $\mathbb{R}^n$

Convex sets and convex functions play an important role in many optimization problems. The study of convex sets in $\mathbb{R}^n$, in turn, involves the notions of line segments, lines, and hyperplanes in $\mathbb{R}^n$. We therefore begin with a discussion of these objects.

In $\mathbb{R}^2$ and $\mathbb{R}^3$ the vector equation of a line and in $\mathbb{R}^n$ the vector equation of a plane are obtained from geometric considerations. In $\mathbb{R}^n$ we shall reverse this process. We shall define a line to be the set of points satisfying an equation which in vector notation is the same as that of a line in $\mathbb{R}^2$ or $\mathbb{R}^3$. We shall define a hyperplane in $\mathbb{R}^n$ in similar fashion.

In $\mathbb{R}^2$ and $\mathbb{R}^3$ a vector equation of the line through two points $\mathbf{x}_1$ and $\mathbf{x}_2$ is given by (see Figure 2.1)

$$\mathbf{x} = \mathbf{x}_1 + t(\mathbf{x}_2 - \mathbf{x}_1), \quad -\infty < t < \infty. \quad (1)$$

The closed-oriented line segment $[\mathbf{x}_1, \mathbf{x}_2]$ with orientation from initial point $\mathbf{x}_1$ to terminal point $\mathbf{x}_2$ corresponds to values of t in $[0, 1]$. The positive ray from $\mathbf{x}_1$ corresponds to values of $t \geq 0$, and the negative ray from $\mathbf{x}_1$ corresponds to values of $t \leq 0$.

In $\mathbb{R}^n$ we define the line through two points $\mathbf{x}_1$ and $\mathbf{x}_2$ to be the set of points $\mathbf{x}$ in $\mathbb{R}^n$ that satisfy (1). Thus

Definition 1.6. The line in $\mathbb{R}^n$ through two points $\mathbf{x}_1$ and $\mathbf{x}_2$ is defined to be the set of points $\mathbf{x}$ such that $\mathbf{x} = \mathbf{x}_1 + t(\mathbf{x}_2 - \mathbf{x}_1)$, where t is any real number, or in set notation

$$\{\mathbf{x} \mid \mathbf{x} = \mathbf{x}_1 + t(\mathbf{x}_2 - \mathbf{x}_1), \quad -\infty < t < \infty\}. \quad (2)$$

The set in (2) can also be written as

$$\{\mathbf{x} \mid \mathbf{x} = (1 - t)\mathbf{x}_1 + t\mathbf{x}_2, \quad -\infty < t < \infty\}. \quad (3)$$



Figure 1.1.

and, recalling $\alpha = (j - i) \cdot d = 0$,

$$[n_0, n_2] = [n_0, n_1] + [n_1, n_2] \quad (i, j = 0). \quad (4)$$

Definition 1.2. The closed line segment joining the points n_1 and n_2 is denoted by $[n_1, n_2]$ and is defined by

$$[n_1, n_2] = \{x \mid x = (1 - \alpha)n_1 + \alpha n_2, 0 \leq \alpha \leq 1\}$$

or equivalently,

$$[n_1, n_2] = \{x \mid x = \alpha n_1 + (1 - \alpha)n_2, \alpha \in \mathbb{R}, \alpha + (1 - \alpha) = 1\}$$

Definition 1.3. The open line segment joining the points n_1 and n_2 is denoted by (n_1, n_2) and is defined by

$$(n_1, n_2) = \{x \mid x = (1 - \alpha)n_1 + \alpha n_2, 0 < \alpha < 1\},$$

or equivalently,

$$(n_1, n_2) = \{x \mid x = \alpha n_1 + (1 - \alpha)n_2, \alpha \in \mathbb{R}, \alpha + (1 - \alpha) = 1\}.$$

The half-open segment which includes n_1 but not n_2 will be denoted by $[n_1, n_2)$ and is obtained by substituting α for the half-open interval $0 \leq \alpha < 1$, or equivalently by substituting α and β for which $\alpha \in \mathbb{R}, \alpha + \beta \in \mathbb{R}, \alpha + \beta = 1$. Similarly, the half-open line segment which includes n_2 but not n_1 will be denoted by $(n_1, n_2]$ and is obtained by substituting α for the half-open interval $0 < \alpha \leq 1$, or equivalently by substituting α and β for which $\alpha \in \mathbb{R}, \alpha + \beta = 1$.

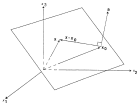


Figure 2.1

Let y be a point in the open line segment (x_1, x_2) . Then $y = \alpha x_1 + \beta x_2$, $\alpha > 0$, $\beta > 0$, $\alpha + \beta = 1$, and $y - x_1 = \beta(x_2 - x_1) = \beta(x_2 - x_1)$. Similarly $y - x_2 = \alpha(x_1 - x_2)$. Hence

$$\frac{\|y - x_1\|}{\|y - x_2\|} = \frac{\beta}{\alpha}. \quad (5)$$

This observation will be useful in the sequel.

In $\mathbb{R}^3$ a plane through the point x_0 with normal n consists of all points x such that $x - x_0$ is orthogonal to n . Thus, the plane is the set defined by (see Figure 2.2)

$$\langle x, \langle n, x - x_0 \rangle \rangle = 0.$$

If $n = (n_1, n_2, n_3)$ and $x_0 = (x_{01}, x_{02}, x_{03})$, we can describe the plane as the set of $x = (x_1, x_2, x_3)$ satisfying

$$n_1(x_1 - x_{01}) + n_2(x_2 - x_{02}) + n_3(x_3 - x_{03}) = 0$$

or

$$n_1x_1 + n_2x_2 + n_3x_3 = b,$$

where $\gamma = a_1x_{n+1} + a_2x_{n+2} + \cdots + a_nx_{2n}$. These are the familiar forms of the equation of a plane in $\mathbb{R}^3$. In vector notation these equations can be written as

$$\langle \mathbf{a}, \mathbf{x} - \mathbf{x}_0 \rangle = 0 \quad \text{or} \quad \langle \mathbf{a}, \mathbf{x} \rangle = \gamma, \quad (5)$$

where $\gamma = \langle \mathbf{a}, \mathbf{x}_0 \rangle$. Conversely, every equation of the form (5) is the equation of a plane with normal vector $\mathbf{a}$. To see this, let $\mathbf{x}_0$ be a point that satisfies (5). Then $\langle \mathbf{a}, \mathbf{x}_0 \rangle = \gamma$, so if $\mathbf{x}$ is any other point satisfying (5), $\langle \mathbf{a}, \mathbf{x} \rangle = \langle \mathbf{a}, \mathbf{x}_0 \rangle$. Hence $\langle \mathbf{a}, \mathbf{x} - \mathbf{x}_0 \rangle = 0$ for all $\mathbf{x}$ satisfying (5). But $\langle \mathbf{a}, \mathbf{x} - \mathbf{x}_0 \rangle = 0$ is an equation of the plane through $\mathbf{x}_0$ with normal $\mathbf{a}$.

In $\mathbb{R}^n$ we take (6) to define a hyperplane.

Definition 1.4. A hyperplane $H_{\mathbf{a}}^c$ in $\mathbb{R}^n$ is defined to be the set of points that satisfy the equation $\langle \mathbf{a}, \mathbf{x} \rangle = c$. Thus

$$H_{\mathbf{a}}^c = \{\mathbf{x} : \langle \mathbf{a}, \mathbf{x} \rangle = c\}.$$

The vector $\mathbf{a}$ is said to be a normal to the hyperplane.

As was the case in $\mathbb{R}^3$, $H_{\mathbf{a}}^c = \{\mathbf{x} : \langle \mathbf{a}, \mathbf{x} \rangle = c\}$ can be written as

$$\langle \mathbf{a}, (\mathbf{x}_0 + (\mathbf{x} - \mathbf{x}_0)) \rangle = c$$

for any $\mathbf{x}_0$ satisfying $\langle \mathbf{a}, \mathbf{x}_0 \rangle = c$, or equivalently for any $\mathbf{x}_0$ in $H_{\mathbf{a}}^c$. Thus we may say that an equation of the hyperplane in $\mathbb{R}^n$ through the point $\mathbf{x}_0$ with normal $\mathbf{a}$ is

$$\langle \mathbf{a}, \mathbf{x} - \mathbf{x}_0 \rangle = 0. \quad (7)$$

Note that in $\mathbb{R}^2$ a hyperplane is a line.

Recall that an $(n - 1)$ -dimensional subspace of $\mathbb{R}^n$ can be represented as the set of vectors $\mathbf{x} = (x_1, \dots, x_n)$ whose coordinates satisfy $a_1x_1 + \cdots + a_nx_n = 0$ for some nonzero vector $\mathbf{a} = (a_1, \dots, a_n)$. Thus, if $c = 0$ and $\mathbf{a} \neq \mathbf{0}$, then the hyperplane $H_{\mathbf{a}}^0$ passes through the origin and is an $(n - 1)$ -dimensional subspace of $\mathbb{R}^n$.

In Section 10 of Chapter I we saw that for a given linear functional L on $\mathbb{R}^n$ there exists a unique vector $\mathbf{a}$ such that $L(\mathbf{x}) = \langle \mathbf{a}, \mathbf{x} \rangle$ for all $\mathbf{x}$ in $\mathbb{R}^n$. This representation of linear functionals and the definition of the hyperplane $H_{\mathbf{a}}^c$ show that hyperplanes are level surfaces of linear functionals. This interpretation will prove to be very useful.

Let A be any set in $\mathbb{R}^n$ and let $\mathbf{x}_0$ be any vector in $\mathbb{R}^n$. The set $A + \mathbf{x}_0$ is defined to be

$$A + \mathbf{x}_0 = \{\mathbf{y} : \mathbf{y} = \mathbf{a} + \mathbf{x}_0, \mathbf{a} \in A\}.$$

and is called the *normal* of A by $\mathbf{x}_0$.

Lemma 1.1. For any $\mathbf{a}$ in $\mathbb{R}^n$ and scalar α , the hyperplane $H_{\mathbf{a}}^{\alpha}$ is the translate of $H_{\mathbf{a}}^0$, the hyperplane through the origin with normal $\mathbf{a}$, by any vector $\mathbf{x}_0$ in $H_{\mathbf{a}}^{\alpha}$.

Proof. Fix an $\mathbf{x}_0$ in $H_{\mathbf{a}}^{\alpha}$. Let $\mathbf{x}$ be an arbitrary point in $H_{\mathbf{a}}^{\alpha}$ and let $\mathbf{u} = \mathbf{x} - \mathbf{x}_0$. Then

$$\alpha = \langle \mathbf{a}, \mathbf{x} \rangle = \langle \mathbf{a}, \mathbf{u} + \mathbf{x}_0 \rangle = \langle \mathbf{a}, \mathbf{u} \rangle + \alpha.$$

Thus $\langle \mathbf{a}, \mathbf{u} \rangle = 0$ and so $\mathbf{u} \in H_{\mathbf{a}}^0$. Since $\mathbf{x} = \mathbf{x}_0 + \mathbf{u}$, we have shown that

$$H_{\mathbf{a}}^{\alpha} \subset H_{\mathbf{a}}^0 + \mathbf{x}_0 \quad \mathbf{x}_0 \in H_{\mathbf{a}}^{\alpha}. \quad (1)$$

Now let $\mathbf{u}$ be an arbitrary vector in $H_{\mathbf{a}}^0$. Then $\mathbf{u} + \mathbf{x}_0 \in H_{\mathbf{a}}^{\alpha} + \mathbf{x}_0$ and

$$\langle \mathbf{a}, \mathbf{u} + \mathbf{x}_0 \rangle = \langle \mathbf{a}, \mathbf{x}_0 \rangle = \alpha.$$

Thus we get the inclusion opposite to the one in (1), and the lemma is proved.

Two hyperplanes $H_{\mathbf{a}}^{\alpha}$ and $H_{\mathbf{b}}^{\beta}$ are said to be *parallel* if their normals are scalar multiples of each other. Thus two hyperplanes are parallel if and only if they are translates of the same hyperplane through the origin. The parallel hyperplane through the origin is called the *parallel subspace*.

Exercise 1.1. (a) Draw a sketch in $\mathbb{R}^2$ and $\mathbb{R}^3$ that illustrates Lemma 1.1.

(b) In $\mathbb{R}^2$ sketch the hyperplane $x + 2y = 1$ and the parallel subspace.

Exercise 1.2. (a) In $\mathbb{R}^2$ find an equation of the hyperplane through the points $(1, 1, 1)$, $(2, 0, 1)$, $(0, 1, 0)$, $(0, 2, 0, 1)$, and $(1, 1, -1, 0)$.

(b) Let $\mathbf{x}_1, \dots, \mathbf{x}_n$ be n points in $\mathbb{R}^n$. Find a sufficient condition for the existence of a unique hyperplane containing these points.

Exercise 1.3. Show that a hyperplane is a closed set.

Exercise 1.4. Given a hyperplane $H_{\mathbf{a}}^{\alpha}$ show that $\mathbf{a}$ is indeed normal to $H_{\mathbf{a}}^{\alpha}$ in the sense that if $\mathbf{x}_1$ and $\mathbf{x}_2$ are any two points in $H_{\mathbf{a}}^{\alpha}$, then $\mathbf{a}$ is orthogonal to $\mathbf{x}_2 - \mathbf{x}_1$.

Exercise 1.5 (Linear Algebra Review). Let V be an $(n - 1)$ -dimensional subspace of $\mathbb{R}^n$. Let $\mathbf{p}$ be a vector in $\mathbb{R}^n$ not in V . Show that every $\mathbf{x}$ in $\mathbb{R}^n$ has a unique representation $\mathbf{x} = \mathbf{v} + \alpha \mathbf{p}$ where $\mathbf{v} \in V$ and $\alpha \in \mathbb{R}$.

2. PROPERTIES OF CONVEX SETS

A subset C of $\mathbb{R}^2$ is said to be *convex* if given any two points x_1, x_2 in C the line segment $[x_1, x_2]$ joining these points is contained in C . We now formulate this definition analytically.

Definition 2.1. A subset C of $\mathbb{R}^2$ is *convex* if for every pair of points x_1, x_2 in C the line segment

$$[x_1, x_2] = \{x: x = \alpha x_1 + \beta x_2, \alpha \geq 0, \beta \geq 0, \alpha + \beta = 1\}$$

belongs to C .

At this point the reader may find it helpful to do the following exercise.

Exercise 2.1. Sketch the following sets in $\mathbb{R}^2$ and determine from your figures which sets are convex and which are not:

- (a) $\{(x, y): x^2 + y^2 \leq 1\}$,
- (b) $\{(x, y): x < x^2 + y^2 \leq 1\}$,
- (c) $\{(x, y): x \geq x^2\}$,
- (d) $\{(x, y): |x| + |y| \leq 1\}$, and
- (e) $\{(x, y): x \geq 1/(1 + y^2)\}$.

In $\mathbb{R}^2$ a plane determines two sets, one on each side of the plane. We extend this notion to $\mathbb{R}^n$. The *positive half space determined by* H_1^n and denoted by H_1^{n+} is defined to be

$$H_1^{n+} = \{x: \phi(n, x) > 0\}. \quad (1)$$

The *negative half space* H_1^{n-} is defined to be

$$H_1^{n-} = \{x: \phi(n, x) < 0\}. \quad (2)$$

The *closed positive half space*, denoted by $\bar{H}_1^{n+}$, is defined to be the closure of H_1^{n+} . The *closed negative half space* $\bar{H}_1^{n-}$ is defined to be the closure of H_1^{n-} .

Note that a half space is not a subspace of $\mathbb{R}^n$, since there exist points x in a half space such that not all scalar multiples of x are in the half space.

Exercise 2.2. Show that

$$\bar{H}_1^{n+} = \{x: \phi(n, x) \geq 0\}, \quad \bar{H}_1^{n-} = \{x: \phi(n, x) \leq 0\}.$$

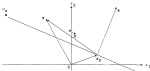


Figure 1.1

If $x_0 \in H_{x_0}^{n+}$, then $\langle n, x_0 \rangle = 0$. Hence, for $x \in H_{x_0}^{n+}$, we have $\langle n, x \rangle = c = \langle n, x_0 \rangle$. Therefore

$$\langle n, x - x_0 \rangle \geq 0 \quad \text{for all } x \in H_{x_0}^{n+}.$$

Thus, vectors n in $H_{x_0}^{n+}$ are such that the angle between n , the normal to $H_{x_0}^{n+}$, and $x - x_0$ lies between 0 and $\pi/2$. In $\mathbb{R}^2$ and $\mathbb{R}^3$ this says that n and $x - x_0$ point in the same direction or have the same orientation. See Figure 1.1.

Half-spaces are convex sets. Although this and other properties of convex sets are "geometrically obvious" in $\mathbb{R}^2$ and $\mathbb{R}^3$, they require proof in $\mathbb{R}^n$. To show that $H_{x_0}^{n+}$ is convex, we must show that if x_1 and x_2 are in $H_{x_0}^{n+}$, then so is $[x_1, x_2]$. Let $x \in [x_1, x_2]$. Then there exist numbers $\lambda \geq 0, \mu \geq 0$ with $\lambda + \mu = 1$ such that $x = \lambda x_1 + \mu x_2$. Hence

$$\langle n, x \rangle = \langle n, \lambda x_1 + \mu x_2 \rangle = \lambda \langle n, x_1 \rangle + \mu \langle n, x_2 \rangle \geq \lambda c + \mu c = c,$$

and so $x \in H_{x_0}^{n+}$. Similar arguments show that $H_{x_0}^{n-}, \hat{H}_{x_0}^{n-}, \hat{H}_{x_0}^{n+}$ and the hyperplane $H_{x_0}^{n0}$ are convex sets.

In the series of lemmas that follow we list some elementary properties of convex sets. We shall provide proofs to show the reader how one "proves the obvious." When speaking of a convex set, we shall always assume that the convex set is not empty.

Lemma 1.1. Let $\{C_\alpha\}$ be a collection of convex sets such that $C = \bigcap_\alpha C_\alpha$ is not empty. Then C is convex.

Proof. Let x_1 and x_2 belong to C . Then for each α the points x_1 and x_2 belong to C_α . Since for each α the set C_α is convex, for each α the line segment $[x_1, x_2]$ belongs to C_α . Hence $[x_1, x_2] \subset C_\alpha$ so C is convex.

Let A be an $m \times n$ matrix with its row $a_i = (a_{i1}, \dots, a_{in})$ and let $b = (b_1, \dots, b_m)$ be a vector in $\mathbb{R}^m$. Let S denote the solution set of $Ax \leq b$. Then S is the intersection of half spaces $\{x : (a_i, x) \leq b_i\}$, so if S is nonempty, then S is convex.

Note that if A and B are convex sets, then $\lambda A + \mu B$ need not be convex.

If A and B are two sets in $\mathbb{R}^n$ and if λ and μ are scalars, we define

$$\lambda A + \mu B = \{x : x = \lambda a + \mu b, a \in A, b \in B\}.$$

Lemma 1.2. *If A and B are convex sets in $\mathbb{R}^n$ and λ and μ are scalars, then $\lambda A + \mu B$ is convex.*

Proof. Let x_1 and x_2 belong to $\lambda A + \mu B$. We must show that $[x_1, x_2]$ belongs to $\lambda A + \mu B$. We have $x_1 = \lambda a_1 + \mu b_1$ for some $a_1 \in A$ and $b_1 \in B$ and $x_2 = \lambda a_2 + \mu b_2$ for some $a_2 \in A$ and $b_2 \in B$. If $x \in [x_1, x_2]$, then there exist scalars $\alpha \geq 0$, $\beta \geq 0$ with $\alpha + \beta = 1$ such that $x = \alpha x_1 + \beta x_2$. Hence

$$\begin{aligned} x &= \alpha x_1 + \beta x_2 = \alpha(\lambda a_1 + \mu b_1) + \beta(\lambda a_2 + \mu b_2) \\ &= \lambda(\alpha a_1 + \beta a_2) + \mu(\alpha b_1 + \beta b_2). \end{aligned}$$

Since A and B are convex, the first term in parentheses on the right is an element a_3 in A and the second term in parentheses on the right is an element b_3 in B . Hence for arbitrary $x \in [x_1, x_2]$,

$$x = \lambda a_3 + \mu b_3$$

so $[x_1, x_2] \subset \lambda A + \mu B$ as required.

Lemma 1.3. *Let $A_1 \subset \mathbb{R}^{n_1}$, $A_2 \subset \mathbb{R}^{n_2}$, ..., $A_k \subset \mathbb{R}^{n_k}$ and let $A_1, \dots, A_k$ be convex. Then $A_1 \times \dots \times A_k$ is a convex set in $\mathbb{R}^{n_1 + \dots + n_k}$.*

We leave the proof as an exercise for the reader.

Lemma 1.4. *If A is convex, then so is $\bar{A}$, the closure of A .*

Proof. Let x and y be points of $\bar{A}$. Let $x \in [x_1, x_2]$, so $x = \alpha x_1 + \beta x_2$ for some $\alpha \geq 0$, $\beta \geq 0$, $\alpha + \beta = 1$. By Remark 1.3.2, there exist sequences $\{x_1\}$ and $\{x_2\}$ of points in A such that $\lim x_1 = x$ and $\lim x_2 = y$. Since A is convex, $x_i = \alpha x_i + \beta x_i \in A$ for each i . Letting $i \rightarrow \infty$, we get that $\lim x_i = x$. From Remark 1.3.1 we have that $x \in \bar{A}$.

The next lemma and its corollaries are fundamental results.

Lemma 1.5. *Let C be a convex set with nonempty interior. Let x_0 and x_1 be two points with $x_0 \in \text{int}(C)$ and $x_1 \in \bar{C}$. Then the line segment $[x_0, x_1]$ is contained in the interior of C .*

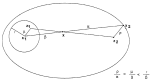


Figure 13

Since $x_1 \in \text{int}(C)$, there exists an $r > 0$ such that $B(x_1, r) \subset C$. To prove the lemma, we must show that if $x \in (x_1, x_2)$, then $x \in \text{int}(C)$. If $x \in (x_1, x_2)$, there exist $\alpha > 0$, $\beta > 0$, $\alpha + \beta = 1$ such that $x = \alpha x_1 + \beta x_2$. Also by relation (1) in Section 1, $\|x - x_1\|/\|x - x_2\| = \beta/\alpha$. Since $x_2 \in C$, there exists a x_3 in C such that $\|x_2 - x_3\| = r/\beta$. Define

$$x_4 = x_1 + \frac{\beta}{\alpha}(x_2 - x_3).$$

Hence $\|x_1 - x_4\| = r$, and so $x_4 \in \text{int}(C)$. Also,

$$\begin{aligned} x &= \alpha x_1 + \beta x_2 = \alpha \left(x_1 + \frac{\beta}{\alpha}(x_2 - x_3) \right) + \beta x_2 \\ &= \alpha x_4 + \beta x_3. \end{aligned}$$

Since $x_4 \in \text{int}(C)$ and $x_3 \in C$, by the case already proved, $x \in \text{int}(C)$.

COROLLARY 1. If C is convex, then $\text{int}(C)$ is either empty or convex.

COROLLARY 2. Let C be convex and let $\text{int}(C) \neq \emptyset$. Then (i) $\overline{\text{int}(C)} = C$ and (ii) $\text{int}(C) = \text{int}(\overline{C})$.

Before proving Corollary 2, we note that its conclusion need not be true if C is not convex. To see this, take $C = \{\text{rational points in } [0, 1] \times [1, 2]\}$.

Proof. Since for any set C , $\text{int}(C) \subseteq C$, it follows that $\overline{\text{int}(C)} \subseteq C$. Let $x \in C$. Then for any y in $\text{int}(C)$, the line segment $[y, x]$ belongs to $\text{int}(C)$. Hence $x \in \text{int}(C)$, and so $C = \text{int}(C)$.

Now, for any set C , $C \subseteq \mathbb{R}^2$, it follows that $\text{int}(C) \subseteq \text{int}(C)$. Let $x \in \text{int}(C)$. Then there exists an $r > 0$ such that $B(x, r) \subseteq C$. Now let $y \in \text{int}(C)$ and consider the line segment $[y, x]$, which is the extension of the segment $[y, x]$ a distance less than r beyond x . Then $x \in C$, and since $y \in \text{int}(C)$, we get that $[x, y] \subseteq \text{int}(C)$. In particular, $x \in \text{int}(C)$. Thus (ii) is proved.

For each positive integer n , define

$$P_n = \left\{ p = (p_1, \dots, p_n) : p_i \geq 0, \sum_{i=1}^n p_i = 1 \right\}.$$

For $n = 1$, P_1 is the point 1. For $n = 2$, P_2 is the closed line segment joining $(0, 1)$ and $(1, 0)$. For $n = 3$, P_3 is the closed triangle with vertices $(1, 0, 0)$, $(0, 1, 0)$, and $(0, 0, 1)$. It is easy to verify that, for each n , P_n is a closed convex set.

Definition 2.2. A point x in $\mathbb{R}^2$ is a *convex combination* of points $x_1, \dots, x_n$ if there exists a $p = (p_1, \dots, p_n)$ in P_n such that

$$x = p_1 x_1 + p_2 x_2 + \dots + p_n x_n.$$

Lemma 2.6. A set C in $\mathbb{R}^2$ is convex if and only if every convex combination of points in C is also in C .

Proof. If every convex combination of points in C is in C , then every convex combination of every two points in C is in C . But for any fixed pair of points x_1 and x_2 in C , the set of convex combinations of x_1 and x_2 is just $[x_1, x_2]$. In other words, for every pair of points x_1 and x_2 in C , the line segment $[x_1, x_2]$ belongs to C , and so C is convex.

We shall prove the "only if" statement by induction on k , the number of points in the convex combination. If C is convex, then the definition of convexity says that every convex combination of two points is in C . Suppose that we have shown that if C is convex, then every convex combination of k points is in C . We shall show that every convex combination of $k + 1$ points is in C . Let

$$x = p_1 x_1 + \dots + p_k x_k + p_{k+1} x_{k+1}, \quad p \in P_{k+1}.$$

If $p_{k+1} = 0$, then there is nothing to prove. If $p_{k+1} < 1$, then $\sum_{i=1}^k p_i > 0$. Hence,

$$x = \left(\sum_{i=1}^k p_i \right) \left(\frac{p_1}{\sum_{i=1}^k p_i} x_1 + \dots + \frac{p_k}{\sum_{i=1}^k p_i} x_k \right) + p_{k+1} x_{k+1}.$$

The term in square brackets is a convex combination of k points in C and so by the inductive hypothesis is in C . Therefore x is a convex combination of two points in C , and so since C is convex, x belongs to C . This proves the lemma.

For a given set A , let $K(A)$ denote the set of all convex combinations of points in A . It is easy to verify that $K(A)$ is convex. Clearly, $K(A) \supseteq A$.

Definition 1.4. The convex hull of a set A , denoted by $\text{co}(A)$, is the intersection of all convex sets containing A .

Since $\mathbb{R}^n$ is convex, $\text{co}(A)$ is nonempty, and since the intersection of convex sets is convex, $\text{co}(A)$ is convex. Also, since we have $\text{co}(A) \subset C$ for any convex set C containing A , it follows that $\text{co}(A)$ is the smallest convex set containing A .

THEOREM 2.1. The convex hull of a set A is the set of all convex combinations of points in A ; that is, $\text{co}(A) = K(A)$.

Proof. Let $\{C_\alpha\}$ denote the family of convex sets containing A . Since $\text{co}(A) = \bigcap C_\alpha$ and $K(A)$ is a convex set containing A , $\text{co}(A) \subset K(A)$. To prove the opposite inclusion, note that $A \subseteq C_\alpha$ for each α and Lemma 2.3 imply that for each α all convex combinations of points in A belong to C_α . That is, for each α , $K(A) \subset C_\alpha$. Hence $K(A) \subset \bigcap C_\alpha = \text{co}(A)$.

Lemma 2.5 states that a set C is convex if and only if $K(C) \subset C$. Since $C \subseteq K(C)$ for any set C , Lemma 2.8 can be restated in the equivalent form that a set C is convex if and only if $K(C) = C$. Since, for any set A , $K(A) = \text{co}(A)$, we have the following corollary to Theorem 2.1.

COROLLARY 1. A set A is convex if and only if $A = \text{co}(A)$.

THEOREM 2.2 (CARATHÉODORY). Let A be a subset of $\mathbb{R}^n$ and let $x \in \text{co}(A)$. Then there exist $n + 1$ points $x_1, \dots, x_{n+1}$ in A and a point p in P_{n+1} such that

$$x = p_1 x_1 + \cdots + p_{n+1} x_{n+1}.$$

Note that since some of the p_i 's may be zero, the result is often stated that a point $x \in \text{co}(A)$ can be written as a convex combination of at most $n + 1$ points of A .

Proof. If $x \in \text{co}(A)$, then by Theorem 2.1,

$$x = \sum_{j=1}^m q_j A_j, \quad A_j \in A, \quad j = 1, \dots, m, \quad q = (q_1, \dots, q_m) \in P_m, \quad (2)$$

for some positive integer m and some q in P_m . If $m \leq n + 1$, then there is

nothing to prove. If $m > n + 1$, we shall express x as a convex combination of at most $m - 1$ points. Repeating this argument, a finite number of times gives the desired result.

Let us then suppose that in (3) $m > n + 1$. Consider the $m - 1 > n$ vectors in $\mathbb{R}^n$, $(x_1 - x_m), (x_2 - x_m), \dots, (x_{m-1} - x_m)$. Since $m - 1 > n$, these vectors are linearly dependent. Hence there exist scalars $\lambda_1, \dots, \lambda_{m-1}$, not all zero, such that

$$\lambda_1(x_1 - x_m) + \lambda_2(x_2 - x_m) + \dots + \lambda_{m-1}(x_{m-1} - x_m) = 0.$$

Let $\lambda_m = -(\lambda_1 + \dots + \lambda_{m-1})$. Then

$$\lambda_1 x_1 + \dots + \lambda_{m-1} x_{m-1} + \lambda_m x_m = 0 \quad (4)$$

and

$$\sum_{i=1}^m \lambda_i = 0. \quad (5)$$

From (1) and (4) it follows that for any i

$$x = \sum_{i=1}^m (x_i - \lambda_i) x_i. \quad (6)$$

We shall show that there is a value of i such that the coefficients of the x_i are nonnegative, have sum equal to 1, and at least one of the coefficients is zero. This will prove the theorem.

Let $I = \{i : i = 1, \dots, m, \lambda_i > 0\}$. Since $\lambda_1, \dots, \lambda_m$ are not all zero, it follows from (5) that $I \neq \emptyset$. Let i_0 denote an index in I such that

$$\frac{q_{i_0}}{\lambda_{i_0}} = \min\{i \in I : q_i/\lambda_i\}.$$

Now take $i = q_{i_0}/\lambda_{i_0}$. Then $i \geq 0$, and for $i \in I$

$$q_i - \lambda_i i = \lambda_i \left(\frac{q_i}{\lambda_i} - \frac{q_{i_0}}{\lambda_{i_0}} \right) \geq 0,$$

with equality holding for $i = i_0$. If $i \notin I$, then $\lambda_i \leq 0$, so $q_i - \lambda_i i \geq 0$ whenever $i \geq 0$. Thus, if $i = q_{i_0}/\lambda_{i_0}$, then

$$q_i - \lambda_i i \geq 0, \quad i = 1, \dots, m, \quad \text{with } q_{i_0} - \lambda_{i_0} i = 0.$$

Upon comparing this statement with (6), we see that we have written x as a nonnegative linear combination of at most $m - 1$ points. This combination is

also a convex combination since, using (2), we have

$$\sum_{i=1}^n (\lambda_i - \lambda_i') = \sum_{i=1}^n \lambda_i - 1 = \sum_{i=1}^n \lambda_i - \sum_{i=1}^n \lambda_i' = 0.$$

Lemma 1.7. *If A is a compact subset of $\mathbb{R}^n$, then so is $\text{co}(A)$.*

If A is closed, then $\text{co}(A)$ need not be closed, as the following example shows. Let

$$A_1 = \{(x_1, x_2) : x_1 > 0, x_2 > 0, x_1 x_2 \geq 1\}$$

and let

$$A_2 = \{(x_1, x_2) : (x_1 > 0, x_2 < 0, x_1 x_2 \leq -1)\}.$$

Let $A = A_1 \cup A_2$. Then A is closed, but $\text{co}(A) = \{(x_1, x_2) : x_1 > 0\}$ is open.

We now prove the lemma. By Theorem 1.2

$$\text{co}(A) = \left\{ \mathbf{x} : \mathbf{x} = \sum_{i=1}^{n+1} \lambda_i \mathbf{x}_i, \quad \mathbf{x}_i \in A, \quad \mathbf{0} = (\lambda_1, \dots, \lambda_{n+1}) \in P_{n+1} \right\}.$$

Let

$$\pi = P_{n+1} \cap \bigcap_{\mathbf{x} \in \text{co}(A)} \lambda_{n+1}^{-1} \mathbf{x}.$$

Since A is compact and P_{n+1} is compact, by Theorem 1.1 the set π is compact in $\mathbb{R}^n$, where $n = -(n+1) + (n+1) = (n+1)^2$. The mapping $\varphi : \mathbb{R}^n \rightarrow \mathbb{R}^n$ defined by

$$\varphi(\lambda_1 \mathbf{x}_1, \dots, \lambda_n \mathbf{x}_n) = \sum_{i=1}^{n+1} \lambda_i \mathbf{x}_i$$

is continuous. Since π is compact, it follows from Theorem 1.6.2 that $\varphi(\pi)$ is compact. By Theorem 1.2, $\varphi(\pi) = \text{co}(A)$ so $\text{co}(A)$ is compact.

Lemma 1.8. *If $\mathcal{C}$ is an open subset of $\mathbb{R}^n$, then $\text{co}(\mathcal{C})$ is also open.*

Proof. Since $\mathcal{O} = \text{co}(\mathcal{C})$, the set $\text{co}(\mathcal{C})$ has nonempty interior. Therefore, by Corollary 1 of Lemma 1.5 $\text{int}(\text{co}(\mathcal{C}))$ is convex and $\mathcal{O} \subset \text{int}(\text{co}(\mathcal{C}))$. (Why?) Since $\text{co}(\mathcal{C})$ is the intersection of all convex sets containing $\mathcal{C}$, we have $\text{co}(\mathcal{C}) \subset \text{int}(\text{co}(\mathcal{C}))$. Since we always have the reverse inclusion, we have that $\text{co}(\mathcal{C}) = \text{int}(\text{co}(\mathcal{C}))$.

Exercise 2.3. Using the definition of a convex set, show that (a) the non-negative orthant in $\mathbb{R}^n = \{x = (x_1, \dots, x_n), x_i \geq 0, i = 1, \dots, n\}$ is convex and (b) a hyperplane H_0^n is convex.

Exercise 2.4. A mapping S from $\mathbb{R}^n$ to $\mathbb{R}^m$ is said to be affine if $Sx = Tx + k$, where T is a linear map from $\mathbb{R}^n$ to $\mathbb{R}^m$ and k is a fixed vector. (If $k = 0$, then S is linear.) Show that if C is a convex set in $\mathbb{R}^n$, then $Sp(C) = \{x, y = Tx + k, x \in C\}$ is a convex set in $\mathbb{R}^m$. In mathematical jargon, you are about to show that under an affine map convex sets are mapped into convex sets or under an affine map the image of a convex set is convex.

Exercise 2.5. Show that the sets P_i are compact and convex.

Exercise 2.6. Show that, for any set A the set $K(A)$ of all convex combinations of points in A is convex.

Exercise 2.7. Show that the open ball $B(0, r)$ is convex.

Exercise 2.8. Consider the linear programming (LP) problem: Minimize $\langle c, x \rangle$ subject to $Ax = b, x \geq 0$. Let $S = \{x, x \text{ is a solution of the problem LP}\}$. Show that if S is not empty, then it is convex.

Exercise 2.9. Let $C \subset \mathbb{R}^n$. Show that C is convex if and only if $\lambda C + \mu C = C$ or λC for all $\lambda \geq 0, \mu \geq 0$.

Exercise 2.10. A set C is said to be a cone with vertex at the origin, or simply a cone, if whenever $x \in C$, all vectors $\lambda x, \lambda \geq 0$, belong to C . If C is also convex, C is said to be a convex cone.

- Give an example of a cone that is not convex.
- Give an example of a cone that is convex.
- Let C be a nonempty set in $\mathbb{R}^n$. Show that C is a convex cone if and only if x_1 and $x_2 \in C$ implies that $\lambda_1 x_1 + \lambda_2 x_2 \in C$ for all $\lambda_1 \geq 0, \lambda_2 \geq 0$.

Exercise 2.11. Show that if C_1 and C_2 are convex cones, then so is $C_1 + C_2$ and that $C_1 + C_2 = \text{co}(C_1 \cup C_2)$.

Exercise 2.12. Show that if A is a bounded set in $\mathbb{R}^n$, then so is $\text{co}(A)$.

Exercise 2.13. Let X be a convex set in $\mathbb{R}^n$ and let p be any vector not in X . Let $[p, X]$ denote the union of all line segments $[p, x]$ with x in X .

- Sketch a set $[p, X]$ in $\mathbb{R}^2$.
- Show that $[p, X]$ is convex.

Exercise 2.14. Let C be a proper subset of $\mathbb{R}^n$. For each $\epsilon > 0$ define a set $\eta_\epsilon(C)$, the ϵ -neighborhood of C , by $\eta_\epsilon(C) = \bigcup_{x \in C} B(x, \epsilon)$. Show that

$$\begin{aligned}\eta_\epsilon(C) &= \{y : \|y - x\| < \epsilon \text{ for some } x \in C\} \\ &= \bigcup_{x \in C} (B(x, \epsilon) + x).\end{aligned}$$

Show that $\eta_\epsilon(C)$ is open. Show that if C is convex, then so is $\eta_\epsilon(C)$.

Exercise 2.15. Show that if a vector x in $\mathbb{R}^n$ has distinct representations as a convex combination of a set of vectors $x_0, x_1, \dots, x_n$, then the vectors $x_0 - x_1, x_1 - x_2, \dots, x_{n-1} - x_n$ are linearly dependent.

Exercise 2.16. Let X be a set contained in the closed negative half space determined by the hyperplane H_0^n . Show that if $x \in \text{int}(X)$, then $\langle a, x \rangle < \alpha$.

Exercise 2.17. Let X be a convex set and let H_0^n be a hyperplane such that $X \cap H_0^n = \emptyset$. Show that X is contained in one of the open half spaces determined by H_0^n .

3. SEPARATION THEOREMS

This section is devoted to the establishment of separation theorems. In some sense, these theorems are the fundamental theorems of optimization theory. The validity of this assertion will be made evident in subsequent chapters of this book.

We have seen that a hyperplane H_0^n divides $\mathbb{R}^n$ into two half spaces, one on each side of H_0^n . It therefore seems natural to say that two sets X and Y are separated by a hyperplane H_0^n if they are contained in different half spaces determined by H_0^n . There are various types of separation, which we will now define precisely and discuss. The reader should draw sketches in $\mathbb{R}^2$. (Recall that in $\mathbb{R}^2$ a hyperplane is a line.)

Definition 3.1. Two sets X and Y are separated by the hyperplane H_0^n if for every $x \in X$, $\langle a, x \rangle \geq \alpha$ and, for every $y \in Y$, $\langle a, y \rangle \leq \alpha$.

This type of separation does not always correspond to what one intuitively considers a separation. For example, in $\mathbb{R}^2$, let

$$X_1 = \{(x_1, x_2) : x_2 = 0, 0 \leq x_1 \leq 2\}$$

and let

$$Y_1 = \{(x_1, x_2) : x_2 = 0, 1 \leq x_1 \leq 3\}.$$

According to Definition 3.1, these sets are separated by the hyperplane $x_n = 0$. On the other hand, one would not consider these sets as being separated in any reasonable way. This example illustrates that two sets which intuitively should not be considered as being separated can be separated by a hyperplane H_n^c according to Definition 3.1 if the sets both lie in H_n^c . To rule out the possibility just discussed, the notion of proper separation is introduced.

Definition 3.2. Two sets X and Y are properly separated by a hyperplane H_n^c if for every x in X , $\langle a, x \rangle = \alpha$, for every y in Y , $\langle a, y \rangle < \alpha$, and at least one of the sets is not contained in H_n^c .

Geometrically, the definition requires that X and Y lie in opposite closed half spaces and at least one of the sets not be contained in H_n^c . Note that the sets X_0 and Y_0 are not properly separated by the hyperplane $x_n = 0$ and that they cannot be properly separated by any hyperplane.

Proper separation of two sets X and Y does not require that the sets be disjoint. The sets

$$X_0 = \{(x_1, x_2) : 0 \leq x_1 \leq 1, 0 \leq x_2 \leq 1\}$$

and

$$Y_0 = \{(x_1, x_2) : 0 \leq x_1 \leq 1, -1 \leq x_2 \leq 0\}$$

are not disjoint but are properly separated by the hyperplane $x_2 = 0$. A set may be a subset of another set, yet the two sets can be properly separated. To see this, let $X_0^* = X_0$ and let $X_1^* = \{(x_1, x_2) : 0 \leq x_1 \leq 1, x_2 = 0\}$. Then $X_1^* \subseteq X_0^*$ and the sets are properly separated by the hyperplane $x_2 = 0$. To rule out the possibility just described, the notion of strict separation is introduced.

Definition 3.3. Two sets X and Y are strictly separated by a hyperplane H_n^c if for every x in X , $\langle a, x \rangle > \alpha$ and, for every y in Y , $\langle a, y \rangle < \alpha$.

Geometrically, the definition requires that X and Y lie in opposite open half spaces determined by H_n^c . Note that the sets X_0 and Y_0 cannot be strictly separated by any hyperplane. The sets

$$X_0 = \{(x_1, x_2) : 0 \leq x_1 \leq 1, 0 < x_2 \leq 1\}$$

and

$$Y_0 = \{(x_1, x_2) : 0 \leq x_1 \leq 1, -1 < x_2 \leq 0\}$$

are strictly separated by the hyperplane $x_2 = 0$. Note that X_0 and Y_0 are properly separated but are not strictly separated by the hyperplane $x_2 = 0$.

Finally, note that the sets X_1 and F_1 are not strictly separated by the hyperplane $x_1 = 0$.

Definition 1.4. Two sets X and F are *strongly separated* by a hyperplane H_0^c if there exists an $\varepsilon > 0$ such that the sets $X + \varepsilon \partial B(1)$ and $F + \varepsilon \partial B(1)$ are strictly separated by H_0^c .

The sets X_0 and F_0 are not strongly separated by the hyperplane $x_0 = 0$ and cannot be strongly separated by any hyperplane. Let $0 < \eta < 1$ be fixed. Then for each such η the sets

$$X_{\eta 0} = \{(x_0, x_1) \in \mathbb{R}^2 : x_0 \leq 1, \eta \leq x_1 \leq 1\}$$

and

$$F_{\eta 0} = \{(x_0, x_1) \in \mathbb{R}^2 : x_0 \leq 1, -1 \leq x_1 \leq -\eta\}$$

are strongly separated by the hyperplane $x_1 = 0$.

It is clear from the definitions that strong separation implies strict separation, which implies proper separation.

Exercise 1.1. Sketch the pairs of sets (X_1, Y_1) , (X_0, Y_0) , (X_1^c, Y_1^c) , (X_0^c, Y_0^c) and $(X_{\eta 0}, Y_{\eta 0})$.

The next lemma gives two conditions, each of which is equivalent to strong separation as given in Definition 1.4.

Lemma 1.1. *The following statements are equivalent:*

- (i) The sets X and F are strongly separated by a hyperplane H_0^c .
- (ii) There exists an $\eta > 0$ such that $\langle a, x \rangle > \alpha + \eta$ for all x in X and $\langle a, y \rangle < \alpha - \eta$ for all y in F .
- (iii) There exists an $\eta' > 0$ such that

$$\inf\{\langle a, x \rangle : x \in X\} \geq \alpha + \eta'$$

and

$$\sup\{\langle a, y \rangle : y \in F\} \leq \alpha - \eta'.$$

Proof. The equivalence of (i) and (iii) is clear, so we need only prove the equivalence of (i) and (ii). We note that $a \in \partial B(1)$ if and only if $-a \in \partial B(1)$. Suppose X and F are strongly separated by H_0^c . Then there exists an $\varepsilon > 0$ such that, for all x in X and all v in $\partial B(1)$, $\langle a, x - \varepsilon v \rangle > \alpha$. Thus for all x in X and

all $x \in B(0, 1)$,

$$-x(a, x) \geq x - \langle a, x \rangle.$$

If we take $x = (R\alpha)/\|a\|$, $0 < R < 1$, and then let $R \rightarrow 1$, we get $-x(a) \geq x - \langle a, x \rangle$, or

$$\langle a, x \rangle \geq x + x(a),$$

for all x in X . A similar argument shows that

$$\langle a, y \rangle \leq x - x(a)$$

for all y in Y . Hence (i) holds for any positive number η strictly less than $x(a)$.

We now suppose that there exists an $\eta > 0$ such that $\langle a, x \rangle \geq x + \eta$ for all x in X . Let x be an arbitrary element of $B(0, 1)$. Since $\|x\| = 1$, from the Cauchy-Schwarz inequality we get that $\langle a, x \rangle \leq \|a\|$, and consequently

$$\langle a, x - a \rangle = \langle a, x \rangle - x(a, a) \geq x + \eta - x(a).$$

Similarly, for all y in Y and $a \in B(0, 1)$,

$$\langle a, y + a \rangle \leq x - \eta + x(a).$$

Hence if we take $x < \eta/\|a\|$, we have that $\langle a, u \rangle \geq x$ for all u in $X + aB(0, 1) = X + aB(0, 1)$ and that $\langle a, v \rangle \leq x$ for all v in $Y + aB(0, 1)$. This proves (i).

Remark 3.2. From the preceding lemma, it should be clear that a necessary and sufficient condition for two sets X and Y to be strongly separated is that there exists a vector a such that

$$\inf\{\langle a, x \rangle : x \in X\} > \sup\{\langle a, y \rangle : y \in Y\}.$$

Our principal objective in this section is to prove Theorem 3.4. Although this theorem is not the best possible separation theorem in $\mathbb{R}^n$, it is the one that is valid in infinite-dimensional spaces and it does cover the situations that occur in optimization problems. The best possible separation theorem in $\mathbb{R}^n$ will be discussed in Section 5. Our proof of Theorem 3.4 will not be the most economical one but will be one that is suggested by rather obvious geometric considerations.

The principal step in the proof is to show that if C is a convex set and p is a point not in C , then p and C can be properly separated. If C is not convex, then it may not be possible to separate p and C , as can be seen by taking, in $\mathbb{R}^2$, $p = 0$ and C to be the disjunct union of any circle with center at the origin.



Figure 2.6.

We first show that if C is closed and convex, then strong separation is possible.

THEOREM 2.1. Let C be a closed convex set and y a vector such that $y \notin C$. Then there exists a hyperplane H_C^y that strongly separates y and C .

To motivate the proof, we argue heuristically from Figure 2.6, in which C is assumed to have a tangent at each boundary point.

Draw a line from y to x_0 , the point of C that is closest to y . The vector $p = x_0 - y$ will be perpendicular to C in the sense that $p = x_0$ is the normal to the tangent line at x_0 . The tangent line, whose equation is $\langle p - x_0, x - x_0 \rangle = 0$, properly separates y and C . The point x_0 is characterized by the fact that $\langle p - x_0, x - x_0 \rangle \leq 0$ for all $x \in C$. To obtain strong separation, we merely move the line parallel to itself so as to pass through a point $x_0 + \alpha(p - x_0)$. We now justify these steps in a series of lemmas, some of which are of interest in their own right.

If X is a set and y is a vector, then we define the distance from y to X , denoted by $d(y, X)$, to be

$$d(y, X) = \inf\{\|y - x\| : x \in X\}.$$

A point x_0 in X is said to attain the distance from y to X , or to be closest to y , if $d(y, X) = \|y - x_0\|$.

LEMMA 2.2. Let C be a convex subset of $\mathbb{R}^n$ and let $y \notin C$. If there exists a point in C that is closest to y , then it is unique.

Proof. Suppose that there were two points x_1 and x_2 of C that were closest to y . Then since C is convex, $(x_1 + x_2)/2$ belongs to C , and so

$$\begin{aligned} d(y, C) &\leq \|(x_1 + x_2)/2 - y\| = \frac{1}{2}(\|x_1 - y\| + \|x_2 - y\|) \\ &\leq \frac{1}{2}(\|x_1 - y\| + \|x_2 - y\|) = d(y, C). \end{aligned}$$

Hence equality holds in the application of the triangle inequality, and so by (5) of Section 2 in Chapter I, we have that, for some $\alpha \geq 0$, $\|x_1 - y\| = \alpha\|x_2 - y\|$. Suppose $\alpha > 0$. Since $\|x_1 - y\| = \|x_2 - y\| = d(y, C)$, we get that $\alpha = 1$, and so $x_1 = x_2$. If $\alpha = 0$, we get $x_1 = y$, contradicting the assumption that $y \notin C$.

Lemma 3.3. *Let C be a closed subset of $\mathbb{R}^n$ and let $y \notin C$. Then there exists a point x_0 in C that is closest to y .*

Proof. Let $x_0 \in C$ and let $r > \|x_0 - y\|$. Then $C_1 = \overline{B(y, r)} \cap C$ is nonempty, closed, and bounded and hence is compact. The function $x \mapsto \|x - y\|$ is continuous on C_1 and so attains its minimum at some point x_0 in C_1 . Thus, for all $x \in C_1$, $\|x - y\| \geq \|x_0 - y\|$. For $x \in C$ and $x \notin C_1$, we have

$$\|x - y\| > r > \|x_0 - y\| \geq \|x_0 - y\|,$$

since $x_0 \in \overline{B(y, r)}$.

Lemmas 3.2 and 3.3 show that if C is a closed convex set and $y \notin C$, then there is a unique closest point in C to y .

The next lemma characterizes closest points.

Lemma 3.4. *Let C be a convex set and let $y \notin C$. Then $x_0 \in C$ is a closest point in C to y if and only if*

$$\langle y - x_0, x - x_0 \rangle \leq 0 \quad \text{for all } x \in C. \quad (1)$$

Note that if C is not convex, then the above characterization of closest point need not hold. To see this, let y be the origin in $\mathbb{R}^2$ and let C be the circumference of the unit circle. Let x_0 be a fixed point in C . Then it is not true that $\langle 0 - x_0, x - x_0 \rangle \leq 0$ for all x in C .

Proof. Let x_0 be a closest point to y and let x be any point in C . Since C is convex, the line segment $[x_0, x] = \{t(x) + (1-t)x_0 : 0 \leq t \leq 1\}$ belongs to C . Let

$$\varphi(t) = \|t(x) - y\|^2 = (x_0 + t(x - x_0) - y, x_0 + t(x - x_0) - y). \quad (2)$$

For $0 \leq t \leq 1$, $\varphi(t)$ is the square of the distance between the point $t(x) \in [x_0, x]$ and y . If $t = 0$, then $x = x_0$. Since φ is continuously differentiable on $(0, 1)$ and

x_0 is a point in C closest to y , we have $\rho(0+) \geq 0$. Calculating $\rho(t)$ from (2) gives

$$\rho(t) = \|y - (y - x_0, x - x_0) + t(x - x_0)\|^2. \quad (3)$$

If we now let $t \rightarrow 0+$ and use $\rho(0+) \geq 0$, we get (1).

We now suppose that (1) holds. Let x be any other point in C . It follows from (3) that $\rho(t) \geq 0$ for $0 \leq t \leq 1$. This ρ is strictly increasing function on $[0, 1]$, and so for any point $x(t)$ in $(x_0, x]$. We have $\|x(t) - y\| \geq \|x_0 - y\|$. In particular, this is true for $x = x_0$, and so x_0 is a closest point to y .

We can now complete the proof of Theorem 3.1. Let x_0 be the closest point in C to y and let $u = y - x_0$. Then for all $x \in C$, $\langle u, x - x_0 \rangle \leq 0$, and so $\langle u, x \rangle \leq \langle u, x_0 \rangle$, with equality occurring when $x = x_0$. Therefore $\sup\{\langle u, x \rangle : x \in C\} = \langle u, x_0 \rangle$. On the other hand, $\langle u, y - x_0 \rangle = \|u\|^2 > 0$, so

$$\langle u, y \rangle = \langle u, x_0 \rangle + \|u\|^2 > \sup\{\langle u, x \rangle : x \in C\}.$$

The conclusion of the theorem now follows from Remark 3.1 with $K = \{y\}$ and $T = C$.

We now take up the separation of a point y and an arbitrary convex set that is, we shall no longer assume that C is closed. We shall obtain separation, but shall not be able to guarantee proper separation unless we assume that C has nonempty interior.

Theorem 3.2. *Let C be a convex set and let $y \notin C$. Then there exists a hyperplane H_C^y such that for all x in C*

$$\langle u, x \rangle \leq \alpha \quad \text{and} \quad \langle u, y \rangle = \alpha, \quad (4)$$

if $\text{int}(C) \neq \emptyset$, then for all x in $\text{int}(C)$

$$\langle u, x \rangle < \alpha. \quad (5)$$

Proof. We first suppose that $y = 0$. Then to establish (4), we must find an $u \neq 0$ such that

$$\langle u, x \rangle \leq 0 \quad \text{for all } x \text{ in } C. \quad (6)$$

For each x in C define

$$N_x = \{u : \|u\| = 1, \langle u, x \rangle \leq 0\}.$$

The set N_x is nonempty since it contains the element $-x/\|x\|$.

To establish (4), it suffices to show that

$$\bigcap_{\alpha \in \mathcal{A}} N_\alpha \neq \emptyset.$$

Each N_α is closed and is contained in $S(\mathbf{B}, 1) = \{\mathbf{u} : \|\mathbf{u}\| = 1\}$. Thus we may write the last relation as

$$\bigcap_{\alpha \in \mathcal{A}} N_\alpha = \left(\bigcap_{\alpha \in \mathcal{A}} N_\alpha \right) \cap S(\mathbf{B}, 1) \neq \emptyset. \quad (4)$$

Since $S(\mathbf{B}, 1)$ is compact, it has the finite-intersection property. Thus, if the intersection in (4) were empty, there would exist a finite subcollection $N_{\alpha_1}, \dots, N_{\alpha_k}$ that is empty. Therefore, we may establish (4) by showing that every finite subcollection $N_{\alpha_1}, \dots, N_{\alpha_k}$ has nonempty intersection.

Now consider any finite collection of sets $N_{\alpha_1}, \dots, N_{\alpha_k}$ and the corresponding k points $\mathbf{x}_1, \dots, \mathbf{x}_k$ in C . Let $\text{co}\{\mathbf{x}_1, \dots, \mathbf{x}_k\}$ denote the convex hull of $\mathbf{x}_1, \dots, \mathbf{x}_k$. Then $\text{co}\{\mathbf{x}_1, \dots, \mathbf{x}_k\}$ is compact and convex. Since C is convex, $\text{co}\{\mathbf{x}_1, \dots, \mathbf{x}_k\} \subset C$. Thus, since $\mathbf{0} \notin C$, we have $\mathbf{0} \notin \text{co}\{\mathbf{x}_1, \dots, \mathbf{x}_k\}$. Hence, by Theorem 3.1, there exists a vector $\mathbf{u} \neq \mathbf{0}$ such that

$$0 = \langle \mathbf{u}, \mathbf{0} \rangle > \langle \mathbf{u}, \mathbf{x}_i \rangle, \quad \mathbf{x}_i \in \text{co}\{\mathbf{x}_1, \dots, \mathbf{x}_k\}. \quad (5)$$

We may divide through by $\|\mathbf{u}\| \neq 0$ in (5) and so assume that $\|\mathbf{u}\| = 1$. Thus, (5) says that

$$\mathbf{u} \in \left(\bigcap_{i=1}^k N_{\alpha_i} \right) \cap S(\mathbf{B}, 1),$$

and we have shown that the intersection of any finite subcollection $N_{\alpha_1}, \dots, N_{\alpha_k}$ of the closed sets $\{N_\alpha\}_{\alpha \in \mathcal{A}}$ has nonempty intersection with $S(\mathbf{B}, 1)$. Hence (4) holds.

We now remove the restriction that $\mathbf{p} = \mathbf{0}$. We have

$$y \in C \quad \text{if and only if} \quad \mathbf{0} \in C - y = \{\mathbf{x}' : \mathbf{x}' = \mathbf{x} - y, \mathbf{x} \in C\}.$$

Therefore, there exists an $\mathbf{u} \neq \mathbf{0}$ such that $\langle \mathbf{u}, \mathbf{x}' \rangle \leq 0$ for all $\mathbf{x}'$ in $C - y$. Hence for all $\mathbf{x}$ in C ,

$$\langle \mathbf{u}, \mathbf{x} - y \rangle \leq 0,$$

and so

$$\langle \mathbf{u}, \mathbf{x} \rangle \leq \langle \mathbf{u}, y \rangle \quad \text{for all } \mathbf{x} \text{ in } C.$$

If we now let $x = (x, y)$, we get the first conclusion of the theorem. The second follows from Exercise 3.16.

Our first separation result for convex sets will follow from Theorem 3.2.

THEOREM 3.3. Let X_0 and Y be two disjoint convex sets. Then there exists a hyperplane H_0^* that separates them.

Note that the theorem does not assert that proper separation can be achieved. Theorem 3.2 in the sequel will allow us to conclude that the separation is proper. Our proof of Theorem 3.3 does not yield this fact.

To prove the theorem, we first note that

$$X_0 \cap Y = \emptyset \quad \text{if and only if } 0 \notin X_0 - Y.$$

Let

$$A = X_0 - Y = X_0 + (-1)Y.$$

Since X_0 and Y are convex, by Lemma 2.2, so is A . Since X_0 and Y are disjoint, $0 \notin A$. Hence by Theorem 3.2, there exists an $\alpha \neq 0$ such that

$$\langle \alpha, x \rangle \leq 0 \quad \text{all } x \text{ in } A.$$

Let x be an arbitrary element of X_0 and let y be an arbitrary element of Y . Then $x - y = y$ is in A and

$$\langle \alpha, x \rangle \leq \langle \alpha, y \rangle. \quad (7)$$

Let $\beta = \sup\{\langle \alpha, x \rangle : x \in X_0\}$ and let $\gamma = \inf\{\langle \alpha, y \rangle : y \in Y\}$. Then from (7) we get that β and γ are finite, and for any number α such that $\beta \leq \alpha \leq \gamma$,

$$\langle \alpha, x \rangle \leq \alpha \leq \langle \alpha, y \rangle$$

for all x in X_0 and all y in Y . Thus the hyperplane H_0^* separates X_0 and Y .

THEOREM 3.4. Let X and Y be two convex sets such that $\text{int}(X)$ is not empty and $\text{int}(Y)$ is disjoint from X . Then there exists a hyperplane H^* that properly separates X and Y .

To illustrate the theorem, let $X = \{(x_1, x_2) : 0 \leq x_1 \leq 1, -1 \leq x_2 \leq 1\}$ and let $Y = \{(x_1, x_2) : x_1 = 0, -1 \leq x_2 \leq 1\}$. The hypotheses of the theorem are fulfilled and $x_1 = 0$ properly separates X and Y and hence X and Y . Note that strict separation of X and Y is not possible.

Proof. Since $\text{int}(X)$ is convex and disjoint from S , we can apply Theorem 3.3 with $K_0 = \text{int}(X)$ and obtain the existence of an $\alpha \neq 0$ and an r such that

$$\langle \alpha, x \rangle < r < \langle \alpha, y \rangle \quad (8)$$

for all x in $\text{int}(X)$ and all y in S .

Let $x \in \bar{X}$. By Corollary 2 to Lemma 2.5, $\bar{X} = \overline{\text{int}(X)}$. Hence $x \in \overline{\text{int}(X)}$. By Remark 2.2 there exists a sequence of points $\{x_k\}$ in $\text{int}(X)$ such that $x_k \rightarrow x$. It then follows from (8) and the continuity of the inner product that (8) holds for all x in $\bar{X}$ and all $y \in S$. Thus, the hyperplane H_α^r separates $\bar{X}$ and S by Corollary 2 to Lemma 2.5, $\text{int}(\bar{X}) = \text{int}(X)$. By Exercise 2.16, for $x \in \text{int}(\bar{X})$, $\langle \alpha, x \rangle < r$, so the separation is proper.

Theorem 3.5. *Let K be a compact convex set and let C be a closed convex set such that K and C are disjoint. Then K and C can be strongly separated.*

Proof. Let

$$K_\epsilon = K + B(0, \epsilon) = \{x + \alpha \mid x \in K, \|\alpha\| < \epsilon\},$$

$$C_\epsilon = C + B(0, \epsilon) = \{y + \alpha \mid y \in C, \|\alpha\| < \epsilon\}.$$

Since $K_\epsilon = \bigcup_{\alpha \in B(0, \epsilon)} (K + \alpha)$, it is a union of open sets and hence is open. It is also readily verified that K_ϵ is convex. Similarly, the set C_ϵ is open and convex. (See Exercise 2.14.)

We now show that there exists an $\epsilon > 0$ such that $K_\epsilon \cap C_\epsilon = \emptyset$. If the assertion were false, there would exist a sequence $\{\epsilon_k\}$ with $\epsilon_k \rightarrow 0$ and $\epsilon_k > 0$ and a sequence $\{\alpha_k\}$ such that, for each k , $\alpha_k \in K_{\epsilon_k} \cap C_{\epsilon_k}$. Since $\alpha_k = x_k + y_k$ with $x_k \in K$, $\|y_k\| < \epsilon_k$ and $\alpha_k = z_k + v_k$ with $z_k \in C$ and $\|v_k\| < \epsilon_k$, we have a sequence $\{x_k\}$ in K and a sequence $\{z_k\}$ in C such that

$$\|x_k - \alpha_k\| < \epsilon_k, \quad \|z_k - \alpha_k\| < \epsilon_k.$$

Hence

$$\|x_k - z_k\| = \|(x_k - \alpha_k) + (\alpha_k - z_k)\| \leq \|x_k - \alpha_k\| + \|\alpha_k - z_k\| < 2\epsilon_k. \quad (9)$$

Since K is compact, there exists a subsequence $\{\alpha_{k_j}\}$ that converges to an element α_0 in K . It follows from (9) that $\alpha_{k_j} \rightarrow \alpha_0$. Since C is closed, $\alpha_0 \in C$. This contradicts the assumption that C and K are disjoint, and so the assertion is true.

We have shown that there exists an $\epsilon > 0$ such that K_ϵ and C_ϵ are disjoint open convex sets. Hence by Theorem 3.4 there is a hyperplane that properly separates K_ϵ and C_ϵ . Since both K_ϵ and C_ϵ are open sets, it follows from Exercise 2.18 that K_ϵ and C_ϵ are strictly separated. According to Definition 3.4, this says that K and C are strongly separated.

We now apply the separation theorem to obtain an additional characterization of convex sets.

THEOREM 3.6. Let A be a set contained in some half space. Then the closure of the convex hull of A , $\text{co}(A)$, is the intersection of all closed half spaces containing A .

Proof. We select the hyperplanes that determine the closed half spaces containing A to be the negative closed half spaces. If H_0^{+} is a closed half space containing A , then $\text{co}(A)$ is contained in H_0^{+} , and hence in the intersection of all such half spaces. Now let x be a point in the intersection. If x were not in $\text{co}(A)$, then by Theorem 3.1, there would exist a hyperplane H_1^0 such that

$$\langle \mathbf{a}, \mathbf{x} \rangle < \alpha < \langle \mathbf{a}, \mathbf{a}_i \rangle, \quad \mathbf{x} \in \overline{\text{co}(A)}.$$

Hence x would not be in the closed half space by H_1^{+} which contains A . Therefore x is not in the intersection of all negative closed half spaces containing A , which contradicts the assumption that x is in the intersection. Thus the intersection is contained in $\text{co}(A)$, and the theorem is proved.

The following corollary follows from the theorem and the fact that a closed convex set not equal to $\mathbb{R}^n$ is contained in some closed half space.

COROLLARY 1. If A is a closed convex set not equal to $\mathbb{R}^n$, then A is the intersection of all closed half spaces containing A .

Note that the theorem is false if we replace $\overline{\text{co}(A)}$ by $\text{co}(A)$. To see this, consider the set A in $\mathbb{R}^2$ defined by

$$A = \{(x_1, x_2) : (x_1 \geq 0, x_1 x_2 \geq 1) \cup \{(x_1, x_2) : (x_1 \leq 0, x_1 x_2 \leq -1)\}.$$

Then $\text{co}(A) = \{(x_1, x_2) : x_1 \geq 0\}$, and the intersection of all closed half spaces is $\{(x_1, x_2) : x_1 \geq 0\}$.

Exercise 3.2. Let V be a linear subspace and let $y \notin V$. Show that $\mathbf{x}_0$ in V is the closest point in V to y if and only if $y - \mathbf{x}_0$ is orthogonal to V ; that is, for every $\mathbf{w}$ in V , $y - \mathbf{x}_0$ is orthogonal to $\mathbf{w}$.

Exercise 3.3. Let C be a closed convex set and let $p \notin C$. Show that $\mathbf{x}_0$ in C is closest to p if and only if $\langle \mathbf{x} - \mathbf{x}_0, \mathbf{x}_0 - p \rangle \geq \|\mathbf{x}_0 - p\|^2$ for all $\mathbf{x}$ in C .

Exercise 3.4. In reference to Theorem 3.5, show that if K is assumed to be convex, and closed, then it may not be possible to separate K and C strictly, let alone strongly.

Exercise 3.8. Show that if F is closed and K is compact, then $K + F$ is closed. Give an example in which F and K are closed yet $F + K$ is not.

Exercise 3.9. Let A and B be two compact sets. Show that A and B can be strongly separated if and only if $\text{co}(A) \cap \text{co}(B) = \emptyset$.

Exercise 3.10. Prove Theorem 3.5 using Theorem 3.1 and Exercise 3.9.

Exercise 3.11. Let A be a bounded set. Show that $\text{co}(A)$ is the intersection of all closed half spaces containing A . Show that the statement is false if A is not bounded and a proper subset of $\mathbb{R}^n$.

Exercise 3.12. Let A be a closed convex set such that $\text{c}A$ (the complement of A) is convex. Show that A is a closed half space.

Exercise 3.13. Let C_1 and C_2 be two convex subsets of $\mathbb{R}^n$. Show that there exists a hyperplane that strongly separates C_1 and C_2 if and only if

$$\inf\{\|x - y\| : x \in C_1, y \in C_2\} > 0.$$

4. SUPPORTING HYPERPLANES: EXTREME POINTS

In studying convex sets we do not wish to restrict ourselves to sets with smooth boundaries, that is, sets such that every boundary point has a tangent plane. The concept that replaces that of tangent plane is that of supporting hyperplane. It is a concept that has many important consequences, as we shall see.

We first define what is meant by a boundary point of a set. A point x is said to be a boundary point of a set S if for every $\epsilon > 0$ the ball $B(x, \epsilon)$ contains a point $u \in S$ and a point $y \notin S$. Note that the definition allows us to have $y = x$ or $x = u$. For example, in $\mathbb{R}^2$ let $S = \{x = (x_1, x_2) : 0 \leq |x| \leq 1\}$. Then 0 is a boundary point of S , and for every $0 < \epsilon < 1$ the only point in $B(0, \epsilon)$ that is not in S is 0 . The other boundary points of S are the points x with $\|x\| = 1$. If there exists an $\epsilon > 0$ such that the only point in $B(x, \epsilon)$ that is in S is x itself, then we say that x is an isolated point of S .

Definition 4.1. A hyperplane H_c^0 is said to be a supporting hyperplane to a set S if for every $x \in S$, $\langle x, x^0 \rangle \leq \alpha$ and there exists at least one point $u_0 \in S$ such that $\langle u_0, u_0^0 \rangle = \alpha$. The hyperplane H_c^0 is said to support S at u_0 . The hyperplane H_c^0 is a nontrivial supporting hyperplane if there exists a point $x_1 \in S$ such that $\langle x_1, u_0^0 \rangle < \alpha$.

To illustrate this definition, consider the set S in $\mathbb{R}^2$ defined by

$$S = \{x \in \mathbb{R}^2 : x_1^2 + x_2^2 \leq 1, x_2 \geq 0\}.$$

As a set in $\mathbb{R}^2$, every point of S is a boundary point. The only supporting hyperplane at a point of S is the trivial one, $\alpha_1 = 0$. Every point in the set $E = \{(x, x) + x_0 \mid x_0 = 0\}$ is also a boundary point of S . At every point of E there is a nontrivial supporting hyperplane with equation $ax_1 + bx_2 = 1$ for appropriate a and b .

Remark 4.1. Let K be a compact set in E^n . Then for each α in E^n there exists an ϵ such that the hyperplane H_α^ϵ with equation $\langle \alpha, x \rangle = \epsilon$ is a supporting hyperplane to K . To see this, note that the linear functional C defined by $C(x) = \langle \alpha, x \rangle$ is continuous on K and so attains its maximum at a point x_0 in K . Then $\langle \alpha, x \rangle \leq \langle \alpha, x_0 \rangle$ for all $x \in K$. The desired hyperplane is obtained by taking $\epsilon = \langle \alpha, x_0 \rangle$.

Theorem 4.1. Let C be a convex set and let x be a boundary point of C . Then there exists a supporting hyperplane H_x^ϵ to C such that $x \in H_x^\epsilon$, $H_x^\epsilon \cap \text{int}(C) \neq \emptyset$, the supporting hyperplane is nontrivial.

Proof. If $x \in C$, then by Theorem 3.2 with $p = x$ there exists a hyperplane H_x^ϵ such that for all y in C

$$\langle y, x \rangle \leq \epsilon \quad \text{and} \quad \langle x, x \rangle = \epsilon.$$

Since $x \in C$ and x is a boundary point of C , it follows that $x \in \bar{C}$. Thus, H_x^ϵ is a supporting hyperplane to C with $x \in H_x^\epsilon$.

If $x \notin C$, then from the definition of boundary point it follows that there exists a sequence of points $\{x_n\}$ with $x_n \in C$ such that $x_n \rightarrow x$. It follows from Theorem 3.2 that for each positive integer n there exists a vector $\alpha_n \neq 0$ such that

$$\langle \alpha_n, x \rangle \leq \langle \alpha_n, x_n \rangle \quad \text{for all } x \in C. \quad (1)$$

If we divide through by $\|\alpha_n\| \neq 0$ in (1), we see that we may assume that $\|\alpha_n\| = 1$. Since $\mathbb{R}^n \times \mathbb{R}$ is compact, there exists a subsequence of $\{\alpha_n\}$ that we again label as $\{\alpha_n\}$ and a vector α with $\|\alpha\| = 1$ such that $\alpha_n \rightarrow \alpha$. If we now let $n \rightarrow \infty$ in (1), we get that

$$\langle \alpha, x \rangle \leq \langle \alpha, x \rangle \quad \text{for all } x \in C.$$

If we set $\epsilon = \langle \alpha, x \rangle$, we see that the hyperplane H_x^ϵ contains x and is a supporting hyperplane to C .

It follows from Exercise 2.16 that the supporting hyperplane is nontrivial if $\text{int}(C) \neq \emptyset$.

Definition 4.2. A point x_0 belonging to a convex set C is an *extreme point* of C if, for no pair of distinct points x_1, x_2 in C and for no $\alpha > 0, \beta > 0, \alpha + \beta = 1$, is it true that $x_0 = \alpha x_1 + \beta x_2$.

In other words, a point x_0 in a convex set C is an extreme point of C if x_0 is interior to no closed interval $[x_1, x_2]$ in C . We leave it as an exercise for the reader to show that an extreme point is a boundary point.

Let

$$C_1 = \{(x_1, x_2) : x_1^2 + x_2^2 \leq 1\}.$$

The set of extreme points of C_1 is $\{(x_1, x_2) : x_1^2 + x_2^2 = 1\}$.

Let

$$C_1 = \{(x_1, x_2) : x_1^2 + x_2^2 \leq 1\}, \quad C_2 = \{(x_1, x_2) : x_1^2 + x_2^2 = 1, x_1 \geq 0\},$$

and let $C_3 = C_1 \cup C_2$. Then the set of extreme points of C_3 is C_2 .

A convex set need not have any extreme points, as evidenced by C_1 . Another example of a convex set without any extreme points is $\{(x_1, x_2) : x_1 = 0\}$.

THEOREM 4.2. Let C be a compact convex set. Then the set of extreme points of C is not empty, and every supporting hyperplane of C contains at least one extreme point of C .

The following observation will be used in the proof of the theorem.

LEMMA 4.1. Let C be a convex set and let H be a supporting hyperplane to C . A point x_0 in $H \cap C$ is an extreme point of C if and only if x_0 is an extreme point of $H \cap C$.

Proof. Let x_0 in $H \cap C$ be an extreme point of C . If x_0 were not an extreme point of $H \cap C$, then x_0 would be an interior point of some closed interval $[x_1, x_2]$ contained in $H \cap C$. But then $[x_1, x_2]$ would be contained in C , and x_0 would not be an extreme point of C . Thus, if x_0 is an extreme point of C , then it is also an extreme point of $H \cap C$.

Let x_0 be an extreme point of $H \cap C$. If x_0 were not an extreme point of C , then there would exist points x_1 and x_2 in C with at least one of x_1 and x_2 not in H such that for some $0 < t < 1$

$$x_0 = tx_1 + (1 - t)x_2.$$

Suppose that $H = H_{x_0}^C$ that $\langle x, n \rangle \leq c$ for x in C , and that $x_1 \notin H_{x_0}^C$. Then $\langle x_1, n \rangle > c$ and

$$c = \langle x, n \rangle = t\langle x_1, n \rangle + (1 - t)\langle x_2, n \rangle < tc + (1 - t)c = c.$$

This contradiction shows that x_0 must be an extreme point of C .

We now show that the set of extreme points of a compact convex set is not empty.

Lemma 4.2. *The set of extreme points of a compact convex set C is not empty.*

Proof. If C is a singleton, then there is nothing to prove. Therefore we may suppose that C is not a singleton. Since C is compact and the norm is a continuous function, there exists a point x^* in C of maximum norm. Thus $\|x\| \leq \|x^*\|$ for all x in C and $\|x^*\| > 0$.

We assert that x^* is an extreme point of C . If x^* were not an extreme point, there would exist distinct points x_1 and x_2 in C different from x^* and scalars $\alpha > 0, \beta > 0, \alpha + \beta = 1$ such that

$$x^* = \alpha x_1 + \beta x_2.$$

Since

$$\|x^*\| = \|\alpha x_1 + \beta x_2\| \leq \alpha\|x_1\| + \beta\|x_2\|,$$

we see that if either $\|x_1\|$ or $\|x_2\|$ was strictly less than $\|x^*\|$, we would have

$$\|x^*\| < (\alpha + \beta)\|x^*\| = \|x^*\|.$$

This contradiction implies that

$$\|x_1\| = \|x_2\| = \|x^*\|. \quad (2)$$

But then,

$$\begin{aligned} \|x^*\|^2 &= \|\alpha x_1 + \beta x_2\|^2 = \alpha^2\|x_1\|^2 + 2\alpha\beta\langle x_1, x_2 \rangle + \beta^2\|x_2\|^2 \\ &\leq \alpha^2\|x_1\|^2 + 2\alpha\beta\|x_1\|\|x_2\| + \beta^2\|x_1\|^2 \\ &= (\alpha + \beta)^2\|x^*\|^2 = \|x^*\|^2. \end{aligned}$$

Hence we must have

$$2\alpha\beta\langle x_1, x_2 \rangle = 2\alpha\beta\|x_1\|\|x_2\|.$$

Since $\alpha > 0, \beta > 0$, we have $\langle x_1, x_2 \rangle = \|x_1\|\|x_2\|$. Therefore $x_2 = \alpha x_1$ for some $\alpha > 0$. The relation (2) implies that $\alpha = 1$, and so $x_2 = x_1$. This contradicts the assumption that x_1 and x_2 are distinct. Thus x^* is an extreme point.

We now complete the proof of the theorem. Let M_1^C be a supporting hyperplane to C . Since C is compact, the point x_0 in C such that $\langle x_0, x_0 \rangle = r$ is also in C . Hence $M_1^C \cap C$ is nonempty and is compact and convex. By Lemma 4.2, $M_1^C \cap C$ contains an extreme point x_1 . By Lemma 4.1, x_1 is also an extreme point of C .

THEOREM 4.3. *Let C be a compact convex set. Let C_e denote the set of extreme points of C . Then $C = \text{co}(C_e)$.*

We first note that by Theorem 4.2 the set C_e is not empty.

Since $C_e \subseteq C$ and C is convex, we have that $\text{co}(C_e) \subseteq \text{co}(C) = C$. Since C is compact, we have that $\overline{\text{co}(C_e)} \subseteq C = C$.

We now assert that $C_e \subseteq \text{co}(C_e)$. If the last assertion were not true, then there would exist an $\mathbf{x}_0 \notin \text{co}(C_e)$ and in C . Since $\text{co}(C_e)$ is closed and convex, it follows from Theorem 3.1 that there exists a hyperplane H_1^* such that for all $\mathbf{y} \in \text{co}(C_e)$

$$\langle \mathbf{a}, \mathbf{x} \rangle = \alpha = \langle \mathbf{a}, \mathbf{y}_0 \rangle, \quad (3)$$

Let

$$\beta = \max\{\langle \mathbf{a}, \mathbf{x} \rangle : \mathbf{x} \in C\}. \quad (4)$$

Since C is compact, the maximum is attained at some point $\mathbf{x}_0$ in C , and thus H_1^* is a supporting hyperplane to C at $\mathbf{x}_0$. By Theorem 4.2, H_1^* contains an extreme point $\mathbf{x}_1$ of C . Thus, $\mathbf{x}_1 \in H_1^* \cap \text{co}(C_e)$, and so by (3)

$$\beta = \langle \mathbf{a}, \mathbf{x}_1 \rangle = \alpha.$$

Since $\mathbf{x}_0 \in C$, it follows from (3) and (4) that

$$\alpha < \langle \mathbf{a}, \mathbf{x}_0 \rangle \leq \beta.$$

Combining the last two inequalities gives $\beta < \beta$. This contradiction shows that $C \subseteq \text{co}(C_e)$.

In more general contexts Theorem 4.3 is known as the *Krein-Milman theorem*. We shall see in Section 7 that in $\mathbb{R}^n$ the conclusion of Theorem 4.3 can be strengthened to $C = \text{co}(C_e)$.

We conclude this section with an example of a compact convex set whose set of extreme points is not closed. In $\mathbb{R}^2$ let

$$C_1 = \{\mathbf{x} : |x_1 - 1|^2 + x_2^2 \leq 1, x_1 = 0\},$$

$$C_2 = \{\mathbf{x} : x_1 = 0, x_2 = 0, -1 \leq x_3 \leq 1\},$$

and let $C = \text{co}(C_1 \cup C_2)$. Then C_e is the union of the points $(0, 0, 1)$, $(0, 0, -1)$ and all points other than the origin that are on the circumference of the circle in the plane $x_3 = 0$ whose center is at $(1, 0, 0)$ and whose radius is 1. The origin is a limit point of this set yet is not in the set.

Exercise 4.1. Let S be a convex set. Show that a point x_0 in S is an extreme point of S if and only if the set $\bar{S} = \{x_0\} = \{x \in \text{co } S : x \neq x_0\}$ is convex.

Exercise 4.2. Show by examples that Theorem 4.3 fails if the convex set C has a nonempty set of extreme points and is merely assumed to be closed or is merely assumed to be bounded rather than closed and bounded.

Exercise 4.3. Let S be a set in $\mathbb{R}^n$. Show that every point of S is either an interior point or a boundary point. Show that if S is compact, then the set of boundary points of S is not empty. Show that if x_0 is an extreme point of S , then it is a boundary point of S .

Exercise 4.4. Let C be a convex subset of $\mathbb{R}^n$ with nonempty interior. Show that $C_+ = \text{int}(C)$ is open.

Exercise 4.5. Show directly, without using Lemma 4.2, that the set of extreme points of a closed ball is the set of its boundary points. (Hint: To simplify notation, take the center of the ball to be the origin.)

Exercise 4.6. Let C be a closed convex set in $\mathbb{R}^n$ not equal to $\mathbb{R}^n$. Show that C is the intersection of all closed half spaces containing C that are determined by the supporting hyperplanes of C .

Exercise 4.7. Let x_0 be a point in $\mathbb{R}^n$. Let C_n denote the closed cube with center at x_0 and length of side equal to 2δ . Thus

$$C_n = \{x : |x_i - x_{0i}| \leq \delta, i = 1, \dots, n\}.$$

The vertices of the cube are the 2^n points of the form $(x_{01}, x_{02}, \dots, x_{0n})$ where the x_i take on the values ± 1 .

- Show that the set of 2^n vertices is the set of extreme points of C_n .
- Show directly, without using Theorem 4.3, that C_n is the convex hull of the set of vertices. (Hint: Use induction on n .)

5. SYSTEMS OF LINEAR INEQUALITIES: THEOREMS OF THE ALTERNATIVE

In this section we shall use the separation theorems to obtain so-called theorems of the alternative. Theorems of the alternative have the following form. Given two systems of inequalities, I and II, either system I has a solution or system II has a solution but never both. Necessary conditions for many optimization problems follow from these theorems.

The following definitions and lemma are needed in our discussion. A vector $\mathbf{x} = (x_1, \dots, x_n)$ in $\mathbb{R}^n$ is said to be nonnegative, or $\mathbf{x} \geq \mathbf{0}$, if, for every $j = 1, 2, \dots, n$, $x_j \geq 0$. A vector $\mathbf{x}$ in $\mathbb{R}^n$ is said to be positive, or $\mathbf{x} > \mathbf{0}$, if, for every $j = 1, 2, \dots, n$, $x_j > 0$.

Lemma 3.1. Let A be an $m \times n$ matrix and let

$$C = \{\mathbf{w} | \mathbf{w} = A\mathbf{x}, \mathbf{x} \in \mathbb{R}^n, \mathbf{x} \geq \mathbf{0}\}.$$

Then C is a closed convex cone.

Proof. A straightforward calculation verifies that C is a convex cone.

To show that C is closed, we first consider the case in which the columns of A are linearly independent. Let $\mathbf{w}_0$ be a limit point of C , and let $\{\mathbf{w}_k\}$ be a sequence of points in C converging to $\mathbf{w}_0$. Then there exists a sequence of points $\{\mathbf{x}_k\}$ in $\mathbb{R}^n$ with $\mathbf{x}_k \geq \mathbf{0}$ such that $A\mathbf{x}_k = \mathbf{w}_k$.

We show that $\{\mathbf{x}_k\}$ is bounded. If $\{\mathbf{x}_k\}$ were unbounded, there would exist a subsequence, again denoted by $\{\mathbf{x}_k\}$, such that $\|\mathbf{x}_k\| \rightarrow \infty$. All points in the sequence $\{\mathbf{x}_k/\|\mathbf{x}_k\|\}$ have norm 1, so there exists a subsequence and a point $\mathbf{u}_0$ of norm 1 such that $\mathbf{u}_k/\|\mathbf{x}_k\| \rightarrow \mathbf{u}_0$. From

$$\|\mathbf{x}_k\| A(\mathbf{x}_k/\|\mathbf{x}_k\|) = \mathbf{w}_k \quad \|\mathbf{x}_k\| \rightarrow \infty$$

we see that we must have $A(\mathbf{u}_0/\|\mathbf{u}_0\|) = \mathbf{0}$. On the other hand, $\mathbf{u}_0/\|\mathbf{u}_0\| = \mathbf{u}_0$ implies that $A\mathbf{u}_0 = \mathbf{0}$. Hence $A\mathbf{u}_0 = \mathbf{0}$. Recall that $\|\mathbf{u}_0\| = 1$, so that $\mathbf{u}_0 \neq \mathbf{0}$. Therefore the relation $A\mathbf{u}_0 = \mathbf{0}$ contradicts the linear independence of the columns of A .

Since $\{\mathbf{x}_k\}$ is bounded, $\{\mathbf{x}_k\}$ has a convergent subsequence, which we denote again by $\{\mathbf{x}_k\}$. Let $\mathbf{x}_0 = \lim \mathbf{x}_k$. Then $\mathbf{x}_0 \geq \mathbf{0}$. Also, $A\mathbf{x}_k \rightarrow A\mathbf{x}_0$, so that $A\mathbf{x}_0 = \mathbf{w}_0$. Thus $\mathbf{w}_0 \in C$, and C is closed.

If the columns of A are not linearly independent, we first show that any point $\mathbf{x}$ in C can be expressed as a nonnegative linear combination of $p < n$ linearly independent columns of A . Let A_j , $j = 1, \dots, n$, denote the j th column of A . Then there exists an $\mathbf{x} \neq \mathbf{0}$ such that

$$\mathbf{x} = \sum_{j=1}^n x_j A_j, \quad x_j \geq 0.$$

Since the columns of A are linearly dependent, there exists a $\mu = (\mu_1, \dots, \mu_p) \neq \mathbf{0}$ such that

$$\mu_1 A_1 + \mu_2 A_2 + \dots + \mu_p A_p = \mathbf{0}.$$

Therefore, for any $\mu \in \mathbb{R}$,

$$\mathbf{x} = \sum_{j=1}^n (x_j - \mu \mu_j) A_j = \sum_{j \neq 1}^n (x_j - \mu \mu_j) A_j + \sum_{j=1}^p x_j A_j.$$

Since $\mu \neq 0$, the first case on the right exists. If $\mu_j < 0$, then $x_j - \mu_j x_1 \geq 0$ whenever $\mu \geq x_1/x_1$. In particular, since $x_1 \geq 0$, we have $x_j - \mu_j x_1 \geq 0$ whenever $\mu \geq 0$. If $\mu_j > 0$, then $x_1/x_1 \geq 0$ and $x_j - \mu_j x_1 \geq 0$ if and only if $\mu \leq x_1/x_1$. If there exist indices j such that $\mu_j > 0$, set

$$\mu = \min \left\{ \frac{x_1}{\mu_j} : \mu_j > 0 \right\}.$$

If $\mu_j \leq 0$ for all indices j for which $\mu_j \neq 0$, set

$$\mu = \max \left\{ \frac{x_1}{\mu_j} : \mu_j \neq 0 \right\}.$$

Then $x_j - \mu_j x_1 \geq 0$ for all $j = 1, \dots, n$ and $x_j - \mu_j x_1 = 0$ for at least one value of j . We have now expressed x as a nonnegative linear combination of $q = n$ columns of A . We continue the process until we express x as a nonnegative linear combination of $p < n$ linearly independent columns of A .

Let x denote a choice of $p < n$ linearly independent columns of A and let $A(x)$ denote the corresponding $m \times p$ matrix. There are a finite number of such choices. Let

$$C(x) = \{w : w = A(x)y, w \in \mathbb{R}^m, y \in \mathbb{R}^p, y \geq 0\}.$$

We have shown that each set $C(x)$ is closed and that

$$C \subseteq \bigcup_x C(x).$$

We now show the opposite inclusion. Let $w \in C(x)$ for some x . By relabeling columns of A , we can assume without loss of generality that x selects the first p columns of A . Then there exists a $y \in \mathbb{R}^p, y \geq 0$ such that

$$w = \sum_{j=1}^p y_j A_j = \sum_{j=1}^p y_j A_j + \sum_{j=p+1}^n 0 A_j, \quad y_j \geq 0, \quad j = 1, \dots, p.$$

Let $x = (x_1, \dots, x_p, 0, \dots, 0)$. Then $w = Ax$, and so $w \in C$.

We have expressed C as the union of a finite number of closed sets. Hence C is closed and Lemma 3.1 is proved.

THEOREM 3.1 (FARKAS'S LEMMA). Let A be an $m \times n$ matrix and let b be a vector in $\mathbb{R}^m$. Then one and only one of the systems

$$\begin{aligned} \text{I. } Ax &\leq b, & \langle b, x \rangle &> 0, & x \in \mathbb{R}^n, \\ \text{II. } Ax &= b, & x &\geq 0, & x \in \mathbb{R}^n \end{aligned}$$

has a solution.

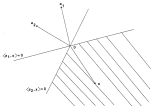


Figure 2.1

The conclusion of the theorem is sometimes stated in the logically equivalent form: Either I has a solution or II has a solution but never both. It can also be stated in the following form: System I has a solution if and only if system II has no solution.

Figure 2.1 illustrates Farkas's lemma when A is a 2×2 matrix with rows a_1 and a_2 . Any vector x in the acute-bounded region satisfies $\langle x, a_i \rangle \leq b$. Thus, if h does not lie in the acute angle determined by a_1 and a_2 , then I has a solution. (Draw the hyperplane $\langle h, x \rangle = 0$ for such a h .) If we write II in the equivalent form $y'A = b$, $y' \geq 0$, we see that for II to have a solution h must lie in the angle determined by a_1, a_2 . The vector h is given, so its position is fixed. It either lies in the angle or does not. Thus, precisely one of the systems I and II has a solution.

We now prove the theorem. The reader will get an appreciation of the efficiency and subtle appeal of convexity arguments by looking up the original proof [Farkas, 1954].

We first suppose that II has a solution y_0 . We shall show that I has no solution. If $\Delta x > 0$ for all x in $\mathbb{R}^n$, then I has no solution. Suppose there exists an u_0 in $\mathbb{R}^n$ such that $\Delta u_0 \leq 0$. Then

$$\langle h, u_0 \rangle = b'y_0 = (A'y_0)'u_0 = f_0'u_0 = \langle f_0, \Delta u_0 \rangle.$$

Since $p_j > 0$ and $\lambda x_j < 0$, it follows that $\langle \hat{b}_j, x_j \rangle < 0$. Hence I has no solution.

We now suppose that II has no solution. Let

$$C = \{x \in \mathcal{A}(p, y) \in \mathbb{R}^n, y \geq 0, x \in \mathbb{R}_+^n\}.$$

The statement that II has no solution is equivalent to the statement that $\hat{b} \notin C$. By Lemma 5.1, C is closed and convex. Therefore, by Theorem 3.1 there exists a hyperplane H_C^0 that strongly separates $\hat{b}$ and C . Thus $\alpha_0 \neq 0$ and

$$\langle \alpha_0, \hat{b} \rangle > \alpha > \langle \alpha_0, x \rangle$$

for all x in C . Since $\hat{b} \in C$, we get that $\alpha > 0$. Thus

$$\langle \alpha_0, \hat{b} \rangle > 0. \quad (1)$$

For all x in C we have

$$\langle \alpha_0, \hat{b} \rangle = \langle \alpha_0, \mathcal{A}(p) \rangle = p' \lambda x_0 = \alpha,$$

the last inequality now holding for all y in $\mathbb{R}^n$ satisfying $y \geq 0$.

We claim that the last inequality implies that

$$\lambda x_0 \leq 0. \quad (2)$$

If (2) were false, there would exist at least one component of λx_0 , say the j th component, that is positive. Let i_j denote this component. Let $y_j = \lambda x_j = \lambda p_1, \dots, \lambda p_n, 0, \dots, 0$, where $\lambda > 0$. Then $\langle y_j, \lambda x_0 \rangle = \lambda i_j < \alpha$. Since $i_j > 0$, if we take λ to be sufficiently large, we get a contradiction. Thus (2) is true. But (1) and (2) say that α_0 is a solution of I.

THEOREM 5.2 [GONZALEZ 1813]. *Let A be an $m \times n$ matrix. Then one and only one of the following systems has a solution:*

- I. $\lambda x = 0, \quad x \in \mathbb{R}_+^n$
 II. $\lambda' y = 0, \quad y \in \mathbb{R}^m, \quad y \neq 0, \quad y \geq 0.$

The reader should draw a figure which illustrates this theorem when A is a 3×2 matrix with rows a_1, a_2, a_3 .

Proof. We shall show that the system I has a solution if and only if the system

$$I': \lambda x \leq 0, \quad \sum x_i = 1$$

has a solution, where e is the vector in $\mathbb{R}^n$ all of whose entries equal 1. Let

(x_0, ζ_0) be a solution of Γ . Then

$$\lambda x_0 + \zeta_0 e \leq 0,$$

so I has a solution. Now let x_0 be a solution of I and let $\zeta = \lambda x_0$. If $\zeta = (\zeta_1, \dots, \zeta_m)$, then $\zeta_i \leq 0$ for $i = 1, \dots, m$. Let $\zeta_0 = \max\{\zeta_1, \dots, \zeta_m\}$. Then (x_0, ζ_0) is a solution of Γ .

The system Γ has a solution if and only if the system

$$IV: \quad (A_i - e) \begin{pmatrix} x \\ \zeta \end{pmatrix} \leq 0, \quad (B_m - I)(x, \zeta) \geq 0,$$

where 0 is the zero vector in $\mathbb{R}^n$, has a solution. By Farkas's lemma, either Γ (and hence I) has a solution or

$$IV: \quad \begin{pmatrix} A' \\ -e' \end{pmatrix} y = \begin{pmatrix} B \\ -1 \end{pmatrix}, \quad y \geq 0, \quad y \in \mathbb{R}^m,$$

has a solution, but never both. System IV can be written as

$$IV: \quad A'y = B, \quad (a, y) = 1, \quad y \geq 0.$$

From $(a, y) = 1$ it follows that $y \neq 0$. Hence if IV has a solution, so does

$$II: \quad A'y = B, \quad y \geq 0, \quad y \neq 0.$$

If II has a solution, then $(a, y) \neq 0$. If we set $y' = y/(a, y)$, then IV will have a solution. Hence either I has a solution or II has a solution, but never both.

Theorem 5.3 [MORSE 1966]. Let A, B, C be given matrices with a column and let $A \neq 0$. Then one and only one of the following systems has a solution:

$$I: \quad Ax \geq 0, \quad Bx \geq 0, \quad Cx = 0, \quad x \in \mathbb{R}^n,$$

$$II: \quad A'y_1 + B'y_2 + C'y_3 = 0, \quad y_1 \geq 0, \quad y_2 \neq 0, \quad y_3 \geq 0.$$

Proof. Let $\alpha = (1, 1, 1, \dots, 1)$ be a vector all of whose entries equal 1 and whose dimension is the row dimension of A . We assert that the system I has a solution if and only if the system Γ has a solution, where

$$\Gamma: \quad Ax \geq \zeta \alpha, \quad Bx \geq 0, \quad Cx = 0, \quad \zeta > 0.$$

Here ζ is a real scalar. It is clear that if Γ has a solution, then I has a solution. If I has a solution x_0 let $\xi = (\xi_1, \dots, \xi_m) = Ax_0$. Let $\zeta = \min\{\xi_1, \dots, \xi_m\}$. Then $\zeta > 0$ and $\lambda x_0 \geq \zeta \alpha$ and so Γ has a solution.

Since I has a solution if and only if $\bar{I}$ has a solution, $\bar{I}$ has no solution if and only if $\bar{I}$ has no solution. Now $\bar{I}$ has no solution if and only if the system

$$I: \begin{pmatrix} -A & B \\ -B & B \\ -C & B \\ C & B \end{pmatrix} \begin{pmatrix} x \\ y \end{pmatrix} \leq \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}, \quad \langle \bar{b}, 1 \rangle \langle \bar{a}, 1 \rangle > 0,$$

where $\bar{b}$ is the zero vector in R , has no solution.

By Farkas's lemma $\bar{I}$ has no solution if and only if the system II' has a solution, where

$$II': \begin{pmatrix} -A' & -B' & -C' & C' \\ A' & B' & B' & 0 \end{pmatrix} \begin{pmatrix} p_1 \\ p_2 \\ q_1 \\ q_2 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 1 \\ 1 \end{pmatrix}, \quad \langle q_2, p_2, q_1, q_2 \rangle \geq 0.$$

But II' having a solution is equivalent to the system

$$II: -Ap_1 + Bp_2 + C(q_1 - q_2) = 0, \quad \langle \bar{a}, p_2 \rangle = 1,$$

having a solution $p_1 \geq 0$, $p_2 \geq 0$, $q_1 \geq 0$, $q_2 \geq 0$. If we let $p_2 = (q_1 - q_2)$, we see that if II' has a solution as indicated, then

$$III: -Ap_1 + Bp_2 + Cq_2 = 0, \quad p_1 \geq 0, \quad p_2 \geq 0, \quad p_2 \geq 0,$$

has a solution. In summary, we have shown that if I has no solution, then II has a solution.

Now suppose that II has a solution (p_1, p_2, q_2) . Then $\langle \bar{a}, p_1 \rangle > 0$. Hence if we set

$$p'_1 = \frac{p_1}{\langle \bar{a}, p_1 \rangle}, \quad p'_2 = \frac{p_2}{\langle \bar{a}, p_1 \rangle}, \quad q'_2 = \frac{q_2}{\langle \bar{a}, p_1 \rangle},$$

then (p'_1, p'_2, q'_2) is a solution of II with $\langle \bar{a}, p'_1 \rangle = 1$. Thus we may assume that II has a solution (p_1, p_2, q_2) with $\langle \bar{a}, p_1 \rangle = 1$.

Let $\bar{q}_1$ be a vector whose dimension is that of p_1 and whose i th component is equal to p_{1i} , the i th component of p_1 , if $p_{1i} > 0$ and zero otherwise. Let $\bar{q}_2$ be a vector whose dimension is that of p_2 and whose i th component is equal to p_{2i} if $p_{2i} > 0$ and zero otherwise. Then $\bar{q}_1 \geq 0$, $\bar{q}_2 \geq 0$, and $p_1 = (\bar{q}_1 - \bar{q}_2)$. Hence if II has a solution, then so does III . Therefore, I has no solution.

Remark 3.1. Theorem 5.2 is valid if either B or C or both are zero. In fact, if both B and C are zero, the lemma becomes Gordan's theorem.

THEOREM 5.4 [Gale 1962]. Let A be an $m \times n$ matrix and let $a \neq 0$ be a vector in $\mathbb{R}^m$. Then one and only one of the following systems has a solution:

$$\begin{aligned} \text{I: } & Ax \leq a, \\ \text{II: } & A^T y = 0, \quad (a, y) = -1, \quad y \geq 0 \end{aligned}$$

Proof. The system $Ax \leq a$ has a solution if and only if the system $\xi \geq 0$, $A\xi \leq 0$ has a solution. Thus the system I is equivalent to the system

$$\text{I': } (A, -a) \begin{pmatrix} x \\ \xi \end{pmatrix} \leq 0, \quad (0, 1, 0, \xi) \geq 0,$$

where 0 is the zero vector in $\mathbb{R}^n$. By Farkas's lemma I' has a solution if and only if the system II'

$$\begin{pmatrix} A^T \\ -a \end{pmatrix} y' = \begin{pmatrix} 0 \\ 1 \end{pmatrix}, \quad y \geq 0, y \in \mathbb{R}^m$$

has no solution. But system II' is the same as

$$(a, y) = -1, \quad A^T y = 0, \quad y \geq 0,$$

which is system (II).

COROLLARY 1. Let A be an $m \times n$ matrix and let $a \neq 0$ be in $\mathbb{R}^m$ for each that, for all $x \in \mathbb{R}^n$, $Ax - a \geq 0$. Then there exists a vector $y \in \mathbb{R}^m$ such that $y \geq 0$, $A^T y = 0$, and $(a, y) = -1$.

To prove the corollary, we need only note that if $Ax - a \geq 0$ for all x in $\mathbb{R}^n$, then $Ax \leq a$ has no solution in $\mathbb{R}^n$.

EXERCISE 5.1. Show that the set C of Lemma 5.1 is a convex cone.

EXERCISE 5.2. Let A be an $m \times n$ matrix and b a vector in $\mathbb{R}^m$. Show that one and only one of the following systems has a solution:

$$\begin{aligned} \text{I: } & Ax \geq b, \quad x \geq 0, \quad (b, u) \geq 0 \\ \text{II: } & A^T y \geq b, \quad y \leq 0 \end{aligned}$$

EXERCISE 5.3. Let A be an $m \times n$ matrix and let $a \neq 0$ be a vector in $\mathbb{R}^m$. Show that one and only one of the following systems has a solution:

$$\begin{aligned} \text{I: } & Ax = a, \quad x \in \mathbb{R}^n, \\ \text{II: } & A^T y = 0, \quad (a, y) = 1. \end{aligned}$$

4. AFFINE GEOMETRY

Consider a plane Π in $\mathbb{R}^3$ and a line A in $\mathbb{R}^3$. Let C be a subset of Π and let S be a subset of A . Although C and S are subsets of $\mathbb{R}^3$, it is natural to consider C as a two-dimensional object and S as a one-dimensional object. In this section we shall make these ideas precise for subsets of $\mathbb{R}^n$.

Let V_r denote an r -dimensional vector subspace of $\mathbb{R}^n$. When we do not wish to emphasize the dimension of the subspace, we shall write V . The reader will recall from linear algebra that an $(n - 1)$ -dimensional subspace V_{n-1} of $\mathbb{R}^n$ can be expressed as

$$V_{n-1} = \left\{ \mathbf{x} : \sum_{i=1}^n a_i x_i = 0 \right\} = \{ \mathbf{x} : (\mathbf{a}, \mathbf{x}) = 0 \},$$

where $\mathbf{a}$ is a nonzero vector. Thus V_{n-1} is a hyperplane through the origin.

Also, given a subspace V_{n-k} , there exist k linearly independent vectors $\mathbf{a}_1, \dots, \mathbf{a}_k$ such that

$$V_{n-k} = \{ \mathbf{x} : (\mathbf{a}_i, \mathbf{x}) = 0, i = 1, \dots, k \}.$$

Thus V_{n-k} , a subspace of dimension $n - k$, is the intersection of k hyperplanes through the origin.

Definition 4.1. A linear manifold M is a subset of $\mathbb{R}^n$ of the form

$$M = V + \mathbf{h} = \{ \mathbf{x} : \mathbf{x} = \mathbf{v} + \mathbf{h}, \mathbf{v} \in V \},$$

where V is a subspace and $\mathbf{h}$ is a vector in $\mathbb{R}^n$.

Note that since $\mathbf{h} = \mathbf{0} + \mathbf{h}$ and $\mathbf{0} \in V$, it follows that $\mathbf{h} \in M$.

Thus a linear manifold is a translate of a subspace. Since subspaces are closed sets in $\mathbb{R}^n$, linear manifolds are closed sets. Note that $\mathbb{R}^n$ is a linear manifold, since $\mathbb{R}^n = \mathbb{R}^n + \mathbf{0}$, and that a subspace V is a linear manifold, since $V = V + \mathbf{0}$. In Section 1 we saw that a hyperplane is the translate of a hyperplane through the origin and thus the translate of an $(n - 1)$ -dimensional vector space. Therefore, hyperplanes are linear manifolds. In $\mathbb{R}^3$ the linear manifolds are planes, lines, and points. (A point is the translate of the zero-dimensional subspace consisting of the origin.)

Linear manifolds are also called *flat* or *affine varieties* or *affine sets*.

Lemma 4.1. Let M be a linear manifold, $M = V + \mathbf{h}$. Then the subspace V is unique.

To prove the lemma, we suppose that M' has another representation $M' = V' + \mathbf{k}'$ and show that $V' = V$. Since $\mathbf{k}_0 \in M$, from $M' = V' + \mathbf{k}'$, we have that $\mathbf{k} = \mathbf{v}' + \mathbf{k}'$ for some $\mathbf{v}' \in V'$. Hence $\mathbf{k}' = \mathbf{k} - \mathbf{v}' \in V'$. Now let $\mathbf{x} \in V$. Then $\mathbf{x} = \mathbf{m} - \mathbf{k}$ for some $\mathbf{m} \in M$. But $\mathbf{m} = \mathbf{y} + \mathbf{k}$ for some $\mathbf{y} \in V'$, and so

$$\mathbf{x} = (\mathbf{y} + \mathbf{k}') - \mathbf{k} = \mathbf{y} + (\mathbf{k}' - \mathbf{k}) \in V'.$$

Thus $V \subseteq V'$. A similar argument shows that $V' \subseteq V$.

If $M = V + \mathbf{k}$, then the vector $\mathbf{k}$ is not unique. In fact, $M = V + \mathbf{k}'$ for any $\mathbf{k}' \in M$. To see this, let $\mathbf{k}'$ be any vector in M and let $M' = V' + \mathbf{k}'$. Let $\mathbf{m} \in M$, so $\mathbf{m} = \mathbf{v} + \mathbf{k}$ for some $\mathbf{v} \in V$. Since $\mathbf{k}' \in M$, we have $\mathbf{k}' = \mathbf{v}' + \mathbf{k}$ for some $\mathbf{v}' \in V$. Hence $\mathbf{m} = (\mathbf{v} - \mathbf{v}') + \mathbf{k}'$, which is an element of M' . Thus, $M' = M$. A similar argument shows that $M' \subseteq M$.

On the other hand, if $\mathbf{k}$ is fixed and $M = V + \mathbf{k}$, then the representation of an element $\mathbf{x}$ in M as $\mathbf{x} = \mathbf{v} + \mathbf{k}$ is unique. For if $\mathbf{x} = \mathbf{v}' + \mathbf{k}$, then $\mathbf{v} = \mathbf{x} - \mathbf{k}$ and $\mathbf{v}' = \mathbf{x} - \mathbf{k}$.

If M is a linear manifold, then the unique subspace V such that $M = V + \mathbf{k}$ for some $\mathbf{k}$ is called the *parallel subspace* of M .

Definition 4.2. The *dimension* of a linear manifold M is defined to be the dimension of the parallel subspace V .

If M_1 and M_2 are two linear manifolds with the same parallel subspace, then M_1 and M_2 are said to be *parallel*.

The next lemma characterizes linear manifolds as sets having the property that if $\mathbf{x}_1$ and $\mathbf{x}_2$ belong to the set, then the entire line determined by these points lies in the set.

Lemma 4.1. A set M is a linear manifold if and only if, for every $\mathbf{x}_1, \mathbf{x}_2$ in M and λ_1, λ_2 real with $\lambda_1 + \lambda_2 = 1$, the point $\lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2$ belongs to M .

Proof. Let M be a linear manifold, $M = V + \mathbf{k}$. Let $\mathbf{x}_1$ and $\mathbf{x}_2$ be in M . Then $\mathbf{x}_1 = \mathbf{v}_1 + \mathbf{k}$ and $\mathbf{x}_2 = \mathbf{v}_2 + \mathbf{k}$, where $\mathbf{v}_1$ and $\mathbf{v}_2$ are in V . Hence for λ_1, λ_2 with $\lambda_1 + \lambda_2 = 1$,

$$\lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2 = \lambda_1 (\mathbf{v}_1 + \mathbf{k}) + \lambda_2 (\mathbf{v}_2 + \mathbf{k}) = (\lambda_1 \mathbf{v}_1 + \lambda_2 \mathbf{v}_2) + \mathbf{k}.$$

Since V is a subspace, $\lambda_1 \mathbf{v}_1 + \lambda_2 \mathbf{v}_2 \in V$, and so $\lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2 \in M$.

We now prove the reverse implication. Let M be a set such that if $\mathbf{x}_1$ and $\mathbf{x}_2$ are in M and λ_1, λ_2 are scalars such that $\lambda_1 + \lambda_2 = 1$, then $\lambda_1 \mathbf{x}_1 + \lambda_2 \mathbf{x}_2$ is in M . Let $\mathbf{x}_0 \in M$ and let $M' = M - \mathbf{x}_0$. We shall show that M' is a subspace, and then show that M is a linear manifold, since $M = M' + \mathbf{x}_0$. Let $\mathbf{y}_1$ and $\mathbf{y}_2$

belong to W and let α and β be arbitrary scalars. Then

$$\alpha y_1 + \beta y_2 + x_3 = \{\alpha(y_1 + x_3) + (1 - \alpha)x_3\} + \{\beta(y_2 + x_3) + (1 - \beta)x_3\} = x_3.$$

Since $y_1 + x_3$ and $y_2 + x_3$ belong to M , each of the expressions in square brackets belongs to M . If we denote the expression in the first bracket by m_1 and the expression in the second bracket by m_2 , we may write

$$\alpha y_1 + \beta y_2 + x_3 = 1(m_1 + m_2) = x_3 = 0(m_1 + m_2)$$

where m_1 denotes the expression in curly braces. Since m_1 and x_3 are in M and the sum of the coefficients on the extreme right is 1, we get that $\alpha y_1 + \beta y_2 + x_3 \in M$. Thus, $\alpha y_1 + \beta y_2 \in W$ and so W is a subspace.

COROLLARY 1. *If M_1 and M_2 are linear manifolds, then so is $M_1 + M_2 = \{x + y : x \in M_1 + M_2, y \in M_1, z \in M_2\}$.*

COROLLARY 2. *If $\{M_\alpha\}$ is a family of linear manifolds, then so is $\bigcup_\alpha M_\alpha$.*

We leave the proofs of these corollaries as exercises for the reader.

Definition 4.3. A point x is said to be an *affine combination* of points $x_1, \dots, x_k$ if there exists a vector $q = (q_1, \dots, q_k)$ in $\mathbb{R}^k$ such that

$$x = q_1 x_1 + \dots + q_k x_k \quad \text{and} \quad \sum_{i=1}^k q_i = 1.$$

COROLLARY 3. *A set X is a linear manifold if and only if every affine combination of points in X is in X .*

The proof is similar to that of Lemma 2.6.

Definition 4.4. A finite set of points $x_1, \dots, x_k$ is *affinely dependent* if there exist real numbers $\lambda_1, \dots, \lambda_k$, not all zero such that $\lambda_1 + \dots + \lambda_k = 0$ and $\lambda_1 x_1 + \dots + \lambda_k x_k = 0$. If the points $x_1, \dots, x_k$ are not affinely dependent, then they are *affinely independent*.

Lemma 4.5. *A set of points $x_1, x_2, \dots, x_k$ is affinely dependent if and only if the vectors $(x_2 - x_1), (x_3 - x_1), \dots, (x_k - x_1)$ are linearly dependent.*

Proof. Let $x_1, \dots, x_k$ be affinely dependent. Then there exist real numbers $\lambda_1, \dots, \lambda_k$ not all zero such that

$$\lambda_1 x_1 + \lambda_2 x_2 + \dots + \lambda_k x_k = 0, \quad \sum_{i=1}^k \lambda_i = 0. \quad (1)$$

Hence

$$\lambda_1 = -(\lambda_2 + \cdots + \lambda_r). \quad (2)$$

Not all of $\lambda_2, \dots, \lambda_r$ are zero, for if they were, λ_1 would also be zero, contradicting the assumption that not all of $\lambda_1, \dots, \lambda_r$ are zero. Substituting (2) into the first relation in (1) gives

$$\lambda_2(\mathbf{x}_1 - \mathbf{x}_2) + \lambda_3(\mathbf{x}_1 - \mathbf{x}_3) + \cdots + \lambda_r(\mathbf{x}_1 - \mathbf{x}_r) = \mathbf{0},$$

where $\lambda_2, \dots, \lambda_r$ are not all zero. Thus the vectors $(\mathbf{x}_1 - \mathbf{x}_2), \dots, (\mathbf{x}_1 - \mathbf{x}_r)$ are linearly dependent.

To prove the reverse implication, we reverse the argument.

Remark 6.3. Clearly, an equivalent statement of Lemma 6.3 is the following. A set of points $\mathbf{x}_1, \dots, \mathbf{x}_r$ is affinely independent if and only if the vectors $\mathbf{x}_2 - \mathbf{x}_1, \dots, \mathbf{x}_r - \mathbf{x}_1$ are linearly independent.

Lemma 6.4. Let M be a linear manifold. Then the following statements are equivalent:

- The dimension of M equals r .
- There exist $r+1$ points $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_r$ in M that are affinely independent and any set of $r+2$ points in M is affinely dependent.
- There exist points $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_r$ that are affinely independent and such that each point $\mathbf{x}$ in M has a unique representation

$$\mathbf{x} = \sum_{i=1}^r \lambda_i \mathbf{x}_i, \quad \sum_{i=1}^r \lambda_i = 1 \quad (3)$$

in terms of the $\mathbf{x}_i$'s.

Proof. We first show that (i) implies (ii). If (i) holds, then we may write $M = V + \mathbf{x}_0$, where V is a subspace of dimension r and $\mathbf{x}_0$ is an arbitrary point in M . Let $\mathbf{v}_1, \dots, \mathbf{v}_r$ be a basis for V and let

$$\mathbf{x}_i = \mathbf{x}_0 + \mathbf{v}_i, \quad i = 1, \dots, r.$$

Then the vectors $\mathbf{x}_i - \mathbf{x}_0$, $i = 1, \dots, r$, are linearly independent and the $r+1$ points $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_r$ are affinely independent.

Let $\mathbf{y}_0, \mathbf{y}_1, \dots, \mathbf{y}_{r+1}$ be any set of $r+2$ points in M . If we write $M = V + \mathbf{x}_0$, then each of the vectors $\mathbf{y}_i - \mathbf{x}_0$, $i = 1, \dots, r+1$, is in V . Therefore the $r+1$ vectors in this set are linearly dependent, and hence $\mathbf{y}_0, \mathbf{y}_1, \dots, \mathbf{y}_{r+1}$ are affinely dependent.

Next we show that (ii) implies (iii). Since M is a linear manifold and $\mathbf{x}_0 \in M$, we may write $M = V + \mathbf{x}_0$, where V is a subspace. We assert that V has dimension r . To see this, let

$$\mathbf{v}_i = \mathbf{x}_i - \mathbf{x}_0, \quad i = 1, \dots, r.$$

Since the points $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_r$ are affinely independent, the vectors $\mathbf{v}_1, \dots, \mathbf{v}_r$ are linearly independent. Now let $\mathbf{w}_1, \dots, \mathbf{w}_{r+1}$ be a set of $r+1$ distinct nonzero vectors in V . Then there exist points $\mathbf{p}_1, \dots, \mathbf{p}_{r+1}$ in M such that

$$\mathbf{w}_j = \mathbf{p}_j - \mathbf{x}_0, \quad j = 1, \dots, r+1.$$

Since the $\mathbf{w}_j$ are distinct and none of the vectors $\mathbf{w}_j$ is zero, the $r+2$ points $\mathbf{x}_0, \mathbf{p}_1, \dots, \mathbf{p}_{r+1}$ are distinct. By hypothesis, they are affinely dependent. Hence the vectors $\mathbf{w}_1, \dots, \mathbf{w}_{r+1}$ are linearly dependent. Thus, V has dimension r ; we may write $V = U$, and the vectors $\mathbf{v}_1, \dots, \mathbf{v}_r$ are a basis for U .

For arbitrary $\mathbf{x}$ in M we have $\mathbf{x} = \mathbf{x}_0 + \mathbf{u}$, uniquely for some $\mathbf{u}$ in U . Hence there exist unique real numbers $\lambda_1, \dots, \lambda_r$ such that

$$\mathbf{x} = \sum_{i=1}^r \lambda_i \mathbf{v}_i + \mathbf{x}_0 = \sum_{i=1}^r \lambda_i (\mathbf{x}_i - \mathbf{x}_0) + \mathbf{x}_0 = \left(1 - \sum_{i=1}^r \lambda_i\right) \mathbf{x}_0 + \sum_{i=1}^r \lambda_i \mathbf{x}_i.$$

We obtain the representation (1) by setting $\lambda_0 = (1 - \lambda_1 - \dots - \lambda_r)$.

To show that the representation is unique in terms of the $\mathbf{x}_i$'s, we suppose that there were another representation

$$\mathbf{x} = \sum_{i=0}^r \lambda'_i \mathbf{x}_i, \quad \sum_{i=0}^r \lambda'_i = 1.$$

Then

$$0 = \sum_{i=1}^r (\lambda'_i - \lambda_i) \mathbf{x}_i, \quad \sum_{i=1}^r (\lambda'_i - \lambda_i) = 0$$

with not all of the coefficients $(\lambda'_i - \lambda_i)$ equal to zero. But then the points $\mathbf{x}_i$, $i = 1, \dots, r$, would be affinely dependent, contrary to assumption.

We conclude the proof by showing that (iii) implies (i). Let $\mathbf{x}$ be an arbitrary point in M . Then by virtue of (1)

$$\mathbf{x} = \sum_{i=1}^r \lambda_i \mathbf{x}_i + \left(1 - \sum_{i=1}^r \lambda_i\right) \mathbf{x}_0 + \sum_{i=1}^r \lambda_i \mathbf{x}_i = \sum_{i=1}^r \lambda_i (\mathbf{x}_i - \mathbf{x}_0) + \mathbf{x}_0. \quad (2)$$

Since $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_r$ are affinely independent, the vectors

$$\mathbf{u}_1 = \mathbf{x}_1 - \mathbf{x}_0, \mathbf{u}_2 = \mathbf{x}_2 - \mathbf{x}_0, \dots, \mathbf{u}_r = \mathbf{x}_r - \mathbf{x}_0$$

are linearly independent.

Let V_1 denote the r -dimensional vector space spanned by $x_1, \dots, x_r = x_p$. It follows from (4) that $x_{p+1} \in V_1 + N_p$. Thus $W \subseteq V_1 + N_p$.

To prove the opposite inclusion, we let $W' = W - N_p$. We saw in the proof of Lemma 6.2 that W' is a subspace. Moreover, since $x_1 = x_p, \dots, x_r = x_p$ belong to W we have that $V_1 \subseteq W'$. Hence

$$V_1 + N_p \subseteq W' + N_p = W$$

and so $W = V_1 + N_p$. This concludes the proof of Lemma 6.4.

Points $x_0, x_1, \dots, x_r$ as in Lemma 6.4 are said to be an *affine basis* for M . The numbers $\lambda_i, i = 0, 1, \dots, r$, are called the *barycentric coordinates* of x relative to the affine basis $x_0, x_1, \dots, x_r$.

Definition 6.5. Let A be a subset of $\mathbb{R}^n$. The *affine hull* of A , denoted by $A[A]$, is defined to be the intersection of all linear manifolds containing A . The *dimension* of A is the dimension of $A[A]$.

Since $A \subset \mathbb{R}^n$, since $\mathbb{R}^n$ is a linear manifold, and since the intersection of linear manifolds is a linear manifold, it follows that, for any nonempty set A , $A[A]$ is not empty. Thus the dimension of A is well defined.

For any linear manifold M , the relations $M = A[M]$ and $A[M] \subset M$ hold, so $M = A[M]$. Thus, for a linear manifold M , the dimension of M as defined in Definition 6.2 coincides with that of Definition 6.5.

We shall often write $\dim A$ to refer to the dimension of A .

Lemma 6.6. Let A be a subset of $\mathbb{R}^n$ and let $M(A)$ denote the set of all affine combinations of points in A . Then

- $A[A] = M(A)$ and
- if there exist $p+1$ points $x_0, x_1, \dots, x_p$ in A that are affinely independent and if any set of $p+2$ points in A is affinely dependent, then $\dim A = p$.

Remark 6.1. Statement (i) is the analog for affine hulls of Theorem 2.1 for convex hulls.

Proof. It is readily verified that if x and p are two points in $M(A)$, then each affine combination of these points is in $M(A)$. Hence, by Lemma 6.2, $M(A)$ is a linear manifold. Since $M(A) \supset A$, it follows from the definition of $A[A]$ that $A[A] \subset M(A)$. On the other hand, it follows from Corollary 3 to Lemma 6.2 that $M(A)$ is contained in every linear manifold containing A . Thus $A[A] = M(A)$. Hence $A[A] = M(A)$.

We now establish (ii). Let $x_0, x_1, \dots, x_p$ be $p+1$ affinely independent points in A and suppose that any set of $p+2$ points is affinely dependent. Since the points $x_0, x_1, \dots, x_p$ are in $A[A]$, it follows from Lemma 6.4 that $\dim A[A] \geq p$.

On the other hand, for arbitrary x in A the vectors

$$x_1 - x_0, x_2 - x_0, \dots, x_p - x_0, x - x_0$$

are linearly dependent, and the first p vectors are linearly independent. Hence for arbitrary x in A there exist unique scalars $\mu_1, \dots, \mu_p$ such that

$$x - x_0 = \sum_{j=1}^p \mu_j (x_j - x_0).$$

On setting $\mu_0 = 1 - (\mu_1 + \mu_2 + \dots + \mu_p)$, we get, uniquely,

$$x = \sum_{j=0}^p \mu_j x_j, \quad \sum_{j=0}^p \mu_j = 1.$$

Hence A is contained in the linear manifold $N[x_0, x_1, \dots, x_p]$ and so $\dim A \leq \dim N[x_0, x_1, \dots, x_p]$.

To complete the proof of (ii), it suffices to show that $N[x_0, x_1, \dots, x_p]$ has dimension p . For then $\dim A \leq p$, and since we have already shown that $\dim A \geq p$, statement (ii) will follow.

We now show that $\dim N[x_0, x_1, \dots, x_p] = p$. Since $x_0, x_1, \dots, x_p$ are affinely independent to show that $\dim N[x_0, x_1, \dots, x_p] = p$, we must show that any set $x_0, x_1, \dots, x_p, x_{p+1}$ of $p+2$ vectors in $N[x_0, x_1, \dots, x_p]$ is affinely dependent. From (i), we have that $N[x_0, \dots, x_p] = A[x_0, \dots, x_p]$. Thus for $j = 0, 1, \dots, p+1$ we have

$$x_j = \sum_{i=0}^p q_{ij} x_i, \quad \sum_{i=0}^p q_{ij} = 1.$$

Since $q_{00} = 1 - q_{01} - q_{02} - \dots - q_{0p}$, we have

$$x_j = \sum_{i=1}^p q_{ij} (x_i - x_0) + x_0, \quad j = 0, 1, 2, \dots, p+1.$$

Hence,

$$x_j - x_0 = \sum_{i=1}^p (q_{ij} - q_{0i})(x_i - x_0), \quad j = 1, \dots, p+1. \quad (5)$$

Thus, the $p+1$ vectors $x_1 - x_0$ in (5) are in the vector space spanned by the p linearly independent vectors $x_1 - x_0, \dots, x_p - x_0$. Hence the vectors $x_1 - x_0$ in (5) are linearly dependent, and the corresponding points $x_0, x_1, \dots, x_{p+1}$ are affinely dependent.

COROLLARY 4. The points $x_0, x_1, \dots, x_n$ are an affine basis for $A[X]$.

Proof. In the proof of Lemma 6.5 we showed that $A \subset A[x_0, x_1, \dots, x_n]$. Hence $A[A] \subset A[x_0, x_1, \dots, x_n]$. Since the opposite inclusion clearly holds, we have $A[A] = A[x_0, x_1, \dots, x_n]$. The corollary then follows from the uniqueness of the representation (ii) of a point x in $A[x_0, x_1, \dots, x_n]$.

Let $x_0, x_1, \dots, x_n$ be affinely independent. We call $\text{co}[x_0, x_1, \dots, x_n]$ the simplex spanned by $x_0, x_1, \dots, x_n$. The points $x_0, x_1, \dots, x_n$ are called the vertices of the simplex.

We leave it as an exercise for the reader to show that any $n+2$ points in the simplex are affinely dependent. Thus the dimension of the simplex $\text{co}[x_0, x_1, \dots, x_n]$ is n .

COROLLARY 5. The dimension of a convex set C is the maximum dimension of simplices in C .

This corollary is a restatement of Lemma 6.5 when C is a convex set.

Intuitively, in $\mathbb{R}^2$ one would expect that a convex set that is not contained in a plane or in a line to have nonempty interior. In terms of our definition of dimension of a set, we would expect a three-dimensional convex set to have nonempty interior. This is indeed the case, not only in $\mathbb{R}^3$, but also in $\mathbb{R}^n$, as the following theorem shows.

THEOREM 6.1. A convex set C in $\mathbb{R}^n$ has dimension n if and only if C has nonempty interior.

Suppose C has nonempty interior. Let x_0 be an interior point of C and let $e_1, \dots, e_n$ denote the standard basis vectors in $\mathbb{R}^n$. Then there exists an $\epsilon > 0$ such that each of the $n+1$ points $x_0, x_0 + \epsilon e_1, \dots, x_0 + \epsilon e_n$ lies in C . These points are affinely independent. Since in $\mathbb{R}^n$ any $n+2$ points are affinely dependent, it follows from (ii) of Lemma 6.5 that C has dimension n .

Now suppose that C has dimension n . Then $A[C] = \mathbb{R}^n$, and by Corollary 4 of Lemma 6.5 there exist points $x_0, x_1, \dots, x_n$ in C that form an affine basis for $A[C] = \mathbb{R}^n$. Thus each x in $\mathbb{R}^n$ can be written uniquely as

$$x = \sum_{i=0}^n \lambda_i x_i, \quad \sum_{i=0}^n \lambda_i = 1. \quad (6)$$

Since C is convex, all points x in $\mathbb{R}^n$ with $\lambda_i \geq 0$, $i = 0, 1, \dots, n$, lie in C . In particular,

$$x_0 = \frac{1}{n+1}(x_0 + x_1 + \dots + x_n)$$

the centroid of the simplex $\text{co}[x_0, x_1, \dots, x_n]$, is in C .

We can also write $\mathbf{x}$ uniquely as

$$\mathbf{x} = \sum_{i=1}^n \lambda_i (\mathbf{u}_i - \mathbf{u}_0),$$

with the λ_i as in (6). Since the vectors $\mathbf{u}_1 - \mathbf{u}_0, \dots, \mathbf{u}_n - \mathbf{u}_0$ are linearly independent, they form a basis for $\mathbb{R}^n$, and the linear transformation T from $\mathbb{R}^n$ to $\mathbb{R}^n$ defined by

$$T(\mathbf{x}) = (\lambda_1, \dots, \lambda_n)$$

is nonsingular and continuous. Let f be the mapping from $\mathbb{R}^n$ to $\mathbb{R}^{n+1}$ defined by

$$f(\mathbf{x}) = f\left(\sum_{i=1}^n \lambda_i \mathbf{u}_i\right) = (\lambda_1, \lambda_2, \dots, \lambda_n),$$

where the λ_i are as in (6). The mapping f can be viewed as a composite map $\mathbf{h} \circ T$, where

$$\mathbf{h}(\lambda_1, \dots, \lambda_n) = (\lambda_1, \lambda_2, \dots, \lambda_n), \quad \lambda_n = 1 - \sum_{i=1}^{n-1} \lambda_i.$$

Thus, f is continuous.

Since f is continuous and since

$$f(\mathbf{x}_0) = \left(\frac{1}{n+1}, \frac{1}{n+1}, \dots, \frac{1}{n+1}\right),$$

there exists a $\delta > 0$ such that for $\mathbf{x} \in B(\mathbf{x}_0, \delta)$ all components λ_i of $f(\mathbf{x})$ will be strictly positive. Thus, all $\mathbf{x} \in B(\mathbf{x}_0, \delta)$ lie in C , and so $\mathbf{x}_0$ is an interior point of C .

By experimenting with a sheet of cardboard, if necessary, the reader should find it plausible that an arbitrary plane in $\mathbb{R}^3$ can be put into coincidence with the plane $x_3 = 0$ by means of an appropriate sequence of rotations followed by a possible translation. The plane $x_3 = 0$ consists of those points in $\mathbb{R}^3$ of the form $(x_1, x_2, 0)$ and thus is essentially $\mathbb{R}^2$. Thus, in some sense a plane in $\mathbb{R}^3$ is equivalent to $\mathbb{R}^2$. Similarly, any line in $\mathbb{R}^3$ can be made to coincide with any one of the coordinate axes by means of an appropriate sequence of rotations followed by a possible translation.

We shall now show that any n -dimensional manifold M in $\mathbb{R}^n$ is equivalent, in a sense to be made precise, to $\mathbb{R}^n$.

Definition 6.6. An *affine transformation* T from $\mathbb{R}^n$ to $\mathbb{R}^n$ is a mapping of the form $Tx = Sx + a$, where a is a fixed vector in $\mathbb{R}^n$ and S is a linear transformation from $\mathbb{R}^n$ to $\mathbb{R}^n$.

Since relative to appropriate bases in $\mathbb{R}^n$ a linear transformation S from $\mathbb{R}^n$ to $\mathbb{R}^n$ is given by $Sx = Ax$, where A is an $n \times n$ matrix, we have $Tx = Ax + a$.

Lemma 6.6. A mapping T from $\mathbb{R}^n$ to $\mathbb{R}^n$ is affine if and only if for each $\lambda \in \mathbb{R}$ and each pair of points x_1 and x_2 in $\mathbb{R}^n$

$$T[(1 - \lambda)x_1 + \lambda x_2] = (1 - \lambda)Tx_1 + \lambda Tx_2. \quad (7)$$

It is a straightforward calculation to show that if T is affine, then (7) holds. Now suppose that (7) holds. Let $a = T(0)$, and let $Sx = Tx - a$. To show that T is affine, we must show that S is linear.

For any α in $\mathbb{R}$ and x in $\mathbb{R}^n$,

$$\begin{aligned} S(\alpha x) &= T(\alpha x + (1 - \alpha)(0)) - a = \alpha T(x) + (1 - \alpha)T(0) - a \\ &= \alpha(T(x) - a) = \alpha S(x). \end{aligned}$$

For any x_1 and x_2 in $\mathbb{R}^n$ we have

$$\begin{aligned} S(x_1 + x_2) &= S[\frac{1}{2}(x_1 + x_2) + \frac{1}{2}(x_1 + x_2)] = 2S[\frac{1}{2}(x_1 + x_2)] = 2T[\frac{1}{2}(x_1 + x_2)] - 2a \\ &= [T(x_1) - a] + [T(x_2) - a] = S(x_1) + S(x_2). \end{aligned}$$

Thus S is linear.

From (7) we see that under an affine map lines are mapped onto lines. Thus under an affine map linear manifolds are mapped onto linear manifolds. By taking $0 \leq \lambda \leq 1$ in (7), we see that under an affine map convex sets are mapped onto convex sets.

THEOREM 6.7. Let $x_0, x_1, \dots, x_r$ and $y_0, y_1, \dots, y_r$ be two affinely independent sets in $\mathbb{R}^n$. Then there exists a one-to-one affine transformation T from $\mathbb{R}^n$ onto itself such that $Tx_i = y_i$, $i = 0, 1, \dots, r$. If $r = n$, then T is unique.

If $r < n$, obtain two affinely independent sets $x_0, \dots, x_n$, $x_{n+1}, \dots, x_n$ and $y_0, \dots, y_n$, $y_{n+1}, \dots, y_n$ of $n + 1$ points by adjoining appropriate points to the original sets. The vectors $x_1 - x_0, \dots, x_n - x_0$ constitute a basis for $\mathbb{R}^n$, as do the vectors $y_1 - y_0, \dots, y_n - y_0$. Thus, there exists a unique one-to-one linear transformation S of $\mathbb{R}^n$ onto itself such that $S(x_i - x_0) = y_i - y_0$, $i = 1, \dots, n$. Let $Tx = Sx + a$ for a fixed a to be determined. Thus for $i = 0, 1, \dots, n$

$$T(x_i) = S(x_i) + a = S(x_i - x_0) + S(x_0) + a = y_i - y_0 + S(x_0) + a.$$

Therefore, if we take $\mathbf{x} = \mathbf{y}_i - \beta\mathbf{y}_n$, we will have $T\mathbf{x}_i = \mathbf{y}_i$, $i = 0, 1, \dots, n$.

COROLLARY 6. *Let M_1 and M_2 be two linear manifolds of the same dimension. Then there exists a one-to-one affine transformation T of $\mathbb{R}^n$ onto itself such that $TM_1 = M_2$.*

By Lemma 6.4, M_1 has an affine basis $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_n$ and M_2 has an affine basis $\mathbf{y}_0, \mathbf{y}_1, \dots, \mathbf{y}_n$. By Theorem 6.1 there exists a one-to-one transformation T from $\mathbb{R}^n$ onto itself such that $T\mathbf{x}_i = \mathbf{y}_i$, $i = 1, \dots, n$. The corollary now follows from the facts that each $\mathbf{x}$ in M_1 is a unique affine combination of $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_n$, each $\mathbf{y}$ in M_2 is a unique affine combination of $\mathbf{y}_0, \mathbf{y}_1, \dots, \mathbf{y}_n$, and T maps affine combinations onto affine combinations.

REMARK 6.1. If M is a linear manifold of dimension r , then there exists a one-to-one affine transformation T of $\mathbb{R}^n$ onto itself that maps M onto the manifold

$$U_r^0 = \{\mathbf{x} \in \mathbb{R}^n : (x_1, \dots, x_n, 0, \dots, 0, 0) = (x_1, \dots, x_n, 0, \dots, 0, 0)\}.$$

The manifold U_r^0 can be identified with $\mathbb{R}^r$ by means of the one-to-one mapping Π from U_r^0 onto $\mathbb{R}^r$ given by

$$\Pi(\mathbf{x}) = (x_1, \dots, x_n, 0, \dots, 0) \mapsto (x_1, \dots, x_r).$$

Denote the restriction of T to M by T_M , and define a mapping τ from M to $\mathbb{R}^r$ by the formula

$$\mathbf{x} \in M : \tau(\mathbf{x}) = \Pi(T_M \mathbf{x}).$$

Then τ is one-to-one onto, is continuous, and has a continuous inverse, τ^{-1} . Note that τ and τ^{-1} map convex sets onto convex sets and linear manifolds onto linear manifolds of the same dimension.

EXERCISE 6.1. Let k be a positive integer, $1 \leq k \leq n - 1$. Show that for a given subspace U_{n-k} of $\mathbb{R}^n$, there exist k linearly independent vectors $\mathbf{u}_1, \dots, \mathbf{u}_k$ such that

$$U_{n-k} = \{\mathbf{x} \in \mathbb{R}^n : \langle \mathbf{u}_i, \mathbf{x} \rangle = 0, i = 1, \dots, k\}.$$

EXERCISE 6.2. Show that a set of points $\mathbf{x}_0, \dots, \mathbf{x}_n$ is affinely dependent if and only if one of the points is an affine combination of the other points.

EXERCISE 6.3. In $\mathbb{R}^3$ consider the plane $2x + y + z = 5$.

- (a) Find an affine basis for this plane that does not include the point $\mathbf{x} = (-1, -2, 9)$.
- (b) Express $\mathbf{x} = (-1, -2, 9)$ in barycentric coordinates relative to your basis.

Exercise 6.6. Let $\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_r$ be affinely independent. Show that any $r + 2$ points $\mathbf{p}_0, \mathbf{p}_1, \dots, \mathbf{p}_r, \mathbf{x}_{r+1}$ in $\text{co}\{\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_r\}$ are affinely dependent.

Exercise 6.7. Show that an affine map T maps affine combinations of points into affine combinations of points.

Exercise 6.8. In Theorem 6.2, why is T unique if $r = n$?

7. MORE ON SEPARATION AND SUPPORT

Although the separation and support theorems of Sections 3 and 4 are adequate for most applications to optimization problems, they are not applicable in some situations where proper separation and nontrivial support of convex sets do occur. To illustrate this, let

$$C_1 = \{\mathbf{x} \in \mathbb{R}^2 : x_1 \geq 0, x_2 \geq 0, x_3 = 0\}$$

$$C_2 = \{\mathbf{x} \in \mathbb{R}^2 : x_1 \geq 0, x_2 \geq 0, x_3 = 0\}.$$

These two sets are properly separated by the plane (hyperplane in $\mathbb{R}^3$) $x_3 = x_4 = 0$. These sets are convex, are not disjoint, and have empty interior so that neither Theorem 3.3 nor Theorem 4.4 can be applied to deduce separation.

Let $C = \{\mathbf{x} \in \mathbb{R}^2 : x_1^2 + x_2^2 \leq 1, x_3 = 0\}$. At each boundary point $\mathbf{x}_0$ with $x_{01}^2 + x_{02}^2 = 1$ there exist a nontrivial supporting hyperplane, namely the plane that is orthogonal to the plane $x_3 = 0$ and whose intersection with this plane is the tangent line to the circle $x_1^2 + x_2^2 = 1, x_3 = 0$ at $\mathbf{x}_0$. Again, since the convex set C has empty interior, Theorem 4.4 does not give us the existence of this nontrivial supporting hyperplane.

The sets C_1, C_2 , and C , considered as sets in $\mathbb{R}^3$, do have interiors. It turns out that this observation is the key to theorems that are applicable to examples such as those described above. We introduce the notion of relative interior of a convex set C of dimension r . The relative interior is the interior of a convex set C considered as a set in $\mathcal{N}(C)$. In other words we consider the set C as a set in $\mathbb{R}^r$ rather than a set in $\mathbb{R}^n$.

Definition 7.1. Let A be a set of dimension $r \leq n$ and let $M = \mathcal{N}(A)$. A point $\mathbf{x}_0$ in A is said to be a relative interior point of A if there exists an $\varepsilon > 0$ such that $\mathbf{B}(\mathbf{x}_0, \varepsilon) \cap M \subset A$. The set of relative interior points of A is called the

relative interior of A and will be denoted by $\text{ri}(A)$. A set A is said to be relatively open if $\text{ri}(A) = A$, that is, if every point of A is a relative interior point of A .

Note that if $\dim M = n$, then a relative interior point is an interior point and a relatively open set is an open set.

Since a linear manifold M is a closed set, any limit point of A will also be in M . Thus the closure of A as a subset of M is the same as the closure of A as a subset of $\mathbb{R}^n$.

If A is contained in M and $\dim M < n$, then every point of A is a boundary point of A . We define the relative boundary of A to be the set $\bar{A} - \text{ri}(A)$, that is, those points of $\bar{A}$ that are not relative interior points of A . If $\dim M = n$, then the relative boundary of A is the boundary of A .

Let M be a linear manifold of dimension $r < n$ and let τ be the mapping defined in Remark 3.1. We noted there that τ and τ^{-1} are one-to-one and continuous.

LEMMA 3.1. *If O is an open set in $\mathbb{R}^n$, then $\tau^{-1}(O)$ is a relatively open set in M . If O is a relatively open set in M , then $\tau(O)$ is an open set in $\mathbb{R}^n$.*

Proof. Let O be an open set in $\mathbb{R}^n$ and let $x_0 \in \tau^{-1}(O)$. Then there is a unique $y_0 \in O$ such that $\tau^{-1}(y_0) = x_0$, or equivalently, $y_0 = \tau(x_0)$. Since $x_0 \in O$, there exists an $\epsilon > 0$ such that the open ball $B_\epsilon(x_0, \epsilon)$ in $\mathbb{R}^r$ is contained in O . Since τ is continuous, there exists a $\delta > 0$ such that for all $B(x_0, \delta) \cap M$ we have $\tau(B(x_0, \delta) \cap M) \subset O$. But this says that $B(x_0, \delta) \cap M \subset \tau^{-1}(O)$. Hence x_0 is a relative interior point of $\tau^{-1}(O)$, and so $\tau^{-1}(O)$ is relatively open. Since $\tau = (\tau^{-1})^{-1}$, a similar argument with τ replaced by τ^{-1} shows that if O is relatively open, then $\tau(O)$ is open in $\mathbb{R}^n$.

Let C be a convex set in $\mathbb{R}^n$ of dimension $r < n$. As a consequence of the one-to-one relationship between $M = A(C)$ and $\mathbb{R}^r$ that preserves open sets, convexity, dimension of convex sets, and linear manifolds, we can obtain many properties of C by establishing these properties for convex sets of full dimension. These properties then hold for the set $\tau(C)$ considered as a set of full dimension in $\mathbb{R}^n$ and hence for $C = \tau^{-1}(\tau(C))$. In particular, we obtain the following results in this fashion.

THEOREM 3.1. *A convex set of dimension r has nonempty relative interior.*

This theorem follows from Theorem 6.1 and Lemma 3.1.

LEMMA 3.2. *Let C be a convex set and let x_1 and x_2 be two points with $x_1 \in \text{ri}(C)$ and $x_2 \in C$. Then the line segment $[x_1, x_2]$ is contained in $\text{ri}(C)$.*

COROLLARY 1. *If C is convex, then $\text{ri}(C)$ is convex.*

COROLLARY 1. *Let C be convex. Then*

- (i) $\overline{\text{ri}(C)} = \bar{C}$ and
- (ii) $\text{ri}(C) = \text{int}(C)$.

LEMMA 7.2. Corollary 1 and (i) of Corollary 2 follow from Lemma 7.1 and the corresponding conclusion. In proving (ii) of Corollary 2 to Lemma 7.1, we argued that since $C \subseteq \bar{C}$, it follows that $\text{int}(C) \subseteq \text{int}(\bar{C})$. In general, it is not true that if A and B are convex sets with $A \subseteq B$ that $\text{ri}(A) = \text{ri}(B)$, as the following example shows. Take B to be a closed unit cube in $\mathbb{R}^n$. Let A be any closed face of B . Then $A \subseteq B$, but it is not true that $\text{ri}(A) \subseteq \text{ri}(B)$. In fact, $\text{ri}(A)$ and $\text{ri}(B)$ are disjoint. Since $A(C) = A(\bar{C})$, it does follow that $\text{ri}(C) = \text{ri}(\bar{C})$. The proof of the opposite inclusion given in Section 2 is valid if we replace interiors by relative interiors.

We now state and prove the first possible separation theorem in $\mathbb{R}^n$.

THEOREM 7.1. *Let X and Y be two convex subsets of $\mathbb{R}^n$. Then X and Y can be properly separated if and only if $\text{ri}(X)$ and $\text{ri}(Y)$ are disjoint.*

To prove the "only if" part of the theorem, we shall need to define what is meant by hyperplane cutting a convex set and to develop conditions under which cutting occurs.

DEFINITION 7.1. A hyperplane H_α^c is said to cut a convex set C if there exist points x_1 and x_2 in C such that $\langle a, x_1 \rangle < \alpha$ and $\langle a, x_2 \rangle > \alpha$. A hyperplane that does not cut C is said to bound C .

LEMMA 7.1. A hyperplane H_α^c cuts a convex set C if and only if the following two conditions hold:

- (i) H_α^c does not contain C and
- (ii) $A_\alpha^c \cap \text{ri}(C) \neq \emptyset$.

Proof. Let H_α^c cut C . Then there exist points x_1 and x_2 in C such that $\langle a, x_1 \rangle < \alpha$ and $\langle a, x_2 \rangle > \alpha$. Since for $y \in H_\alpha^c$ we have $\langle a, y \rangle = \alpha$, conclusion (i) follows. Now let $x \in \text{ri}(C)$. Then by Lemma 7.2 the line segments $[x, x_1]$ and $[x, x_2]$ belong to $\text{ri}(C)$. Let

$$p(t) = \begin{cases} x + t(x_1 - x), & -1 \leq t \leq 0, \\ x + t(x_2 - x), & 0 \leq t \leq 1. \end{cases}$$

Then $p(t) \in C$ for $-1 \leq t \leq 1$ and $p(t) \in \text{ri}(C)$ for $-1 < t < 1$. Let $q(t) = \langle a, p(t) \rangle$. Since q is continuous on $[-1, 1]$ with $q(-1) < \alpha$ and $q(1) > \alpha$, there exists a t_0 with $-1 < t_0 < 1$ such that $q(t_0) = \alpha$. But then $p(t_0) \in \text{ri}(C)$ and $\langle a, p(t_0) \rangle = \alpha$, which contradicts (i).

Now suppose that (i) and (ii) hold. Then there exist a point $x_0 \in \text{ri}(C) \cap A_1^0$ and a point $x_1 \in C$ and not in A_1^0 . Suppose that $\langle a, x_1 \rangle = \alpha$. Since $x_0 \in \text{ri}(C)$, there exists a $\delta > 0$ such that $x_2 = x_0 - \delta x_1 = x_0 \pm C$. Then

$$\langle a, x_2 \rangle = \alpha - \delta \langle a, x_1 \rangle = \alpha - \delta \alpha = \alpha.$$

If $\langle a, x_0 \rangle > \alpha$, then $\langle a, x_2 \rangle < \alpha$.

We now prove Theorem 3.2.

If a hyperplane properly separates the convex sets X and Y then it certainly properly separates $\text{ri}(X)$ and $\text{ri}(Y)$. Conversely, if a hyperplane properly separates $\text{ri}(X)$ and $\text{ri}(Y)$, then since $\text{ri}(X) = X$ and $\text{ri}(Y) = Y$ it follows that the hyperplane properly separates X and Y . Thus a hyperplane properly separates X and Y if and only if it properly separates $\text{ri}(X)$ and $\text{ri}(Y)$. Thus we may assume without loss of generality that X and Y are disjoint relatively open convex sets. The proof of separation will be carried out by a descending induction on $d = \max\{\dim X, \dim Y\}$.

If $\dim X = n$, then the separation follows from Theorem 3.1. We now assume that proper separation covers if $d \geq k$, where $1 \leq k \leq n$. We now suppose that $\dim X = k - 1$ and $\dim Y \leq k - 1$. To simplify the notation, we assume that $0 \in Y$. There is no loss of generality in assuming this, since this can always be achieved by a translation of coordinates. Since $0 \in Y$, the manifold $A(X)$ is a subspace of dimension $k \leq n - 1$. Let w be a point in $\mathbb{R}^n$ that is not in $A(X)$. Let

$$A = \text{co}\{X \cup (X + w)\}, \quad B = \text{co}\{X \cup (X - w)\}.$$

We assert that Y cannot intersect both A and B . See Figure 23.

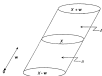


Figure 23.

By Theorem 2.2 point $\bar{q}$ in A can be written

$$\bar{q} = \sum_{i=1}^{n+1} \bar{q}_i \bar{x}_i + \bar{q}_0 \bar{w},$$

where $\bar{x}_i \in X$ for $i = 1, \dots, n+1$, the vector $\bar{z} = (\bar{q}_0, \dots, \bar{q}_{n+1}) \in F_{n+1,1}$, and $\bar{q}_i$ is either zero or 1. Thus, if $\bar{q} \in A$, then there exists an $x \in X$ such that

$$\bar{q} = x + \theta w, \quad 0 \leq \theta \leq 1. \quad (1)$$

Since

$$\bar{q} = x + \theta w = (1 - \theta)x + \theta(w + x),$$

every point $\bar{q}$ given by (1) is in A . A similar argument shows that

$$\bar{q} = \{\bar{q}; \bar{q} = x + \theta w, x \in X, 0 \leq \theta \leq 1\}.$$

Thus if $A \cap F \neq \emptyset$ and $B \cap F \neq \emptyset$, then there exist points y_1 and y_2 in F such that

$$y_1 = x_1 + \theta_1 w, \quad y_2 = x_2 + \theta_2 w,$$

where $x_1, x_2 \in X$, $0 \leq \theta_1 \leq 1$, $0 \leq \theta_2 \leq 1$. The inequalities $\theta_1 = 0$ and $\theta_2 = 0$ follow from $X \cap F = \emptyset$.

Let $\alpha = \theta_2/\theta_1 + \theta_2$ and let $\beta = \theta_2/\theta_1 + \theta_2$. Then $\alpha > 0$, $\beta > 0$, $\alpha + \beta = 1$, and

$$\alpha y_1 + \beta y_2 = \alpha x_1 + \beta x_2. \quad (2)$$

Since X and F are convex, the left-hand side of (2) belongs to X and the right-hand side belongs to X . This contradicts $X \cap F = \emptyset$, and so F cannot intersect both A and B .

For the sake of definiteness, suppose that $A \cap F = \emptyset$. Since $\dim A = k$ (why?) and $\dim F \leq k-1$, by the inductive hypothesis there exists a hyperplane H that properly separates A and F . Since $X \subset A$, the hyperplane H properly separates X and F .

To prove the "only if" assertion, suppose that there exists a hyperplane H that properly separates X and F and that there exists a point x in $\text{ri}(X) \cap \text{ri}(F)$. Then $x \in H$, and thus $H \cap \text{ri}(X) \neq \emptyset$ and $H \cap \text{ri}(F) \neq \emptyset$. Since the separation is proper, at least one of the sets, say X , is not contained in H . Hence by Lemma 7.3, H cuts X , which contradicts the assumption that H separates X and F .

We conclude this section with results that are stronger than Theorems 4.1 and 4.3.

THEOREM 7.1. *Let C be a convex set and let D be a convex subset of the relative boundary of C . Then there exists a nontrivial supporting hyperplane to C that contains D .*

Proof. Since $\text{ri}(C) \cap D = \emptyset$, we have $\text{ri}(C) \cap \text{ri}(D) = \emptyset$. Hence by Theorem 7.2 there exists a hyperplane H_0^C that properly separates C and D . Thus

$$(x, x_0) < \alpha \quad \text{for } x \in C \quad \text{and} \quad (x, y_0) \geq \alpha \quad \text{for } y \in D. \quad (3)$$

Since each y in D is a relative boundary point of C , for each y in D there exists a sequence $\{x_n\}$ of points in C such that $x_n \rightarrow y$. For this sequence we have $(x, x_0) < \alpha$ and $(x, x_0) \rightarrow (x, y_0)$. Thus $(x, y_0) < \alpha$. From this and from (3) we get that $(x, y_0) = \alpha$. Thus, each y in D is in H_0^C . Since $D \subset \bar{C}$, it follows that H_0^C is a supporting hyperplane to C that contains D . Since the separation of C and D is proper and since $D \subset H_0^C$, there is an x_0 in C such that $(x, x_0) < \alpha$. Thus the supporting hyperplane is nontrivial.

The next theorem is stronger than Theorem 4.3 but is only valid in finite dimensions.

THEOREM 7.4. *Let C be a compact convex subset of $\mathbb{R}^k$. Then C is the convex hull of its extreme points.*

Proof. Let C_e denote the set of extreme points of C . By Theorem 4.2 the set $C_e \neq \emptyset$. Since $C_e \subset C$ and $\text{co}(C_e) \subset \text{co}(C) = C$, to prove the theorem, we must show that

$$C = \text{co}(C_e). \quad (4)$$

The proof of (4) proceeds by induction on the dimension k of C .

If $k = 0$, then C is a point, and if $k = 1$, then C is a closed line segment. Thus (4) holds in these cases. Now suppose that (4) holds if $\dim C = k - 1$, where $1 \leq k \leq n$.

Let $\dim C = k$ and let $x \in C$. By Theorem 7.1, either $x \in \text{ri}(C)$ or x belongs to the relative boundary of C . If x belongs to the relative boundary of C , then by Theorem 7.3, there exists a supporting hyperplane to C that contains x . This hyperplane will intersect $\text{ri}(C)$ in a manifold M_{k-1} that is a hyperplane of dimension $k - 1$ in $\mathbb{R}^k$ and that supports C at x . The set $M_{k-1} \cap C$ is compact and convex and has dimension $\leq k - 1$. By the inductive hypothesis x is a convex combination of the extreme points of $M_{k-1} \cap C$. By Lemma 4.1 the extreme points of $M_{k-1} \cap C$ are extreme points of C . Thus $x \in \text{co}(C_e)$ whenever x is a relative boundary point of C .

If $x \in \partial C$, let L be a line in $A[C]$ through the point x . Then $L \cap \partial C$ is a line segment $[p, q]$ whose end points p and q are in the relative boundary of C . Since C is compact, $y \in C$ and $xy \subset C$. Since x is a convex combination of p and q and since p and q are in $\text{co}(C_{\geq 0})$, it follows that $x \in \text{co}(C_{\geq 0})$ and (8) is established.

III

CONVEX FUNCTIONS

1. DEFINITION AND ELEMENTARY PROPERTIES

The simplest class of functions are real-valued functions that are translations of linear functions, or affine functions, f defined by the formula $f(x) = (x, x_0) + b$. One could consider the next class of functions in order of complexity to be the translation of the functions $y = |x|^p$, $p \geq 1$. In $\mathbb{R}^2 = (x, y)$, where $x \in \mathbb{R}^1$, the graph of $y = |x|$ is the upper nappe of a right circular cone with vertex at the origin and the graph of $y = |x|^2$ is a paraboloid of revolution. The graphs of such functions are actually symmetric about a line parallel to the y -axis. Intuitively, it appears that all of the above functions have the property that the set of points (x, y) in $\mathbb{R}^{n+1}$ that lie on or above the graph is a convex set. It is therefore reasonable to consider the next level of functions to be the functions with this property. Such functions are called convex functions. It turns out that convex functions are important in optimization problems as well as in other areas of mathematics. We shall define convex functions using a different property and then show in Lemma 1.1 below that our definition is equivalent to the one suggested.

Definition 1.1. A real-valued function f defined in a convex set C in $\mathbb{R}^n$ is said to be convex if for every x_1 and x_2 in C and every $\alpha \geq 0, \beta \geq 0, \alpha + \beta = 1$,

$$f(\alpha x_1 + \beta x_2) \leq \alpha f(x_1) + \beta f(x_2). \quad (1)$$

The function f is said to be strictly convex if strict inequality holds in (1).

The geometric interpretation of this definition is that a convex function is a function such that if x_1 and x_2 are any two points in $\mathbb{R}^{n+1}$ on the graph of f then points of the line segment $[x_1, x_2]$ joining x_1 and x_2 lie on or above the graph of f .

Figure 1.1 illustrates the definition. Let $C = (x_\alpha, x_\beta)$, where $x_\alpha = \alpha x_1 + \beta x_2$, $\alpha \geq 0, \beta \geq 0, \alpha + \beta = 1$, be a point on the line segment joining the points $A = (x_1, f(x_1))$ and $B = (x_2, f(x_2))$ on the graph of f . Let $D = (x_\alpha, \alpha)$ and let

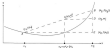


Figure 1.1

$E = (x_{i-1}, x_i]$. Then

$$\frac{f(x_i) - w}{(x_i - x_{i-1})} = \frac{(f(x_i) - f(x_{i-1}))}{(x_i - x_{i-1})} = \frac{f(x_i) - f(x_{i-1})}{(x_i - x_{i-1})}.$$

From the extreme terms of these equalities we get that

$$w \leq f(x_i) \leq w + (f(x_i) - f(x_{i-1})) = f(x_{i-1}).$$

Thus $f(x_{i-1})$ lies above w on the graph (and only $f(x_{i-1})$ lies).

Note that if the domain of f is not convex, then the left-hand side of (1.1) need not be defined. Therefore, when defining convex functions, it will suffice to use $\mathbb{R}$ -valued functions defined on $\mathbb{R}$ -valued sets. (Incidentally we shall not repeat the same statement when $\mathbb{C}$ -valued convex functions; it shall be tacitly assumed that f is real valued.)

Definition 1.1. A function f defined on a convex set C is said to be convex if $-f$ is concave.

Definition 1.2. Let f be a real valued function defined on a set A . The epigraph of f , denoted by $\text{epi}(f)$ is the set

$$\text{epi}(f) = \{(x, y) \in \text{dom}(f) \mid y \geq f(x)\}.$$

In Figure 1.2 the bounded curves represent the epigraphs of $y = x^2$, $y = 1/x$, $y = \ln x$, and $y = x^2 + 1$, $y = 1/x + 1$, $y = \ln x + 1$.

The epigraph of a function is the set of points (x, y) in $\mathbb{R}^2$ that lie on or above the graph of f .



Figure 3.1.

Lemma 3.1. *A function f defined on a convex set C is convex if and only if $\text{epi } f$ is convex.*

Proof. Suppose that $\text{epi } f$ is convex. Let x_1 and x_2 belong to C . Then $(x_1, f(x_1))$ and $(x_2, f(x_2))$ belong to $\text{epi } f$. Since $\text{epi } f$ is convex, for any $(\alpha, \beta) \in P_1$

$$\alpha(x_1, f(x_1)) + \beta(x_2, f(x_2)) = (\alpha x_1 + \beta x_2, \alpha f(x_1) + \beta f(x_2))$$

belongs to $\text{epi } f$. By the definition of $\text{epi } f$,

$$\alpha f(x_1) + \beta f(x_2) \geq f(\alpha x_1 + \beta x_2),$$

and so f is convex.

Now suppose f is convex. Let (x_1, y_1) and (x_2, y_2) be any two points in $\text{epi } f$. Then $y_1 \geq f(x_1)$ and $y_2 \geq f(x_2)$. For any $(\alpha, \beta) \in P_1$ let

$$(x, y) = \alpha(x_1, y_1) + \beta(x_2, y_2) = (\alpha x_1 + \beta x_2, \alpha y_1 + \beta y_2).$$

Since (x_1, y_1) and (x_2, y_2) are in $\text{epi } f$, and since f is convex, we have

$$y = \alpha y_1 + \beta y_2 \geq \alpha f(x_1) + \beta f(x_2) \geq f(\alpha x_1 + \beta x_2) = f(x).$$

Thus, $(x, y) \in \text{epi } f$, and so $\text{epi } f$ is convex.

Lemma 3.2. *If f is a convex function defined on a convex set C in $\mathbb{R}^n$, then for any real γ the set $\{x \in C \mid f(x) \leq \gamma\}$ is either empty or convex.*

We leave the proof of this lemma as an exercise for the reader.

The converse to this lemma is false, as can be seen from the function $f(x) = x^2$.

Lemma 1.3. Let $\{f_\alpha\}$ be a family of convex functions defined on a convex set C in $\mathbb{R}^n$ and let there exist a real-valued function M with domain C such that, for each α , $f_\alpha(x) \leq M(x)$ for all x in C . Then the function f defined by $f(x) = \sup\{f_\alpha(x) : \alpha\}$ is convex.

Proof. Since for each α in C , $f_\alpha(x) \leq M(x)$ for all x , it follows that $f(x) \leq M(x)$, and so f is real valued. Also $\text{epi}(f) = \bigcap_\alpha \text{epi}(f_\alpha)$, so $\text{epi}(f)$, being the intersection of convex sets, is convex. Therefore f is convex.

Theorem 1.1 (Jensen's Inequality). Let f be a real-valued function defined on a convex set C in $\mathbb{R}^n$. Then f is convex if and only if for every finite set of points $x_1, \dots, x_k$ in C and every $\lambda = (\lambda_1, \dots, \lambda_k)$ in P_k

$$f(\lambda_1 x_1 + \dots + \lambda_k x_k) \leq \lambda_1 f(x_1) + \dots + \lambda_k f(x_k) \quad (2)$$

Proof. If (2) holds for all positive integers k , then it holds for $k = 2$. But for $k = 2$, the relation (2) is the definition of convex function.

Now suppose that f is convex. We shall prove (2) by induction on k . For $k = 1$ the relation is trivial. For $k = 2$ the relation follows from the definition of a convex function. Suppose that $k \geq 2$ and that (2) has been established for $k - 1$. We shall show that (2) holds for k . If $\lambda_k = 1$, then there is nothing to prove. If $\lambda_k \neq 1$, set $A = \sum_{i=1}^{k-1} \lambda_i$. Then

$$A > 0, \quad A + \lambda_k = 1, \quad \sum_{i=1}^{k-1} \frac{\lambda_i}{A} = 1,$$

and we may write (compare the proof of Lemma B.2.6)

$$\begin{aligned} f\left(\sum_{i=1}^k \lambda_i x_i\right) &= f\left(A \left[\sum_{i=1}^{k-1} \left(\frac{\lambda_i}{A}\right) x_i\right] + \lambda_k x_k\right) \\ &\leq \lambda_k f\left(\sum_{i=1}^{k-1} \left(\frac{\lambda_i}{A}\right) x_i\right) + \lambda_k f(x_k) \\ &\leq \lambda_k f(x_1) + \dots + \lambda_{k-1} f(x_{k-1}) + \lambda_k f(x_k) \end{aligned}$$

where the first inequality follows from the convexity of f and the second inequality follows from the inductive hypothesis.

The next lemma in conjunction with other criteria is sometimes useful in determining whether a given function is convex.

Lemma 1.4. Let f be a convex function defined on a convex set C in $\mathbb{R}^n$. Let g be a nondecreasing convex function defined on an interval I in $\mathbb{R}$. Let $f(C) \subseteq I$. Then the composite function $g \circ f$ defined by $(g \circ f)(x) = g(f(x))$ is convex on C .

We leave the proof as an exercise.

Let C be a convex set. Choose an $\mathbf{x}$ in C and a $\mathbf{v}$ in $\mathbb{R}^p$. Let

$$A = \{\lambda : \lambda \in \mathbb{R}, \mathbf{x} + \lambda \mathbf{v} \in C\}.$$

Let $\bar{\lambda}_1 = \sup\{\lambda : \lambda \in A\}$ and let $\bar{\lambda}_2 = \inf\{\lambda : \lambda \in A\}$. Then $\bar{\lambda}_1 \leq \bar{\lambda}_2 \leq \bar{\lambda}_1$. If $\bar{\lambda}_1 \neq \bar{\lambda}_2$ then since C is convex, $\mathbf{x} + \lambda \mathbf{v} \in C$ for all λ in the open interval $(\bar{\lambda}_2, \bar{\lambda}_1)$. Note that, depending on the structure of the set C , $\mathbf{x} + \lambda \mathbf{v}$ can belong to C for all λ in $[\bar{\lambda}_2, \bar{\lambda}_1]$ or $(\bar{\lambda}_2, \bar{\lambda}_1]$ or $[\bar{\lambda}_2, \bar{\lambda}_1)$.

Let f be a convex function defined on a convex set C . From the discussion in the preceding paragraph it follows that for each $\mathbf{x}$ in C and each $\mathbf{v}$ in $\mathbb{R}^p$ there is an interval $I \subset \mathbb{R}$ that depends on $\mathbf{x}$ and $\mathbf{v}$ such that the function $\varphi(\cdot; \mathbf{x}, \mathbf{v})$ defined by

$$\varphi(\lambda; \mathbf{x}, \mathbf{v}) = f(\mathbf{x} + \lambda \mathbf{v}) \quad (3)$$

has I as its domain of definition. The interval I always includes the origin and can degenerate to the single point $\{0\}$; it can be open, half open, or closed. The function $\varphi(\cdot; \mathbf{x}, \mathbf{v})$ gives the one-dimensional section of the function f through the point $\mathbf{x}$ in the direction of $\mathbf{v}$. The next lemma states that f is convex if and only if all the nondegenerate one-dimensional functions $\varphi(\cdot; \mathbf{x}, \mathbf{v})$ are convex.

Lemma 1.5. *A function f defined on a convex set C is convex if and only if for each $\mathbf{x}$ in C and $\mathbf{v}$ in $\mathbb{R}^p$ such that $\varphi(\cdot; \mathbf{x}, \mathbf{v})$ is defined on a nondegenerate interval I the function $\varphi(\cdot; \mathbf{x}, \mathbf{v})$ is convex on I .*

Proof. Suppose that f is convex. Then for any λ_1, λ_2 in I and (α, β) in P_2

$$\begin{aligned} \varphi(\alpha \lambda_1 + \beta \lambda_2; \mathbf{x}, \mathbf{v}) &= f(\mathbf{x} + (\alpha \lambda_1 + \beta \lambda_2) \mathbf{v}) \\ &= f(\alpha(\mathbf{x} + \lambda_1 \mathbf{v}) + \beta(\mathbf{x} + \lambda_2 \mathbf{v})) \\ &\stackrel{\text{conv. } f}{\leq} \alpha f(\mathbf{x} + \lambda_1 \mathbf{v}) + \beta f(\mathbf{x} + \lambda_2 \mathbf{v}) \\ &\stackrel{\text{conv. } \varphi}{\leq} \alpha \varphi(\lambda_1; \mathbf{x}, \mathbf{v}) + \beta \varphi(\lambda_2; \mathbf{x}, \mathbf{v}) \\ &\stackrel{\text{conv. } \varphi}{=} \varphi(\alpha \lambda_1 + \beta \lambda_2; \mathbf{x}, \mathbf{v}). \end{aligned}$$

This proves the convexity of $\varphi(\cdot; \mathbf{x}, \mathbf{v})$.

Now suppose that for each $\mathbf{x}$ in C and $\mathbf{v}$ in $\mathbb{R}^p$ such that $I \neq \{0\}$, $\varphi(\cdot; \mathbf{x}, \mathbf{v})$ is convex. Let $\mathbf{x}_1$ and $\mathbf{x}_2$ be in C and let (α, β) be in P_2 . Then

$$\begin{aligned} f(\alpha \mathbf{x}_1 + \beta \mathbf{x}_2) &= f(\bar{\mathbf{x}}) = \varphi(\bar{\beta}; \mathbf{x}_1, \mathbf{x}_2 - \mathbf{x}_1) \\ &= f(\mathbf{x}_1 + \beta(\mathbf{x}_2 - \mathbf{x}_1)) = \varphi(\bar{\beta}; \mathbf{x}_1, \mathbf{x}_2 - \mathbf{x}_1) \\ &= \varphi(\alpha + \beta; \mathbf{x}_1, \mathbf{x}_2 - \mathbf{x}_1). \end{aligned} \quad (4)$$

Since

$$\varphi(\lambda \mathbf{x}_0 + \mathbf{x}_1 - \mathbf{x}_0) = f(\mathbf{x}_0) + \lambda(\mathbf{x}_1 - \mathbf{x}_0)$$

and since C is convex, $\varphi(\lambda \mathbf{x}_0 + \mathbf{x}_1 - \mathbf{x}_0)$ is defined for $0 \leq \lambda \leq 1$ and

$$\varphi(0 \mathbf{x}_0 + \mathbf{x}_1 - \mathbf{x}_0) = f(\mathbf{x}_1), \quad \varphi(1 \mathbf{x}_0 + \mathbf{x}_1 - \mathbf{x}_0) = f(\mathbf{x}_0). \quad (5)$$

From the convexity of φ ($\mathbf{x}_0, \mathbf{x}_1 - \mathbf{x}_0$) and (5) we get that the rightmost side of (4) is less than or equal to

$$\alpha \varphi(0 \mathbf{x}_0 + \mathbf{x}_1 - \mathbf{x}_0) + \beta \varphi(1 \mathbf{x}_0 + \mathbf{x}_1 - \mathbf{x}_0) = \alpha f(\mathbf{x}_1) + \beta f(\mathbf{x}_0).$$

Hence f is convex.

Remark 1.1. An examination of the proof shows that f is strictly convex if and only if, for each $\mathbf{x}$ and $\mathbf{x}_1$ w/ $(\mathbf{x}, \mathbf{x}_1)$ is strictly convex.

We now present some results for convex functions defined on real intervals. The first result is the so-called three-chord property.

Theorem 1.2. Let g be a convex function defined on an interval $I \subset \mathbb{R}$. Then if $x < y < z$ are three points in I ,

$$\frac{g(y) - g(x)}{y - x} \leq \frac{g(z) - g(x)}{z - x} \leq \frac{g(z) - g(y)}{z - y}. \quad (6)$$

If g is strictly convex, then strict inequality holds.

Figure 1.1 illustrates the theorem and shows why the conclusion of the theorem is called the three-chord property.

Proof. Since $x < y < z$, there exists a α ($0 < \alpha < 1$) such that $y = x + \alpha(z - x)$. Since g is convex,

$$\begin{aligned} g(y) - g(x) &= g(1 - \alpha)x + \alpha z - g(x) \\ &\leq (1 - \alpha)g(x) + \alpha(g(z) - g(x)) \\ &= \alpha(g(z) - g(x)). \end{aligned} \quad (7)$$

If in (7) we now divide through by $\alpha(z - x) > 0$ and note that $y - x = \alpha(z - x)$, we get that

$$\frac{g(y) - g(x)}{y - x} \leq \frac{g(z) - g(x)}{z - x}.$$

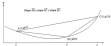


Figure 1.4.

To obtain the second inequality, we write

$$\begin{aligned} g(x) - g(y) &= g(x) - g(y) + g(y) - g(z) + g(z) - g(y) \\ &= g(x) - (z - y)g(y) + g(y) \\ &= 0 + (y)g(y) - g(y)g. \end{aligned}$$

If we now divide through by $z - y$ (it will note that $z - y > 0$) we obtain the second inequality.

LEMMA 1.1. For each rule f the slope function f^m is defined by

$$f^m(x, y) = \frac{f(y) - f(x)}{y - x}, \quad \text{if } y > x, \quad \text{and}$$

is an increasing, concave (if f is strictly convex, then f^m is strictly decreasing).

Proof. We must show that if $x_1 < x_2$, then

$$f^m(x_1, y) \leq f^m(x_2, y) \quad (1)$$

For two points x_1 and x_2 such that $x_1 < x_2$, inequality (1) follows from (4) by taking $z = x_1$, $y = x_2$, and $x = x_1$. For two points x_1 and x_2 such that $x_1 < x_2$, i.e., inequality (1) follows from (4) by taking $z = x_1$, $y = x_2$, and $x = x_2$. Lastly, if $x_1 < x_2$, the inequality (1) follows from (4) by taking $z = x_1$, $y = x_2$, and $x = x_1$.

Let f be a real-valued function defined on an interval I . If at a point c in I

$$\lim_{h \rightarrow 0^+} \frac{f(c+h) - f(c)}{h}$$

exists, then f is said to have a *right-hand derivative* at c whose value $f'_+(c)$ is the value of the above limit. Similarly, the *left-hand derivative* of f at c , denoted by $f'_-(c)$, is

$$\lim_{h \rightarrow 0^-} \frac{f(c+h) - f(c)}{h}$$

provided the limit exists.

COROLLARY 1. At each point c in $\text{int}(I)$, the right- and left-hand derivatives $f'_+(c)$ and $f'_-(c)$ exist and $f'_-(c) \leq f'_+(c)$. Moreover, f'_+ and f'_- are increasing functions.

Proof. In (5) take $y = c$. Then for $x < c < z$ we have

$$\frac{g(x) - g(c)}{x - c} \leq \frac{g(z) - g(c)}{z - c}.$$

It follows from this relation, from Corollary 1, and from Exercise 16.6, that $g'_-(c)$ and $g'_+(c)$ exist and $g'_-(c) \leq g'_+(c)$.

Now let $c = d$ be any point in $\text{int}(I)$ and let $h > 0$ be such that $c + h < d$ and $d - h > c$. Then by Theorem 1.2

$$\frac{g(c+h) - g(c)}{h} \leq \frac{g(d) - g(c)}{d - c} \leq \frac{g(d - h) - g(d)}{-h}.$$

If we now let $h \rightarrow 0^+$, we get $g'_+(c) \leq g'_-(d)$. Combining this with $g'_-(c) \leq g'_+(c)$ for any c in $\text{int}(I)$, we get

$$g'_-(c) \leq g'_+(c) \leq g'_-(d) \leq g'_+(d)$$

which gives the desired monotonicity property.

Remark 1.1. If c is a left-hand end point of I and g is defined at c , then from the monotonicity of $g(x)$ it follows that $g'_-(c)$ always exists if we permit the value $-\infty$.

The next corollary is an immediate consequence of Corollary 2 and is a result that the reader encountered in elementary calculus.

COROLLARY 1. *If g is convex and twice differentiable on I , then g' is an increasing function on I and $g''(x) \geq 0$ for all x in I .*

We can now return to functions defined on $\mathbb{R}^n$.

Definition 1.4. Let f be a real-valued function defined on a set S in $\mathbb{R}^n$. Let $\mathbf{x}$ be a point in S and let $\mathbf{v}$ be a vector in $\mathbb{R}^n$ such that $\mathbf{x} + \lambda\mathbf{v}$ belongs to S for λ in some interval (δ, δ) . Then f is said to have a directional derivative at $\mathbf{x}$ in the direction $\mathbf{v}$ if

$$\lim_{\lambda \rightarrow 0^+} \frac{f(\mathbf{x} + \lambda\mathbf{v}) - f(\mathbf{x})}{\lambda}$$

exists. We denote the directional derivative by $f'(\mathbf{x}; \mathbf{v})$.

Lemma 1.6. *Let f be a convex function defined on a convex set C and let $\mathbf{x}$ be a point of C . If $\mathbf{x}$ is an interior point, then $f'(\mathbf{x}; \mathbf{v})$ exists for every $\mathbf{v}$ in $\mathbb{R}^n$. If $\mathbf{x}$ is a boundary point and $\mathbf{v}$ is a vector in $\mathbb{R}^n$ such that $\mathbf{x} + \lambda\mathbf{v}$ belongs to C for sufficiently small $\lambda > 0$, then $f'(\mathbf{x}; \mathbf{v})$ exists if we allow the value $-\infty$.*

Proof. If $\mathbf{x}$ is an interior point of C , then for each $\mathbf{v}$ in $\mathbb{R}^n$ there exists a positive number $\lambda_0 = \lambda_0(\mathbf{v})$ such that for $|\lambda| \leq \lambda_0$ the point $\mathbf{x} + \lambda\mathbf{v}$ is in C . Hence the function $\varphi: (\lambda, \mathbf{v})$ of (3) is defined for $|\lambda| \leq \lambda_0$. By Lemma 1.3, $\varphi: (\lambda, \mathbf{v})$ is convex, and so by Corollary 2 of Theorem 1.1, $\varphi'(\delta; \mathbf{x}, \mathbf{v})$ exists. But

$$\frac{\varphi(\delta; \mathbf{x}, \mathbf{v}) - \varphi(0; \mathbf{x}, \mathbf{v})}{\delta} = \frac{f(\mathbf{x} + \delta\mathbf{v}) - f(\mathbf{x})}{\delta},$$

and so $f'(\mathbf{x}; \mathbf{v})$ exists and equals $\varphi'(\delta; \mathbf{x}, \mathbf{v})$. If $\mathbf{x}$ is a boundary point of C , the assertion follows from similar arguments and Remark 1.1.

In the next lemma we record a useful inequality.

Lemma 1.7. *Let f be a convex function defined on a convex set C . Let $\mathbf{x}$ and $\mathbf{y}$ be points in C and let $\mathbf{z} = \alpha\mathbf{x} + (1 - \alpha)\mathbf{y}$, $0 \leq \alpha \leq 1$. Then*

$$f(\mathbf{z}) - f(\mathbf{x}) \geq \frac{f(\mathbf{y}) - f(\mathbf{x})}{1} \alpha. \quad (4)$$

Proof. The function $\varphi: (\lambda, \mathbf{y} - \mathbf{x})$ in (1) is defined for $0 \leq \lambda \leq 1$, with $\varphi(0; \mathbf{x}, \mathbf{y} - \mathbf{x}) = f(\mathbf{x})$, $\varphi(1; \mathbf{x}, \mathbf{y} - \mathbf{x}) = f(\mathbf{y})$ and $\varphi(\alpha; \mathbf{x}, \mathbf{y} - \mathbf{x}) = f(\mathbf{z})$ for $0 \leq$

$t < 1$. By Corollary 1 to Theorem 1.2,

$$\varphi(t)x_0 + (1-t)x_1 = \varphi(t)x_0 + (1-t)\frac{\varphi(1)x_0 + (1-t)x_1}{1},$$

which is the same as (9).

We note that (8) can also be established directly from the definition of a convex function.

We next show that a convex function is continuous on the interior of its domain of definition.

THEOREM 1.3. *Let f be a convex function defined on a convex set C in $\mathbb{R}^n$. If x is an interior point of C , then f is continuous at x .*

If x is not an interior point of C , then f need not be continuous at x , as the following example shows. Let $C = [-1, 1] \subset \mathbb{R}$ and let $f(x) = x^2$ for $-1 < x < 1$ and $f(-1) = f(1) = 2$. The function f is convex on $[-1, 1]$ since $\text{epi } f$ is convex, but f is discontinuous at ± 1 .

We now prove the theorem.

If x is an interior point of C , then there exists a $\rho > 0$ such that $\overline{B(x, \rho)} \subset C$. Let $\delta = \rho/\sqrt{n}$. Let C_δ denote the cube with center at x and length of side equal to 2δ ; that is,

$$C_\delta = \{y \mid |y_i - x_i| \leq \delta, i = 1, \dots, n\}.$$

Since

$$\|y - x\| = \left[\sum_{i=1}^n (y_i - x_i)^2 \right]^{1/2},$$

it follows that $\|x_i - y_i\| \leq \rho/\sqrt{n}$ for each i , then $\|y - x\| \leq \rho$. Thus $C_\delta \subset \overline{B(x, \rho)}$. On the other hand, if $\|y - x\| \leq \delta$, then we cannot have $|y_i - x_i| > \delta$ for any index i in set $\{1, \dots, n\}$. Hence $B(x, \delta) \subset C_\delta$.

The first step in the proof is to show that f is bounded above on $B(x, \delta)$. It is easier to show that f is bounded above on the larger set C_δ than to work directly with $B(x, \delta)$. Let $C_{2\delta}$ denote the set of 2^n vertices of C_δ . By Exercise 11.4.7(a), the set $C_{2\delta}$ is the set of extreme points of C_δ . By Theorem 11.4.1, $C_\delta = \text{co}(C_{2\delta})$. But $C_{2\delta}$ is compact, so by Lemma 11.1.1, $\text{co}(C_{2\delta})$ is compact. Hence

$$C_\delta = \overline{\text{co}(C_{2\delta})} = \text{co}(C_{2\delta}).$$

[We could also conclude this directly from Exercise 11.4.7 (b).] Thus, for each

p is in C_0 , then there is a p in $\bar{C}_0$ such that

$$f \equiv \sum_{j=1}^n \alpha_j \phi_j,$$

where the α_j are the vertices of C_0 .

Let

$$M = \max\{|\phi_j| : j \in C_0\}.$$

Then by the convexity of f and Theorem 1.1 (Jensen's inequality), for each p in C_0

$$f(p) \leq \sum_{j=1}^n p_j f(\alpha_j) \leq \sum_{j=1}^n p_j M = M.$$

Thus, since $\overline{B(x, r)} \subset C_0$,

$$f(x) \leq M \quad \text{for } x \in \overline{B(x, r)}. \quad (3.6)$$

Let x be an arbitrary point in the disk, and consider the line determined by x and z . This line will intersect the surface of $\overline{B(x, r)}$ at two points, $u = u_1$ and $v = u_2$, where $u_1 < x < u_2$ (the converse is 1 and $u_2 = 0$ in Figure 1a).

We shall use the bound on f and the convexity of f to estimate

$$f(x) = f(u_1)$$

in terms of $|z - x|$ and thus prove the continuity of f at x . To this end, we



Figure 1a.

first express x as a convex combination of x and $x + u$. Thus

$$x = (1 - \lambda)(x) + \lambda(x + u) = x + \lambda u, \quad 0 < \lambda < 1. \quad (11)$$

Therefore, $\lambda = \|x - u\|^{-1}$. We next express x as a convex combination of x and $x - u$. Thus

$$x = (1 - \alpha)(x - u) + \alpha x, \quad 0 < \alpha < 1,$$

resulting relation (1) in Section 1, of Chapter II and using the relation $\lambda = \|x - u\|^{-1}$, we see that

$$\lambda(1 - \alpha)^{-1} = \lambda\|x - u\|^{-1} = \lambda^{-1},$$

and so $\alpha = (1 + \lambda)^{-1}$. Therefore,

$$x = \frac{1}{1 + \lambda}x + \frac{\lambda}{1 + \lambda}(x + u). \quad (12)$$

Since f is convex, we have, using (10), (11), and (12),

$$f(x) \leq (1 - \lambda)f(x) + \lambda f(x + u) \leq (1 - \lambda)f(x) + \lambda M,$$

$$f(x) \leq \frac{1}{1 + \lambda}f(x) + \frac{\lambda}{1 + \lambda}f(x + u) \leq \frac{1}{1 + \lambda}f(x) + \frac{\lambda}{1 + \lambda}M.$$

From these inequalities we get

$$-K[M - f(x)] \leq f(x) - f(x) \leq \lambda[M - f(x)].$$

Hence, for $x \in B(x, \delta)$

$$|f(x) - f(x)| \leq \delta^{-1}[M - f(x)]\|x - x\|, \quad (13)$$

from which the continuity of f at x follows.

Remark 1.1. A function f defined on a set D in $\mathbb{R}^n$ with range in $\mathbb{R}^m$ is said to be *Lipschitz continuous* on D if there exists a constant $K > 0$ such that

$$\|f(x) - f(y)\| \leq K\|x - y\|$$

for all x, y in D . Since a real-valued continuous function defined on a compact set is bounded above and below, relation (13) implies that a convex function defined on an open convex set C is Lipschitz continuous on every compact subset contained in C .

Remark 1.1. The proof of Theorem 1.3 can easily be modified to show that if f is convex on a convex set C , then f is continuous on $\text{ri}(C)$.

In Section 11 of Chapter I we saw that the existence of partial derivatives at a point does not imply differentiability. For convex functions, however, the existence of partial derivatives at a point does imply differentiability.

THEOREM 1.4. Let f be a convex function defined on a convex set C . If the partial derivatives of f with respect to each variable exist at an interior point $\mathbf{x}$ of C , then f is differentiable at $\mathbf{x}$.

Proof. Define the function η on $\mathbb{R}^n$ by

$$\eta(\mathbf{h}) = f(\mathbf{x} + \mathbf{h}) - f(\mathbf{x}) - \sum_{j=1}^n \frac{\partial f}{\partial x_j}(\mathbf{x})h_j.$$

To prove that f is differentiable at $\mathbf{x}$, we must show that

$$\frac{\eta(\mathbf{h})}{\|\mathbf{h}\|} \rightarrow 0 \quad \text{as } \|\mathbf{h}\| \rightarrow 0. \quad (1.6)$$

Let

$$\eta_i = \frac{f(\mathbf{x} + \alpha h_i \mathbf{e}_i) - f(\mathbf{x})}{\alpha h_i} - \frac{\partial f}{\partial x_i}(\mathbf{x}).$$

Since f is convex and a linear function is convex and concave, the function η is concave. Thus

$$\begin{aligned} \eta(\mathbf{h}) &= \eta\left(\sum_{i=1}^n \frac{\alpha h_i \mathbf{e}_i}{\alpha} \right) \leq \frac{1}{\alpha} \sum_{i=1}^n \eta(\alpha h_i \mathbf{e}_i) \\ &= \sum_{i=1}^n \eta_i h_i \leq \sum_{i=1}^n |h_i| |\eta_i| \leq \left(\sum_{i=1}^n |\eta_i| \right) \|\mathbf{h}\| \end{aligned} \quad (1.7)$$

for $i = 1, \dots, n$. The existence of the partial derivatives implies that each of the $\eta_i \rightarrow 0$ as $\|\mathbf{h}\| \rightarrow 0$. Hence,

$$0 = \eta(\mathbf{h}) = \eta\left(\frac{\mathbf{h} + (-\mathbf{h})}{2}\right) \leq \frac{\eta(\mathbf{h}) + \eta(-\mathbf{h})}{2}.$$

Hence

$$\eta(\mathbf{h}) \geq -\eta(-\mathbf{h}) \geq -\left(\sum_{i=1}^n |\eta_i|\right) \|\mathbf{h}\|. \quad (1.8)$$

where, for each $i = 1, \dots, n$, $x_i^j \rightarrow 0$ as $\|h\| \rightarrow 0$. We now obtain (14) from (15) and (16).

Exercise 1.1. Show that if f and g are convex functions defined on a convex set C , then for any $\lambda > 0$, $\mu > 0$, the function $\lambda f + \mu g$ is convex.

Exercise 1.2. Let $\{f_n\}$ be a sequence of convex functions defined on a convex set C . Show that if $f(x) = \lim_{n \rightarrow \infty} f_n(x)$ exists for all $x \in C$, then f is convex.

Exercise 1.3. Prove Lemma 1.2.

Exercise 1.4. Prove Lemma 1.4.

Exercise 1.5. Let $J = [a, b]$ be a real interval and let C be a convex set in $\mathbb{R}^n$. Let $f(x, t) = f(x, t)$ be a real-valued function defined on $J \times C$ that is continuous on J for each fixed x in C and is convex on C for each fixed t in J . Show that

$$F(x) = \int_a^b f(x, t) dt$$

is defined for each x in C and that F is convex on C .

Exercise 1.6. (a) Show that any norm $\|\cdot\|$ on $\mathbb{R}^n$ is convex on $\mathbb{R}^n$.

(b) Let S be a nonempty convex set, and let $\|\cdot\|$ denote the euclidean norm in $\mathbb{R}^n$. Let f be the distance function to S defined by

$$f(x) = \inf\{\|y - x\| : y \in S\}.$$

Show that f is convex on $\mathbb{R}^n$.

(c) Show that f is Lipschitz continuous on compact subsets of $\mathbb{R}^n$.

Exercise 1.7. Let f be a convex function defined on $\mathbb{R}^n$. Let x_0 be a fixed point in $\mathbb{R}^n$. Show that the function $\lambda f(\cdot) + f(x_0)$ defined by $\lambda f(\cdot) + f(x_0)$ is (a) positively homogeneous [i.e., for $\lambda > 0$, $\lambda f(x) = \lambda f(x)$] and (b) convex.

Exercise 1.8. Let g be a convex function defined on $\mathbb{R}^n$. Let $S = \{x : g(x) > 0\}$ be nonempty.

(a) Show that S is convex.

(b) Show that if $f(x) = 1/g(x)$ for x in S , then f is convex on S .

Exercise 1.8. Let C be a convex cone (see Exercise II.2.10 for the definition). A real-valued function g defined on C is said to be a *gauge function* on C if

- (a) $g(px) = pg(x)$ for all x in C and $p \geq 0$ and
- (b) $g(x + y) \leq g(x) + g(y)$ for all x, y in C .

Show that g is convex.

Exercise 1.9. Let C be a convex set. The *support function* s of C is the real-valued function defined by

$$s(\mathbf{n}) = \sup\{(\mathbf{n}, \mathbf{x}) : \mathbf{x} \in C\}$$

for all $\mathbf{n}$ such that the supremum is finite.

- (a) Show that if the domain D of s is nonempty, then D is in a convex cone.
- (b) Show that s is a gauge function on D .

Exercise 1.11. Let C be a compact convex set with $\dim(C) \geq 1$. The *Minkowski distance function* p determined by C is defined by

$$p(\mathbf{x}) = \inf\{\lambda : \lambda \geq 0, \mathbf{x} \in \lambda C\}.$$

Show that p is a gauge function.

Exercise 1.12. Let C be a compact convex set, and let s be its support function. Show that the following statements hold:

- (a) The domain of s is $\mathbb{R}^n$.
- (b) For each $\mathbf{n} \neq \mathbf{0}$ there exists a point $\mathbf{x}_n$ in C such that $s(\mathbf{n}) = (\mathbf{n}, \mathbf{x}_n)$.
- (c) The hyperplane $H_n = \{\mathbf{x} : (\mathbf{n}, \mathbf{x}) = s(\mathbf{n})\}$ supports C at $\mathbf{x}_n$.
- (d) The distance from H_n to $\mathbf{0}$ is equal to $s(\mathbf{n})/|\mathbf{n}|$.

Exercise 1.13. Let C be a nonempty closed convex set $C \neq \mathbb{R}^n$. Let s be the support function of C with nonempty domain D . Show that

$$C = \{\mathbf{x} : (\mathbf{n}, \mathbf{x}) \leq s(\mathbf{n}) \text{ for all } \mathbf{n} \text{ in } D\}.$$

Hint: An alternate definition of C is

$$C = \bigcap_{\mathbf{n} \in D} \{\mathbf{x} : (\mathbf{n}, \mathbf{x}) \leq s(\mathbf{n}), \mathbf{x} \text{ fixed in } B_1\},$$

and see Exercise II.6.6.

2. SUBGRADIENTS

If f is a real-valued function defined on a set C in $\mathbb{R}^n$, then the graph of f is the set of points (x, z) in $\mathbb{R}^{n+1}$ of the form $(x, f(x))$, where $x \in C$. If f is a differentiable function, then given a point x_0 in the interior of C , the graph of f has a tangent plane at $(x_0, f(x_0))$ whose normal vector is $(\nabla f(x_0), -1)$.

If C is convex and f is convex, f need not be differentiable at an interior point of C , as the function f defined by $f(x) = |x|$ shows. In this section we shall introduce a generalization of the gradient vector, called the subgradient, that exists at points $(x, f(x))$ of the graph of a convex function f , where x is an interior point of C . We shall show that if f is a convex function defined on C and x is an interior point of C , then $\text{epi } f$ has a supporting hyperplane at the point $(x, f(x))$. The supporting hyperplane need not be unique. In the next section we shall show that f is differentiable at x if and only if f has a unique subgradient at x . This subgradient turns out to be $\nabla f(x)$. Thus, if f is differentiable at x , there is a unique supporting hyperplane to $\text{epi } f$ at $(x, f(x))$. This hyperplane is the tangent plane to the graph at $(x, f(x))$. If $(\bar{q}, -1)$ is a normal vector to a supporting hyperplane, then $\bar{q}$ has some of the properties that the gradient vector would have.

We now define the notion of subgradient.

Definition 2.1. A function f defined on a set C is said to have a subgradient at a point x_0 in C if there exists a vector $\bar{q}$ in $\mathbb{R}^n$ such that

$$f(x) \geq f(x_0) + \langle \bar{q}, x - x_0 \rangle$$

for all $x \in C$. The vector $\bar{q}$ is called a subgradient.

Geometrically, $\bar{q}$ is a subgradient of f at x_0 if the graph of f in $\mathbb{R}^{n+1}$ lies on or above the graph of the hyperplane $z = f(x_0) + \langle \bar{q}, x - x_0 \rangle$. Since $(x_0, f(x_0))$ is in this hyperplane, this hyperplane will be a supporting hyperplane to $\text{epi } f$ at $(x_0, f(x_0))$. Thus, the existence of a subgradient is a statement about the existence of a nonvertical supporting hyperplane to $\text{epi } f$. In Figure 2.1 we illustrate the definition with the function $f(x) = |x|$ at $x = 0$.

Any line through the origin with slope $\bar{c}$ satisfying $-1 \leq \bar{c} \leq 1$ has the property that

$$|x| \geq \bar{c}x = \bar{c}|x| + \bar{c}|x| - \bar{c}|x|$$

for all x . Thus any $\bar{c}$ in the interval $[-1, 1]$ is a subgradient of $|x|$ at $x = 0$. If $x_0 > 0$, then $\bar{c} = 1$ is the only subgradient of $|x|$ at x_0 ; if $x_0 < 0$, then $\bar{c} = -1$ is the only subgradient of $|x|$ at x_0 .

Definition 2.2. The set of subgradients of a function f at a point x_0 is called the subdifferential of f at x_0 and is denoted by $\partial f(x_0)$.

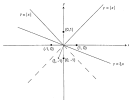


Figure 1b.

The subdifferential of $|x|$ at $x = 0$ is the interval $[-1, 1]$. The subdifferential of $|x|$ at $x_0 > 0$ is the singleton $\{1\}$, and the subdifferential at $x_0 < 0$ is $\{-1\}$.

Example 1f. The subdifferential will be the empty set if no subgradient exists. This may happen, as the following example shows. Let $C = [-1, 1]$ and let $f(x) = -\sqrt{1-x^2}$. (Sketch the graph.) The graph of f has a nonvertical tangent line at any point x_0 in $(-1, 1)$. From the graph it appears that for such x_0 the graph lies above the tangent line, and so $\partial f(x_0) = \{f'(x_0)\}$. Although this can be verified in this example by relatively straightforward calculations, we shall not do so. The validity of the statement will follow from general theorems that we will prove. At the points $x = \pm 1$, the graph has a vertical tangent line, but the tangent line does not lie below the graph. Thus it appears that f does not have a subgradient at these points. We now show this analytically.

A number ξ is a subgradient of f at $x = -1$ if and only if

$$-\sqrt{1-x^2} \geq \xi - \xi(x+1) \quad \text{for all } x \text{ in } [-1, 1].$$

For x in $(-1, 1)$ this inequality is equivalent to

$$-\left(\frac{1-x^2}{1+x}\right)^{1/2} \geq \xi - \xi(x+1), \quad x \text{ in } (-1, 1).$$

If we let x in $(-1, 1]$ tend to -1 , the left-hand side of this inequality tends to $-\infty$. Hence there is no $\ell \in \mathbb{R}$ such that the subgradient inequality holds for all x in $(-1, 1]$.

A similar argument shows that $\partial f(1) = \emptyset$.

THEOREM 2.1. Let f be a convex function defined on a convex set C . Then at each point $x_0 \in \text{int}(C)$ there exists a vector $\bar{q}$ in $\mathbb{R}^n$ such that

$$f(x) \geq f(x_0) + \langle \bar{q}, x - x_0 \rangle \quad (1)$$

for all x in C . In other words, $\partial f(x_0) \neq \emptyset$ at every interior point x_0 of C .

Example 2.1 shows that the requirement $x_0 \in \text{int}(C)$ is essential.

Proof. We shall show that the conclusion of the theorem is a consequence of the existence of a supporting hyperplane to $\text{epi}(f)$ at $(x_0, f(x_0))$.

The point $(x_0, f(x_0))$ is a boundary point of $\text{epi}(f)$. Since $\text{epi}(f)$ is convex, we may apply Theorem (1.4.1) with $C = \text{epi}(f)$ and get that there exists a vector $(\bar{\alpha}, \bar{\beta}) \neq (0, 0)$ in $\mathbb{R}^{n+1}$ such that

$$\langle \bar{\alpha}, \bar{\beta}, (x, y) \rangle \leq \langle \bar{\alpha}, \bar{\beta}, (x_0, f(x_0)) \rangle$$

for all (x, y) in $\text{epi}(f)$. This relation is equivalent to

$$\langle \bar{\alpha}, x \rangle + \bar{\beta} y \leq \langle \bar{\alpha}, x_0 \rangle + \bar{\beta} f(x_0) \quad (2)$$

for all (x, y) in $\text{epi}(f)$, that is, for all x in C and $y = f(x)$.

Since the right-hand side of (2) is a constant, if $\bar{\beta}$ were strictly positive, we could obtain a contradiction by letting y become arbitrarily large. Hence $\bar{\beta} \leq 0$.

If $\bar{\beta} = 0$, then $\bar{\alpha} \neq 0$ and (2) becomes

$$\langle \bar{\alpha}, x \rangle \leq \langle \bar{\alpha}, x_0 \rangle \quad \text{all } x \in C. \quad (3)$$

Since $x_0 \in \text{int}(C)$, there exists an $\varepsilon > 0$ such that $x_1 = x_0 + \varepsilon \bar{\alpha}$ belongs to C . Setting $x = x_1$ in (3) gives $\varepsilon \langle \bar{\alpha}, \bar{\alpha} \rangle \leq 0$, which is a contradiction. Thus, $\bar{\beta} < 0$.

In (2) we now divide through by $\bar{\beta} < 0$ and set $y = f(x)$ to get

$$f(x) \geq f(x_0) + \left\langle \frac{\bar{\alpha}}{\bar{\beta}}, x_0 - x \right\rangle$$

for all $x \in C$. Now set $\bar{q} = -\bar{\alpha}/\bar{\beta}$ to get (1).

COROLLARY 1. If f is strictly convex, then strict inequality holds in (1).

We leave the proof as an exercise.

The converse of Theorem 2.1 is false. To see this, let C be the square in $\mathbb{R}^2$, $\{x \mid 0 \leq x_1 \leq 1, 0 \leq x_2 \leq 1\}$. Let $f(x) = 0$ for x such that $0 \leq x_1 \leq 1$ and $0 < x_2 < 1$. Let $f(x) = \frac{1}{2} - |x_1 - \frac{1}{2}|$ for x such that $0 \leq x_1 \leq 1$ and $x_2 = 0$. The function f has a unique subgradient $\xi = 0$ at every point in $\text{int}(C)$, yet f is not convex on C . The function ξ however, is convex on $\text{int}(C)$, and this phenomenon illustrates the general state of affairs.

THEOREM 2.2. Let C be a convex set in $\mathbb{R}^n$. If at each point x_0 in C f has a subgradient, then f is convex on C .

REMARK 2.1. If we merely assume that at each point in $\text{int}(C)$ the subdifferential is not empty, then we can only conclude that f is convex on $\text{int}(C)$.

Proof. Let x and y be two points in C and let $x = \alpha x_1 + \beta y$, $\alpha > 0$, $\beta > 0$, $\alpha + \beta = 1$. Since f has a subgradient at x , we get

$$f(x) \geq f(x) + \langle \xi_x, x - x \rangle = f(x) + \beta \langle \xi_x, x - y \rangle,$$

$$f(y) \geq f(y) + \langle \xi_y, y - x \rangle = f(x) - \alpha \langle \xi_y, x - y \rangle.$$

If we now multiply the first inequality by $\alpha > 0$, the second inequality by $\beta > 0$, and add the results, we get

$$\alpha f(x) + \beta f(y) \geq f(x) = f(\alpha x + \beta y).$$

Hence f is convex.

REMARK 2.2. The proof of the theorem shows that if at each x_0 in C there exists a subgradient ξ such that $f(x) \geq f(x_0) + \langle \xi, x - x_0 \rangle$, for $x \neq x_0$, x in C , then f is strictly convex on C .

EXERCISE 2.1. Show that if $\partial f(x_0)$ is not empty, then $\partial f(x_0)$ is closed and convex.

EXERCISE 2.2. Prove Corollary 1 to Theorem 2.1.

EXERCISE 2.3. Let f be a convex function defined on a convex set C and let $x_0 \in \text{int}(C)$. Show that $\xi \in \partial f(x_0)$ if and only if $f(x_0 + v) \geq \langle \xi, v \rangle$ for all v in $\mathbb{R}^n$.

EXERCISE 2.4. Let f be convex and concave on $\mathbb{R}^n$. Show that f is affine, that is, $f(x) = \langle a, x \rangle + b$.

EXERCISE 2.5. Let $f(x) = |x|$. Show that $\partial f(0) = \{\xi \mid \|\xi\| \leq 1\}$. (From Theorem 2.1 we will get that, for $x \neq 0$, $\partial f(x) = \{\lambda(x)/|x|\}$.)

Exercise 14. Let $f(x_1, x_2) = |x_1 + x_2|$. Show that f is convex and calculate $\partial f(x, y)$.

3. DIFFERENTIABLE CONVEX FUNCTIONS

We first show that if f is a convex function and differentiable [in the sense of relation (5) in Section 11 of Chapter I] at an interior point of its domain, then there is a unique subgradient at the point and the subgradient is the gradient vector, and conversely, if f has a unique subgradient at x , then f is differentiable at x .

THEOREM 3.1. Let f be a convex function defined on a convex set C in E^n , and let x be an interior point of C . Then f is differentiable at x if and only if f has a unique subgradient at x . Moreover, the unique element in the subgradient is $\nabla f(x)$.

Proof. Let f be differentiable at x . Since f is convex and $x \in \text{int}(C)$, by Theorem 2.1 the subgradient $\partial f(x)$ is not empty. Let $\xi \in \partial f(x)$. Then for any h in E^n and sufficiently small $\varepsilon > 0$

$$f(x + \varepsilon h) \leq f(x) + \varepsilon \langle \xi, h \rangle.$$

Since f is differentiable at x , we have that for $h \neq 0$

$$f(x + \varepsilon h) = f(x) + \varepsilon \langle \nabla f(x), h \rangle + o(\varepsilon),$$

where $o(\varepsilon)/\varepsilon \rightarrow 0$ as $\varepsilon \rightarrow 0$. If we subtract the second relation from the first, we get

$$0 \leq \varepsilon \langle \xi, h \rangle - \varepsilon \langle \nabla f(x), h \rangle + o(\varepsilon) = \varepsilon [\langle \xi - \nabla f(x), h \rangle + o(1)].$$

If we now divide by $\varepsilon > 0$ and then let $\varepsilon \rightarrow 0$, we get that $\langle \xi - \nabla f(x), h \rangle \leq 0$ for all $h \neq 0$ in E^n . In particular, if $\xi - \nabla f(x) \neq 0$, then this inequality holds for $h = \xi - \nabla f(x)$, whence $\langle \xi - \nabla f(x), \xi - \nabla f(x) \rangle \leq 0$. Thus, $\xi = \nabla f(x)$.

We now suppose that the subgradient $\partial f(x)$ consists of a unique element ξ . We will show that all the first-order partial derivatives of f exist at x . The differentiability of f at x will then follow from Theorem 1.4.

To simplify the notation, we assume that $x = 0$ and that $f(x) = 0$. There is no loss of generality in making this assumption, since this amounts to a translation of the origin in $E^n \times \mathbb{R}$ to $(x, f(x))$. Let us now suppose that at least one partial derivative of f at 0 does not exist. For definiteness let us assume that it is the partial with respect to x_1 . Let

$$g(t) = f(0 + tx_1).$$

where e_1 is the first standard basis vector $(1, 0, \dots, 0)$. Since f is convex and $\Phi_1 \in C_1$, the function g is defined and convex on some maximal interval I containing $\bar{x}$ in its interior. Moreover, $g'(x)$ does not exist at $\bar{x}$ if and only if $g''(\bar{x})$ does not exist. Since g is convex, it follows from Corollary 2 to Theorem 1.2 that $g'_-(\bar{x})$ and $g'_+(\bar{x})$ exist and that $g'_-(\bar{x}) \leq g'_+(\bar{x})$. Therefore, if $g'(\bar{x})$ does not exist, then $g'_-(\bar{x}) < g'_+(\bar{x})$.

Let ζ_1 be any number satisfying

$$g'_-(\bar{x}) < \zeta_1 < g'_+(\bar{x}). \quad (1)$$

We shall show that corresponding to each such ζ_1 we can find a $\bar{h}$ in $J \cap \Phi_1$ with distinct ζ_1 determining distinct $\bar{h}$. This will contradict the hypothesis that $J \cap \Phi_1$ consists of a unique element, and the proof will be accomplished.

It follows from (1) and Corollary 1.1 that for $h > 0$

$$\frac{g(\bar{x} + h) - g(\bar{x})}{h} < g'_-(\bar{x}) < \zeta_1 < g'_+(\bar{x}) < \frac{g(\bar{x}) - g(\bar{x} - h)}{h}.$$

Since $g(\bar{x}) = f(\bar{x}) = 1$, we get from the preceding inequalities that for all h in J_+

$$f(\bar{x} + h) = g(h) > \zeta_1 h, \quad (2)$$

with equality holding only if $h = 0$. Let V_1 denote the one-dimensional subspace spanned by e_1 . Then $V_1 = \{x|x = \alpha e_1, \alpha \in \mathbb{R}\}$. Define a linear functional L_1 on V_1 by the formula

$$L_1(x) = L_1(\alpha e_1) = \zeta_1 \alpha. \quad (3)$$

Since the representation $x = \alpha e_1$ is unique for x in V_1 , the linear functional L_1 is well defined. Note that L_1 depends on ζ_1 . Combining (2) and (3), we get that for x in $V_1 \cap C$

$$f(x) \geq L_1(x) = \zeta_1 \alpha. \quad (4)$$

If $n = 1$, we are through. Otherwise, let V_2 denote the two-dimensional subspace spanned by e_1 and e_2 . Then

$$V_2 = \{y|y = \alpha_1 e_1 + \alpha_2 e_2, \alpha_i \in \mathbb{R}\} = \{y|y = x + \alpha_2 e_2, x \in V_1, \alpha_2 \in \mathbb{R}\}.$$

We shall extend L_1 to a linear functional L_2 defined on V_2 and shall find a number ζ_2 such that for all y in $V_2 \cap C$

$$f(y) \geq L_2(y) = \zeta_2 \alpha + \zeta_2 \alpha_2. \quad (5)$$

Let x and y belong to $V_1 \cap C$ and let $\alpha > 0$, $\beta > 0$. It then follows from the linearity of L_α , from the convexity of $V_1 \cap C$, from (4), and from the convexity of f that

$$\begin{aligned} \frac{\alpha}{\alpha+\beta} L_\alpha(x) + \frac{\beta}{\alpha+\beta} L_\beta(y) &= L_{\alpha+\beta} \left(\frac{\alpha}{\alpha+\beta} x + \frac{\beta}{\alpha+\beta} y \right) \\ &\leq_{(2)} f \left(\frac{\alpha}{\alpha+\beta} x + \frac{\beta}{\alpha+\beta} y \right) \\ &= f \left(\frac{\alpha}{\alpha+\beta} (x - \beta w) + \frac{\beta}{\alpha+\beta} (y + \alpha w) \right) \\ &\leq_{(3)} \frac{\alpha}{\alpha+\beta} f(x - \beta w) + \frac{\beta}{\alpha+\beta} f(y + \alpha w). \end{aligned}$$

Multiplying through by $\alpha + \beta$ gives

$$\alpha L_\alpha(x) + \beta L_\beta(y) \leq \alpha f(x - \beta w) + \beta f(y + \alpha w)$$

and so

$$\frac{L_\alpha(x) - f(x - \beta w)}{\beta} \leq \frac{f(y + \alpha w) - L_\beta(y)}{\alpha}. \quad (5)$$

Denote the left-hand side of (5) by $g(x, \beta)$ and denote the right-hand side of (5) by $h(y, \alpha)$. Hence

$$\sup\{g(x, \beta) : x \in V_1 \cap C, \beta > 0\} \leq \inf\{h(y, \alpha) : y \in V_1 \cap C, \alpha > 0\}$$

and both the supremum and infimum are finite. Hence there exists a number L_∞ such that

$$\frac{L_\alpha(x) - f(x - \beta w)}{\beta} \leq L_\infty \leq \frac{f(y + \alpha w) - L_\beta(y)}{\alpha} \quad (7)$$

for all $x \in V_1 \cap C$, $\alpha > 0$, $\beta > 0$.

Now let y be an element in $V_1 \cap C$ that is not in V_∞ . Then $y = x + \alpha w$, with $x \in w_\perp$ and $\alpha > 0$. If $\alpha > 0$, take $\alpha = w$ in (7). If $\alpha < 0$, take $\beta = -\alpha$ in (7). Then from (7) and (5) we get that for all y in $V_0 \cap C$

$$f(y) - f(x + \alpha w) \geq L_\infty(x) + L_\infty \alpha = L_\infty x + L_\infty \alpha,$$

which is (2).

Proceeding inductively, we get that for each i_1 satisfying (1) there exists a vector $\xi = (\xi_0, \xi_1, \dots, \xi_n)$ with the fact that if $\mathbf{x} = (x_1, \dots, x_n)$ is in C , then

$$f(\mathbf{x}) \geq \xi_0 + \xi_1 x_1$$

Since $f(\mathbf{0}) = 0$, this says that $\xi_0 \leq f'(\mathbf{0})$.

The next theorem is an immediate consequence of Theorems 3.1, 3.3, and 3.4.

THEOREM 3.2. *Let C be an open convex set in $\mathbb{R}^n$ and let f be real valued and differentiable on C . Then f is convex if and only if for each $\mathbf{x}_0$ in C*

$$f(\mathbf{x}) \geq f(\mathbf{x}_0) + \langle \nabla f(\mathbf{x}_0), \mathbf{x} - \mathbf{x}_0 \rangle$$

for all $\mathbf{x}$ in C . Also, f is strictly convex if and only if for each $\mathbf{x}_0$ in C and all $\mathbf{x} \neq \mathbf{x}_0$ in C

$$f(\mathbf{x}) > f(\mathbf{x}_0) + \langle \nabla f(\mathbf{x}_0), \mathbf{x} - \mathbf{x}_0 \rangle.$$

Remark 3.2. If $n = 2$ and f is differentiable at $\mathbf{x}_0$, then the surface in $\mathbb{R}^3$ defined by $y = f(\mathbf{x})$ has a tangent plane at $\mathbf{x}_0$ whose equation can be written as

$$y - f(\mathbf{x}_0) = \langle \nabla f(\mathbf{x}_0), \mathbf{x} - \mathbf{x}_0 \rangle.$$

Thus, the geometric interpretation of Theorem 3.2 is that a function f that is differentiable on C is convex on C if and only if for each point $\mathbf{x}_0$ in C the surface defined by $y = f(\mathbf{x})$ lies above the tangent plane to the surface at $\mathbf{x}_0$. This is not surprising since we obtained the result in Theorem 3.2 by considering the supporting hyperplanes to $\text{epi}(f)$, and for differentiable f these are the tangent planes to the surface.

Let f be a real-valued function defined and of class $C^{(2)}$ on an open set B in $\mathbb{R}^n$. (See Section 11 in Chapter I for the definition of class $C^{(2)}$.) For such functions we have a computationally handy criterion for convexity. In the ensuing discussion recall that, by Theorem 3.1.1, f is differentiable at each point $\mathbf{x}_0$ in B .

We first develop our criterion if B is an interval $I \subset \mathbb{R}$.

LEMMA 3.1. *Let f be of class $C^{(2)}$ on an open interval $I \subset \mathbb{R}$. Then f is convex on I if and only if $f''(x) \geq 0$ for all x in I . If $f''(x) > 0$ or all x in I , then f is strictly convex.*

Note that if f is strictly convex, then we cannot conclude that $f''(x) > 0$ for all x , as the example $f(x) = x^4$ shows. The function f is strictly convex, yet $f''(0) = 0$.

Proof. If f is convex on I , then by Corollary 1 to Theorem 3.2, $f''(x) \geq 0$ on I . Conversely, suppose $f''(x) \geq 0$ on I . Let x_0 be a fixed point in I and let x be any other point in I . Then by Taylor's theorem

$$f(x) - f(x_0) - f'(x_0)(x - x_0) = \frac{1}{2}f''(x_1)(x - x_0)^2 \geq 0,$$

where x_1 is a point between x_0 and x . It now follows from Theorem 3.2 that f is convex on I . If $f''(x) > 0$ on I it follows from Theorem 3.2 that f is strictly convex.

Recall that a symmetric $n \times n$ matrix A with entries (a_{ij}) is said to be positive semidefinite if the quadratic form

$$\mathbf{x}'A\mathbf{x} = (\mathbf{x}, A\mathbf{x}) = \sum_{j=1}^n a_{ij}x_j$$

is nonnegative for all $\mathbf{x}$ in $\mathbb{R}^n$. The matrix A is said to be positive definite if $(\mathbf{x}, A\mathbf{x})$ is positive for all $\mathbf{x} \neq \mathbf{0}$ in $\mathbb{R}^n$.

THEOREM 3.3. Let f be of class $C^{(2)}$ on an open convex set D . Then f is convex on D if and only if the Hessian matrix

$$H(\mathbf{x}) = \left(\frac{\partial^2 f}{\partial x_i \partial x_j}(\mathbf{x}) \right)$$

is positive semidefinite at each point $\mathbf{x}$ in D . If $H(\mathbf{x})$ is positive definite at each $\mathbf{x}$, then f is strictly convex.

Proof. Let $\mathbf{x}$ be a point in D . Since D is open, for each vector $\mathbf{v}$ in $\mathbb{R}^n$, $\mathbf{v} \neq \mathbf{0}$, there exists an $\epsilon > 0$ such that the line segment $(\mathbf{x} - \epsilon\mathbf{v}, \mathbf{x} + \epsilon\mathbf{v})$ is in D . Hence the function $\varphi: (s, t) \mapsto f(\mathbf{x} + s\mathbf{v})$ in relation (3) in Section 1 is defined on $(-\epsilon, \epsilon)$. Also, since f is of class $C^{(2)}$ on D , $\partial^2 f(\mathbf{x})/\partial x_i \partial x_j = \partial^2 f(\mathbf{x})/\partial x_j \partial x_i$ at each $\mathbf{x}$ in D , and so $H(\mathbf{x})$ is symmetric.

If f is convex, then by Lemma 3.1, the function $\varphi: (s, t) \mapsto f(\mathbf{x} + s\mathbf{v})$ is convex on the interval $(-\epsilon, \epsilon)$. By Lemma 3.1, $\varphi''(t) = \varphi''(0) \geq 0$. If $\mathbf{v} = (v_1, \dots, v_n)$, then straightforward calculations as in relations (12) and (11) in Section 11 of Chapter 1 give

$$\begin{aligned} \varphi''(0) &= \frac{d}{ds} f(\mathbf{x} + s\mathbf{v}) = \sum_{j=1}^n \frac{\partial f}{\partial x_j}(\mathbf{x} + s\mathbf{v})v_j \\ \varphi''(0) &= \sum_{i=1}^n \sum_{j=1}^n \frac{\partial^2 f}{\partial x_i \partial x_j} f(\mathbf{x} + s\mathbf{v})v_i v_j = (\mathbf{v}, H(\mathbf{x} + s\mathbf{v})\mathbf{v}). \end{aligned} \quad (4)$$

Setting $v = 0$ gives

$$(x, H(x, 0)) = \varphi''(x, 0) \geq 0. \quad (9)$$

Since (9) holds for all x , $H(x)$ is positive semidefinite.

Conversely, suppose now that $H(x)$ is positive semidefinite for each x in D . Then, from (8), for each x in D and for each v in $\mathbb{R}^n$, $\varphi''(x, v) \geq 0$ for all v in $\{-v, v\}$. Hence by Lemma 3.1, φ' is concave. Since x and v are arbitrary, it follows from Lemma 1.5 that f is concave on D . The argument also shows that if $H(x)$ is positive definite for all x in D , then f is strictly concave.

We assume that the reader is familiar with the following facts. Given a symmetric matrix A , there exists an orthogonal matrix P such that $P^TAP = D$, where D is a diagonal matrix whose diagonal elements are the eigenvalues of A . Thus A is positive definite if and only if all the diagonal entries of D (eigenvalues of A) are positive and A is positive semidefinite if and only if they are nonnegative. For a survey of numerical methods for obtaining the eigenvalues of a symmetric matrix see Golub and Van Loan [1996].

Let A be an $n \times n$ matrix with entries a_{ij} . Let

$$A_k = A_{kk}, \quad A_k = \det \begin{pmatrix} a_{11} & \cdots & a_{kk} \\ \vdots & \ddots & \vdots \\ a_{k1} & \cdots & a_{kk} \end{pmatrix}, \quad k = 1, 2, 3, \dots, n.$$

The determinants A_k are called the principal minors of A . Another criterion for positive definiteness of a symmetric matrix A with entries a_{ij} is the following.

Lemma 3.2. *The symmetric matrix A is positive definite if and only if $A_k > 0$ for $k = 1, \dots, n$ and is negative definite if and only if $(-1)^k A_k > 0$ for $k = 1, \dots, n$.*

We shall prove this lemma in Exercise IV.2.16.

If we assume that $A_k \geq 0$ for all $k = 1, 2, \dots, n$ and not all the $A_k = 0$, then there is difference between $n \geq 1$ and $n = 2$.

For $n \geq 3$, if $A_k \geq 0$, $k = 1, \dots, n$, and not all $A_k = 0$, then A need not be positive semidefinite. To see this, consider the matrix

$$A = \begin{pmatrix} 1 & 0 & \cdots & 0 \\ 0 & 0 & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & -1 \end{pmatrix}.$$

Then $(x, Ax) = x_1^2 - x_n^2$ which is not nonnegative for all x .

For $n = 2$, however, if $A_{11} > 0$, $A_{22} > 0$, and not both A_{11} and A_{22} are zero, then the symmetric matrix A is positive semidefinite. To see this, let

$$A = \begin{pmatrix} a_{11} & a \\ a & a_{22} \end{pmatrix}.$$

If $A_{11} = 0$, then $A_{22} = -a^2$. Then to have $A_{22} > 0$, we must have $a = 0$. But then $A_{22} = 0$. Hence we must have $a_{11} > 0$. Thus $A_{11} > 0$. Since $A_{22} > 0$,

$$a_{11}a_{22} > a^2. \quad (15)$$

Since $a_{22} > 0$, it follows that $a_{11} > 0$. For $\mathbf{x} = (x_1, x_2)$

$$\langle \mathbf{x}, A\mathbf{x} \rangle = a_{11}x_1^2 + 2ax_1x_2 + a_{22}x_2^2.$$

If $\mathbf{x} = (x_1, 0)$ with $x_1 \neq 0$, we get

$$\langle \mathbf{x}, A\mathbf{x} \rangle = a_{11}x_1^2 > 0.$$

If $\mathbf{x} = (x_1, x_2)$ with $x_2 \neq 0$, we have $x_2 = cx_1$ for some $-\infty < c < \infty$ and

$$\langle \mathbf{x}, A\mathbf{x} \rangle = [a_{11}c^2 + 2ac + a_{22}]x_1^2.$$

Denote the term in square brackets by $q(c)$. From (15) we get that the discriminant of q , which equals $4a^2 - 4a_{11}a_{22}$, is nonpositive. Thus the quadratic q either has a double root or has no real roots. Since $a_{11} > 0$, it follows that $q(c) \geq 0$ for all c . Hence, $\langle \mathbf{x}, A\mathbf{x} \rangle \geq 0$ for all $\mathbf{x}$.

Exercise 3.1. Determine whether or not the following functions are convex on $\mathbb{R}^2$:

(i) $f(x, y, z) = x^2 + 2xy + 4xz + 3y^2 + yz + 3z^2$,

(ii) $f(x, y, z) = x^2 + 4xy + 4y^2 + 2xz + 4yz$.

Exercise 3.2. Determine for convex set in $\mathbb{R}^2$ on which the function

$$f(x, y) = x^2 - 2xy + \frac{1}{2}y^2 - 3y$$

is (i) convex and (ii) strictly convex.

Exercise 3.3. For what values of r is the function $f(x) = x^r - r \ln x$ convex on $(0, \infty)$.

Exercise 3.4. Show that $f(\mathbf{x}) = x_1^2 + \cdots + x_n^2$, $n \geq 1$, is convex on $\{\mathbf{x} : \mathbf{x} \geq \mathbf{0}\}$.

Exercise 3.4. Show that the n -dimensional ellipsoid

$$\left\{ \mathbf{x} : \sum_{i=1}^n \left(\frac{x_i}{a_i} \right)^2 \leq 1 \right\}$$

is convex. *Hint:* There is a very easy solution.

Exercise 3.6. For what values of p is $f(\mathbf{x}) = \|\mathbf{x}\|^p$ convex on $\mathbb{R}^n$. *Hint:* Use Lemma 3.4.

Exercise 3.7. Show that if $c_i > 0$, $i = 1, \dots, 4$, the function

$$f(\mathbf{x}) = \exp(c_1 x_1^2) + c_2 x_1^{-1} + c_3 x_1^2 + c_4 x_1^{-1}$$

is convex on $\bar{S} = \{(\mathbf{x}_1, \mathbf{x}_2) : x_1 > 0, x_2 > 0\}$.

Exercise 3.8. Show that the function $f(\mathbf{x}) = (1 + \|\mathbf{x}\|^2)^{p/2}$, $p \geq 1$, is convex on $\mathbb{R}^n$.

Exercise 3.9. Show that the function $f(\mathbf{x}) = -(x_1 + x_2 + \dots + x_n)^{1/p}$, $n \geq 1$, is convex on $\bar{S} = \{\mathbf{x} : \mathbf{x} \geq \mathbf{0}\}$.

Exercise 3.10. A function f that is positive on a convex set C in $\mathbb{R}^n$ is said to be *logarithmically convex* if the function $g(\mathbf{x}) = \ln f(\mathbf{x})$ is convex on C .

(a) Show that if f is logarithmically convex, then f is convex.

(b) Is it true that every positive convex function is logarithmically convex?

Exercise 3.11. Let f be differentiable on an open convex set X . Show that if f is strictly convex, then $\nabla f(\mathbf{x}) \neq \nabla f(\mathbf{y})$ for all distinct pairs of points $\mathbf{x}, \mathbf{y}$ in X .

Exercise 3.12. Let C be an open convex set and let f be real valued and differentiable on C . Prove that f is convex on C if and only if for each pair of points $\mathbf{x}_0$ and $\mathbf{x}_1$ in C

$$\langle \nabla f(\mathbf{x}_1) - \nabla f(\mathbf{x}_0), \mathbf{x}_1 - \mathbf{x}_0 \rangle \geq 0.$$

Show that f is strictly convex if and only if strict inequality holds.

4. ALTERNATIVE THEOREMS FOR CONVEX FUNCTIONS

The theorems of this section are alternative theorems for systems of inequalities involving convex and affine functions. They find application in optimization problems involving such functions.

Definition 4.1. A vector-valued function $\mathbf{f} = (f_1, \dots, f_m)$ with domain a convex set C in $\mathbb{R}^n$ and range in $\mathbb{R}^m$ is said to be convex if each of the real-valued component functions f_i , $i = 1, \dots, m$, is convex.

Affine functions or affine transformations were defined in Definition B.3.6. It was shown there that an affine function $\mathbf{h}$ with domain $\mathbb{R}^n$ and range $\mathbb{R}^m$ has the form

$$\mathbf{h}(\mathbf{x}) = A\mathbf{x} + \mathbf{b},$$

where A is an $m \times n$ matrix and $\mathbf{b}$ is a vector in $\mathbb{R}^m$.

Note that $\mathbf{h} = (h_1, \dots, h_m)$, where $h_i = (\mathbf{a}_i, \mathbf{b}_i) \cdot \mathbf{x} + b_i$ and $\mathbf{a}_i$ is the i th row of A .

Theorem 4.1. Let C be a convex set in $\mathbb{R}^n$ and let $\mathbf{f} = (f_1, \dots, f_m)$ be defined and convex on C . Let $\mathbf{h} = (h_1, \dots, h_m)$ be an affine mapping from $\mathbb{R}^n$ to $\mathbb{R}^m$. If

$$f(\mathbf{x}) < 0, \quad \mathbf{h}(\mathbf{x}) = \mathbf{0} \quad (1)$$

has no solution $\mathbf{x}$ in C , then there exist a vector $\mathbf{p}$ in $\mathbb{R}^m$ and a vector $\mathbf{q}$ in $\mathbb{R}^m$ such that $\mathbf{p} \geq \mathbf{0}$, $(\mathbf{p}, \mathbf{q}) \neq \mathbf{0}$, and

$$(\mathbf{p}, \mathbf{f}(\mathbf{x})) + (\mathbf{q}, \mathbf{h}(\mathbf{x})) \geq 0 \quad (2)$$

for all $\mathbf{x}$ in C .

Proof. For each $\mathbf{x}$ in C define the set

$$S(\mathbf{x}) = \{(\mathbf{p}, \mathbf{q}) : \mathbf{p} \in \mathbb{R}^m, \mathbf{q} \in \mathbb{R}^m, \mathbf{p} \geq \mathbf{f}(\mathbf{x}), \mathbf{q} = \mathbf{h}(\mathbf{x})\}.$$

For each $\mathbf{x}$ in C , the set $S(\mathbf{x})$ is the set of solutions $(\mathbf{p}, \mathbf{q})$ in $\mathbb{R}^{2m}$ of the system

$$\mathbf{p} \geq \mathbf{f}(\mathbf{x}), \quad \mathbf{q} = \mathbf{h}(\mathbf{x}). \quad (3)$$

Thus for each $\mathbf{x}$, the origin in $\mathbb{R}^{2m}$ is not in $S(\mathbf{x})$. Let

$$S = \bigcup_{\mathbf{x} \in C} S(\mathbf{x}).$$

Since the origin in $\mathbb{R}^{2m}$ belongs to none of the sets $S(\mathbf{x})$, it does not belong to S . It is readily verified that S is convex. Hence by Theorem B.3.2 there is a hyperplane that separates S and $\mathbf{0}$ in $\mathbb{R}^{2m}$. That is, there exist vectors $\mathbf{p}$ in $\mathbb{R}^m$ and $\mathbf{q}$ in $\mathbb{R}^m$ such that $(\mathbf{p}, \mathbf{q}) \neq \mathbf{0}$ and

$$(\mathbf{p}, \mathbf{f}(\mathbf{x})) + (\mathbf{q}, \mathbf{h}(\mathbf{x})) \geq 0 \quad (4)$$

for all $(\mathbf{p}, \mathbf{q})$ in S .

If (p, q) belongs to S , then $p > 0$ for some x in C . Hence if some component of q , say q_i , were negative, we could take x_i to be arbitrarily large and obtain a contradiction to (4). Hence $q \geq 0$.

Now let $\varepsilon > 0$ and let ε denote the vector in $\mathbb{R}^n$ all of whose components are equal to ε . Then for $x \in C$ the vector $(f(x) + \varepsilon, h(x))$ belongs to S , and so

$$\langle p, f(x) \rangle + \langle q, h(x) \rangle \geq -\varepsilon \langle p, \varepsilon \rangle.$$

We obtain (2) by letting $\varepsilon \rightarrow 0$.

Remark 4.1. From the proof it is clear that if the affine function h is absent, then the conclusion is that there exists a $p \neq 0$, $p \geq 0$ such that $\langle p, f(x) \rangle \geq 0$ for all x in C .

Remark 4.2. This theorem is not a true alternative theorem, as it does not assert that either the inequality (1) has a solution or the inequality (2) has a solution, but never both. If we strengthen the hypotheses of the theorem by assuming $C = \mathbb{R}^n$ and the rows of A to be linearly independent, then we do get a true alternative theorem. This is taken up in Exercise 4.2.

The next two theorems are corollaries of Theorem 4.1.

Theorem 4.2. Let f be a vector-valued convex function defined on a convex set C in $\mathbb{R}^n$ and with range in $\mathbb{R}^m$. Then either

$$f - f(x_0) = 0$$

has a solution x_0 in C or there exists a $p \neq 0$, $p \geq 0$, such that

$$(B) \quad \langle p, f(x) \rangle \geq 0$$

for all x in C , but never both.

Proof. Suppose f has a solution x_0 . Then for all $p \neq 0$, $p \geq 0$, we have $\langle p, f(x_0) \rangle \leq 0$, so (B) cannot hold.

Suppose now that f has no solution. Then by Theorem 4.1 and Remark 4.1 following the theorem, there exists a $p \neq 0$, $p \geq 0$ such that (B) holds for all x .

Theorem 4.2 is a true generalization of Gordan's theorem, Theorem II.5.1. To see that this is so, let $C = \mathbb{R}^n$ and let $f(x) = Ax$. Then $Ax = 0$ has no solution x_0 if and only if there exists a vector p in $\mathbb{R}^m$, $p \neq 0$, $p \geq 0$, such that $\langle p, Ax \rangle \geq 0$ for all x . If we take $x = Ap$, we get $\langle Ap, Ap \rangle \geq 0$. If we take $x = -Ap$, we get $\langle Ap, Ap \rangle \leq 0$. Hence $Ap = 0$. In other words we have shown that $Ax = 0$ has no solution if and only if there exists a $p \neq 0$, $p \geq 0$ such that $Ap = 0$, which is Gordan's theorem.

THEOREM 4.1. Let C be a compact convex set in $\mathbb{R}^n$ and let $\{f_a\}_{a \in A}$ be a possibly infinite family of real-valued convex functions that are continuous on C . If the system

$$f_a(x) \leq 0, \quad a \in A, \quad (5)$$

has no solution in C , then there exists a finite subfamily $f_{a_1}, \dots, f_{a_m}$ and a vector $p = (p_1, \dots, p_m)$ such that $p \neq 0$, $p \geq 0$ and such that the real-valued function F defined by

$$F(x) = \sum_{i=1}^m p_i f_{a_i}(x) \quad (6)$$

satisfies

$$\min_{x \in C} F(x) > 0. \quad (7)$$

The proof will be carried out by translating the statement of the nonexistence of solutions to the system (5) to a statement about the empty intersection of a collection of closed subsets of a compact set. The finite-intersection property will then be used to obtain a statement about the empty intersection of a finite collection of sets. The last statement will then be translated to a statement about the nonexistence of a solution to the system (1). An application of Theorem 4.1 will then yield (7).

Proof. Since for each $a \in A$ the function f_a is continuous, it follows that for each $a \in A$ and $\varepsilon > 0$ the set

$$C_{a,\varepsilon} = \{x \in C, f_a(x) \leq \varepsilon\}$$

is closed and is contained in C . If the intersection of the closed sets $\{C_{a,\varepsilon}\}$, where $a \in A$ and $\varepsilon > 0$, were not empty, there would exist an x_0 in C such that $f_a(x_0) \leq \varepsilon$ for all $a \in A$ and all $\varepsilon > 0$. Hence $f_a(x_0) \leq 0$ for all $a \in A$, contradicting the hypothesis that $\{f_a\}_{a \in A}$ has no solution in C . Thus the intersection of the sets $C_{a,\varepsilon}$ is empty. Since C is compact, by Remark 1.1.1 and Theorem 1.1.1 there exists a finite subcollection $C_{a_1,\varepsilon_1}, C_{a_2,\varepsilon_2}, \dots, C_{a_m,\varepsilon_m}$ of $\{C_{a,\varepsilon}\}$ whose intersection is empty. That is, the system $f_{a_i}(x) - \varepsilon_i \leq 0$, $i = 1, \dots, m$, has no solution x in C . Therefore the system

$$f_{a_i}(x) - \varepsilon_i \leq 0, \quad i = 1, \dots, m,$$

has no solution x in C . Therefore, by Theorem 4.1 and Remark 1, there exists a vector $p = (p_1, \dots, p_m)$ with $p \neq 0$ and $p_i \geq 0$, $i = 1, \dots, m$, such that

$$\sum_{i=1}^m p_i f_{a_i}(x) \geq \sum_{i=1}^m p_i \varepsilon_i > 0 \quad (8)$$

for all x in C . The function J defined in (5) is continuous on the compact set C and hence attains a minimum on C . From this and from (5) the conclusion follows.

Remark 4.1. The reader who is familiar with the notion of semicontinuity will see that the requirement that the functions f_i be continuous can be replaced by the less restrictive requirement that they be lower semicontinuous.

Exercise 4.1. Verify that the set S in Theorem 4.1 is convex.

Exercise 4.2. Prove the following corollary to Theorem 4.1. Let $f = (f_1, \dots, f_n)$ be defined and convex on $\mathbb{R}^n$. Let $h = Ax - b$ be an affine mapping from $\mathbb{R}^n$ to $\mathbb{R}^m$ and let the rows of A be linearly independent. Then either

$$1. \quad f(x) \leq 0, \quad Ax - b = 0$$

has a solution x in $\mathbb{R}^n$ or

$$2. \quad \exists p, (p, q) \in \mathbb{R}_+^n \times \mathbb{R}_+^m, p + q = 0 \quad \text{all } x \in \mathbb{R}^n$$

has a solution $(p, q) \neq 0, p \in \mathbb{R}_+^n, q \in \mathbb{R}_+^m$, but never both.

5. APPLICATION TO GAME THEORY

In this section we use Theorem 4.1 to prove the von Neumann minimum theorem and then apply the minimum theorem to prove the existence of value and saddle points in mixed strategies for matrix games. Our proof will require the concept of uniform continuity and a sufficient condition for uniform continuity.

Definition 5.1. A function f defined on a set S in $\mathbb{R}^n$ with range in $\mathbb{R}^m$ is said to be *uniformly continuous* on S if for each $\epsilon > 0$ there exists a $\delta > 0$ such that if x and x' are points in S satisfying $\|x - x'\| < \delta$, then $\|f(x) - f(x')\| < \epsilon$.

Theorem 5.1. Let S be a compact set in $\mathbb{R}^n$ and let f be a mapping defined and continuous on S with range in $\mathbb{R}^m$. Then f is uniformly continuous on S .

For a proof we refer the reader to Bartle and Sherbert [1989] or Rudin [1976].

THEOREM 5.2 (VON NEUMANN MINIMUM THEOREM). Let $X \subset \mathbb{R}^n$ and $Y \subset \mathbb{R}^m$ be compact convex sets. Let $f: X \times Y \rightarrow \mathbb{R}$ be continuous on $X \times Y$. For each x in X let $f(x, \cdot): Y \rightarrow \mathbb{R}$ be convex and for each y in Y let $f(\cdot, y): X \rightarrow \mathbb{R}$ be concave.

Then

$$\max_{x \in X} \min_{y \in Y} f(x, y) = \min_{y \in Y} \max_{x \in X} f(x, y), \quad (1)$$

and there exists a pair (x_0, y_0) with $x_0 \in X$ and $y_0 \in Y$ such that for all x in X and all y in Y

$$f(x, y_0) \leq f(x_0, y_0) \leq f(x_0, y). \quad (2)$$

Moreover, if v denotes the common value in (1), then $f(x_0, y_0) = v$.

A pair (x_0, y_0) as in (2) is called a *saddle point* of the function f .

The relation (1) is not true in general for continuous functions f defined in $X \times Y$ where X and Y are compact and convex. To see this, let $X = [0, 1]$, let $Y = [0, 1]$, and let $f(x, y) = (x - y)^2$. Then for fixed x , $\min\{f(x, y) : y \in Y\}$ is achieved at $y = x$ and equals zero. Hence in this example the left-hand side of (1) equals zero. On the other hand, for fixed y in the interval $[0, \frac{1}{2}]$

$$\max_{x \in X} \{f(x, y) : x \in X\} = (1 - y)^2,$$

and for fixed y in the interval $[\frac{1}{2}, 1]$

$$\max_{x \in X} \{f(x, y) : x \in X\} = y^2.$$

Hence in this example the right-hand side of (1) equals $\frac{1}{4}$.

The convexity in X and concavity in Y is a sufficient condition but not a necessary condition, as the example $X = [0, 1]$, $Y = [0, 1]$, and $f(x, y) = x^2 y^2$ shows.

Before we take up the proof of Theorem 5.2, which will be carried out in several steps, we observe the following. If A and B are arbitrary sets and g is a real-valued bounded function defined on $A \times B$, then the quantities

$$\sup_{x \in A} \inf_{y \in B} g(x, y) \quad \text{and} \quad \inf_{y \in B} \sup_{x \in A} g(x, y)$$

are well defined, while the quantities

$$\max_{x \in A} \min_{y \in B} g(x, y) \quad \text{and} \quad \min_{y \in B} \max_{x \in A} g(x, y)$$

need not be. For a simple example, take $A = [0, 1]$, $B = [0, 1]$, and $g(x, y) = xy$. Thus, implicit in (1) is the statement, which must be proved, that the left- and right-hand sides of (1) exist. The proof of this statement will constitute the second step of our proof. The first step is the general observation given in Lemma 5.1.

LEMMA 5.1. Let A and B be arbitrary sets and let $g: A \times B \rightarrow \mathbb{R}$ be bounded. Then

$$\sup_{a \in A} \inf_{b \in B} g(a, b) \leq \inf_{b \in B} \sup_{a \in A} g(a, b).$$

Proof. Choose an element $b_1 \in B$. Then for each $a \in A$,

$$g(a, b_1) \leq \sup_{b \in B} g(a, b)$$

Since the choice of b_1 was arbitrary, this relationship holds for each $a \in A$ and each $b \in B$. Hence, for each $a \in A$

$$\inf_{b \in B} g(a, b) \leq \inf_{b \in B} \sup_{a \in A} g(a, b).$$

The right-hand side is a fixed real number, so

$$\sup_{a \in A} \inf_{b \in B} g(a, b) \leq \inf_{b \in B} \sup_{a \in A} g(a, b).$$

We now proceed to the second step of the proof and show that the left- and right-hand sides of (1) exist.

Fix fixed x in X , the function $f(x, \cdot): Y \rightarrow \mathbb{R}$ defined by $f(x, y)$ is continuous on the compact set Y . Hence

$$\phi(x) = \min_{y \in Y} f(x, y)$$

is defined for each x in X . We assert that the function ϕ is continuous on X . To prove the assertion, let $\varepsilon > 0$ be given. Since f is continuous on the compact set $X \times Y$, it is uniformly continuous on $X \times Y$. Hence there exists a $\delta > 0$ such that if $\|x_1 - x_2\| = \delta$, then

$$f(x_1, y) - \frac{1}{2}\varepsilon < f(x_2, y) < f(x_1, y) + \frac{1}{2}\varepsilon$$

for all y in Y . Upon taking the minimum over Y we get that

$$\phi(x_1) - \frac{1}{2}\varepsilon \leq \phi(x_2) \leq \phi(x_1) + \frac{1}{2}\varepsilon$$

or $|\phi(x_1) - \phi(x_2)| \leq \varepsilon$ whenever $\|x_1 - x_2\| = \delta$. This proves the continuity of ϕ .

Since ϕ is continuous on the compact set X , there exists and x_0 in X such that

$$\phi(x_0) = \max_{x \in X} \phi(x) = \max_{x \in X} \left[\min_{y \in Y} f(x, y) \right]. \quad (3)$$

It follows from (3) that the left-hand side of (1) is well defined.

From the definition of ϕ we have

$$\phi(x_0) = \min_{y \in Y} f(x_0, y). \quad (4)$$

Similarly, for each y in Y we can set

$$\psi(y) = \max_{x \in X} f(x, y)$$

and then obtain the existence of a y_0 in Y such that

$$\psi(y_0) = \min_{y \in Y} \psi(y) = \min_{y \in Y} \left[\max_{x \in X} f(x, y) \right] \quad (5)$$

and

$$\psi(y_0) = \max_{x \in X} f(x, y_0). \quad (6)$$

Thus the right-hand side of (1) is well defined.

The third step in the proof is to establish the equality in (1). Let

$$\alpha = \max_{x \in X} \min_{y \in Y} f(x, y), \quad \beta = \min_{y \in Y} \max_{x \in X} f(x, y).$$

By Lemma 5.1, with $X = A$, $Y = B$, $f = \varphi$ and by the second step, we have that $\alpha \leq \beta$. Hence to prove (1), we must show that $\alpha \geq \beta$.

Let $\varepsilon > 0$ be arbitrary. We assert that there does not exist a y_0 in Y such that for all x in X

$$f(x, y_0) \leq \beta - \varepsilon.$$

For if such a y_0 did exist, we would have

$$\beta = \min_{y \in Y} \max_{x \in X} f(x, y) \leq \max_{x \in X} f(x, y_0) \leq \beta - \varepsilon,$$

which is not possible.

Let $\{p_i\}_{i=1}^n$ be the family of functions defined on Y by the formula

$$p_i(x) = f(x, p_i) \quad x \in X.$$

Then, for each $x \in X$, the function p_i is continuous and convex on Y . Moreover, for arbitrary $\varepsilon > 0$ the system

$$p_i(y) = \beta + \varepsilon < 0, \quad x \in X, \quad (7)$$

has no solution y in Y .

Since (7) has no solution, it follows from Theorem 4.3 that there exist points $x_1, \dots, x_n$ in X , a vector $p \in E$, $p \geq 0$, $p \neq 0$ and a $\beta > 0$ such that

$$G(y) = \sum_{i=1}^n p_i f(x_i, y) = \beta + \varepsilon > \beta > 0 \quad (8)$$

for each y in Y . Since $\sum_{i=1}^n p_i > 0$, we may divide through by this quantity in (8) and assume that $p \in E_1$. Therefore, since X is convex, $\sum_{i=1}^n p_i x_i$ is in X . Using (8) and the concavity of $f(\cdot, y)$ for each fixed y , we may write

$$f\left(\sum_{i=1}^n p_i x_i, y\right) \geq \sum_{i=1}^n p_i f(x_i, y) > \beta + \varepsilon$$

for each y in Y . Therefore,

$$\begin{aligned} \alpha &= \max_{x \in X} \min_{y \in Y} f(x, y) \geq \min_{y \in Y} f\left(\sum_{i=1}^n p_i x_i, y\right) \\ &\geq \min_{y \in Y} \sum_{i=1}^n p_i f(x_i, y) = \beta + \varepsilon. \end{aligned}$$

Since $\varepsilon > 0$ is arbitrary, we get that $\alpha \geq \beta$, and (1) is established.

The final step in the proof is to show that $\alpha = f(x_0, y_0)$ and that (2) holds. It follows from (1), (3), and (5) that

$$\beta(x_0) = \beta(y_0) = \alpha.$$

It now follows from (6) and (8) that

$$f(x, x_0) \leq \alpha \leq f(x_0, y) \quad (9)$$

for all x in X and y in Y . In particular, (9) holds for $x = x_0$ and $y = y_0$, so that $\alpha = f(x_0, y_0)$. The relation (2) now follows from (9).

The simplest two-person zero-sum games are matrix games. These games are played by two players that we shall call Blue and Red. Blue selects a positive integer i from a finite set I of positive integers $\{1, \dots, m\}$. Red, simultaneously and in ignorance of Blue's choice, selects a positive integer j from a finite set J of positive integers $\{1, \dots, n\}$. The "payoff" or "reward" to Blue in the event that he chooses i and Red chooses j is a predetermined real number a_{ij} . The payoff to Red is $-a_{ij}$. Hence the name zero-sum game. A strategy for Blue is a rule which tells Blue which positive integer i in I to select. A strategy for Red is a rule which tells Red which positive integer j in J to select. Blue's objective is to choose i in I so as to maximize a_{ij} . Red's objective is to choose j in J so as to minimize a_{ij} , that is, to maximize $-a_{ij}$. If the strategy for Blue is to choose the number i_0 in I , then we call the strategy a pure strategy for Blue. Similarly, if the strategy for Red is to choose the number j_0 in J , then we call the strategy a pure strategy for Red. Later we shall consider strategies that are not pure.

Another way of describing a matrix game is as follows. An $m \times n$ matrix A with entries a_{ij} , $i = 1, \dots, m$, $j = 1, \dots, n$, is given. Blue selects a row i and Red selects a column j . The choices are made simultaneously and without knowledge of the opponent's choice. If Blue selects row i and Red selects column j , then the payoff to Blue is a_{ij} and the payoff to Red is $-a_{ij}$. Blue's objective is to maximize the payoff, while Red's objective is to minimize the payoff.

At this point, the formulation of a matrix game may strike the reader as rather artificial and uninteresting. It turns out, however, that many decision problems in conflict situations can be modeled as matrix games. As an example, we shall formulate a tactical naval problem as a matrix game. The reader will find many additional interesting examples elsewhere [Dresher, 1981]. The naval game appears there as a land conflict game called Colonel Blotto.

Two Red cargo ships are protected by three escort frigates under the command of Captain Photo. The entire convoy is subject to attack by a pack of four Blue submarines under the command of Captain Nemo. The shipwreckary a hazardous cargo and thus maintains a separation such that one submarine cannot attack both ships. Captain Nemo's problem is to allocate his submarines to attack the two ships. Thus, he may allocate all four of his submarines to attack one of the ships and none to attack the second ship. Or, he may allocate three submarines to attack one ship and one submarine to attack the second ship. Or, he may allocate two submarines to attack each of the ships. We may summarize the choices, or strategies, available to Captain Nemo by five ordered pairs, (α, β) , $\alpha, \beta \in \{1, 2, 3, 4\}$, $\alpha + \beta = 4$, where (α, β) indicates that he has allocated α submarines to attack ship 1 and β submarines to attack ship 2. Captain Photo's problem is to allocate his frigates to defend the ships. He has four choices, or strategies, each of which can be denoted by an ordered pair (γ, δ) , $\gamma \in \{1, 2, 3, 4\}$, $0 \leq \delta \leq 3$, $\gamma + \delta = 3$, where γ frigates are allocated to defend ship 1 and δ frigates are allocated to defend ship 2. We assume that conditions

an attack that when Captain Nemo allocates his submarines, he cannot determine the disposition of the escort frigates, and that when Captain Plato allocates his escort frigates, he cannot determine the disposition of the attacking submarines.

If at a given ship the number of attacking submarines exceeds the number of defending frigates, then the payoff to Blue at that ship equals the number of Red frigates at that ship plus 1. The "plus 1" is awarded for the cargo ship that is assumed to be sunk. If at a given ship the number of attacking submarines is less than the number of defending frigates, then the payoff to Blue at the ship is

$$-(\text{Number of attacking submarines} + 1).$$

The "plus 1" is deducted for allowing the cargo ship to proceed to its destination. If at a given ship the number of attacking submarines equals the number of defending ships, then the encounter is a draw, and the payoff to Blue at that ship is zero. The total payoff to Blue resulting from given choices by the two antagonists is the sum of the payoffs at each ship. The payoff to Red is the negative of the payoff to Blue.

The game just described can be cast as a matrix game with payoff matrix

		Captain Plato's strategies			
		(3, 0)	(0, 3)	(2, 1)	(1, 2)
Captain Nemo's strategies	(4, 0)	$\begin{pmatrix} -4 & 0 & 2 & 1 \\ 0 & 4 & 1 & 2 \\ 1 & -1 & 3 & 0 \\ -1 & 1 & 0 & 3 \\ -2 & -2 & 2 & 2 \end{pmatrix}$			
	(0, 4)				
	(1, 1)				
	(1, 3)				
	(2, 2)				

A strategy for Blue (Captain Nemo) is a choice of row; a strategy for Red (Captain Plato) is a choice of column.

If Captain Nemo chooses row i , he is guaranteed a payoff of at least

$$\min\{a_{ij} : j = 1, \dots, 4\}.$$

If he picks the row that will give the payoff

$$v'' = \max_i \min_j a_{ij},$$

then he will be guaranteed to obtain v'' . Captain Nemo can thus pick either row 1 or row 2 and be guaranteed a payoff of at least zero. Thus $v'' = 0$. On the other hand, if Captain Plato chooses column j , he will lose at most

$\max\{a_i, f_i\}$, $i = 1, \dots, 5$. If Captain Photo picks the columns that gives Blue the payoff

$$v^* = \min_i \max_j a(i, j),$$

then he is guaranteed to lose no more than v^* . Captain Photo can thus pick either column 3 or 4 and be guaranteed to lose no more than three. Thus $v^* = 3$.

Note that if Captain Nemo knows beforehand that Captain Photo will pick column 3, then he can pick row 3 and obtain a payoff of 5, which is greater than $v^* = 3$. On the other hand, if Captain Photo knows beforehand that Captain Nemo will pick row 3, then he can pick column 1 and lose zero units, which is less of a loss than three units ($v^* = 3$). An examination of the payoff matrix shows that if one of the commander's choices is known in advance by the second commander, then the second commander can usually improve his payoff if he bases his choice on this information. Thus, in the submarine-convooy game secrecy of choice is essential. The open question, however, in the submarine-convooy game is, "What is the best strategy for each player?"

In contrast to the submarine-convooy game, consider the game with payoff matrix

$$\begin{pmatrix} 3 & 7 & 4 & 3 \\ 12 & 10 & 7 & 8 \\ 18 & 4 & 1 & 6 \\ 2 & 5 & 6 & 9 \end{pmatrix}.$$

Here

$$\max_i \min_j a(i, j) = \min_j \max_i a(i, j) = 7.$$

The strategies $i = i_0$, where $i_0 = 2$, and $j = j_0$, where $j_0 = 3$, have the property that $a(i_0, j_0) = 7$ and

$$a(i, j_0) \leq a(i_0, j_0) \leq a(i_0, j)$$

for all $i = 1, \dots, 4$ and $j = 1, \dots, 4$. We say that the value of the game is 7, that $i_0 = 2$ is an optimal (pure) strategy for Blue, and that $j_0 = 3$ is an optimal (pure) strategy for Red. If Blue chooses his optimal strategy $i_0 = 2$, then the best that Red can do is to also choose his optimal strategy $j_0 = 3$. Also, if Red chooses his optimal strategy $j_0 = 3$, then the best that Blue can do is to choose his optimal strategy $i_0 = 2$. Thus, if either player announces beforehand that he will choose an optimal strategy, the opponent cannot take advantage of this information, other than to choose an optimal strategy.

We now consider the general matrix game with $m \times n$ payoff matrix A with entries a_{ij} , $i = 1, \dots, m$, $j = 1, \dots, n$. We always have

$$v^* = \max_i \min_j a_{ij} \leq \min_j \max_i a_{ij} = v^*,$$

If the reverse inequality $v^* \geq v^*$ holds, then we set

$$v = \min_j \max_i a_{ij} = \max_i \min_j a_{ij}$$

saying that the matrix game has a value v in pure strategies. It is easy to show that if the matrix game has a value v in pure strategies, then there exists a pair (i_0, j_0) such that

$$v = a(i_0, j_0), \quad (10)$$

and for all $i = 1, \dots, m$, $j = 1, \dots, n$

$$a(i, j_0) \leq a(i_0, j_0) \leq a(i_0, j). \quad (11)$$

The strategies i_0 and j_0 are called *optimal* (pure) strategies. We point out that optimal strategies need not be unique. A pair of strategies (i_0, j_0) such that (11) holds is called a *saddle point* in pure strategies. It is also easy to show that if a matrix game has a saddle point in pure strategies, then the game has value v and (10) holds.

If the matrix game has a value, then each player should use an optimal strategy. For then, Blue, the maximizer, is guaranteed a payoff of v and Red, the minimizer, is guaranteed to lose at most $-v$. If one player, say Blue, uses an optimal strategy and Red does not, then the payoff to Blue, who plays optimally, will, in general, be more favorable than it would be had Red played optimally. This is the content of (1.5).

The question that still thralls itself upon us is, "How does one play if the game does not have a saddle point, as in the information-concealment game?" To answer this question, we introduce the notion of mixed strategies.

Definition 1.1. A mixed strategy for Blue in the game with $m \times n$ payoff matrix A is a probability distribution on the rows of A . Thus, a mixed strategy is a vector $x = (x_1, \dots, x_m)$ with $x_i \geq 0$, $i = 1, \dots, m$, and $\sum_{i=1}^m x_i = 1$.

A mixed strategy for Red is a probability distribution on the columns of A . Thus, a mixed strategy for Red is a vector $y = (y_1, \dots, y_n)$ with $y_j \geq 0$, $j = 1, \dots, n$, and $\sum_{j=1}^n y_j = 1$.

The set of pure strategies for a player is a subset of the set of mixed strategies. For example, the pure strategy for Blue in which row i is always

chosen is also the mixed strategy in which row i is chosen with probability 1, that is, the vector $a_i = (0, \dots, 0, 1, 0, \dots, 0)$, where the i th entry is 1 and all other entries are zero.

In actual play, Blue will select his row using a random device which produces the value i with probability a_i , $i = 1, \dots, m$. Blue chooses row i if the random device produces the value i . Red chooses column j in similar fashion. If mixed strategies are used, neither player tries to outguess his opponent and both tend to lose. Moreover, if Blue uses a mixed strategy, Red cannot determine Blue's choice beforehand and take advantage of this information. A similar statement holds if Red uses a mixed strategy.

If mixed strategies are allowed, then it is appropriate to define the payoff of the matrix game with mixed strategies as the expected value of the payoff. Thus, if we denote the payoff resulting from the choice of strategies (x, y) by $P(x, y)$, we have

$$P(x, y) = x^T Ay = (x, Ay). \quad (12)$$

We define the matrix game with matrix A in which mixed strategies are used as follows. Blue selects a vector x in P_m and Red selects a vector y in P_n . The payoff to Blue is given by (12), and the payoff to Red is $-P(x, y)$. We say that the matrix game with mixed strategies has a value if

$$\max_{x \in P_m} \min_{y \in P_n} P(x, y) = \min_{y \in P_n} \max_{x \in P_m} P(x, y). \quad (13)$$

We denote the number in (13) by v and call it the value of the game. We say that the game has a saddle point (x_0, y_0) , that x_0 is an optimal mixed strategy for Blue, and that y_0 is an optimal mixed strategy for Red if

$$P(x, y_0) \leq P(x_0, y_0) \leq P(x_0, y) \quad (14)$$

for all x in P_m and y in P_n .

THEOREM 3.1. *A matrix game has a value and a saddle point in mixed strategies.*

This theorem is a corollary of Theorem 3.2. Take $N = P_m$, $F = P_n$, and $f = P$. The function P is continuous on $P_m \times P_n$. The sets P_m and P_n are compact and convex. Since P is bilinear, it is concave in x and convex in y . Theorem 3.1 follows by applying Theorem 3.2 to P .

Note that the saddle point condition (14) implies that

$$v = P(x_0, y_0).$$

Exercise 8.1. Verify that, in the subgame-perfect game, the value in mixed strategies is $\frac{1}{2}$ and that

$$u_{\mathbf{a}} = \left(\frac{1}{2}, \frac{1}{2}, 0, 0, \frac{1}{2} \right), \quad v_{\mathbf{a}} = \left(\frac{1}{2}, \frac{1}{2}, \frac{1}{2}, \frac{1}{2} \right)$$

is a saddle point.

Exercise 8.2. Show that if a matrix game has a value in pure strategies, then (i) and (iii) hold. Conversely, show that if a matrix game has a saddle point in pure strategies, then the game has a value v and (ii) holds.

Exercise 8.3. Let $X_{\mathbf{a}}$ denote the set of optimal mixed strategies for Blue and let $Y_{\mathbf{a}}$ denote the set of optimal mixed strategies for Red. Show that $X_{\mathbf{a}}$ and $Y_{\mathbf{a}}$ are compact convex sets.

IV

OPTIMIZATION PROBLEMS

1. INTRODUCTION

In Section 9 of Chapter I we stated the basic problem in optimization theory as follows: Given a set S (not necessarily in $\mathbb{R}^n$) and a real-valued function f defined on S , does there exist an element x_0 in S such that $f(x_0) \leq f(x)$ for all x in S , and if so find it. We also showed that the problem of maximizing f over S is equivalent to the problem of minimizing $-f$ over S . Thus, there is no loss of generality in only considering minimization problems, as we shall do henceforth.

If S is a set in $\mathbb{R}^n$ and f is a real-valued function on S , we say that a point x_0 in S is a *local minimizer* or that f has a *local minimum* at x_0 if there exists a $\delta > 0$ such that $f(x_0) \leq f(x)$ for all x in $B(x_0, \delta) \cap S$. If the strict inequality $f(x_0) < f(x)$ holds, then x_0 is said to be a *strict local minimizer*. If S is open, then $B(x_0, \delta) \cap S$ is an open set contained in S . Hence there exists a $\delta > 0$ such that for all x in $B(x_0, \delta)$, $f(x_0) \leq f(x)$. Therefore, if S is open, for x_0 to be a local minimizer, we need to find a $\delta > 0$ such that $f(x_0) \leq f(x)$ for all x in $B(x_0, \delta)$ and can omit the intersection with S . A point x_0 that is a minimizer is also a local minimizer. A point x_0 that is a local minimizer need not be a minimizer.

In this chapter we shall develop necessary conditions and sufficient conditions for a point x_0 to be a minimizer or local minimizer for various classes of problems. The determination of the set E of points that satisfy the necessary conditions does not, of course, solve the optimization problem. All we learn is that the set of solutions is to be found in this set. The determination of those points in E that are solutions to the problem requires further analysis. Sufficient conditions, if applicable to the problem at hand, are useful in this connection.

If E is an open set in $\mathbb{R}^n$, the optimization problem is often called an *unconstrained problem*.

A *programming problem* is a problem of the following type. A set X_0 in $\mathbb{R}^n$ is given as are functions f and g with domains X_0 and ranges in $\mathbb{R}^1$ and $\mathbb{R}^m$

respectively. Let

$$X = \{x: x \in X_0, g(x) \leq 0\}.$$

Problem P

Minimize f over the set X . That is, find an x_0 in X such that $f(x_0) \leq f(x)$ for all x in X .

The set X is called the *feasible set*, and elements of X are called *feasible points* or *feasible vectors*.

The problem is sometimes stated as follows: Minimize $f(x)$ subject to $x \in X_0$ and $g(x) \leq 0$.

If $X_0 = \mathbb{R}^n$, A is an $m \times n$ matrix, and

$$f(x) = c - k^T x, \quad g(x) = \begin{pmatrix} Ax - r \\ -x \end{pmatrix}, \quad (1)$$

the problem is said to be a *linear programming problem*. We have written $-k$ rather than k because in the linear programming literature the problem is usually stated as follows: Minimize (k, x) subject to $Ax \leq r$ and $x \geq 0$.

If X_0 is convex and the functions f and g are convex, the problem is said to be a *convex programming problem*.

2. DIFFERENTIABLE UNCONSTRAINED PROBLEMS

We will now discuss the set on which the function f is defined by X rather than $\mathbb{R}^n$. We first assume that X is an open interval in $\mathbb{R}^1$ and that f is differentiable on X . We will derive necessary conditions and sufficient conditions for a point x_0 in X to be a local minimum. Although the reader surely knows these conditions from elementary calculus, it is instructive to review them.

Let x_0 in X be a local minimum of f and let f be differentiable at x_0 . Then there exists a $\delta > 0$ such that $f(x) \geq f(x_0)$ for x in $(x_0 - \delta, x_0 + \delta)$. Therefore, for $x_0 - \delta < x < x_0 + \delta$

$$\frac{f(x) - f(x_0)}{x - x_0} \geq 0, \quad (1)$$

and for $x_0 - \delta < x < x_0$

$$\frac{f(x) - f(x_0)}{x - x_0} \leq 0. \quad (2)$$

Since f is differentiable at x_0 ,

$$\lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} = \lim_{x \rightarrow x_0} \frac{f(x) - f(x_0)}{x - x_0} = f'(x_0).$$

From this and from (1) we get $f'(x_0) \geq 0$, while from (2) we get $f'(x_0) \leq 0$. Hence $f'(x_0) = 0$.

Let us now assume further that f is of class C^2 on $(x_0 - \delta, x_0 + \delta)$. By Taylor's theorem, for each $x \neq x_0$ in $(x_0 - \delta, x_0 + \delta)$ there exists a point $\bar{x}$ that lies between x_0 and x such that

$$f(x) = f(x_0) + f'(x_0)(x - x_0) + \frac{1}{2}f''(\bar{x})(x - x_0)^2. \quad (3)$$

Since x_0 is a local minimum, $f(x_0) = 0$ and therefore

$$\frac{1}{2}f''(\bar{x})(x - x_0)^2 = f(x) - f(x_0) \geq 0.$$

Thus

$$f''(\bar{x}) = 2 \frac{f(x) - f(x_0)}{(x - x_0)^2} \geq 0.$$

If we now let $x \rightarrow x_0$, then $\bar{x} \rightarrow x_0$ and $f''(\bar{x}) \rightarrow f''(x_0)$. Hence $f''(x_0) \geq 0$.

We have proved the following theorem.

THEOREM 2.1. Let X be an open interval in $\mathbb{R}$, let f be a real-valued function defined on X , and let x_0 in X be a local minimum for f . If f is differentiable at x_0 , then $f'(x_0) = 0$. If f is of class C^2 on some interval about x_0 , then $f''(x_0) \geq 0$.

A point c in X such that $f'(c) = 0$ is called a *critical point* of f . We have shown that for a differentiable function a local minimum is a critical point.

If we replace the necessary condition $f'(x_0) = 0$ by the stronger condition $f'(x_0) > 0$, we obtain a sufficient condition for a critical point to be a local minimum.

THEOREM 2.2. Let X be an interval in $\mathbb{R}$ and let f be real-valued and of class C^2 on some interval about a point x_0 in X . If $f'(x_0) = 0$ and $f''(x_0) > 0$, then x_0 is a strict local minimum for f .

Proof. Since f'' is continuous on some interval about x_0 and since $f''(x_0) > 0$, there exists a $\delta > 0$ such that $f''(x) > 0$ for all x in $(x_0 - \delta, x_0 + \delta)$ (See Exercise 15.4.) By Taylor's Theorem, for each $x \neq x_0$ in $(x_0 - \delta, x_0 + \delta)$ there exists an $\bar{x}$ lying between x_0 and x such that (3) holds. Since $f'(x_0) = 0$, we get that for each $x \neq x_0$ in $(x_0 - \delta, x_0 + \delta)$

$$f(x) - f(x_0) = \frac{1}{2}f''(\bar{x})(x - x_0)^2 > 0. \quad (4)$$

Thus x_0 is a strict local minimum.

In elementary calculus texts Theorem 2.2 is often called the *second derivative test*.

We must have $f''(x_0) > 0$ rather than $f''(x_0) \neq 0$ for the conclusion of Theorem 2.2 to hold. To see this, take $X = \mathbb{R}$ and take $f(x) = x^4$. Then $f'(0) = 0$, $f''(0) = 0$, but zero is not a local minimum. If, however, we assume that $f''(x) \neq 0$ for all x in X , then all the statements in the proof of Theorem 2.2 hold for each x in X , except that (4) now reads

$$f(x) - f(x_0) = \frac{1}{2}f''(\xi)(x - x_0)^2 \neq 0.$$

Thus x_0 is a minimum. If we assume that $f''(x) \geq 0$ on some interval about x_0 , then x_0 is a local minimum. Also, if we assume that $f''(x) > 0$ on all of X , then x_0 is a strict minimum.

We summarize this discussion in the following corollary to Theorem 2.2.

COROLLARY 1. Let $f(x_0) = 0$. If $f''(x) \geq 0$ on X , then x_0 is a minimum. If $f''(x) > 0$ on X , then x_0 is a strict minimum. If $f''(x) \geq 0$ on some interval about x_0 , then x_0 is a local minimum.

We now take X to be an open set in $\mathbb{R}^n$. The next two theorems are the n -dimensional generalizations of Theorems 2.1 and 2.2.

THEOREM 2.3. Let X be an open set in $\mathbb{R}^n$. Let f be a real-valued function defined on X , and let x_0 be a local minimum. If f has first-order partial derivatives at x_0 , then

$$\frac{\partial f}{\partial x_i}(x_0) = 0, \quad i = 1, \dots, n. \quad (2)$$

If f is of class C^2 on some open ball $B(x_0, \delta)$ centered at x_0 , then

$$B(x_0) = \left(\frac{\partial^2 f}{\partial x_i \partial x_j}(x_0) \right)$$

is positive semidefinite.

Proof. Since X is open, there exists a $\delta > 0$ such that for each $i = 1, \dots, n$ the point

$$x_0 + \theta e_i = (x_{0,1}, \dots, x_{0,i-1}, x_{0,i} + \theta, x_{0,i+1}, \dots, x_{0,n})$$

lies in X whenever $|\theta| < \delta$. Therefore for each $i = 1, \dots, n$ the function φ_i defined by

$$\varphi_i(\theta) = f(x_0 + \theta e_i), \quad |\theta| < \delta,$$

is well defined and has a local minimum at $t = 0$. Since the partial derivatives of f exist at $\mathbf{x}_0$, the chain rule implies that for each $i = 1, \dots, n$ the derivative $\varphi'_i(t)$ exists and is given by

$$\varphi'_i(t) = \frac{df}{dx_i}(\mathbf{x}_0).$$

Since each φ_i has a local minimum at $t = 0$, $\varphi'_i(0) = 0$. Thus (2) holds.

Now suppose that f is of class C^{2n} on some open ball $B(\mathbf{x}_0, \delta)$. Then for each unit vector $\mathbf{u}$ in $\mathbb{R}^n$ the function $\varphi(\cdot; \mathbf{u})$ defined by

$$\varphi(t; \mathbf{u}) = f(\mathbf{x}_0 + t\mathbf{u}), \quad |t| < \delta,$$

is well defined, is of class C^{2n} , and has a local minimum at $t = 0$. By Theorem 2.1, $\varphi'(0; \mathbf{u}) = 0$ and $\varphi''(0; \mathbf{u}) \geq 0$. By a straightforward calculation as in relation (11) in Section 1.1 of Chapter I, we get

$$\varphi''(0; \mathbf{u}) = \langle \mathbf{u}, H(\mathbf{x}_0) \mathbf{u} \rangle.$$

Setting $t = 0$, we get that $\langle \mathbf{u}, H(\mathbf{x}_0) \mathbf{u} \rangle \geq 0$. Since $\mathbf{u}$ is an arbitrary unit vector, we get that $H(\mathbf{x}_0)$ is positive semidefinite.

In $\mathbb{R}^n$ points $\mathbf{x}$ such that $\nabla f(\mathbf{x}) = \mathbf{0}$ are called *critical points* of f .

THEOREM 2.6. *Let X be an open set in $\mathbb{R}^n$ and let f be a real-valued function of class C^{2n} on some ball $B(\mathbf{x}_0, \delta_0)$. If $H(\mathbf{x}_0)$ is positive definite and $\nabla f(\mathbf{x}_0) = \mathbf{0}$, then $\mathbf{x}_0$ is a strict local minimizer. If X is convex, f is of class C^{2n} on X , and $H(\mathbf{x})$ is positive semidefinite for all $\mathbf{x}$ in X , then $\mathbf{x}_0$ is a minimizer for f . If $H(\mathbf{x})$ is positive definite for all $\mathbf{x}$ in X , then $\mathbf{x}_0$ is a strict minimizer.*

Proof. We shall only prove the local statement and leave the proof of the last two to the reader.

If $H(\mathbf{x}_0)$ is positive definite at $\mathbf{x}_0$, then it follows from the continuity of the second partial derivatives of f and Lemma III.1.2 that there exists a $\delta > 0$ such that $H(\mathbf{x})$ is positive definite for all $\mathbf{x}$ in $B(\mathbf{x}_0, \delta)$. Let $\mathbf{x}$ be any point in $B(\mathbf{x}_0, \delta)$ and let $\mathbf{v} = \mathbf{x} - \mathbf{x}_0$. Then the function $\varphi(\cdot; \mathbf{v})$ defined by

$$\varphi(t; \mathbf{v}) = f(\mathbf{x}_0 + t\mathbf{v})$$

is well defined and is C^{2n} on some real interval $(-1 + \eta, 1 + \eta)$, where $\eta > 0$. Moreover

$$\varphi(0; \mathbf{v}) = f(\mathbf{x}_0), \quad \varphi(1; \mathbf{v}) = f(\mathbf{x}). \quad (3)$$

By Taylor's theorem, there exists a $\bar{t}$ in the interval $(0, 1)$ such that

$$\varphi(t; \mathbf{x}) = \varphi(0; \mathbf{x}) + \varphi'(0; \mathbf{x}) + \frac{1}{2}\varphi''(\bar{t}; \mathbf{x}). \quad (7)$$

Again, straightforward calculations as in relations (12) and (13) in Section 1.1 of Chapter 1 give

$$\begin{aligned} \varphi(0; \mathbf{x}) &= (\nabla f)(\mathbf{x}_0 + \alpha(\mathbf{x}), \mathbf{x}), \\ \varphi'(0; \mathbf{x}) &= (\mathbf{x}, H(\mathbf{x}_0 + \alpha(\mathbf{x})), \quad -1 - \eta < t < 1 + \eta. \end{aligned}$$

If we set $t = 0$ in the first of these relations, we get that $\varphi'(0; \mathbf{x}) = (\nabla f)(\mathbf{x}_0), \mathbf{x}$. If we then set $t = \bar{t}$ in the second of these relations, use the hypothesis $\nabla f(\mathbf{x}_0) = \mathbf{0}$, and substitute into equation (7), we get

$$\varphi(1; \mathbf{x}) - \varphi(0; \mathbf{x}) = \frac{1}{2}(\mathbf{x}, H(\mathbf{x}_0 + \bar{t}(\mathbf{x})).$$

From this relation, from (5), and from the fact that $H(\mathbf{x})$ is positive-definite for all $\mathbf{x}$ in $B(\mathbf{x}_0, \delta)$, we get that

$$f(\mathbf{x}) - f(\mathbf{x}_0) = \frac{1}{2}(\mathbf{x}, H(\mathbf{x}_0 + \bar{t}(\mathbf{x})) > 0$$

for all $\mathbf{x}$ in $B(\mathbf{x}_0, \delta)$. Hence $\mathbf{x}_0$ is a strict local minimizer.

Remark 2.1. Theorems 2.1 and 2.3 suggest the well-known strategy of solving unconstrained optimization problems by first finding the critical points. Some of these critical points can then be eliminated by considering the second derivative or Hessian at these points. By then considering strengthened conditions on the second derivative or the Hessian as in Theorems 2.2 and 2.4, one can show in some cases that a critical point is indeed a local minimizer or maximizer.

Exercise 2.1. Let K be a closed interval $[a, b]$ in $\mathbb{R}$ and let f be a real-valued function defined on K . Show that if $x = a$ is a local minimizer for f and if the right-hand derivative $f'_+(a)$ exists, then $f'_+(a) \geq 0$. Show that if $x = b$ is a local minimizer and the left-hand derivative $f'_-(b)$ exists, then $f'_-(b) \leq 0$.

Exercise 2.2. Find local minimizers, local maximizers, minimizers, and maximizers, if they exist, for the following functions. Exercise 15.5 may be useful in some cases.

- $f(x_1, x_2) = x_1^2 - 6x_1 + 2x_2^2 + 7$
- $f(x_1, x_2, x_3) = 3x_1^2 + 2x_2^2 + 2x_3^2 + 2x_1x_2 + 2x_1x_3 + 2x_2x_3$
- $f(x_1, x_2) = x_1x_2 + 1/x_1 + 1/x_2, (x_1, x_2) \neq (0, 0)$, and
- $f(x_1, x_2) = (x_1 + x_2)(x_1, x_2 + 1)$

Exercise 2.3. Let

$$f(x_1, x_2) = x_1^2 + x_2^2 \quad g(x_1, x_2) = x_1^2 + x_2^3.$$

Show that $\nabla f(\mathbf{0}) = \nabla g(\mathbf{0}) = \mathbf{0}$ and that both f and g have positive semidefinite Hessians at $\mathbf{0}$, but $\mathbf{0}$ is a local minimum for g but not for f . This shows that the condition of positive definiteness of $H(x_0)$ in Theorem 2.4 cannot be replaced by positive semidefiniteness.

Exercise 2.4. Complete the proof of Theorem 2.4 by proving the last two statements.

Exercise 2.5. Let

$$f(x, y) = x^2 + y^2 - 3xy^2.$$

Minimize f on $\mathbb{R}^2$. (Hint: see Exercise 18.1.)

Exercise 2.6. (a) Show that for no value of a does the function

$$f(x, y) = x^2 - 3axy + y^2$$

have a minimum.

(b) For each value of a find local minimizers and local maximizers, if any exist, of this function.

Exercise 2.7. Let f be defined on $\mathbb{R}^2$ by

$$f(x, y) = x^2 + e^{2y} - 3xy^2.$$

Show that there is a unique point (x_0, y_0) at which $\nabla f(x, y) = \mathbf{0}$ and that this point is a local minimum but not a minimizer.

Exercise 2.8 (Linear Regression). In the linear regression problem n points $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ are given in the xy -plane and it is required to “fit” a straight line $y = ax + b$ to these points in such a way that the sum of the squares of the vertical distances of the given points from the line is minimized. That is, a and b are to be chosen so that

$$f(a, b) = \sum_{i=1}^n (ax_i + b - y_i)^2$$

is minimized. The resulting line is called the linear regression line for the given points.

Show that the coefficients a and b of the linear regression line are given by

$$b = \bar{y} - a\bar{x}, \quad a = \frac{\sum_{i=1}^n x_i y_i - \sum_{i=1}^n x_i \sum_{i=1}^n y_i}{\sum_{i=1}^n x_i^2 - \sum_{i=1}^n x_i \sum_{i=1}^n x_i},$$

where

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i, \quad \bar{y} = \frac{1}{n} \sum_{i=1}^n y_i.$$

Exercise 2.5 (Best Mean-Square Approximation by Orthogonal Functions). Let $\varphi_1, \dots, \varphi_n$ be continuous and not identically zero on an interval $[a, b]$ and such that

$$\int_a^b \varphi_j(x) \varphi_k(x) dx = 0 \quad \text{if } j \neq k.$$

Given a piecewise continuous function f on $[a, b]$ show that the values of $c_1, \dots, c_n$ that minimize

$$F(c) = \int_a^b \left| f - \sum_{j=1}^n c_j \varphi_j \right|^2 dx$$

are given by the formulas

$$c_i = \frac{\langle \varphi_i, f \rangle}{\langle \varphi_i, \varphi_i \rangle}, \quad i = 1, \dots, n,$$

where

$$\langle \varphi, \psi \rangle = \int_a^b \varphi(x) \psi(x) dx.$$

Exercise 2.10. In this exercise we shall establish a series of results, some of which are of interest in their own right, that lead to a proof of Lemma II.1.2.

(a) Let $Q(x) = \langle x, Ax \rangle$, where A is a symmetric matrix. Show that if Q is positive definite, then A is nonsingular. (Hint: If A were singular, there would exist an $x_0 \neq 0$ such that $Ax_0 = 0$.)

(b) Show that if Q is positive semidefinite and A is nonsingular, then Q is positive definite. (Hint: If A is positive semidefinite, then there exists an $x_0 \neq 0$ that minimizes Q .)

(c) If Q is positive definite, then $\det A > 0$. (Hint: Let I denote the identity matrix and let $0 < \lambda < 1$. For each $0 < \lambda < 1$ consider the quadratic form

$$P(x, y) = (1 - y)(x)^2 + y(g(x)) = (x, y)(1 - y)I + yA(x).$$

Let $d(y) = \det[(1 - y)I + yA]$. Show that d is a continuous function, that $d(0) \neq 0$, and that $d(1) = 1$.

(c) Let the function f with domain $\mathbb{R}^n$ be defined by

$$f(x) = (x, Ax) + 2(b, x) + c,$$

where A is an $n \times n$ positive-definite matrix and b is an n -vector. Show that $x_0 = -A^{-1}b$ is the unique minimizer of f and that

$$f(x_0) = (b, x_0) + c = -\frac{\det \begin{pmatrix} A & b \\ b^T & c \end{pmatrix}}{\det A}.$$

Hint:

$$\det \begin{pmatrix} A & b \\ b^T & c \end{pmatrix} = \det \begin{pmatrix} A & Ax_0 + b \\ (b, x_0) + c & \end{pmatrix}.$$

Why?

(d) Let

$$g(x, y) = (x, Ax) + 2y(b, x) + cy^2 = (x, y)B(x, y),$$

where

$$B = \begin{pmatrix} A & b \\ b^T & c \end{pmatrix}.$$

Show that the quadratic form g is positive definite if and only if A is positive definite and $\det B > 0$. Hint:

(i) Since $g(x, 0) = (x, Ax)$, if g is positive definite, then A is positive definite.

(ii) For $y \neq 0$,

$$g(x, y) = y^2 g\left(\frac{x}{y}, 1\right)$$

(iii) If A is positive definite, then

$$g(x, 0) = f(x) > f(x_0).$$

(c) Prove Lemma III.5.2. (Hint: If A is positive definite, use (i) and induction. If A is negative definite, then $-A$ is positive definite.

3. OPTIMIZATION OF CONVEX FUNCTIONS

The points at which a convex function attains a maximum or minimum have important properties not shared by all functions.

THEOREM 3.1. *Let f be a convex function defined on a convex set C .*

- (i) *If f attains a local minimum at x_0 , then f attains a minimum at x_0 .*
- (ii) *The set of points at which f attains a minimum is either empty or convex.*
- (iii) *If f is strictly convex and f attains a minimum at x_0 , then x_0 is unique.*
- (iv) *If f is not a constant function and if f attains a maximum at some point x in C , then x must be a boundary point of C .*

Proof. (i) Since f attains a local minimum at x_0 , there exists a $\delta > 0$ such that for all x in $B(x_0, \delta) \cap C$

$$f(x) - f(x_0) \geq 0. \quad (1)$$

Let y be an arbitrary point in C . Then $[x_0, y]$ belongs to C . Hence for $0 < t < \delta/(y - x_0)$ the point $x = x_0 + t(y - x_0)$ belongs to $B(x_0, \delta) \cap C$. Hence, from Lemma III.1.7 and (1) we get

$$f(y) - f(x_0) \geq \frac{f(x) - f(x_0)}{t} \geq 0.$$

(ii) If f attains a minimum at a unique point, then there is nothing to prove. Suppose that f attains a minimum at two distinct points x_1 and x_2 . Let $\mu = f(x_1) = f(x_2) = \min\{f(x) : x \in C\}$. Then for any (α, β) in P_2 with $\alpha + \beta = 1$, $\beta > 0$,

$$\mu \leq f(\alpha x_1 + \beta x_2) \leq \alpha f(x_1) + \beta f(x_2) = \alpha\mu + \beta\mu = \mu.$$

Hence f attains a minimum at each point of $[x_1, x_2]$.

(iii) Suppose that f attains a minimum at two distinct points x_1 and x_2 . Let μ be as in (ii). Then for any (α, β) in P_2 ,

$$\mu \leq f(\alpha x_1 + \beta x_2) \leq \alpha f(x_1) + \beta f(x_2) = \mu.$$

Thus, we get a contradiction, $\mu < \mu$, and so x_1 and x_2 cannot be distinct.

(iv) Suppose f attains a maximum at x . If x were not a boundary point, then x would be an interior point. Since f is not a constant function, there exists a point y in C such that $f(y) < f(x)$. Since C is convex, the line segment $[x, y]$ belongs to C . Since x is an interior point, there is a line segment $[x, y]$ such that x is an interior point of $[x, y]$. Then $x = \alpha + \beta y$ for some (α, β) in

P_1 with $\alpha > 0$, $\beta > 0$, and

$$\beta(x) \leq \alpha f(x) + \beta f(y) < \alpha f(x) + \beta f(x) = f(x).$$

Thus we have the contradiction $f(x) < f(x)$.

The next theorem and its corollary characterize points x_0 at which a function attains a minimum.

THEOREM 3.1. *Let f be a real-valued function defined on a set C . Then f attains a minimum at x_0 in C if and only if $\nabla f(x_0) = 0$.*

Proof. If f attains a minimum at x_0 , then for all x in C ,

$$\beta(x) \geq \beta(x_0) = f(x_0) + \langle \nabla f(x_0), x - x_0 \rangle. \quad (2)$$

Hence $\nabla f(x_0) = 0$. Conversely, if $\nabla f(x_0) = 0$, then (2) holds and f attains a minimum at x_0 .

Note that the theorem does not require f or C to be convex. If, however, f is not convex, it can fail to have a subgradient at all points of its domain. Such a function therefore cannot attain a minimum. The logarithm function in x , or any other logarithmic function, is an example. The next corollary does require f and C to be convex.

COROLLARY 1. *Let f be convex and differentiable on an open convex set C . Then f attains a minimum at x_0 in C if and only if $\nabla f(x_0) = 0$.*

Proof. Let f attain a minimum at x_0 . That $\nabla f(x_0) = 0$ follows from Theorem 3.1, whether or not f and C are convex. If f is convex and differentiable on an open convex set C and $\nabla f(x_0) = 0$, then by Theorem III.3.1

$$\beta_0 = \nabla f(x_0) = 0 \in \beta(x_0).$$

Thus, by the theorem, f attains a minimum at x_0 .

Remark 3.1. In geometric terms, Theorem 3.1 states that for a convex function f defined on a convex set C , a point x_0 in C furnishes a minimum if and only if the epigraph of f has a horizontal supporting hyperplane at $(x_0, \beta(x_0))$. The corollary states that for a differentiable convex function defined on an open convex set C , a point x_0 furnishes a minimum if and only if the tangent plane to the graph of f at x_0 is horizontal.

The corollary also points out the true content of Theorem 2.4. By Theorem III.3.3, the positive definiteness or semidefiniteness of $\beta(x)$ on X implies that

f is convex. It is the convexity of f that is the important element in Theorem 2.4. The alert reader may have observed that much of the proof of Theorem 2.4 was a repetition of the proof of Theorem H1.1.3.

Exercise 3.1. Show that $f(x) = x_1^2 + x_2^2 - 2x_1 - 7x_2$ attains a minimum on $\mathbb{R}^2$ at a unique point and find the minimum.

Exercise 3.2. Maximize $f(x) = x_1' + x_2x_3 + x_3'$ subject to

$$-2x_1 - 2x_2 + 4 \leq 0, \quad -x_2 + x_3 - 3 \leq 0, \quad x_3 - 2 \leq 0.$$

- Sketch the feasible set.
- Show that a solution exists.
- Find the solution.

4. LINEAR PROGRAMMING PROBLEMS

Linear programming is used extensively in resource allocation problems. To illustrate, we formulate a simple job-scheduling problem.

An office furniture manufacturer makes a different types of desks. His profit per desk of type j is b_j . Thus, if each week he produces x_j desks of type j , then his profit for the week is

$$b_1x_1 + \cdots + b_sx_s.$$

The manufacturer's objective is to maximize his profit. The number of desks that he can produce per week is not unlimited but is constrained by the capacity of his plant. Suppose that the manufacture of a type 1 desk requires a_{11} hours per week on machine 1, that the manufacture of a type 2 desk requires a_{12} hours of time on machine 1, and the manufacture of a type j desk requires a_{1j} hours per week of time on machine 1. If $x_j, j = 1, \dots, s$, denotes the number of desks of type j manufactured per week, then the number of hours that machine 1 is used per week is

$$a_{11}x_1 + a_{12}x_2 + \cdots + a_{1s}x_s.$$

Because of maintenance requirements, readjustments, and the number of shifts employed, the total time that machine 1 is in operation cannot exceed some fixed preassigned number c_1 . Thus we have a constraint

$$a_{11}x_1 + \cdots + a_{1s}x_s \leq c_1.$$

Similarly, if the manufacture of one desk of type j requires a_{ij} hours per week

on machine i and the total number of hours per week that machine i can operate is c_i , then

$$a_{1j}x_j + \cdots + a_{ij}x_j \leq c_i, \quad i = \bar{1}, \bar{2}, \dots, \bar{n}.$$

Also the number of disks of type j made is nonnegative, so $x_j \geq 0$.

The problem for the disk manufacturer is to manufacture x_j disks of type j so as to maximize

$$\sum_{j=1}^n b_j x_j$$

subject to

$$\sum_{j=1}^n a_{ij} x_j \leq c_i, \quad i = 1, \dots, n, \quad x_j \geq 0, \quad j = 1, \dots, n.$$

This is a linear programming problem in standard form.

For a brief account of the origins and early days of linear programming, by a pioneer in the field, see Dantzig [1963]. Some early applications are also discussed there.

In this section we shall derive a necessary and sufficient condition that a point x_0 be a solution to a linear programming problem. In practice the use of this condition to solve the problem is precluded by the dimension of the space E^n . Instead, a direct method, the simplex method, is used. The simplex method will be discussed in Chapter VI. Nevertheless, we derive the necessary conditions for linear programming problems to illustrate the simplest case of the use of Farkas's theorem (and hence separation theorems) to derive necessary conditions in constrained optimization problems. The basic idea in this proof is a theme that runs throughout much of optimization theory.

The linear programming problem in standard form (SLP) is given as

$$\begin{aligned} & \text{Minimize } \langle -b, x \rangle \\ & \text{subject to } Ax \leq c, \quad x \geq 0, \end{aligned}$$

where $x \in E^n$, A is an $m \times n$ matrix, $b, c \in E^m$ is a vector in E^n , and c is a vector in E^m .

The linear programming problem SLP can also be written as follows: Minimize $\langle -b, x \rangle$ subject to

$$\begin{pmatrix} A \\ -I \end{pmatrix} x \leq \begin{pmatrix} c \\ 0 \end{pmatrix} \quad (1)$$

where I is the $n \times n$ identity matrix and $\mathbf{0}$ is the zero vector in $\mathbb{R}^n$. If we denote the matrix on the left in (2) by $\tilde{A}$ and the vector on the right by $\tilde{c}$ and then drop the tildes, we can state the linear programming problem (LPI) as

$$\begin{aligned} & \text{minimize } \langle -\mathbf{b}, \mathbf{x} \rangle \\ & \text{subject to } A\mathbf{x} \leq \mathbf{c}, \end{aligned}$$

where A is an $m \times n$ matrix, $\mathbf{b} \in \mathbb{R}^n$, and $\mathbf{c} \in \mathbb{R}^m$. Here, we have incorporated the constraint $\mathbf{x} \geq \mathbf{0}$ in the matrix A and the vector $\mathbf{c}$.

THEOREM 4.1. Let $\mathbf{x}_0$ be a solution of LPI. Then there exists a $\mathbf{k} \neq \mathbf{0}$ in $\mathbb{R}^m$ such that $(\mathbf{x}_0, \mathbf{k})$ satisfies the following:

$$\begin{aligned} \text{(i)} \quad & A\mathbf{x}_0 \leq \mathbf{c}, & \text{(ii)} \quad A\mathbf{k} = \mathbf{b}, \\ \text{(iii)} \quad & \mathbf{k} \geq \mathbf{0}, & \text{(iv)} \quad \langle A\mathbf{x}_0 - \mathbf{c}, \mathbf{k} \rangle = 0, \\ \text{(v)} \quad & \langle \mathbf{b}, \mathbf{x}_0 \rangle = \langle \mathbf{c}, \mathbf{k} \rangle. \end{aligned} \quad (2)$$

If there exists an $\mathbf{x}_0 \in \mathbb{R}^n$ and a $\mathbf{k} \neq \mathbf{0}$ in $\mathbb{R}^m$ that satisfy (i)–(v), then $\mathbf{x}_0$ is a solution of LPI.

Proof. The feasible set for LPI is $X = \{\mathbf{x} \mid A\mathbf{x} \leq \mathbf{c}\}$. Since $\mathbf{x}_0$ is a solution to LPI, $\mathbf{x}_0 \in X$, and so (i) holds.

We now prove (ii) and (iii). If $\mathbf{x}_0$ in X is a solution to LPI, then $\langle -\mathbf{b}, \mathbf{x}_0 \rangle \leq \langle -\mathbf{b}, \mathbf{x} \rangle$ for all $\mathbf{x}$ in X . Hence $\langle \mathbf{b}, \mathbf{x} \rangle \leq \langle \mathbf{b}, \mathbf{x}_0 \rangle$ for all $\mathbf{x}$ in X . If we set $\mathbf{v} = \langle \mathbf{b}, \mathbf{x}_0 \rangle$, then we have that X is contained in the closed negative half space determined by the hyperplane $H_{\mathbf{v}}$. By Exercise B.2.15, for $\mathbf{x}$ in $\text{int}(X)$, we have $\langle \mathbf{b}, \mathbf{x} \rangle < v$, and so $\mathbf{x}_0$ cannot be an interior point of X .

Let $\mathbf{a}_i$ denote the i th row of the matrix A . Let

$$E = \{i \mid \langle \mathbf{a}_i, \mathbf{x}_0 \rangle = c_i\}, \quad F = \{i \mid \langle \mathbf{a}_i, \mathbf{x}_0 \rangle < c_i\}.$$

Since $\mathbf{x}_0$ is not an interior point of X , the set $E \neq \emptyset$. Let A_E denote the submatrix of A consisting of those rows $\mathbf{a}_i$ with i in E , and let A_F denote the submatrix of A consisting of those rows $\mathbf{a}_i$ with i in F .

We claim that the system

$$A_E \mathbf{x} \leq \mathbf{0}, \quad \langle \mathbf{b}, \mathbf{x} \rangle > 0, \quad (3)$$

has no solution $\mathbf{x}$ in $\mathbb{R}^n$.

Intuitively this is clear, for if $\mathbf{x}_0$ would be a solution, then we would move, for $\epsilon > 0$, and we would have $A_E(\mathbf{x}_0 + \epsilon \mathbf{x}_0) \leq \mathbf{c}_E$. Also for ϵ sufficiently small, by continuity, we would have $A_F(\mathbf{x}_0 + \epsilon \mathbf{x}_0) \leq \mathbf{c}_F$. Thus, if (3) had a solution, we would be able to move into the feasible set and at the same time increase $\langle \mathbf{b}, \mathbf{x}_0 \rangle$ or decrease $\langle -\mathbf{b}, \mathbf{x}_0 \rangle$. This would contradict the optimality of $\mathbf{x}_0$.

We now make this argument precise. Again, if x_0 were a solution of (1), then x_0 would be a solution of (3) for all $\alpha > 0$, and

$$\zeta(-k, x_0 + \alpha x_0) = \zeta(-k, x_0) - \alpha \langle k, x_0 \rangle < \zeta(-k, x_0) \quad (6)$$

for all $\alpha > 0$. Also,

$$A_0(x_0 + \alpha x_0) = x_0 + \alpha A_0 x_0 < x_0 \quad \alpha > 0, \quad (7)$$

and

$$A_0(x_0 + \alpha x_0) = A_0 x_0 + \alpha A_0 x_0 = (p_1 - p) + \alpha A_0 x_0,$$

where $p > 0$. (Recall that $p > 0$ means that all components of p are strictly positive.) By taking $\alpha > 0$ to be sufficiently small, we can make $\alpha A_0 x_0 - p < 0$. Hence for $\alpha > 0$ and sufficiently small

$$A_0(x_0 + \alpha x_0) = x_0 + (\alpha A_0 x_0 - p) < x_0. \quad (8)$$

In summary, from (5) and (8), we get that for sufficiently small $\alpha > 0$ the vector $x_0 + \alpha x_0$ is in X . From this and from (4), we get that x_0 is not a solution to the linear programming problem. This contradiction shows that x_0 cannot be a solution of (1).

Since the system (2) has no solution, it follows, from Theorem II.5.1 (Farkas's lemma) that there exists a vector y in $\mathbb{R}^m$, where m is the row dimension of A_0 , such that

$$A_0^T y = k, \quad y \geq 0.$$

Therefore, if m is the row dimension of A_0 , then

$$A_0^T y + A_1^T z = k,$$

where z is the zero vector in $\mathbb{R}^n$. If we now take $\lambda = (y, z)$, then λ is in $\mathbb{R}^m$ and satisfies

$$A\lambda = (A_0^T, A_1^T) \begin{pmatrix} y \\ z \end{pmatrix} = k, \quad \lambda \geq 0. \quad (9)$$

Thus, (i) and (ii) are established.

From (7) we see that those components of λ corresponding to indices in J are zero. By the definition of $\bar{L}$, $A_1 x_0 - x_0 = \bar{L}$. Hence

$$\langle Ax_0 - x, \lambda \rangle = \langle A_0 x_0 - x_0, A_1 x_0 - x_0 \rangle = \langle \bar{L}, \bar{L} \rangle = 0, \quad (10)$$

which is (iv).

From (iv) and (ii) we get

$$\langle \mathbf{a}_i, \mathbf{k} \rangle = \langle \mathbf{A}\mathbf{x}_0, \mathbf{k} \rangle = \langle \mathbf{x}_0, \mathbf{A}^T \mathbf{k} \rangle = \langle \mathbf{x}_0, \mathbf{b} \rangle,$$

which is (v).

The proof of sufficiency is elementary and does not involve separation theorems. Let $\mathbf{x}$ be any feasible point. Then

$$\begin{aligned} \langle -\mathbf{b}, \mathbf{x} \rangle - \langle -\mathbf{b}, \mathbf{x}_0 \rangle &= \langle -\mathbf{A}^T \mathbf{k}, \mathbf{x} \rangle + \langle \mathbf{x}_0, \mathbf{k} \rangle \\ &= \langle \mathbf{k}, -\mathbf{A}\mathbf{x} + \mathbf{e} \rangle = -\langle \mathbf{e}, \mathbf{A}\mathbf{x} - \mathbf{e}, \mathbf{k} \rangle \geq 0. \end{aligned}$$

Hence $\mathbf{x}_0$ is a solution to L.P.

Remark 4.1. An alternate proof of (iv) and (v) is to establish (v) first and then use (v) to establish (iv). Thus

$$\begin{aligned} \langle \mathbf{b}, \mathbf{x}_0 \rangle &= \langle \mathbf{A}^T \mathbf{k}, \mathbf{x}_0 \rangle = \langle \mathbf{k}, \mathbf{A}\mathbf{x}_0 \rangle = \left\langle \mathbf{y}, \mathbf{B}_0 \begin{pmatrix} \mathbf{A}_0 \\ \mathbf{A}_1 \end{pmatrix} \mathbf{x}_0 \right\rangle \\ &= \langle \mathbf{y}, \mathbf{A}_0 \mathbf{x}_0 \rangle = \langle \mathbf{y}, \mathbf{e}_0 \rangle = \langle \mathbf{y}, \mathbf{B}_0 (\mathbf{x}_0, \mathbf{e}_1) \rangle \\ &= \langle \mathbf{k}, \mathbf{e} \rangle, \end{aligned}$$

which is (v). To establish (iv), we write

$$\begin{aligned} 0 &= \langle \mathbf{b}, \mathbf{x}_0 \rangle - \langle \mathbf{b}, \mathbf{k} \rangle = \langle \mathbf{A}^T \mathbf{k}, \mathbf{x}_0 \rangle - \langle \mathbf{b}, \mathbf{k} \rangle \\ &= \langle \mathbf{k}, \mathbf{A}\mathbf{x}_0 \rangle - \langle \mathbf{b}, \mathbf{k} \rangle = \langle \mathbf{A}\mathbf{x}_0 - \mathbf{b}, \mathbf{k} \rangle. \end{aligned}$$

Remark 4.2. It follows from (i) and (iii) that (iv) holds if and only if for each $i = 1, \dots, m$

$$\langle \mathbf{a}_i, \mathbf{x}_0 \rangle - c_i b_i = 0, \quad \text{if } i$$

That is, (iv') is equivalent to (iv). The condition (iv') is sometimes called the complementary slackness condition. It states that if there is "slack" in the i th constraint at $\mathbf{x}_0$, that is, if $\langle \mathbf{a}_i, \mathbf{x}_0 \rangle < c_i b_i$, then $\lambda_i = 0$.

Remark 4.3. Let the linear programming problem be reformulated with X_0 taken to be an arbitrary open set instead of $\mathbb{R}^n$. It is clear from the proof that if $\mathbf{x}_0$ is a solution to this problem, then the necessary conditions hold. Conversely, if there is a point $\mathbf{x}_0$ in X_0 at which the necessary conditions hold, then $\mathbf{x}_0$ is feasible and the minimum is achieved at $\mathbf{x}_0$.

The inequality constraints $Ax \leq r$ can be converted to equality constraints by introducing a nonnegative vector $\bar{q}$ whose components are called slack variables and writing

$$Ax + \bar{q} = r, \quad \bar{q} \geq \bar{0}.$$

Thus the linear programming problem formulated in (1) can be written as

$$\begin{aligned} &\text{Minimize } (-b, \bar{0}) \cdot (x, \bar{q}) \\ &\text{subject to } (A, I) \begin{pmatrix} x \\ \bar{q} \end{pmatrix} = r, \quad x \geq \bar{0}, \bar{q} \geq \bar{0}. \end{aligned}$$

That is, the problem LPI can be written as a canonical linear programming (CLP) problem

$$\begin{aligned} &\text{Minimize } (-b, \bar{0}) \cdot x \\ &\text{subject to } Ax = r, \quad x \geq \bar{0}. \end{aligned}$$

The simplex method is applied to problems in canonical form.

Conversely, the problem CLP can be written as a problem LPI by writing the equality constraint $Ax = r$ as two inequality constraints, $Ax \leq r$ and $-Ax \leq -r$, and then incorporating the constraint $x \geq \bar{0}$ in the matrix. That is, we write CLP as

$$\begin{aligned} &\text{Minimize } (-b, \bar{0}) \cdot x \\ &\text{subject to } \begin{pmatrix} A \\ -A \end{pmatrix} x \leq \begin{pmatrix} r \\ -r \end{pmatrix}, \end{aligned}$$

Remark 4.1. We warn the reader that the definitions "standard form" and "canonical form" are not uniform throughout the linear programming literature. Problems that we have designated as being in canonical form are designated as being in standard form by some authors who use the designation "canonical form" differently from ours. In consulting the literature the reader should always check the definitions used.

Exercise 4.1. Prove the following corollaries to Theorem 4.1.

COROLLARY 1. Let x_0 be a solution of the standard linear programming problem SLP. Then there exists a vector $k \geq \bar{0}$ in $\mathbb{R}^n$ such that (x_0, k) satisfies

- (i) $Ax_0 \leq a$, $x_0 \geq 0$
- (ii) $\sum_{j=1}^m \lambda_j a_{ij} \geq b_i, i = 1, \dots, n$, with equality holding if $x_{0j} > 0$
- (iii) $\lambda \geq 0$
- (iv) $c(Ax_0 - a, \lambda) = 0$ and
- (v) $(\lambda, x_0) = (\alpha, \bar{x})$.

If there exist an x_0 in $\mathbb{R}^n$ and a $\lambda, \lambda \neq 0$ in $\mathbb{R}^m$ such that (i)–(v) hold, then x_0 is a solution of CLP.

COROLLARY 2. Let x_0 be a solution of the canonical linear programming problem CLP. Then there exists a vector $\lambda, \lambda \neq 0$ in $\mathbb{R}^m$ such that (x_0, λ) satisfies

- (i) $Ax_0 = a$, $x_0 \geq 0$
- (ii) $\sum_{j=1}^m \lambda_j a_{ij} \geq b_i, i = 1, \dots, n$, with equality holding if $x_{0j} > 0$
- (iii) $c(Ax_0 - a, \lambda) = 0$ and
- (iv) $(\lambda, x_0) = (\alpha, \bar{x})$.

If there exist an x_0 in $\mathbb{R}^n$ and a $\lambda, \lambda \neq 0$ in $\mathbb{R}^m$ such that (i)–(iv) hold, then x_0 is a solution of CLP.

Note that in Corollary 2 the condition $\lambda \geq 0$ is absent.

Example 4.2. Show that the problem LP1 can be written as a problem in standard form.

9. FIRST-ORDER CONDITIONS FOR DIFFERENTIABLE NONLINEAR PROGRAMMING PROBLEMS

In this section we shall derive first-order necessary conditions for differentiable programming problems. As noted in the introduction of this chapter, the set of points that satisfy the necessary conditions need not furnish a solution to the programming problem. If a solution exists, then it belongs to this set. Satisfaction of the necessary conditions corresponds to "bringing in the usual suspects" in a criminal investigation. Further investigation is required to determine which, if any, of these points, or "suspects," furnish a solution. In this vein, we shall show that if the data of the problem are convex, then a point that satisfies the necessary conditions is a solution.

Problem P1

Let X_0 be an open convex set in $\mathbb{R}^n$. Let f, g , and h be C^1 functions with domain X_0 and ranges in $\mathbb{R}$, $\mathbb{R}^m$, and $\mathbb{R}^p$, respectively. Let

$$K = \{x : x \in X_0, g(x) \leq 0, h(x) = 0\}.$$

Minimize f over K .

The problem is often stated as

$$\begin{aligned} & \text{Minimize } f(\mathbf{x}) \\ & \text{subject to } g(\mathbf{x}) \leq 0, \quad h(\mathbf{x}) = 0. \end{aligned}$$

If the constraints $g(\mathbf{x}) \leq 0$ and $h(\mathbf{x}) = 0$ are absent, the problem becomes an unconstrained problem, which was treated in Section 2. We saw there that a necessary condition for a point $\mathbf{x}_0$ to be a solution is that all first-order partial derivatives of f vanish at $\mathbf{x}_0$. For $n = 2$, 3 and $1 \leq k < n$, the problem with equality constraints and no inequality constraints is often usually treated in elementary calculus courses. It is stated there, with or without proof, that if $\mathbf{x}_0$ is a solution of the problem and if the k gradient vectors $\nabla h_j(\mathbf{x}_0)$ are linearly independent, then there exist scalars μ_j , $j = 1, \dots, k$, such that for each $j = 1, \dots, n$

$$\frac{\partial f}{\partial x_j}(\mathbf{x}_0) + \sum_{j=1}^k \mu_j \frac{\partial h_j}{\partial x_j}(\mathbf{x}_0) = 0$$

and

$$h_j(\mathbf{x}_0) = 0, \quad j = 1, \dots, k.$$

This is the Lagrange multiplier rule.

We now give a first-order necessary condition that is satisfied by a solution of problem PL. The condition includes a more general version of the Lagrange multiplier rule in the case that inequality constraints are absent.

We recall the notation established in Section 11 of Chapter 1. If f is a real-valued function and $\mathbf{g} = (g_1, \dots, g_m)$ is a vector-valued function that is differentiable at $\mathbf{x}_0$, then

$$\nabla f(\mathbf{x}_0) = \left(\frac{\partial f}{\partial x_1}(\mathbf{x}_0), \frac{\partial f}{\partial x_2}(\mathbf{x}_0), \dots, \frac{\partial f}{\partial x_n}(\mathbf{x}_0) \right)$$

and

$$\nabla \mathbf{g}(\mathbf{x}_0) = \begin{pmatrix} \nabla g_1(\mathbf{x}_0) \\ \nabla g_2(\mathbf{x}_0) \\ \vdots \\ \nabla g_m(\mathbf{x}_0) \end{pmatrix} = \left(\frac{\partial g_j}{\partial x_i}(\mathbf{x}_0) \right).$$

THEOREM 5.1. Let $\mathbf{x}_0$ be a solution of problem PL. Then there exists a real number $\lambda_0 \geq 0$, a vector $\lambda \geq 0$ in $\mathbb{R}^m$, and a vector μ in $\mathbb{R}^k$ such that

- (i) $(\lambda_0, \lambda, \mu) \neq 0$,
- (ii) $\langle \lambda_0, g(x_0) \rangle = 0$, and
- (iii) $\lambda_0 f'(x_0) + \lambda' \nabla g(x_0) + \mu' \nabla h(x_0) = 0$.

Vectors $(\lambda_0, \lambda, \mu)$ having the properties stated in the theorem are called *multipliers*, as are the components of these vectors.

Definition 3.1. A point x_0 at which the conclusion of Theorem 3.1 holds will be called a *critical point* for problem P1.

Thus, Theorem 3.1 says that a solution x_0 of problem P1 is a critical point.

Remark 3.1. The necessary condition asserts that $\lambda_0 \geq 0$ and not the stronger statement $\lambda_0 > 0$. If $\lambda_0 > 0$, then we may divide through by λ_0 in (i) and (iii) and relabel λ/λ_0 and μ/λ_0 as λ and μ , respectively, and then obtain statements (i)–(iii) with $\lambda_0 = 1$. In the absence of further conditions that would guarantee that $\lambda_0 > 0$, we cannot assume that $\lambda_0 = 1$. We shall return to this point later.

Remark 3.2. If $g = 0$, that is, the inequality constraints are absent, Theorem 3.1 becomes the Lagrange multiplier rule.

Remark 3.3. Since $\lambda \geq 0$ and $g(x_0) \leq 0$, condition (ii) is equivalent to

$$\lambda_j g_j(x_0) = 0, \quad j = 1, \dots, m. \quad (3')$$

Condition (iii) is a system of n equations

$$\lambda_0 \frac{\partial f}{\partial x_j}(x_0) + \sum_{i=1}^m \lambda_i \frac{\partial g_i(x_0)}{\partial x_j} + \sum_{i=1}^p \mu_i \frac{\partial h_i(x_0)}{\partial x_j} = 0, \quad j = 1, \dots, n. \quad (3'')$$

Remark 3.4. The necessary conditions are also necessary conditions for a local minimum. For if x_0 is a local minimum, then there exists a $\delta > 0$ such that $f(x_0) \leq f(x)$ for all x that are in $B(x_0, \delta) \cap X_0$ and that satisfy the constraints $g(x) \leq 0$ and $h(x) = 0$. Thus, x_0 is a global solution to the problem in which X_0 is replaced by $X_0 \cap B(x_0, \delta)$.

Remark 3.5. Theorem 3.1, under additional hypotheses guaranteeing that $\lambda_0 > 0$, was proved in Kuhn and Tucker [1951]. This paper had great impact since interest in nonlinear programming was developing at that time in economics and other areas of applications. The theorem with $\lambda_0 \geq 0$ had been proved in Fiacchi [1948], and prior to that, again under conditions ensuring that $\lambda_0 > 0$, in Karush, [1959]. The results of Fiacchi and of Karush were mentioned when they appeared, probably because interest in programming problems had not yet come about. These papers were looked upon only as mathematical

generalizations of the Lagrange multiplier rule. In the literature, Theorem 5.1, with additional hypotheses guaranteeing that $\lambda_0 = 0$, is often referred to as the Karush–Kuhn–Tucker Theorem. Theorem 5.1, as stated here, is referred to as the Fiacz–John theorem. We shall adopt this terminology and refer to Theorem 5.1 as the FJ theorem and call the multipliers (λ_0, λ, k) FJ multipliers.

The next example, due to Fiacz and McCormick ([198]), illustrates the use of Theorem 5.1 and will be used to illustrate some subsequent theorems.

Example 4.1. Let $X_0 = \mathbb{R}^2$, let

$$f(x) = (x_1 - 1)^2 + x_2^2,$$

and let

$$g(x) = 2kx_1 - x_2^2 \leq 0, \quad k > 0.$$

For each value of $k > 0$ we obtain a programming problem n_k . The feasible set X_k for problem n_k is the set of points on and exterior to the graph of the parabola $x_2^2 = 2kx_1$, and the function to be minimized is the distance from x to the point $(1, 0)$. In Figure 4.1 we have sketched some level curves of f and the parabolas for two values of k . Although the solution can be read from Figure 4.1, we will apply the appropriate theorems in this section and in the next section to solve this problem.

By straightforward calculation

$$\nabla f = (2(x_1 - 1), 2x_2), \quad \nabla g = (2k, -2x_2).$$

From Theorem 5.1, we get that if a point x is a solution to the programming problem n_k , then there exist a $\lambda_0 \geq 0$ and a $\lambda \geq 0$ with $(\lambda_0, \lambda) \neq (0, 0)$ such that (λ_0, λ) and x satisfy

$$\lambda_0(2(x_1 - 1), 2x_2) + \lambda(2k, -2x_2) = (0, 0). \quad (1)$$

$$\lambda(2kx_1 - x_2^2) = 0. \quad (2)$$

If $\lambda_0 = 0$, then since $\lambda \neq 0$, equations (1) would imply that $k = 0$, contradicting the assumption that $k > 0$. Hence we may take $\lambda_0 = 1$. If $\lambda = 0$, then (1) would imply that $x_1 = 1$ and $x_2 = 0$. The point $(1, 0)$ is not feasible for any value of $k > 0$, so we may assume that $\lambda \neq 0$ in (2). It then follows from (2) that $2kx_1 - x_2^2 = 0$; that is, the point x must lie in the parabola. Thus we may replace (1) and (2) by

$$k(x_1 - 1, x_2) + \lambda(k, -x_2) = (0, 0), \quad x_2^2 = 2kx_1. \quad (3)$$

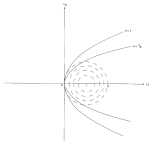


Figure 4.1

If $x_2 \neq 0$, then we must have $k = 1$, and so

$$x_1 = 1 - k, \quad x_2 = \pm \sqrt{(2k) - k^2}, \quad 0 < k < 1.$$

If $x_2 = 0$, then from the second equation in (1) we get that $x_1 = 0$, and consequently $k = 1/2$.

In summary, we have shown that if $0 < k < 1$, then the only points $\mathbf{x} = (x_1, x_2)$ and multipliers (λ, μ) that satisfy the necessary conditions for problem n_k are

$$\lambda = 1, \quad x_1 = 1 - k, \quad x_2 = \pm \sqrt{(2k) - k^2}, \quad (4)$$

$$\lambda = \frac{1}{2}, \quad x_1 = 0, \quad x_2 = 0. \quad (5)$$

If $k \geq 1$, the necessary conditions are only satisfied by (5). Note that if $k = 1$, then (4) and (5) both give $k = 1$ and $x_1 = x_2 = 0$.

We emphasize that Theorem 5.1 is a necessary condition, so that without further arguments we cannot determine which, if any, of the points (4) and (5) furnish a minimum in the problem n_k ; we are minimizing the distance from the point $(1, 0)$ to the closed set X_k . By Lemma B.3.3 a solution to this problem exists. Since for $k \geq 1$ the only point that satisfies the necessary conditions is the origin, this point is the solution to problem n_k for $k \geq 1$. If we denote the two points in (4) by x_1 and x_2 , then

$$f(x_1) = f(x_2) = -k^2 + 2k.$$

The point in (5) is the origin, and $f(0) = 1$. Since $-k^2 + 2k < 1$ for $k \neq 1$, a solution to problem n_k for $k \neq 1$ is given by each of the points x_1 and x_2 .

Before we present our proof of Theorem 5.1, we shall present the principal idea of a proof which, with variations, is popular in the literature. We suppose that the equality constraints have been converted into inequality constraints, $h(x) \leq 0$ and $-h(x) \leq 0$, so that the only constraints are $g(x) \leq 0$. By a translation of coordinates, we can assume that $x_0 = 0$ and that $f(x_0) = f(0) = 0$. We also assume that $\nabla f(0) \neq 0$.

As in the linear programming problem, let

$$K = \{x \mid g(x) = 0\}, \quad I = \{x \mid g(x) < 0\}. \quad (6)$$

The constraints $g(x) = 0$ with i in J are said to be *active or binding* at 0 those with $i \in I$ are said to be *inactive* at 0 .

Since f and g are C^1 , they are differentiable by Theorem B.1.1. Therefore, since $f(0) = 0$ and $g_0(0) = 0$,

$$f(x) = \langle \nabla f(0), x \rangle + o(\|x\|),$$

$$g_i(x) = \nabla g_i(0)x + o(\|x\|),$$

where $o(\|x\|)$ and $o(\|x\|)$ represent terms such that $o(\|x\|)/\|x\| \rightarrow 0$ and $o(\|x\|)/\|x\| \rightarrow 0$ as $\|x\| \rightarrow 0$. Thus, $o(\|x\|)$ and $o(\|x\|)$ represent higher order terms.

Since g is continuous and $g_0(0) < 0$, there exists a $\delta > 0$ such that $g_0(x) < 0$ whenever $\|x\| < \delta$. The origin is, therefore, a solution of the local problem $\text{Minimize } f(x) \text{ subject to } g_0(x) \leq 0 \text{ and } \|x\| = \delta$. If we drop the terms $o(\|x\|)$ and $o(\|x\|)$, that is, if we linearize the problem about the solution 0 and assume that the problem is the linearized problem, then $x_0 = 0$ is the solution of the linear programming problem:

$$\text{Minimize } \langle \nabla f(0), x \rangle = \langle -(-\nabla f(0)), x \rangle,$$

$$\text{subject to } \nabla g_i(0)x \leq 0, \quad \|x\| < \delta. \quad (7)$$

By Remark 4.3, we can now apply Theorem 4.1 with $J = \nabla g_{\mathbf{H}_0}(\bar{\mathbf{x}})$ and $\mathbf{h} = -\nabla f(\bar{\mathbf{x}})$ to the linearized problem and obtain the existence of a vector $\mathbf{v}_1 \in \Phi$ in $\mathbb{R}^n$ such that $\mathbf{v}_1 \in \Phi$ and

$$v_1^T \nabla g_{\mathbf{H}_0}(\bar{\mathbf{x}}) = -\nabla f(\bar{\mathbf{x}}).$$

If we now set $\lambda = (v_1, \bar{\lambda}_1)$, then $\lambda \geq 0$ and we can rewrite this equation as

$$\nabla f(\bar{\mathbf{x}}) + v_1^T \nabla g_{\mathbf{H}_0}(\bar{\mathbf{x}}) + \bar{\lambda}_1^T \nabla g_{\mathbf{H}_1}(\bar{\mathbf{x}}) = \nabla f(\bar{\mathbf{x}}) + \lambda^T \nabla g(\bar{\mathbf{x}}) = \mathbf{0}.$$

The rightmost equation is statement (ii) of Theorem 3.1 with $\lambda_0 = 1$. (Recall that we are taking all the constraints to be inequality-constraints.) Since $\lambda_0 = 1$, $(\lambda_0, \bar{\lambda}) \neq \mathbf{0}$, also

$$\langle \lambda, g(\bar{\mathbf{x}}) \rangle = \langle (v_1, \bar{\lambda}), (g_1(\bar{\mathbf{x}}), g_2(\bar{\mathbf{x}})) \rangle = 0,$$

which is (i). Note that if the linearization is valid, then we can get a set of multipliers $(\lambda_0, \bar{\lambda})$ with $\lambda_0 = 1$.

To complete the proof, one must show that the terms $c(\|\mathbf{x}\|)$ and $w(\|\mathbf{x}\|)$ can indeed be neglected. That is, it must be shown that if the terms represented by $c(\|\mathbf{x}\|)$ and $w(\|\mathbf{x}\|)$ are absent, then $\bar{\mathbf{x}}$ is a solution of (7). This part of the proof becomes quite technical and is carried out in different ways by different authors depending on the additional assumptions that are made. In Exercise 3.13 below the reader will be asked to carry out one such proof, the original proof by Kuhn and Tucker.

Our proof of Theorem 3.1, which we now present, will be a "penalty function" proof and is due to McShane ([1977]).

Proof. As noted in the preceding discussion, we may assume that $\mathbf{x}_0 = \bar{\mathbf{x}}$ and that $f(\mathbf{x}_0) = 0$.

Let E and I be as in (4). To simplify the notation, we assume that $E = \{1, 2, \dots, r\}$ and that $I = \{r+1, \dots, m\}$. Since E or I can be empty, we have $0 \leq r \leq m$.

Let α be a function from $(-\infty, \infty)$ to $\mathbb{R}^1$ such that (i) α is strictly increasing on $(0, \infty)$, (ii) $\alpha(u) = 0$ for $u \leq 0$, (iii) α is C^1 , and (iv) $\alpha(u) \rightarrow 0$ for $u \rightarrow 0$. Note that $\alpha(u) > 0$ for $u > 0$.

Since g is continuous and $N_{\mathbf{H}}$ is open, there exists an $\alpha_0 > 0$ such that $\mathbf{H}\mathbf{H}, \alpha_0 = N_{\mathbf{H}}$ and $g_{\mathbf{H}}(\mathbf{x}) = \bar{\mathbf{g}}$ for $\mathbf{x} \in \mathbf{H}\mathbf{H}, \alpha_0$.

Define a penalty function F as follows:

$$F(\mathbf{x}, \beta) = f(\mathbf{x}) + \|\mathbf{x}\|^p + \beta \left\{ \sum_{i \in E} \alpha(g_i(\mathbf{x})) + \sum_{i \in I} \|h_i(\mathbf{x})\|^q \right\}, \quad (8)$$

where $x \in X_p$ and p is a positive integer. We assert that for each ε satisfying $0 < \varepsilon < \alpha_p$ there exists a positive integer $p(\varepsilon)$ such that for n with $\|x\| = \varepsilon$

$$F(x, p(\varepsilon)) > 0. \quad (9)$$

We prove the assertion by assuming it to be false and arriving at a contradiction. If the assertion were false, then there would exist an ε' with $0 < \varepsilon' < \alpha_p$ such that for each positive integer p there exists a point x_p with $\|x_p\| = \varepsilon'$ and $F(x_p, p) \leq 0$. Hence, from (8)

$$F(x_p) + \|x_p\|^2 \leq -p \left\{ \sum_{i=1}^r \omega g_i(x_p) + \sum_{i=1}^k [h_i(x_p)]^2 \right\}. \quad (10)$$

Since $\|x_p\| = \varepsilon'$ and since $\overline{B(0, \varepsilon')} = \{x : \|x\| = \varepsilon'\}$ is compact, there exist sub-sequences, which we relabel as x_p and p , and a point x_0 with $\|x_0\| = \varepsilon'$ such that $x_p \rightarrow x_0$. Since f, g , and h are continuous,

$$f(x_p) \rightarrow f(x_0), \quad g_i(x_p) \rightarrow g_i(x_0), \quad h_i(x_p) \rightarrow h_i(x_0).$$

Therefore, if in (10) we divide through by $-p$ and then let $p \rightarrow \infty$, we get

$$0 \leq \sum_{i=1}^r \omega g_i(x_0) + \sum_{i=1}^k [h_i(x_0)]^2 \leq 0.$$

Hence for each $i = 1, \dots, r$ we have $g_i(x_0) \leq 0$ and for each $i = 1, \dots, k$ we have $h_i(x_0) = 0$. Since $\|x_0\| = \varepsilon' < \alpha_p$, we have $g_j(x_0) = 0$. Thus x_0 is a feasible point, and so

$$f(x_0) \geq f^*(0) = 0. \quad (11)$$

On the other hand, from (10) and from $\|x_p\| = \varepsilon'$ we get that $f(x_p) \leq -\varepsilon'^2$, and so

$$f(x_0) \leq -\varepsilon'^2 < 0,$$

which contradicts (11). This proves the assertion.

For each $\varepsilon \in [0, \alpha_p]$ the function $F(\cdot, p)(\varepsilon)$ is continuous on the closed ball $\overline{B(0, \varepsilon)}$. Since $\overline{B(0, \varepsilon)}$ is compact, $F(\cdot, p)(\varepsilon)$ attains its minimum on $\overline{B(0, \varepsilon)}$ at some point x_ε with $\|x_\varepsilon\| \leq \varepsilon$.

Since $F(x, p)(\varepsilon) = 0$ for x with $\|x\| = \alpha$, and since $F(0, p)(\varepsilon) = f(0) = 0$, it follows that $F(\cdot, p)(\varepsilon)$ attains its minimum on $\overline{B(0, \varepsilon)}$ at an interior point x_ε of $\overline{B(0, \varepsilon)}$. Hence,

$$\frac{\partial F}{\partial x_j}(x_\varepsilon, p)(\varepsilon) = 0, \quad j = 1, \dots, n.$$

Calculating $\partial F/\partial x_j$ from (8) gives

$$\begin{aligned} \frac{\partial F}{\partial x_j}(x_0) + 2\lambda_{r+1} + \sum_{i=1}^r \mu(i) \partial g_i(x_0) \frac{\partial g_i}{\partial x_j}(x_0) \\ + \sum_{i=1}^s 2\mu(i) h_i(x_0) \frac{\partial h_i}{\partial x_j}(x_0) = 0 \quad \text{for } j = 1, \dots, n. \end{aligned} \quad (12)$$

Define

$$\begin{aligned} \Omega(x) &= 1 + \sum_{i=1}^r [\mu(i) \partial g_i(x_0, x)]^2 + \sum_{i=1}^s [2\mu(i) h_i(x)]^2, \\ \lambda_0(x) &= \frac{1}{\sqrt{\Omega(x)}}, \\ \lambda_i(x) &= \frac{\mu(i) \partial g_i(x_0, x)}{\sqrt{\Omega(x)}}, \quad i = 1, \dots, r, \\ \lambda_i(x) &= 0, \quad i = r+1, \dots, m, \\ \mu_i(x) &= \frac{2\mu(i) h_i(x)}{\sqrt{\Omega(x)}}, \quad i = 1, \dots, s. \end{aligned}$$

Note that

$$\begin{aligned} \text{(I)} \quad & \lambda_0(x) > 0, \\ \text{(II)} \quad & \lambda_i(x) \geq 0, \quad i = 1, \dots, r, \\ \text{(III)} \quad & \lambda_i(x) = 0, \quad i = r+1, \dots, m, \\ \text{(IV)} \quad & \|\lambda_0(x), \lambda(x), \mu(x)\| = 1, \end{aligned} \quad (13)$$

where $\lambda(x) = (\lambda_1(x), \dots, \lambda_m(x))$ and $\mu(x) = (\mu_1(x), \dots, \mu_s(x))$.

If we divide through by $\sqrt{\Omega(x)}$ in (12), we get

$$\lambda_0(x) \frac{\partial F}{\partial x_j}(x_0) + \frac{2\lambda_{r+1}}{\sqrt{\Omega(x)}} + \sum_{i=1}^r \lambda_i(x) \frac{\partial g_i(x_0)}{\partial x_j} + \sum_{i=1}^s \mu_i(x) \frac{\partial h_i(x_0)}{\partial x_j} = 0. \quad (14)$$

Now let $x \rightarrow 0$ through a sequence of values x_k . Then since $\|x_k\| = x_k$ we have that

$$x_k \rightarrow 0. \quad (15)$$

Since the vectors $(\lambda_0(x), \lambda(x), \mu(x))$ are all unit vectors [see (13)], there exists a subsequence, that we again denote as x_k , and a unit vector $(\lambda_0, \lambda, \mu)$ such that

$$(\lambda_0(x_k), \lambda(x_k), \mu(x_k)) \rightarrow (\lambda_0, \lambda, \mu). \quad (16)$$

Since $(\lambda_0, \lambda, \mu)$ is a unit vector, it is different from zero.

From (14), (15), (16), and the continuity of the partials of f , g , and h , we get

$$\lambda_0 \frac{\partial f}{\partial x_j}(\bar{x}) + \sum_{i=1}^r \lambda_i \frac{\partial g_i(\bar{x})}{\partial x_j} + \sum_{i=1}^m \mu_i \frac{\partial h_i(\bar{x})}{\partial x_j} = 0. \quad (17)$$

From (14) and (16) we see that $\lambda_i \geq 0$ for $i = 0, 1, \dots, r$ and $\lambda_i = 0$ for $i = r + 1, \dots, m$. Thus, $\lambda_0 \geq 0$ and $\lambda \geq 0$. Since $g_i(\bar{x}) = 0$ for $i = 1, \dots, r$ and $h_i = 0$ for $i = r + 1, \dots, m$, we have that $\lambda_i g_i(\bar{x}) = 0$ for $i = 1, \dots, m$. Also, we can take the upper limit in the second term in (17) to be m and write

$$\lambda_0 \frac{\partial f}{\partial x_j}(\bar{x}) + \sum_{i=1}^r \lambda_i \frac{\partial g_i(\bar{x})}{\partial x_j} + \sum_{i=1}^m \mu_i \frac{\partial h_i(\bar{x})}{\partial x_j} = 0.$$

This completes the proof of the theorem.

We now present conditions that ensure that $\lambda_0 = 0$ cannot occur in Theorem 5.1. Such conditions are called *constraint qualifications* and are statements about the geometry of the feasible set in a neighborhood of x_0 . There are many constraint qualifications in the literature. We have selected two that are relatively simple and may be apply. The second is easier to apply, but the first, given in Definition 5.2 below, will represent more points at which $\lambda_0 = 0$ cannot occur. All of the constraint qualifications that are appropriate to problem P1 have the defect that, in general, they require knowledge of the point x_0 rather than requiring knowledge of global properties of the functions involved. The reader who is interested in other constraint qualifications and the relationships among them is referred to BAZARAN, BERNIKI, and SHIBUY [1983] and Mangasarian [1985].

Definition 5.2. The functions g and h satisfy the *constraint qualification (CQ)* at a feasible point x_0 if

- (i) the vectors $\nabla g_1(x_0), \dots, \nabla g_r(x_0)$ are linearly independent and
- (ii) the system

$$\nabla g_i(x_0)x \leq 0, \quad \nabla h_i(x_0)x = 0, \quad (18)$$

has a solution x in $\mathbb{R}^n$. Here, $E = \{i \mid g_i(x_0) = 0\}$.

If the equality constraints are absent, then the statements involving $\nabla h_i(x_0)$ are to be deleted.

Condition (i) of the constraint qualification says that "to first order" we can move from x_0 into the feasible set by moving along the surface defined by the equality constraints $h(x) = 0$ and into the interior of the set defined by $g(x) \leq 0$. It turns that this implies that we can actually move from x_0 into the

feasible set in this fashion. We shall prove this in the corollary to Lemma 5.1 below.

Note that CQ does not impose any condition on the point x_0 other than that it be feasible. Also note that $x = 0$ cannot be a solution of (18).

The constrained qualification CQ of Definition 5.1 was introduced by Mangasarian and Fromowitz [1967].

THEOREM 5.2. Let x_0 be a solution of problem P1 and let CQ hold at x_0 . Then $\lambda_0 > 0$ and there exists a vector $\lambda \geq 0$ in $\mathbb{R}^m$ and a vector g in $\mathbb{R}^n$ such that

- (i) $\langle g, g(x_0) \rangle = 0$ and
- (ii) $\nabla f(x_0) + \lambda^* \nabla g(x_0) + g^* \nabla h(x_0) = 0$.

We shall refer to the necessary conditions of Theorem 5.2 as the Karush-Kuhn-Tucker (KKT) conditions and to the multipliers of Theorem 5.2 as the KKT multipliers.

We now prove the theorem. First we will show that if CQ holds at x_0 , then λ_0 in Theorem 5.1 cannot be zero. Hence if (λ_0, λ, g) is as in Theorem 5.1, so is $k(1, \lambda, \lambda_0, g(x_0))$ and Theorem 5.2 follows. Note that since we have $\lambda_0 > 0$, then $(\lambda_0, \lambda, g) \neq (0, 0, 0)$.

We now show that $\lambda_0 = 0$ cannot occur. If $\lambda_0 = 0$ did occur, then by (ii) of Theorem 5.1

$$\lambda^* \nabla g(x_0) + g^* \nabla h(x_0) = 0. \quad (19)$$

If $\lambda = 0$, then $g = 0$, since $(\lambda_0, \lambda, g) \neq (0, 0, 0)$. But then $g^* \nabla h(x_0) = 0$, contradicting the linear independence of $\nabla h_1(x_0), \dots, \nabla h_k(x_0)$. Thus, $\lambda \neq 0$.

Since $\lambda = (\lambda_0, \lambda_1)$ and $\lambda_1 = 0$, we have that $\lambda_0 \neq 0$. We may therefore rewrite (19) as

$$\lambda_0^* \nabla g_0(x_0) + g^* \nabla h(x_0) = 0. \quad (20)$$

Multiplying on the right by a vector x that is a solution of (18) gives

$$\lambda_0^* \nabla g_0(x_0)x = 0.$$

Since $\lambda_0 \neq 0$, $\lambda_0 > 0$, and $\nabla g_0(x_0)x < 0$, we have a contradiction. Thus $\lambda_0 \neq 0$.

COROLLARY 1. Let x_0 be a solution of problem P1 such that the set J defined in (9) is $\{1, \dots, r\}$ and such that the vectors

$$\nabla g_1(x_0), \dots, \nabla g_k(x_0), \nabla h_1(x_0), \dots, \nabla h_r(x_0) \quad (21)$$

are linearly independent. Then the conclusion of Theorem 5.2 holds.

Proof. If $\lambda_0 = 0$ in Theorem 5.1, then (20) holds. Since the vectors in (21) are linearly independent, this implies that $(\lambda_0, \lambda) = 0$. Hence $(\lambda_0, \lambda, \mu) = 0$, which cannot occur, and so $\lambda_0 > 0$.

We leave to the reader the modifications of the arguments in the proofs of Theorems 5.2 and Corollary 1 if the equality constraints are absent.

Remark 5.6. We noted in Remark 5.3 that the necessary condition (i) of Theorem 5.2 is equivalent to the m equations

$$\lambda_i g_i(\mathbf{x}_0) = 0, \quad i = 1, \dots, m. \quad (27)$$

Condition (ii) in component form gives the system of n equations

$$\frac{\partial f}{\partial x_i}(\mathbf{x}_0) + \sum_{j=1}^m \lambda_j \frac{\partial g_j(\mathbf{x}_0)}{\partial x_i} + \sum_{k=1}^k \mu_k \frac{\partial h_k}{\partial x_i} = 0, \quad i = 1, \dots, n. \quad (28)$$

Since $\mathbf{x}_0$ is a solution, it is feasible, and so we obtain an additional k equations

$$h_k(\mathbf{x}_0) = 0, \quad k = 1, \dots, k.$$

Thus we have a system of $m + n + k$ equations in the $m + n + k$ unknowns

$$\lambda = (\lambda_1, \dots, \lambda_m), \quad \mathbf{x}_0 = (x_{01}, \dots, x_{0n}), \quad \mu = (\mu_1, \dots, \mu_k).$$

Remark 5.7. If the inequality constraints are absent and we only have the equality constraints $\mathbf{h}(\mathbf{x}) = 0$, then CQ_1 becomes “ $\nabla h_1(\mathbf{x}_0), \dots, \nabla h_k(\mathbf{x}_0)$ are linearly independent.” The constraint in Corollary 1 reduces to the same statement. Under these circumstances, the multiplier (λ, μ) is unique. This follows from the observation that if $(1, \mu')$ with $\mu \neq \mu'$ is another multiplier, then $(0, \mu - \mu')$ will be a multiplier. This, however, is impossible since $\lambda_0 = 0$ cannot be. A similar argument shows that if all the inequality constraints $g_i(\mathbf{x}) \leq 0$ are inactive at $\mathbf{x}_0$, then the multiplier $(1, \lambda, \mu)$ is unique.

In Remark 5.7 we have established the form of the Lagrange multiplier rule that is usually given in multivariate calculus courses. We state it as a corollary to Theorem 5.1 and 5.2.

COROLLARY 2. Let $\mathbf{x}_0$ be a solution to the problem of minimizing $f(\mathbf{x})$ subject to the equality constraints $\mathbf{h}(\mathbf{x}) = 0$ and such that the vectors $\nabla h_1(\mathbf{x}_0), \dots, \nabla h_k(\mathbf{x}_0)$ are linearly independent. Then there exists a unique vector μ in $\mathbb{R}^k$ such that

$$\nabla f(\mathbf{x}_0) + \mu^T \nabla \mathbf{h}(\mathbf{x}_0) = 0$$

as, in component form,

$$\frac{\partial f}{\partial x_j}(x_0) + \sum_{i=1}^k \alpha_i \frac{\partial g_i(x_0)}{\partial x_j} = 0, \quad j = 1, \dots, n.$$

The following example illustrates CQ by showing what happens when it is not fulfilled. Also, the vectors in (21) are linearly dependent in this example.

Example 3.1. Let $N_0 = \mathbb{R}^2$, let $f(x) = -x_2$ and let

$$g(x) = \begin{pmatrix} -x_1 \\ -x_2 \\ x_1 + (x_2 - 1)^2 \end{pmatrix} \in \begin{pmatrix} \mathbb{R} \\ \mathbb{R} \\ \mathbb{R} \end{pmatrix}.$$

The feasible set N is sketched in Figure 4.2.

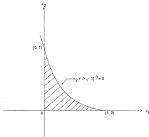


Figure 4.2.

The minimum is attained at $\mathbf{x}_0 = (1, 0)$, and so

$$\mathbf{h}_0(\mathbf{0}) = \begin{pmatrix} -x_2 \\ x_2 + (x_1 - 1)^2 \end{pmatrix}.$$

Hence

$$\nabla \mathbf{h}_0(\mathbf{x}_0)\mathbf{x} = \begin{pmatrix} 0 & -1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = \begin{pmatrix} -x_2 \\ x_2 \end{pmatrix} \neq \begin{pmatrix} 0 \\ 0 \end{pmatrix}$$

has no solution in $\mathbb{R}^2$. Thus neither CQ nor the linear independence of the vectors in (2.1) holds at $\mathbf{x}_0 = (1, 0)$.

The necessary condition with $\lambda_0 = 1$ would imply

$$\nabla f(\mathbf{x}_0) + (\lambda_0, \lambda_2) \begin{pmatrix} 0 & -1 \\ 0 & 1 \end{pmatrix} = (0, 0)$$

or

$$(-1, 0) + (0, -\lambda_2 + \lambda_2) = (0, 0),$$

which is impossible. On the other hand, if we took $\lambda_0 = 0$, then any (λ_2, λ_1) with $\lambda_2 = \lambda_1 \neq 0$ would satisfy the conditions of Theorem 5.1.

Example 4.3. In this example we consider a problem with equality constraints only in which CQ does not hold at the minimum point. Let $X_0 = \mathbb{R}^2$, let $f(\mathbf{x}) = 2x_1^2 - x_2^2$, and let $\mathbf{h}(\mathbf{x}) = x_1^2(x_2 - x_1^2)$. The constraint set X consists of the lines $x_2 = 0$ and $x_2 = \pm x_1$. The level curves of f are the two branches of the hyperbolas $2x_1^2 - x_2^2 = c$ and the asymptotes of this family, the lines $x_2 = \pm x_1/\sqrt{2}$, see Figure 4.3.

From the figure we see that the point $\mathbf{x}_0 = (0, 0)$ furnishes a minimum. Since $\nabla f(\mathbf{x}) = (4x_1, -2x_2)$ and $\nabla \mathbf{h}(\mathbf{x}) = (2x_1x_2, x_1^2 - 2x_2x_1)$, we get that

$$\nabla f(\mathbf{x}_0) = \nabla f(0, 0) = (0, 0), \quad \nabla \mathbf{h}(\mathbf{x}_0) = (0, 0).$$

Since $\nabla \mathbf{h}(\mathbf{x}_0) = \mathbf{0}$, the constraint qualification CQ does not hold at $\mathbf{x}_0$. Any (λ_0, μ) of the form $(0, \mu)$, μ arbitrary, will serve as a Lagrange multiplier, as will any $(\lambda, 0)$, λ arbitrary.

Example 4.4 (Continued). In this example CQ holds at a point $\mathbf{x}_0$ if $g(\mathbf{x}_0) = 0$ and there exists a solution $\mathbf{u}$ to the equation

$$\langle \nabla f(\mathbf{x}_0), \mathbf{u} \rangle = (2\lambda_0 - 2x_0) \langle \mathbf{1}, \mathbf{u} \rangle, \quad x_0 \mathbf{0} = 0.$$

Clearly, $\mathbf{u} = (-1, 0)$ is a solution. Thus CQ holds at all boundary points of the constraint set. This is consistent with our being able to rule out $\lambda_0 = 0$ in the outset in our first treatment of this example.

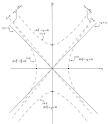


FIGURE 4.6

The next theorem is a sufficiency theorem which states that if f, X_{g^*} , f^* and g are convex and if the KKT necessary conditions hold, then x_0 is a solution.

Theorem 4.6. Let f^* and g^* be as in the statement of problem P1 and let X_{g^*} , f^* and g^* be convex. Let $x_0 \in \mathbb{R}^n$ and $\lambda \in \mathbb{R}^m$ be such that

$$(i) \quad g(x_0) = 0,$$

$$(ii) \quad \lambda \geq 0,$$

- (a) $\langle \mathbf{g}, \mathbf{g}(\mathbf{x}_0) \rangle = 0$, and
 (b) $\nabla f(\mathbf{x}_0) + \lambda^* \nabla g(\mathbf{x}_0) = \mathbf{0}$.

Then $\mathbf{x}_0$ is a solution of problem P1.

Remark 3.8. The sufficiency theorem can also be applied to problems with equality constraints $\mathbf{h}(\mathbf{x}) = \mathbf{0}$ provided the function $\mathbf{h}$ is affine, that is, has the form $\mathbf{h} = A\mathbf{x} + \mathbf{b}$. For then we can write the equality constraint $\mathbf{h}(\mathbf{x}) = \mathbf{0}$ as a pair of inequality constraints $\mathbf{h}(\mathbf{x}) \leq \mathbf{0}$ and $-\mathbf{h}(\mathbf{x}) \leq \mathbf{0}$, with both $\mathbf{h}$ and $-\mathbf{h}$ convex. The function $\mathbf{g}$ in the theorem is to be interpreted as incorporating this pair of inequality constraints.

Remark 3.9. Under the assumptions of the theorem the feasible set X is convex. To see this, let $X_i = \{\mathbf{x} \mid \mathbf{x} \in X_0, g_i(\mathbf{x}) \leq 0\}$. Since each g_i is convex, each set X_i is convex. Since $X = \cap_{i=1}^n X_i$, we get that X is convex.

Proof. Condition (b) states that $\mathbf{x}_0 \in X$. Since f is convex,

$$f(\mathbf{x}) - f(\mathbf{x}_0) \geq \langle \nabla f(\mathbf{x}_0), \mathbf{x} - \mathbf{x}_0 \rangle, \quad \mathbf{x} \in X.$$

Substituting (a) into this inequality gives

$$\begin{aligned} f(\mathbf{x}) - f(\mathbf{x}_0) &\geq \langle -\lambda^* \nabla g(\mathbf{x}_0), \mathbf{x} - \mathbf{x}_0 \rangle \\ &= -\langle \mathbf{g}, \nabla g(\mathbf{x}_0)(\mathbf{x} - \mathbf{x}_0) \rangle, \quad \mathbf{x} \in X. \end{aligned} \quad (22)$$

Since $\mathbf{g}$ is convex,

$$\mathbf{g}(\mathbf{x}) - \mathbf{g}(\mathbf{x}_0) \geq \nabla g(\mathbf{x}_0)(\mathbf{x} - \mathbf{x}_0).$$

Therefore, since $\mathbf{h} \geq \mathbf{0}$,

$$\langle \mathbf{g}, \mathbf{g}(\mathbf{x}) - \mathbf{g}(\mathbf{x}_0) \rangle \geq \langle \mathbf{g}, \nabla g(\mathbf{x}_0)(\mathbf{x} - \mathbf{x}_0) \rangle.$$

Substituting this inequality into (22) and then using (B) give

$$f(\mathbf{x}) - f(\mathbf{x}_0) \geq \langle \mathbf{h}, \mathbf{g}(\mathbf{x}_0) - \mathbf{g}(\mathbf{x}) \rangle = -\langle \mathbf{h}, \mathbf{g}(\mathbf{x}) \rangle.$$

Since $\mathbf{h} \geq \mathbf{0}$ and since for $\mathbf{x}$ in X we have $\mathbf{g}(\mathbf{x}) \leq \mathbf{0}$, it follows that $-\langle \mathbf{h}, \mathbf{g}(\mathbf{x}) \rangle \geq 0$. Hence $f(\mathbf{x}) \geq f(\mathbf{x}_0)$ for all $\mathbf{x}$ in X , and so $\mathbf{x}_0$ is a solution to problem P1.

The function φ in Example 5.1 is not convex, as a calculation of the Hessian matrix of φ shows. Therefore Theorem 3.3 cannot be used to analyze Example 5.1.

Exercise 5.1. State and prove the form of Theorem 5.3 that results when affine equality constraints $h(x) = 0$ are also present.

Exercise 5.2. Use the following to complete the exercise:

- Sketch the feasible set.
- Use an appropriate theorem to show that a solution exists.
- Find all critical points and associated multipliers.
- Without using the second-order criteria of the next section, determine, if possible, whether a given critical point furnishes a solution.

(i) Minimize $f(x) = x_1^2 + (x_2 + 1)^2$ subject to

$$x_1^2 + x_2^2 \leq 4, \quad x_2^2 \leq x_1 - 1, \quad -2x_1 \leq x_1 - 1.$$

(ii) Minimize $f(x) = x_1^2 + x_2^2$ subject to

$$x_1^2 + x_2^2 \leq 4, \quad x_1^2 \geq x_2, \quad -2x_1 \leq x_2 - 1, \quad -2x_2 \geq x_2 - 2.$$

(iii) Minimize $(x_1 - 2)^2 + x_2^2$ subject to

$$x_1^2 + x_2^2 \leq 4, \quad x_2^2 - 2x_1 + x_2^2 \leq 4, \quad x_1 + x_2^2 \leq 2.$$

(iv) Minimize $f(x) = \|x\|^2$, where $x \in \mathbb{R}^2$, subject to

$$x_1 x_2 \geq 1, \quad x_1 \geq 0, \quad x_2 \geq 0.$$

(v) Maximize $f(x) = x_1^2 + x_2^2 + 4x_1 x_2$, $x \in \mathbb{R}$, subject to $x_1^2 + x_2^2 = 1$. What can you say about minimizing $f(x)$ subject to $x_1^2 + x_2^2 = 1$. [In the minimization problem you need not answer questions (i) and (ii).]

Exercise 5.3. Minimize $f(x) = (x_1 - 2)^2 + (x_2 - 1)^2$ subject to

$$x_1^2 - x_1 \leq 0, \quad x_1 + x_2 - 4 \leq 0, \quad -x_1 \leq 0, \quad -x_2 \leq 0.$$

Give a full justification of your answer. You may use geometric considerations to assist you in finding critical points.

Exercise 5.4. Minimize $f(x) = x_1^2 + x_1 x_2 + x_2^2$ subject to

$$-2x_1 - 2x_2 + 6 \leq 0, \quad -x_1 + x_2 - 3 \leq 0, \quad 4x_1 - 2x_2^2 + 4x_2 - 1 \leq 25.$$

Hint: Compare with Exercise 5.2.

Exercise 5.3. Minimise $f(x) = e^{-2x_1 + x_2^2}$ subject to $x^{(2)} + x^{(2)} \leq 20$.

Exercise 5.4. Minimise $f(x) = \sum_{i=1}^n c_i x_i$ subject to

$$\sum_{i=1}^n a_i x_i = b, \quad x_i \geq 0, \quad i = 1, \dots, n,$$

where $c_i, a_i, i = 1, \dots, n$, and b are positive constants.

Exercise 5.5. Let A be a positive-definite $n \times n$ symmetric matrix and let y be a fixed vector in $\mathbb{R}^n$. Show that the maximum value of $f(x) = \langle y, x \rangle$ subject to the constraint $\langle x, Ax \rangle \leq 1$ is $\langle y, A^{-1}y \rangle^{1/2}$. Use this result to establish the generalisation of the Cauchy-Schwarz inequality

$$|\langle x, y \rangle|^2 \leq \langle x, Ax \rangle \langle y, A^{-1}y \rangle.$$

Exercise 5.6. Use the Lagrange multiplier rule to show that the distance $\rho(x_0)$ from the point x_0 to the hyperplane $H_1^c = \{x : \langle x, n \rangle = c\}$ is given by

$$\rho(x_0) = \left| \frac{\langle n, x_0 \rangle - c}{\|n\|} \right|.$$

(Note: In problems involving the minimisation of a norm or a distance it is often simpler to minimise the square of the norm or distance.)

Exercise 5.7. In this problem we will use the Lagrange multiplier rule to deduce important properties of eigenvectors and eigenvalues of a real $n \times n$ symmetric matrix A .

(a) Show that the maximum of $f(x) = \langle x, Ax \rangle$ subject to $\|x\| = 1$ is attained at a unit eigenvector x_1 and that $f(x_1)$ is the largest eigenvalue λ_1 of A .

(b) Suppose $m < n$ mutually orthogonal unit eigenvectors $x_1, \dots, x_m$ of A have been determined. Show that the minimum of $f(x) = \langle x, Ax \rangle$ subject to $\|x\| = 1$ and $\langle x, x_j \rangle = 0, j = 1, \dots, m$, is attained at a point x_{m+1} that is an eigenvector of A and that $f(x_{m+1})$ is the corresponding eigenvalue.

(c) Use (a) and (b) to show that a real $n \times n$ symmetric matrix has a set of n orthonormal eigenvectors $x_1, \dots, x_n$ with corresponding eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_n$.

Exercise 5.10. In this exercise we will use Farkas's lemma and a constraint qualification introduced by Karush and Tucker [1991] that is different from CQ to establish the conclusion of Theorem 5.2 without using Theorem 5.1. The Karush-Tucker constraint qualification KTQ is said to hold at a point x_0 if there

exists a vector a such that

$$\nabla g_{\alpha}(\alpha_0)a < 0, \quad \nabla h_{\beta}(\alpha_0)a = 0$$

and if for each such a there exists a C^2 function $\bar{q}(\cdot)$ defined on an interval $[0, \tau]$ such that

$$\begin{aligned}\bar{q}(0) &= \alpha_0, & \bar{q}'(0) &= a, \\ g(\bar{q}(t)) &< 0, & h(\bar{q}(t)) &= 0 \quad \text{for all } t \text{ in } [0, \tau].\end{aligned}$$

The Karush–Tucker constraint qualification is difficult to check. In Lemma 5.1 we give a verifiable condition that implies KTQ.

Let α_0 be a solution of problem P1 and let KTQ hold at α_0 . Show that the conclusion of Theorem 5.2 holds without invoking Theorem 5.1. *Hint:* What can you say about the right-hand derivatives at $t = 0$ of the mappings $t \mapsto g(\bar{q}(t)$, $t \mapsto g_{\alpha}(\bar{q}(t)$, and $t \mapsto h(\bar{q}(t)$ in relation to Farkas's theorem?

6. SECOND-ORDER CONDITIONS

In this section we first obtain second-order necessary conditions. Points that satisfy the first-order necessary conditions but fail to satisfy the second-order conditions cannot solve the programming problem. We conclude this section with a presentation of some second-order sufficiency conditions. Points that satisfy the first-order necessary conditions and these sufficiency conditions are solutions to the programming problem.

Since second-order conditions involve second derivatives, in this section we assume that the functions f , g , and h that occur in the statement of problem P1 are of class C^2 on X_{opt} .

The derivation of our second-order necessary condition will require an application of the implicit function theorem, which we now discuss.

Consider the system of equations

$$\begin{aligned}s_1 f(x_1, \dots, x_n, y_1, \dots, y_m) &= 0, \\ &\vdots \\ s_m f(x_1, \dots, x_n, y_1, \dots, y_m) &= 0,\end{aligned}\tag{1}$$

where the real-valued functions s_i , $i = 1, \dots, m$, are defined on an open set $D = D_x \times D_y$ where D_x is open in the x -space $\mathbb{R}^n$ and D_y is open in the y -space $\mathbb{R}^m$. Equation (1) asks us to find all (x, y) in D such that (1) is satisfied. For a given x in D_x , the system (1) can therefore be interpreted as a system of m equations for the m unknowns $(y_1, \dots, y_m)$. If for each x in some subset I of D_x we can find a unique solution $y(x) = (y_1(x), \dots, y_m(x))$, then in this section we obtain a function $y(\cdot)$ with domain I . In this event, we say that (1) defines y as a function of x implicitly.

Let (x_0, y_0) be a point in D that satisfies (1). The implicit function theorem gives conditions ensuring that (1) has a unique solution for points x sufficiently close to x_0 . The resulting function $y(x)$ will inherit the regularity properties of the functions $u_1, \dots, u_m$.

To illustrate these ideas, take $B_1 = \mathbb{R}^2$, $B_2 = \mathbb{R}^1$, and $D = \mathbb{R}^2 = \mathbb{R}^2 \times \mathbb{R}^1$. Consider

$$w(x_1, x_2, y) = x_1^2 + x_2^2 + y^2 - 1 = 0. \quad (2)$$

The set of points $(x, y) = (x_1, x_2, y)$ in $\mathbb{R}^3$ that satisfy (2) is the surface of the ball with center at the origin and radius 1.

Given any point $(x_0, y_0) = (x_{01}, x_{02}, y_0)$ with $y_0 > 0$ that satisfies (2), we can find an $\alpha > 0$ and a $\beta > 0$ such that

$$y = \sqrt{1 - x_1^2 - x_2^2}$$

is the unique solution of (2) for x in $B(x_0, \alpha)$ having range in $B(y_0, \beta)$. The point $(x_0, -y_0)$ also satisfies (2), and

$$y = -\sqrt{1 - x_1^2 - x_2^2}$$

is the unique solution of (2) for x in $B(x_0, \alpha)$ having range in $B(-y_0, \beta)$.

Although all points of $\mathcal{F} = \{(x, y) | x_1^2 + x_2^2 + y^2 = 1, y \neq 0\}$ are solutions of (2), equation (2) does not define y as a function of x on any ball about a point x_0 such that $\|x_0\| = 1$. Since $dy/dy = 2y$, we see that at any point (x_0, y_0) that satisfies (2) we have $dy/dy \neq 0$ if $y_0 \neq 0$ and $dy/dy = 0$ if $y_0 = 0$. Thus at all points of $\mathcal{F}$, $dy/dy \neq 0$. The reader should return to this paragraph after reading Theorem 6.1.

Equation (1) in vector form is

$$w(x, y) = 0.$$

THEOREM 6.1 (IMPLICIT FUNCTION THEOREM). Let B_1 be an open set in $\mathbb{R}^n$ and let B_2 be an open set in $\mathbb{R}^m$. Let $D = B_1 \times B_2$. Let

$$w(x, y) \rightarrow w(x, y), \quad x \in B_1, y \in B_2,$$

be a mapping that is of class C^2 on D and with range in $\mathbb{R}^k$. Let (x_0, y_0) , $x_0 \in B_1$, $y_0 \in B_2$, be a solution of

$$w(x, y) = 0$$

and suppose that the $m \times m$ matrix

$$P(x_0, x_0) = \left(\frac{\partial^2 g_i}{\partial x_j^2}(x_0, x_0) \right), \quad i = 1, \dots, m, \quad j = 1, \dots, m, \quad (3)$$

is nonsingular. Then there exist an $\alpha > 0$, a $\beta > 0$, and a function $g(x)$ defined on $B(x_0, \alpha) \cap D_1$ with range in $B(y_0, \beta) \cap D_2$ such that

$$g(\cdot) \text{ is of class } C^2, \quad (4)$$

$$g(x_0, g(x_0)) = 0 \quad \text{for all } x_0 \in B(x_0, \alpha), \quad (5)$$

Moreover, for x in $B(x_0, \alpha)$, $g(x)$ is the only solution of $w(x, y) = 0$ such that $y(x) \in B(y_0, \beta)$.

For a proof of this theorem we refer the reader to Rockafellar [1978] or Fleming [1977].

Definition 4.1. Let x_0 be a point of X_0 . The vector $g \in \mathbb{R}$ in $\mathbb{R}^r$ is said to satisfy the tangential constraints at x_0 if

$$\nabla U_{ij}(x_0)g \leq 0, \quad \forall i(x_0)g = 0, \quad (6)$$

where $\mathcal{J} = \{i(x_0)g = 0\}$.

For definiteness we shall henceforth suppose that $E = \{1, \dots, r\}$. Let $I = \{j = 1, \dots, m\}$. Since either E or I may be empty, $0 \leq r \leq m$.

Let x be a vector that satisfies the tangential constraints. Let

$$E_1 = \{j \in E, \langle \nabla U_{ij}(x_0), x \rangle = 0\},$$

$$E_2 = \{j \in E, \langle \nabla U_{ij}(x_0), x \rangle < 0\}.$$

Then $E = E_1 \cup E_2$.

$$\nabla U_{ij}(x_0)x = 0, \quad \forall i(x_0)x = 0, \quad (7)$$

and

$$\nabla U_{ij}(x_0)x \leq 0. \quad (8)$$

For the sake of definiteness suppose that $E_1 = \{1, \dots, q\}$. Since E_1 may be empty, $0 \leq q \leq r$. If $q = 0$, statements involving E_1 in (7) in the ensuing discussion should be deleted. The reader should keep in mind that the sets E_1 and E_2 depend on x and x_0 .

The next lemma gives a condition under which, given a vector α that satisfies the tangential constraints, we can construct a curve $\xi(\cdot)$ emanating from x_0 and going into the feasible set in the direction given by α .

Lemma 6.1. Let g and h be of class C^{p+1} on E_p , $p \geq 1$. Let x_0 be feasible and let α satisfy the tangential constraints at x_0 . Suppose that the vectors

$$\nabla g(x_0), \dots, \nabla g_j(x_0), \nabla h_1(x_0), \dots, \nabla h_r(x_0) \quad (7)$$

are linearly independent. Then there exists a $\delta > 0$ and a C^p mapping $\xi(\cdot)$ from $(-\delta, \delta)$ to E^n such that

$$\begin{aligned} \xi(0) &= x_0, & \xi'(0) &= \alpha, \\ g_j(\xi(t)) &= 0, & h_i(\xi(t)) &= 0, & g(\xi(t)) &\leq 0 \\ \|\xi(t)\| &\leq \delta, & \text{for } 0 < t < \delta. \end{aligned}$$

Proof. Let $p = q + k$. In view of (7), $p \leq n$. If $p = n$, then again in view of (7) the only solution of (5) is $\alpha = 0$. Hence, $E_q = \{0\}$ and $E_k = E$. The mapping $\xi(\cdot)$ is then the constant map $\xi(t) = x_0$.

We now suppose that $p < n$. Let Γ denote the $p \times n$ matrix whose rows are the vectors in (7). Thus

$$\Gamma = \begin{pmatrix} \nabla g_1(x_0) \\ \vdots \\ \nabla h_r(x_0) \end{pmatrix}$$

and Γ has rank p . By relabeling the coordinates corresponding to p linearly independent columns of Γ , we may suppose without loss of generality that the first p columns of the matrix Γ are linearly independent. Thus the $p \times p$ matrix

$$\Gamma_p = \begin{pmatrix} \frac{\partial g_1(x_0)}{\partial x_1} \\ \vdots \\ \frac{\partial g_p(x_0)}{\partial x_1} \\ \vdots \\ \frac{\partial h_r(x_0)}{\partial x_1} \\ \vdots \\ \frac{\partial h_r(x_0)}{\partial x_p} \end{pmatrix} \quad i = 1, \dots, q, \quad i = 1, \dots, k, \quad j = 1, \dots, p,$$

is nonsingular.

The system of equations

$$\begin{aligned} g_1(x) &= 0, \\ h(x) &= 0, \\ x_{j+1} - x_{0,j+1} - \alpha_{j+1} &= 0, \\ &\vdots \\ x_n - x_{0,n} - \alpha_n &= 0 \end{aligned} \quad (8)$$

where $u_{p+1}, \dots, u_n$ let the last $n - p$ components of the vector u , has a solution $(u, v) = (u_0, 0)$. We shall use the implicit function theorem to show that the system (8) determines u as a function of v near $v = 0$.

For the system (8), viewed as defining u implicitly as a function of v , the matrix J defined in (1) becomes

$$J = \begin{pmatrix} F_p & F_{n-p} \\ 0_{n-p, p} & I_{n-p} \end{pmatrix},$$

where F_{n-p} is the $p \times (n - p)$ matrix whose columns are the last $n - p$ columns of F , the matrix $0_{n-p, p}$ is the $(n - p) \times p$ zero matrix, and I_{n-p} is the $(n - p) \times (n - p)$ identity matrix. Since F_p is nonsingular, so is J . Hence by Theorem 6.1 there exists a $\tau > 0$ and a C^2 function $\tilde{u}(\cdot)$ defined on the interval $(-\tau, \tau)$ and with range in $\mathbb{R}^p$ such that

$$\begin{aligned} \tilde{u}(0) &= u_0 \\ g_{\tilde{u}}(\tilde{u}(t)) &= 0, & \tilde{h}(\tilde{u}(t)) &= 0, \\ \tilde{u}'(t) &= u_{p+1} + \dots + u_n, & t &= p + 1, \dots, n. \end{aligned} \quad (9)$$

Differentiation of the equations in (8) with respect to v and then setting $v = 0$ give

$$\begin{pmatrix} F_p & F_{n-p} \\ 0_{n-p, p} & I_{n-p} \end{pmatrix} \tilde{u}'(0) = \begin{pmatrix} g' \\ \tilde{h}' \end{pmatrix},$$

where $\tilde{h}' = (h_{p+1}, \dots, h_n)$. The matrix on the left is the matrix J for the system (8) and is nonsingular. Therefore the system

$$\begin{pmatrix} F_p & F_{n-p} \\ 0_{n-p, p} & I_{n-p} \end{pmatrix} u = \begin{pmatrix} g' \\ \tilde{h}' \end{pmatrix}$$

has a unique solution. Since $\tilde{u}'(0)$ and u are both solutions, we get that $\tilde{u}'(0) = u$.

Since $g_u(u_0) < 0$ and $\tilde{u}(0) = u_0$, it follows from the continuity of g and $\tilde{u}$ that for τ sufficiently small we have $g_{\tilde{u}}(\tilde{u}(t)) < 0$ for $|t| < \tau$.

Since

$$\frac{d g_{\tilde{u}}(\tilde{u}(t))}{dt} = g_{\tilde{u}}(\tilde{u}(t)) \tilde{u}'(t)$$

and since $\tilde{u}(0) = u_0$ and $\tilde{u}'(0) = u$, it follows from (5) that $d g_{\tilde{u}}(\tilde{u}(t))/dt < 0$. It then follows from continuity that by taking τ to be smaller, if necessary, there is an interval $(0, \tau)$ on which all of the previous conclusions hold and on which

$d\mathbf{g}_{\mathbf{x}_0}(\mathbf{f})(\mathbf{v})\mathbf{v}^T < 0$. Since $\mathbf{g}_{\mathbf{x}_0}(\mathbf{f})(\mathbf{v}) = \mathbf{g}_{\mathbf{x}_0}(\mathbf{v}_0) = 0$, we have $\mathbf{g}_{\mathbf{x}_0}(\mathbf{f})(\mathbf{v}) < 0$ on $(0, \epsilon)$. This completes the proof.

COROLLARY 1. Let CQ hold at a point $\mathbf{x}_0$. Then for every $\mathbf{z}$ satisfying relation (18) in Section 5 there exists a C^0 function $\mathbf{q}(\cdot)$ such that $\mathbf{q}(0) = \mathbf{x}_0$, $\mathbf{q}'(0) = \mathbf{z}$ and $\mathbf{q}(t)$ is feasible for all t in some interval $[0, \epsilon)$. Moreover, $\mathbf{g}(\mathbf{q}(t)) = 0$ on $(0, \epsilon)$.

Proof. A vector $\mathbf{z}$ satisfying relation (18) in Section 5 satisfies the tangential constraints, and the set of indices E_1 corresponding to $\mathbf{z}$ is empty.

Example 5.2 (Continued). At the point $\mathbf{x}_0 = (1, 0)$

$$\nabla \mathbf{h}_0(\mathbf{x}_0)\mathbf{z} = \begin{pmatrix} 0 & -1 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} 1 \\ 0 \end{pmatrix} = \begin{pmatrix} -1 \\ 0 \end{pmatrix}.$$

Therefore, since there are no equality constraints, every vector of the form $(z_1, 0)$ satisfies the tangential constraints. Clearly, there exists no curve $\mathbf{q}(\cdot)$ emanating from $\mathbf{x}_0 = (1, 0)$ with tangent vector $(1, 0)$ at $\mathbf{x}_0$ such that $\mathbf{q}(t)$ is feasible for t in some interval $[0, \epsilon)$. Lemma 6.1 is not contradicted, since we have $E_1 = E$ and the rows of $\nabla \mathbf{h}_0(\mathbf{x}_0)$ are linearly dependent. Also, the conclusion of the corollary is not contradicted since CQ does not hold at $(1, 0)$.

We now present a second-order necessary condition.

Definition 6.2. Let $(\mathbf{x}_0, \mathbf{k}, \mu)$ be as in Theorem 5.2. Let $\mathbf{z}$ be a vector that satisfies the tangential constraints (9) at $\mathbf{x}_0$. Let

$$E_1 = E_1(\mu) = \{i \in E, \nabla \mathbf{g}_i(\mathbf{x}_0, \mu) = 0\},$$

$$E_2 = E_2(\mu) = \{i \in E, \nabla \mathbf{g}_i(\mathbf{x}_0, \mu) \neq 0\}.$$

The vector $\mathbf{z}$ will be called a *second-order test vector* if (i) $\lambda_i = 0$ for $i \in E_1$ and (ii) the rows of the matrix

$$\begin{pmatrix} \nabla \mathbf{h}_0(\mathbf{x}_0, \mathbf{f}) \\ \nabla \mathbf{h}_i(\mathbf{x}_0, \mathbf{f}) \end{pmatrix}$$

are linearly independent.

THEOREM 6.2. Let f , $\mathbf{g}$, and $\mathbf{h}$ be of class $C^{2,1}$ on X_0 . Let $\mathbf{x}_0$ be a solution to problem P1. Let λ and μ be KKT multipliers as in Theorem 5.2 and let $H(\cdot, \mathbf{k}, \mu)$ be defined by

$$\bar{F}(\mathbf{x}, \lambda, \mu) = f(\mathbf{x}) + \langle \lambda, \mathbf{g}(\mathbf{x}) \rangle + \langle \mu, \mathbf{h}(\mathbf{x}) \rangle.$$

Then for every second-order test vector x ,

$$\langle x, F_{xx}(x_0, \lambda, p)x \rangle \geq 0, \quad (3)$$

where

$$F_{ij} = \left\langle \frac{\partial^2 F}{\partial x_i \partial x_j} \right\rangle \quad i = 1, \dots, n, \quad j = 1, \dots, n.$$

We first establish an elementary result.

LEMMA 6.2. Let ϕ be a C^2 function defined on an interval $(-\tau, \tau)$ such that $\phi(0) = 0$ and $\phi(t) \geq \phi(0)$ for all t in $[0, \tau)$. Then $\phi'(0) \geq 0$.

Proof. By Taylor's theorem, for $0 < t < \tau$

$$\phi(t) - \phi(0) = \frac{1}{2} \phi''(\theta) t^2,$$

where $0 < \theta < t$. If $\phi'(0)$ were negative, then by continuity there would exist an interval $[0, \epsilon)$ on which ϕ' would be negative. Hence $\phi(t) < \phi(0)$ on this interval, contradicting the hypothesis.

Proof (Theorem 6.2). Let x be a second-order test vector. Since $\lambda_1 = 0$ for $t \in E_1$ and

$$\lambda = (\lambda_0, \lambda_1) = (\lambda_0, \lambda_{n_1-1}, \lambda_1) = (\lambda_{1^*}, 0, 0),$$

we have

$$F(x, \lambda, p) = f(x) + \langle \lambda_0, g_0(x) \rangle + \langle \lambda_1, h(x) \rangle.$$

Let $\tilde{f}(\cdot)$ be the function corresponding to x as in Lemma 5.1, with $x_0 = x_{1^*}$. Since $g_0(\tilde{q}(t)) = 0$ and $h(\tilde{q}(t)) = 0$ for $|t| < \tau$, we have

$$F(\tilde{q}(t)) = F(\tilde{q}(t), \lambda, p) \quad \text{for } |t| < \tau.$$

Since for $0 < t < \tau$, all points $\tilde{q}(t)$ are feasible, and since $\tilde{q}(0) = x_{1^*}$, the mapping $t \rightarrow F(\tilde{q}(t))$, where $0 < t < \tau$ has a minimum at $t = 0$. Hence, so does the mapping $\phi(\cdot)$ defined on $[0, \tau)$ by

$$\phi(t) = F(\tilde{q}(t), \lambda, p).$$

Weighted-formal calculation gives

$$\begin{aligned}\phi'(0) &= (\nabla F(\bar{x}(0)), \lambda, \mu, \bar{v}(0)) \\ \phi''(0) &= (\bar{K}''(0), F_{xx}(\bar{x}(0)), \lambda, \mu \bar{K}''(0)) + (\nabla^2 F(\bar{x}, \lambda, \mu, \bar{v}(0)).\end{aligned}$$

Setting $v = 0$ in the first equation and using the relations $\bar{K}(0) = x_0$ and $\nabla F(x_0, \lambda, \mu) = 0$ give $\phi'(0) = 0$. Hence by Lemma 6.2 we have $\phi''(0) \geq 0$. Setting $v = 0$ in the second equation and using $\bar{K}(0) = x_0$, $\nabla F(x_0, \lambda, \mu) = 0$, and $\bar{K}''(0) = 0$ give the conclusion of the theorem.

The following corollary requires that (10) holds for *linear* v thus the theorem does and thus is more restrictive than the theorem but easier to apply.

COROLLARY 2. Let the functions f , g , and h be as in Theorem 6.2 and let x_0 be a solution to problem P1. Suppose that the vectors

$$\nabla g_1(x_0), \dots, \nabla g_r(x_0), \nabla h_1(x_0), \dots, \nabla h_s(x_0)$$

are linearly independent. Then there exist KKT multipliers λ, μ as in Theorem 6.2, and (10) holds for every vector v satisfying

$$\nabla g_1(x_0)v = 0, \quad \nabla h_1(x_0)v = 0. \quad (11)$$

Proof. The existence of KKT multipliers follows from Corollary 1 of Theorem 5.2. The second conclusion follows from the theorem by taking $E_1 = E$, in which case $E = \{1, \dots, r\}$ and $E_2 = \emptyset$.

If the inequality constraints are absent, then the linear independence hypothesis of Theorem 6.2 and Corollary 2 becomes "the vectors $\nabla g_1(x_0), \dots, \nabla g_r(x_0)$ are linearly independent," which is also the constraint qualification CQ for this problem. Thus, we get the following corollary.

COROLLARY 3. Let f and h be as in Theorem 6.2, let the inequality constraints be absent, let x_0 be a solution of problem P1 and let CQ hold at x_0 . Then there exists a unique vector μ in $\mathbb{R}^s$ such that the function $\bar{W}(\cdot, \mu)$ defined by

$$\bar{W}(x, \mu) = f(x) + \langle \mu, h(x) \rangle$$

satisfies

$$(i) \quad \nabla f(x_0) + \mu \nabla h(x_0) = 0$$

and

$$(10) \quad \langle u, H_{xx}(x_0, p(x)) \rangle \geq 0$$

for all x in $\mathbb{R}^n$ satisfying

$$V_0(x_0) = 0.$$

Corollary 3 is the classical second-order necessary condition for equality-constrained problems.

Example 5.4 (Continued). The reader will recall that for $0 < k < 1$, the first-order necessary conditions were satisfied by three points and corresponding multipliers given by

$$\begin{array}{lll} x_1 = 1-k, & x_2 = \pm\sqrt{2k(1-k)}, & \lambda = 1 \\ x_1 = 0, & x_2 = 0, & \lambda = 1/k. \end{array}$$

We will use the second-order necessary condition as given in Corollary 2 to rule out the origin as a candidate for optimality when $0 < k < 1$.

Recall that for any x we have $V_0(x) = (2k, -2x_1)$. At $x_0 = (0, 0)$, $V_0(x_0) \neq 0$, so the linear independence requirement of Corollary 2 holds. Thus for x to satisfy (11) we require that

$$\langle V_0(x_0), x \rangle = (2k, 0) \text{ is }_p \text{ } x_1(0) = 0.$$

This is satisfied by all vectors (x_1, x_2) of the form

$$x_1 = 0, \quad x_2 \text{ arbitrary.} \quad (12)$$

The function F with $\lambda = 1/k$ becomes

$$F = (x_1 - 1)^2 + x_2^2 + \frac{1}{k}(2kx_1 - x_2^2).$$

Condition (12) with $x_0 = (0, 0)$ becomes

$$(x_1, x_2) \begin{pmatrix} 2 & 0 \\ 0 & 2(1-k) \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = 2x_1^2 + 2x_2^2 \left(1 - \frac{k}{1}\right) \geq 0$$

for all (x_1, x_2) satisfying (12). Since $x_1 = 0$ and for $0 < k < 1$ the coefficient of $2x_2^2$ is negative, the last relation only holds for $x_2 = 0$. Thus (12) fails to hold for all (x_1, x_2) satisfying (12), and so $x_0 = (0, 0)$ is not optimal if $0 < k < 1$.

For $\lambda \geq 1$ the only critical point is $\mathbf{x}_\lambda = (3, 0)$ with $\lambda = 1/3$. We now have $\langle \mathbf{x}, F_{\mathbf{x}\mathbf{x}}\mathbf{x} \rangle \geq 0$ for all $\mathbf{x}$, so the second-order conditions of Theorem 5.2 and Corollary 2 both hold at $(3, 0)$.

We now show that for $0 < \lambda < 1$ the second-order condition (10) holds at the two points $(1 - \lambda, \pm\sqrt{\lambda(1 - \lambda)})$. At these points $\lambda = 1$. Thus, for a second-order test vector $\mathbf{x}$ we have $\mathbf{E}_1 = \mathbf{E}$ and $\mathbf{F}_1 = \mathbf{F}$. Hence (5) and (6) reduce to $\nabla_{\mathbf{F}\mathbf{F}}(\mathbf{x}_\lambda)\mathbf{x} = 0$ at these points. Thus a test vector $\mathbf{x}$ satisfies

$$\langle 2\mathbf{E}_1 + 2\sqrt{\lambda(1 - \lambda)}\mathbf{E}, \mathbf{x}_\lambda \mathbf{x} \rangle = 2\mathbf{E}_1 + 2\sqrt{\lambda(1 - \lambda)}\mathbf{E} = \mathbf{E}. \quad (11)$$

The function F with $\lambda = 1$ becomes

$$F = (x_1 - 1)^2 + x_2^2 + (2\lambda x_1 - x_1^2)$$

and condition (10) becomes

$$\langle \mathbf{E}_1, \mathbf{x} \rangle \begin{pmatrix} 2 & 0 \\ 0 & 2 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = 2\langle \mathbf{E}_1, \mathbf{x} \rangle \geq 0.$$

This inequality is satisfied for all $\mathbf{x}$ and hence for all $\mathbf{x}$ satisfying (11).

We now state and prove a second-order sufficiency theorem (Hart and Montgomery, 1979). This theorem is more useful and easier to apply than the second-order necessary conditions.

Theorem 5.3. Let $(\mathbf{x}_0, \lambda_0, \mathbf{h}, \mathbf{g})$ satisfy the *FF* necessary conditions of Theorem 5.1. Suppose that every $\mathbf{z} \neq \mathbf{0}$ that satisfies the tangential constraint (4) at $\mathbf{x}_0 = \mathbf{x}_\lambda$ and the inequality $\langle \nabla f(\mathbf{x}_\lambda), \mathbf{z} \rangle \leq 0$ also satisfies

$$\langle \mathbf{z}, F''_{\mathbf{x}\mathbf{x}}(\mathbf{x}_\lambda, \lambda, \mathbf{g})(\mathbf{z}) \rangle > 0, \quad (14)$$

where

$$F''(\mathbf{x}, \lambda, \mathbf{g}) = \lambda_{\mathbf{x}\mathbf{x}}f''(\mathbf{x}) + \langle \lambda, \mathbf{g}(\mathbf{x}) \rangle + \langle \mathbf{g}, \mathbf{h}(\mathbf{x}) \rangle.$$

Then f attains a strict local minimum for problem *P1* at $\mathbf{x}_0$.

Note that strict inequality is required in (14), while equality is allowed in the necessary condition (10).

Proof. To simplify notation, we assume that $\mathbf{x}_0 = \mathbf{0}$ and that $f(\mathbf{0}) = 0$. For any real-valued function $\varphi: \mathbb{R}_n \rightarrow \mathbb{R}$, the symbol $\varphi_{\mathbf{x}\mathbf{x}}$ will denote the Hessian matrix whose (i, j) th entry is $\partial^2 \varphi / \partial x_i \partial x_j$.

If f did not attain a strict local minimum at $\mathbf{0}$, then for every positive integer q there would exist a feasible point $\mathbf{x}_q \neq \mathbf{0}$ such that $\mathbf{x}_q \in \mathbb{R}\mathbf{h}, \lambda/\mathbf{g}$ and $f(\mathbf{x}_q) \leq 0$.

Then there would exist a sequence of points $\{x_j\}$ in X such that

$$x_j \neq \bar{x}, \quad \lim_{j \rightarrow \infty} x_j = \bar{x}, \quad f(x_j) < \bar{v}, \quad g(x_j) < 0, \quad \lim_{j \rightarrow \infty} g(x_j) = 0.$$

Let $E = \{1, \dots, r\}$. It then follows from Taylor's theorem (Corollary 1, Theorem 1.1.2) and our assumption that $x_0 = \bar{x}$ and $f(\bar{x}) = \bar{v}$ that

$$f(x_j) - \langle \nabla f(x_0), x_j \rangle < 0 \quad (15)$$

and that for $i = 1, \dots, r$ and $j = 1, \dots, k$

$$\begin{aligned} a_i(x_j) &= \langle \nabla a_i(x_0), x_j \rangle < 0, \\ b_j(x_j) &= \langle \nabla b_j(x_0), x_j \rangle = 0 \end{aligned} \quad (16)$$

where the x_j are numbers in the open interval $(0, 1)$.

It also follows from Taylor's theorem that

$$f(x_j) = \langle \nabla f(x_0), x_j \rangle + \frac{1}{2} \langle x_j, f_{xx}(x_j) x_j \rangle < 0 \quad (17)$$

and that for $i = 1, \dots, r$ and $j = 1, \dots, k$

$$\begin{aligned} a_i(x_j) &= \langle \nabla a_i(x_0), x_j \rangle + \frac{1}{2} \langle x_j, a_{ixx}(x_j) x_j \rangle < 0, \\ b_j(x_j) &= \langle \nabla b_j(x_0), x_j \rangle + \frac{1}{2} \langle x_j, b_{jxx}(x_j) x_j \rangle = 0, \end{aligned} \quad (18)$$

where the x_j are real numbers in the open interval $(0, 1)$.

Since $x_0 \neq \bar{x}$, the sequence $x_j/|x_0|$ is a sequence of unit vectors. Hence there exists a unit vector v and a subsequence that we relabel as $\{x_j\}$ such that

$$x_j/|x_0| \rightarrow v.$$

Now suppose that the sequence $\{x_j\}$ in (15) and (16) is the subsequence. If we divide through by $|x_0|$ in the rightmost inequalities and the equalities in (15) and then let $j \rightarrow +\infty$, we get that v satisfies

$$\langle \nabla f(x_0), v \rangle < 0, \quad \nabla g_0(v) < 0, \quad \nabla h_0(v) = 0.$$

Thus v satisfies the tangential constraints (4).

Next, multiply the first inequality in (16) by $\lambda_i \geq 0$, the i th inequality in (16) by $\lambda_i \geq 0$ (the i th component of λ), and the j th equation in (16) by μ_j (the j th component of μ) and then add the resulting equations and inequalities. Since

$\lambda = (\lambda_0, \lambda_2)$ and $\lambda_1 = 0$, we get

$$\langle \lambda_0 \nabla f(\bar{x}) + \lambda_2^* \nabla g(\bar{x}) + \mu^* \nabla h(\bar{x}, \bar{v}_q) + \frac{1}{2} \langle \bar{v}_q, P_{\bar{x}}^* (P_{\bar{x}} \bar{v}_q) \rangle \rangle \leq 0,$$

where

$$P_{\bar{x}}^* (P_{\bar{x}} \bar{v}_q) = \lambda_{00} \bar{L}_{xx} (P_{\bar{x}} \bar{v}_q) + \sum_{i=1}^k \lambda_{0i} \bar{L}_{xi} (P_{\bar{x}} \bar{v}_q) + \sum_{j=1}^l \alpha_j \bar{L}_{xj} (P_{\bar{x}} \bar{v}_q).$$

Since $(\bar{x}, \bar{v}^*, \lambda, \mu)$ satisfies the *FJ* necessary conditions of Theorem 3.1, we get that

$$\frac{1}{2} \langle \bar{v}_q, P_{\bar{x}}^* (P_{\bar{x}} \bar{v}_q) \rangle \leq 0.$$

If we now divide both sides of this inequality by $\|\bar{v}_q\|^2$ and then let $q \rightarrow \infty$, we get that $\langle \bar{x}, P_{\bar{x}}^* (P_{\bar{x}} \bar{v}) \rangle \leq 0$ for a nonzero vector that satisfies the tangential constraints (4) and $\langle \nabla f(\bar{x}), \bar{v} \rangle \leq 0$. This contradicts the assumption that (14) holds for all such $\bar{v}$.

If there exists a set of multipliers at $\bar{x}_0$ with $\lambda_0 = 1$, we have a slightly different version of the sufficient conditions.

COROLLARY 4. Let $(\bar{x}_0, \bar{v}_0, \mu)$ satisfy the *KKT* necessary conditions of Theorem 3.2. Suppose that every $\bar{v} \neq 0$ in $\mathbb{R}^n$ that satisfies the tangential constraints (4) at $\bar{x}_0 = \bar{x}_0$ and the equality $\langle \nabla f(\bar{x}_0), \bar{v} \rangle = 0$ also satisfies

$$\langle \bar{x}, P_{\bar{x}_0} (\bar{x}_0, \bar{v}_0, \mu) \bar{v} \rangle > 0, \quad (17)$$

where

$$P(\bar{x}, \bar{v}_0, \mu) = f(\bar{x}) + \langle \bar{v}_0, g(\bar{x}) \rangle + \langle \bar{\mu}, h(\bar{x}) \rangle. \quad (18)$$

Then f attains a strict local minimum for problem *P1* at $\bar{x}_0$.

The corollary differs from the theorem in that the condition $\langle \nabla f(\bar{x}_0), \bar{v} \rangle \leq 0$ is replaced by $\langle \nabla f(\bar{x}_0), \bar{v} \rangle = 0$.

Proof. If $(\bar{x}_0, \bar{v}_0, \mu)$ satisfies the conditions of Theorem 3.2, then

$$\nabla f(\bar{x}_0) + \lambda_2^* \nabla g(\bar{x}_0) + \mu^* \nabla h(\bar{x}_0) = 0.$$

Since $\lambda_1 = 0$,

$$\nabla f(\bar{x}_0) = -\lambda_2^* \nabla g(\bar{x}_0) - \mu^* \nabla h(\bar{x}_0).$$

Hence, for any $x \in E^*$

$$\langle U_j(x_0), x \rangle = -\langle \lambda_0, \nabla u_j(x_0) \rangle - \langle g, \nabla h_j(x) \rangle.$$

Since $\lambda_0 \geq 0$, it follows that for any x satisfying the tangential constraints (8) we have $\langle \nabla u_j(x_0), x \rangle \geq 0$. Hence the condition $\langle U_j(x_0), x \rangle \leq 0$ reduces to $\langle \nabla h_j(x_0), x \rangle = 0$.

If the inequality constraints are absent, then a vector x satisfies the tangential constraints at a point x_0 if and only if $\nabla h(x_0)x = 0$. Thus, if at a point x_0 the necessary conditions of Corollary 2 of Theorem 5.2 are satisfied, the proof of Corollary 4 shows that for every vector x such that $\nabla h(x_0)x = 0$ we have $\langle \nabla f(x_0), x \rangle = 0$. Therefore the following corollary holds.

COROLLARY 3. Let (x_0, λ) satisfy the necessary conditions of Corollary 2 of Theorem 5.2. Suppose that for every x such that $\nabla h(x_0)x = 0$

$$\langle x, F_{xx}(x_0, \lambda)x \rangle \geq 0,$$

where

$$F(x, \mu) = f(x) + \langle \mu, h(x) \rangle.$$

The f attains a strict local minimum for the problem of minimizing f subject to $h(x) = 0$.

There is a sufficient condition due to Peneda [1833] which is not as convenient to apply as Corollary 4 but is equivalent to it. We now state this condition and prove its equivalence to Corollary 4.

COROLLARY 6. Let (x_0, λ, μ) satisfy the KKT necessary conditions of Theorem 5.2. Let $E^* = \{x \in E \text{ and } \lambda_1 > 0\}$. Suppose that for every $x \neq 0$ in E^* that satisfies

$$\begin{aligned} \langle \nabla u_j(x_0), x \rangle &= 0, \quad x \in E^*, & \langle \nabla u_j(x_0), x \rangle &\leq 0, & x \in E, & \quad \forall j \in J^*, \\ \nabla h(x_0)x &= 0, \end{aligned} \tag{19}$$

the inequality

$$\langle x, F_{xx}(x_0, \lambda, \mu)x \rangle > 0$$

holds, where F is as in (18). Then f attains a strict local minimum for problem $P1$ at x_0 .

Proof. To prove the corollary it suffices to show that the set V_1 of vectors that satisfy (19) is the same as the set V_2 of vectors that satisfy the tangential

constraints (4) and the equality $\langle \nabla f(x_a), x \rangle = 0$, and then invoke Corollary 4.

We first show that $V_1 \subseteq V_7$. If $V_1 = \emptyset$, there is nothing to prove, so we assume that $V_1 \neq \emptyset$. From (6) of Theorem 5.2 we get that for x in $\mathbb{R}^2$

$$\langle \nabla f(x_a), x \rangle + \lambda_1 \nabla g_1(x_a)x + \langle p, \nabla h(x_a)x \rangle = 0. \quad (20)$$

Now let $x \in V_1$. Then $\nabla g_1(x_a)x \leq 0$, and $\nabla h(x_a)x = 0$. Thus x satisfies the tangential constraints (6). Moreover, since $\lambda_i = 0$ for $i \neq 1$, we have that $\lambda_i = 0$ for $i \in \mathbb{R}^2$. Thus (20) becomes

$$\langle \nabla f(x_a), x \rangle + \sum_{i \in \mathbb{R}^2} \lambda_i \langle \nabla g_i(x_a), x \rangle = 0$$

for all x in V_1 . For such x and for $v \in \mathbb{R}^2$ we have $\langle \nabla g_i(x_a), v \rangle = 0$. Hence $\langle \nabla f(x_a), v \rangle = 0$, and so $V_1 \subseteq V_7$.

We now show that $V_1 \subseteq V_5$. We assume that $V_1 \neq \emptyset$. Since $\lambda_i = 0$ for $i \in \mathbb{R}^2$ and $\langle \nabla f(x_a), x \rangle = 0$, it follows that, for $x \in V_1$, (20) becomes

$$\sum_{i \in \mathbb{R}^2} \lambda_i \langle \nabla g_i(x_a), x \rangle = 0.$$

Since for any vector x that satisfies the tangential constraints (4) we have $\langle \nabla g_i(x_a), x \rangle \leq 0$, and since $\lambda_i \geq 0$ for $i \in \mathbb{R}^2$, the last equality implies that $\langle \nabla g_i(x_a), x \rangle = 0$ for $i \in \mathbb{R}^2$. Since $\nabla g_i(x_a)x \leq 0$ for any vector satisfying the tangential constraints, we get that $V_1 \subseteq V_7$.

Example 5.1 (Continued). We now apply Corollary 4 to show that the points given in relation (4) in Section 5 furnish local minima for problems π_k when $0 < k < 1$ and that the origin furnishes a local minimum for problems π_k when $k \geq 1$.

We first consider the problem π_k for $0 < k < 1$. At the point $x_{1,k} = (1-k, \sqrt{2k(1-k)})$, the condition $\langle \nabla f(x_a), x \rangle = 0$ becomes

$$(1-2k, 2_k/\sqrt{2k(1-k)})(x_1, x_2) = 0.$$

At the point $x_{1,k} = (1-k, \sqrt{2k(1-k)})$ the condition $\langle \nabla f(x_a), x \rangle = 0$ becomes

$$(1-2k, -2_k/\sqrt{2k(1-k)})(x_1, x_2) = 0.$$

Thus at $x_{1,k}$ and $x_{2,k}$ we require that

$$x_{1,k}x_{2,k} = 2_k \sqrt{\frac{2(1-k)}{k}}, \quad x_{1,k}x_{2,k} = -2_k \sqrt{\frac{2(1-k)}{k}}. \quad (21)$$

The tangential constraints at the points x_{1a} and x_{2a} become

$$\begin{aligned} x_{1a} : (2k_1 - 2\sqrt{2k_1}) - k_1(x_1 - x_2) &\leq 0, \\ x_{2a} : (2k_1 - 2\sqrt{2k_1}) - k_1(x_1 - x_2) &\leq 0 \end{aligned}$$

or

$$x_{1a} : x_1 \leq x_2 \sqrt{\frac{2(1-k)}{k}}, \quad x_{2a} : x_1 \leq -x_2 \sqrt{\frac{2(1-k)}{k}}.$$

Hence at x_{1a} the vectors $x = (x_1, x_2)$ that satisfy $\langle \nabla f(x_{1a}), x \rangle = 0$ and the tangential constraints are given in (21). Similarly, at x_{2a} the vectors that satisfy $\langle \nabla f(x_{2a}), x \rangle = 0$ and the tangential constraints are given in (21). Since $d = 1$ at x_{1a} and x_{2a} , the function F in (18) becomes

$$F = (x_1 - 1)^2 + x_2^2 + (2kx_1 - x_2^2).$$

Thus

$$\langle x, F_{xx}(x, 1, \mu) \rangle = (x_1, x_2) \begin{pmatrix} 2 & 0 \\ 0 & 2(1-k) \end{pmatrix} = 2(x_1^2 - kx_2^2). \quad (22)$$

At x_{1a} and at x_{2a} vectors $x \neq 0$ at which (22) holds have $x_1 \neq 0$. It therefore follows that the hypotheses of Corollary 4 hold at x_{1a} and x_{2a} , and so both of these points furnish local minima. We leave it as an exercise to show that these points furnish minima for problem π_{μ} with $0 < k < 1$.

We now consider the problem π_{μ} , $k > 1$. The origin is the only point that satisfies the conditions of Theorem 5.2. Since $\nabla f(0) = (-2, 0)$, all vectors x of the form $(0, x_2)$, with x_2 arbitrary, satisfy $\langle \nabla f(0), x \rangle = 0$. Since $\nabla g(0) = (2k, 0)$, all vectors x of the form $x = (x_1, x_2)$, with $x_1 \leq 0$, and x_2 arbitrary, satisfy the tangential constraint $\langle \nabla g(0), x \rangle \leq 0$. Hence the vectors that satisfy the tangential constraints and $\langle \nabla f(0), x \rangle = 0$ are vectors of the form $(0, x_2)$, where x_2 is arbitrary.

The function F with $d = 1/k$ becomes

$$F = (x_1 - 1)^2 + x_2^2 + \frac{2kx_1 - x_2^2}{k}.$$

Thus (17) with $a = (0, x_2)$ becomes

$$(0, x_2) \begin{pmatrix} 2 & 0 \\ 0 & 2(1 - \frac{1}{k}) \end{pmatrix} \begin{pmatrix} 0 \\ 0 \end{pmatrix} = 2 \left(1 - \frac{1}{k} \right) (x_2)^2 > 0, \quad x_2 \neq 0.$$

Hence the origin furnishes a strict local minimum for problem π_k , $k > 1$. Again, we leave it as an exercise for the reader to show that the origin furnishes a global minimum.

If $k = 1$, then $\mathbf{x}_k = (0, 0)$ and

$$\langle \nabla f(\mathbf{0}), \mathbf{x} \rangle = \langle (-2, 0), (x_1, x_2) \rangle = -2x_1,$$

$$\langle \nabla g(\mathbf{0}), \mathbf{x} \rangle = \langle (2, 0), (x_1, x_2) \rangle = 2x_1.$$

Hence a vector $\mathbf{x}$ satisfies the tangential constraint $\langle \nabla g(\mathbf{0}), \mathbf{x} \rangle \leq 0$ and the constraint $\langle \nabla f(\mathbf{0}), \mathbf{x} \rangle = 0$ if and only if it has the form $\mathbf{x} = (0, x_2)$, x_2 arbitrary. If $\mathbf{x}_k = \mathbf{0}$, then $\lambda = 1$ and (22) gives $\langle \mathbf{x}, F_{\alpha}(\mathbf{0}, 1, \rho)\mathbf{x} \rangle = 0$. Hence the sufficient condition gives no information in this case.

Exercise 6.1. Show that the points $\mathbf{x}_{1a}$, $\mathbf{x}_{2a}$ and $\mathbf{0}$ in Example 5.1 furnish minima for the relevant problems.

Exercise 6.2. In Exercises 5.1–5.5 use an appropriate second-order sufficiency condition to determine whether the critical points furnish a local minimum.

Exercise 6.3. Given m distinct points $\mathbf{y}_1, \dots, \mathbf{y}_m$ in $\mathbb{R}^n$, find the smallest closed ball $\overline{B(\mathbf{x}, \rho)}$ that encloses these points. This problem can be formulated analytically as follows: Minimize $f(\mathbf{x}, \rho) = \rho^2$ subject to $\|\mathbf{x} - \mathbf{y}_i\|^2 \leq \rho^2$.

- Show that a solution $(\mathbf{x}_0, \rho_0)$ exists.
- Show that at a critical point $(\mathbf{x}_0, \rho_0)$

- (i) $\lambda_0 = 1$,

- (ii) $\lambda_i \geq 0$, $i = 1, \dots, m$,

- (iii) $\sum_{i=1}^m \lambda_i = 1$,

- (iv) $\mathbf{x}_0 = \sum_{i=1}^m \lambda_i \mathbf{y}_i$, and

- (v) $\rho_0^2 + \|\mathbf{x}_0\|^2 = \sum_{i=1}^m \lambda_i \|\mathbf{y}_i\|^2$.

Note that (ii), (iii), and (iv) show that $\mathbf{x}_0$ is in the simplex determined by $\mathbf{y}_1, \dots, \mathbf{y}_m$.

- Show that with the KKT multipliers as in (b) the function F can be written as

$$F(\rho, \mathbf{x}, \lambda) = \|\mathbf{x}\|^2 - 2 \left\langle \mathbf{x}, \sum_{i=1}^m \lambda_i \mathbf{y}_i \right\rangle.$$

- Does the problem have a unique solution?
- Solve the special case $\mathbf{y}_i = \mathbf{a}_i$, $i = 1, 2, 3$, in detail.

V

CONVEX PROGRAMMING AND DUALITY

1. PROBLEM STATEMENT

In the preceding chapter we studied programming problems with differentiable data. In this chapter we shall replace the assumption that the data are differentiable by the assumption that the data are convex. This assumption will enable us to obtain much more detailed information about the behavior of solutions. In Section 1 of Chapter IV we defined a convex programming problem to be a programming problem in which the underlying set X_0 is convex, the functions f and g are convex, and there are no equality constraints.

Convex Programming Problem 1 (CPI)

Let X_0 be a convex set and let $f: X_0 \rightarrow \mathbb{R}$ and $g: X_0 \rightarrow \mathbb{R}^m$ be convex.

$$\text{Minimize } \{f(x) : x \in X\},$$

where

$$X = \{x : g(x) \in \theta\}.$$

The problem is often stated as

$$\text{Minimize } f(x) \quad \text{subject to } g(x) \in \theta.$$

Remark 1.1. A convex programming problem with the additional equality constraint $h(x) = \theta$ can be cast in the format CPI if the function h is affine. One replaces the constraint $h(x) = \theta$ by the two inequality constraints $h(x) \leq \theta$ and $-\mathbf{h}(x) \leq \theta$. Since h is affine, both h and $-\mathbf{h}$ are convex.

In Remark IV.5.8 we pointed out that if the feasible set X in CPI is not empty, then it is convex. Thus in CPI we are minimizing a convex function

defined on the convex set X . Therefore, the properties of solutions listed in Theorems IV.5.1 are valid for CPl.

If the functions f and g are differentiable on N_p , then CPl is a special case of the programming problem Pl. Hence Theorem IV.5.1 (the Fritz John necessary conditions), Theorem IV.5.2 (the KKT necessary conditions), and the sufficiency theorem IV.5.3 are all valid.

It turns out that in many problems involving affine equality constraints $\mathbf{h}(\mathbf{x}) = \mathbf{0}$, sharper results can be obtained if we retain the equality constraints as such rather than replace them by two sets of inequality constraints. It will also be convenient to take $N_p = \mathbb{R}^n$ when affine constraints are present. We therefore formulate the convex programming problem with affine constraints as follows.

Convex Programming Problem II (CPII)

Let $f: \mathbb{R}^n \rightarrow \mathbb{R}$ be convex, let $g: \mathbb{R}^n \rightarrow \mathbb{R}^m$ be convex and let $\mathbf{h}: \mathbb{R}^n \rightarrow \mathbb{R}^k$ be affine. Minimize

$$f(\mathbf{x}) : \mathbf{x} \in X,$$

where

$$X = \{\mathbf{x} : g(\mathbf{x}) \leq \mathbf{0}\} \cap \{\mathbf{x} : \mathbf{h}(\mathbf{x}) = \mathbf{0}\}.$$

If the feasible set X is not empty, the problem is said to be consistent.

Remark 1.1. The components of the affine equality constraints $\mathbf{h}(\mathbf{x}) = \mathbf{0}$ can be written as

$$\langle \mathbf{a}_i, \mathbf{x} \rangle - b_i = 0, \quad i = 1, \dots, k. \quad (1)$$

Thus, if the problem is consistent, then the system

$$A\mathbf{x} = \mathbf{b}, \quad (2)$$

where A is the $k \times n$ matrix whose i th row is $\mathbf{a}_i$ and $\mathbf{b} = (b_1, \dots, b_k)^T$ has solutions. We recall from linear algebra that for (1) to have solutions the rank of A must equal the rank of the augmented matrix $(A, \mathbf{b})$. If the system (2) has solutions and if the rank of $(A, \mathbf{b})$ were less than k , then one of the equations in (1) could be written as a linear combination of the others and thus eliminated from the system. Since we are only interested in problems that are consistent, there is no loss of generality in assuming that A has maximum possible rank, namely k . In that case the ranks of A and $(A, \mathbf{b})$ both equal k . If the feasible set consists of only one point, then the programming problem is trivial. Therefore, we further assume that $k < n$. For if $k = n$ and A has rank

$k = n$, then (2) has a unique solution and the feasible set consists of a single point. The assumptions that A has rank k and that $k = n$ will be in force henceforth.

In the next section we shall develop necessary conditions and a sufficient condition for a point x_0 to be a solution of CPL without assuming differentiability.

2. NECESSARY CONDITIONS AND SUFFICIENT CONDITIONS

Our first necessary condition will involve a function A with domain $\mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}^n \times \mathbb{R}$ and range in $\mathbb{R}$ defined by

$$A(x, \lambda, k, p) = \lambda_0 f(x) + (p, g(x)) + (q, h(x)). \quad (1)$$

The function A is sometimes called a generalized Lagrangian.

THEOREM 2.1. Let f and g be convex on $\mathbb{R}^n$ and let h be affine. Let x_0 be a solution of CPL. Then there exist a real number λ_{0x_0} , a vector k_{x_0} in $\mathbb{R}^n$, and a vector p_{x_0} in $\mathbb{R}^m$ with $(\lambda_{0x_0}, k_{x_0}) \geq (0, 0)$ and $(\lambda_{0x_0}, k_{x_0}, p_{x_0}) \neq (0, 0, 0)$ such that

$$Q_{x_0}(g(x_0)) = 0 \quad (2)$$

and such that for all x in $\mathbb{R}^n$, all $k \geq 0$ in $\mathbb{R}^n$, and all p in $\mathbb{R}^m$

$$A(x_0, \lambda_{0x_0}, k, p) \leq A(x_0, \lambda_{0x_0}, k_{x_0}, p_{x_0}) \leq A(x, \lambda_{0x_0}, k_{x_0}, p_{x_0}). \quad (3)$$

Moreover,

$$A(x_0, \lambda_{0x_0}, k_{x_0}, p_{x_0}) = \lambda_{0x_0} f(x_0). \quad (4)$$

Proof. Since x_0 is a solution of CPL, the system

$$\begin{aligned} f(x) - \beta(x_0) &< 0, \\ g(x) &< 0, \quad h(x) = 0, \end{aligned}$$

has no solution in $\mathbb{R}^n$. Hence by Theorem III.4.1 there exists a vector $(\lambda_{0x_0}, k_{x_0}, p_{x_0}) \neq (0, 0, 0)$ with $(\lambda_{0x_0}, k_{x_0}) \geq 0$ such that

$$\lambda_{0x_0} f(x_0) = f(x_0) + Q_{x_0}(g(x_0)) + (p_{x_0}, h(x_0)) \geq 0$$

for all x in $\mathbb{R}^n$. Hence

$$\lambda_{\text{opt}} f(x_0) \leq \lambda_{\text{opt}} f(x) + \langle \lambda_{\text{opt}}, g(x) \rangle + \langle \mu_{\text{opt}}, h(x) \rangle \quad (3)$$

for all x in $\mathbb{R}^n$. In particular, (3) holds for $x = x_0$. Since $h(x_0) = 0$, we get that $\langle \lambda_{\text{opt}}, g(x_0) \rangle \geq 0$. On the other hand, since x_0 is feasible, $g(x_0) \leq 0$. But $\lambda_{\text{opt}} \geq 0$, so $\langle \lambda_{\text{opt}}, g(x_0) \rangle \leq 0$. Hence, $\langle \lambda_{\text{opt}}, g(x_0) \rangle = 0$, and (2) is established.

From (2), (3), and $h(x_0) = 0$ we now get

$$\begin{aligned} \lambda_{\text{opt}} f(x_0) + \langle \lambda_{\text{opt}}, g(x_0) \rangle + \langle \mu_{\text{opt}}, h(x_0) \rangle \\ = \lambda_{\text{opt}} f(x_0) \leq \lambda_{\text{opt}} f(x) + \langle \lambda_{\text{opt}}, g(x) \rangle + \langle \mu_{\text{opt}}, h(x) \rangle \end{aligned}$$

for all x . This establishes (4) and the rightmost inequality in (3). From (2) and from $g(x_0) \leq 0$ it follows that for $\lambda \geq 0$

$$\langle \lambda, g(x_0) \rangle \leq 0 = \langle \lambda_{\text{opt}}, g(x_0) \rangle.$$

From this and from $h(x_0) = 0$ we get that

$$\lambda_{\text{opt}} f(x_0) + \langle \lambda, g(x_0) \rangle + \langle \mu_{\text{opt}}, h(x_0) \rangle \leq \lambda_{\text{opt}} f(x_0) + \langle \lambda_{\text{opt}}, g(x_0) \rangle + \langle \mu_{\text{opt}}, h(x_0) \rangle,$$

which is the leftmost inequality in (3).

Remark 2.1. Since $\lambda_{\text{opt}} \geq 0$ and $g(x_0) \leq 0$, condition (2) is equivalent to

$$\lambda_{\text{opt}} g(x_0) = 0, \quad i = 1, \dots, m. \quad (5)$$

Remark 2.2. Although (5) is written as a saddle point condition for the function λ , the leftmost inequality does not involve μ since $h(x_0) = 0$.

Remark 2.3. If the functions f and g are differentiable, then since x_0 minimizes the function $\lambda(\cdot, \lambda_{\text{opt}}, \lambda_{\text{opt}}, \mu_{\text{opt}})$ on $\mathbb{R}^n$, where

$$\lambda(\lambda, \lambda_{\text{opt}}, \lambda_{\text{opt}}, \mu_{\text{opt}}) = \lambda_{\text{opt}} f(x) + \langle \lambda_{\text{opt}}, g(x) \rangle + \langle \mu_{\text{opt}}, h(x) \rangle,$$

we have

$$\frac{\partial \lambda}{\partial x_i}(\lambda_{\text{opt}}, \lambda_{\text{opt}}, \lambda_{\text{opt}}, \mu_{\text{opt}}) = 0, \quad i = 1, \dots, n.$$

Thus, if A is the matrix in (1.2),

$$\lambda_{\text{opt}} \nabla f(x_0) + \lambda_{\text{opt}}^T \nabla g(x_0) + \mu_{\text{opt}}^T A = 0,$$

which together with (2) and the condition $\langle \lambda_{\text{opt}}, \lambda_{\text{opt}}, \mu_{\text{opt}} \rangle \in (0, 0, 0)$ constitutes the FJ necessary condition of Theorem IV.3.1.

Remark 1.1. In the proof of Theorem 2.1, the optimality of x_0 is shown to imply the existence of a vector $(\lambda_{00}, \lambda_{01}, p_0) \in (0, \infty) \times \mathbb{R}$ with $\lambda_{00} > 0$ and $\lambda_{01} > 0$ such that relation (2) holds. All the conclusions of the theorem then follow from (2) and the feasibility of x_0 . We shall use this observation in the sequel.

As in the study of differentiable problems, we seek a constraint qualification that guarantees that $\lambda_{00} \neq 0$, for then we can take $\lambda_{00} = 1$. For the convex programming problem the following simple constraint qualification suffices.

Definition 1.1. The problem CPH is said to be *strongly consistent*, or to satisfy the Slater condition, if there exists an x_0 such that $g(x_0) < 0$ and $h(x_0) = 0$.

Theorem 2.1. Let x_0 be a solution of CPH and let the problem be strongly consistent. Then there exists a $\lambda_0 > 0$ in $\mathbb{R}^n$ such that

$$\langle \lambda_0, g(x_0) \rangle = 0 \quad (9)$$

and a vector p_0 in $\mathbb{R}^m$ such that for all x , all $\lambda > 0$ in $\mathbb{R}^n$, and all p in $\mathbb{R}^m$

$$f(x_0, \lambda, p) \leq f(x_0, \lambda_0, p_0) \leq f(x, \lambda_0, p_0), \quad (10)$$

where

$$f(x, \lambda, p) = f(x) + \langle \lambda, g(x) \rangle + \langle p, h(x) \rangle, \quad (11)$$

Moreover,

$$f(x_0) = f(x_0, \lambda_0, p_0). \quad (12)$$

Proof. Since x_0 is a solution of CPH, the conclusion of Theorem 1.1 holds. If $\lambda_{00} = 0$, then since (2) holds and $h(x_0) = 0$, the rightmost inequality in (3) becomes

$$0 \leq \langle \lambda_0, g(x) \rangle + \langle p_0, h(x) \rangle \quad (13)$$

for all x . If we also have $\lambda_0 = 0$, then the condition $(\lambda_{00}, \lambda_{01}, p_0) \neq (0, 0, 0)$ implies that $p_0 \neq 0$. Since $h(x) = Ax - b$, the inequality (13) and $\lambda_0 = 0$ imply that

$$\langle p_0^T A, x \rangle \geq \langle p_0^T A, x_0 \rangle \quad (14)$$

for all x . The standing hypothesis made in Section 1 that A has rank k and the condition $p_0 \neq 0$ imply that $p_0^T A \neq 0$. Hence if we take $x = pp_0^T A$ and then let $p \rightarrow +\infty$, we see that (14) cannot hold for all x . Hence $\lambda_0 \neq 0$.

Since the problem is strongly consistent, there exists an x_0 such that $g(x_0) < 0$ and $h(x_0) = 0$. Thus, from (10), $\langle \lambda_0, g(x_0) \rangle > 0$. On the other hand,

$\lambda_{i_0} \geq 0$ and $\lambda_{i_0} \neq 0$ imply that $\langle \lambda_{i_0}, g(x_{i_0}) \rangle < 0$. This contradiction shows that $\lambda_{i_0} = 0$. Therefore, we can divide through by $\lambda_{i_0} = 0$ in (2), (3), and (4), relabel $\lambda_{i_0}/\lambda_{i_0}$ as λ_{i_0} and relabel μ_{i_0}/λ_{i_0} as μ_{i_0} to get the conclusion of Theorem 2.2.

Remark 2.3. The rightmost inequality in (7) and equation (9) give

$$f(x_{i_0}) = \inf_{x \in X} f(x) + \langle \lambda_{i_0}, g(x_{i_0}) \rangle \leq f(x) + \langle \lambda_{i_0}, g(x) \rangle + \langle \mu_{i_0}, h(x) \rangle$$

for all x . Therefore, if we assume that CPH has a solution x_{i_0} and if we know the multipliers λ_{i_0} and μ_{i_0} , then the solution x_{i_0} of the constrained problem

$$\min\{f(x) : x \in X\}, \quad X = \{x : g(x) \leq 0, h(x) = 0\}$$

will be a solution of the unconstrained problem

$$\min\{f(x) + \langle \lambda_{i_0}, g(x) \rangle + \langle \mu_{i_0}, h(x) \rangle\}. \quad (12)$$

The set S of solutions of the unconstrained problem (12) may contain points that do not satisfy the constraints. To find solutions of the constrained problem, one must determine those points of S that satisfy the constraints. It follows from the rightmost inequality in (7) that if (12) has a unique solution, then it must be a solution of the constrained problem.

One sometimes finds the statement that we may replace the constrained problem by the unconstrained problem of minimizing a Lagrangian. This statement is only true in the sense described in the preceding paragraph.

It is conceivable that there exist vectors (λ, g) with $\lambda \geq 0$ such that the unconstrained problem (12) with $(\lambda_{i_0}, \mu_{i_0}) = (\lambda, \mu)$ has the same infimum as constrained problem. Such a vector will be called a *Kuhn-Tucker vector*. The multiplier vectors $(\lambda_{i_0}, \mu_{i_0})$ of Theorem 2.2 are Kuhn-Tucker vectors. In the next theorem, we show that if CPH has a solution, then a Kuhn-Tucker vector for CPH is a multiplier vector as in Theorem 2.2.

Theorem 2.3. Let CPH have a solution x_{i_0} . Then a Kuhn-Tucker vector $(\bar{\lambda}, \bar{\mu})$ for CPH is a multiplier vector $(\lambda_{i_0}, \mu_{i_0})$ as in Theorem 2.2.

Proof. Since x_{i_0} is a solution of CPH and since $(\bar{\lambda}, \bar{\mu})$ is a Kuhn-Tucker vector,

$$f(x_{i_0}) = \inf\{f(x) + \langle \bar{\lambda}, g(x) \rangle + \langle \bar{\mu}, h(x) \rangle : x \in \mathbb{R}^n\}.$$

Thus (7) holds with $\lambda_{i_0} = 1$. Hence by the italicized statement in Remark 2.4, the vector $(\bar{\lambda}, \bar{\mu})$ is a multiplier vector as in Theorem 2.2.

Theorems 2.2 and 2.3 imply the following:

COROLLARY 1. Let CPT have a solution x_0 . Let $\mathfrak{M}$ denote the set of subgradients associated with x_0 as in Theorem 2.2 and let $\mathcal{K}$ denote the set of Kuhn-Tucker vectors for CPT. Then $\mathfrak{M} = \mathcal{K}$.

Remark 2.6. The hypothesis of strong consistency in Theorem 2.2 guaranteed that the set $\mathfrak{M}$ is not empty. Thus if CPT is strongly consistent and has a solution, then $\mathcal{K}$ is also nonempty.

We now give an example of a problem, taken from Rockafellar [170], for which no Kuhn-Tucker vector, and hence no vector in $\mathfrak{M}$, exists.

Example 2.7. Consider $\mathbb{R}^2$, and let

$$f(x_1, x_2) = x_1, \quad g_1(x_1, x_2) = x_2, \quad g_2(x_1, x_2) = x_1^2 - x_2.$$

Then the only point in $\mathbb{R}^2$ satisfying $g(x_1, x_2) \leq 0$ is $(0, 0)$, and the value of the minimum is zero. Thus, a Kuhn-Tucker vector $\lambda = (\lambda_1, \lambda_2)$ would have to satisfy $\lambda_1 \geq 0, \lambda_2 \geq 0$ and

$$0 \leq x_1 + \lambda_1 x_1 + \lambda_2 (x_1^2 - x_2) = x_1 + (\lambda_1 - \lambda_2)x_1 + \lambda_2 x_1^2$$

for all x . We leave it to the reader to check that this is impossible.

Note that the problem in this example is not strongly consistent.

Condition (I) of Theorem 2.1 and condition (7) of Theorem 2.2 are called *saddle point conditions*. The points $(x_0, \lambda_{x_0}, \lambda_{x_0}, p_0)$ in (I) and $(x_0, \lambda_{x_0}, p_0)$ in (7) are called *saddle points*. We previously encountered saddle point conditions and saddle points in our discussion of game theory in Section 3 of Chapter III. If the convex programming problem CPT has a solution and if the problem is strongly consistent, then the saddle point condition (7) gives us a method of computing $f(x_0)$, the value of the minimum, without finding the optimal point x_0 , as is shown in the following theorem. We shall return to this point in our study of duality in Section 4.

THEOREM 2.8. Let CPT be strongly consistent, let x_0 be a solution, let v denote the value of the minimum, and let λ_{x_0}, p_{x_0} be the subgradients associated with x_0 as in Theorem 2.2. Then

$$\begin{aligned} v &= L(x_0, \lambda_{x_0}, p_{x_0}) = \min_{\lambda, p} L(x_0, \lambda, p) \\ &= \sup_{\lambda, p} \inf_{x \in \mathbb{R}^n} L(x, \lambda, p) = \inf_{\lambda, p} \sup_{x \in \mathbb{R}^n} L(x, \lambda, p), \end{aligned} \quad (2.8)$$

where $\lambda \in \mathbb{R}^m, \lambda \geq 0, p \in \mathbb{R}^r, v \in \mathbb{R}$.

Proof. By a slight modification of the argument in Lemma III.5.1 to take into account the possible value of $+\infty$ for a supremum and $-\infty$ for an infimum, we have that

$$\sup_{k \in K} \inf_{x \in X} L(x, k, p) \leq \inf_{x \in X} \sup_{k \in K} L(x, k, p). \quad (14)$$

From (7) of Theorem 2.2 we get

$$\sup_{k \in K} L(x_{op}, k, p) = V(x_{op}, k_{op}, p_{op}) = \inf_{k \in K} L(x_{op}, k, p). \quad (15)$$

Hence

$$\inf_{k \in K} \sup_{x \in X} L(x, k, p) \leq L(x_{op}, k_{op}, p_{op}) \leq \sup_{k \in K} \inf_{x \in X} L(x, k, p). \quad (16)$$

We also note that by (7) of Theorem 2.2 we can replace the infimum in (15) by a minimum. This observation and (14) and (16) establish all the equalities in (11) except the first, which follows from (9) of Theorem 2.2.

Remark 2.7. It follows from Remark 2.6 and the proof of Theorem 2.4 that the hypothesis of strong consistency in the theorem can be replaced by the hypothesis that W is not empty.

Theorem 2.2 states that if CPT is strongly consistent and x_0 is a solution, then we may take $\lambda_0 = 1$. Hence, if the functions f and g are differentiable, we may take $\lambda_0 = 1$ in Remark 2.3. Theorem IV.5.3 states that, for differentiable convex problems, if the KKT conditions hold at a feasible point x_0 , then x_0 minimizes. Thus, we have the following theorem.

Theorem 2.8. Let f and g be convex and differentiable and let h be affine. Let CPT be strongly consistent. A necessary and sufficient condition that a feasible point x_0 be a solution to CPT is that there exist multipliers $\lambda_0 \in \mathbb{R}^r$, $\lambda_0 \geq 0$, and $\mu_0 \in \mathbb{R}^m$ such that

$$\begin{aligned} \nabla f(x_0) + \lambda_0' \nabla g(x_0) + \mu_0' h &= 0, \\ \langle \lambda_0, g(x_0) \rangle &= 0. \end{aligned}$$

In the next theorem we shall show that, with no restrictions placed on f and g , the saddle point condition (7) is a sufficient condition that x_0 be a solution to the programming problem P: Minimize f over the set X , where $X = \{x: x \in X_0, g(x) \leq 0\}$.

Note that since we place no conditions on the constraints, equality constraints can be replaced by inequality constraints. We assume that this has been done, and so $h = 0$.

THEOREM 2.6. Let X_0 be a set in $\mathbb{R}^n$, let f be a real-valued function defined on X_0 , and let g be a function from X_0 to $\mathbb{R}^m$. Let x_0 be a point in X_0 and let $\lambda_0 \geq \theta$ be a vector in $\mathbb{R}^m$ such that the saddle point condition (7) holds for all x in X_0 and all $\lambda \in \theta$ in $\mathbb{R}^m$. Then x_0 is a solution to problem P.

Proof. We first show that x_0 is feasible. The leftmost inequality in (7) is

$$f(x_0) + \langle \lambda, g(x_0) \rangle \leq f(x_0) + \langle \lambda_0, g(x_0) \rangle, \quad \lambda \geq \theta.$$

Hence

$$\langle \lambda - \lambda_0, g(x_0) \rangle \leq 0. \quad (17)$$

Now fix a value j in the set $\{1, \dots, m\}$. Let

$$\lambda_j = (d_{j1}, \dots, d_{j,j-1}, d_{j,j} + 1, d_{j,j+1}, \dots, d_{jm}).$$

If for each $j = 1, \dots, m$ we take $\lambda = \lambda_j$ in (17), we get

$$g_j(x_0) \leq 0, \quad j = 1, \dots, m.$$

Thus x_0 is feasible.

We next show that (6) holds. If we take $\lambda = \theta$ in (17), we get that $\langle \lambda_0, g(x_0) \rangle \geq 0$. On the other hand, $g(x_0) \leq \theta$ and $\lambda_0 \geq \theta$ imply that $\langle \lambda_0, g(x_0) \rangle \leq 0$. Hence $\langle \lambda_0, g(x_0) \rangle = 0$, which is (6).

From the rightmost inequality in (7) and from (6) we get that

$$f(x_0) \leq f(x) + \langle \lambda_0, g(x) \rangle, \quad x \in X_0.$$

For x feasible, $g(x) \leq \theta$. Hence for x feasible

$$f(x_0) \leq f(x),$$

and so $f(x_0) = \min\{f(x) : x \in X\}$.

REMARK 2.8. Note that the proof of Theorem 2.6 only involves very elementary arguments.

The next theorem emphasizes some of the points implicit in the material presented in this section.

THEOREM 2.7. A necessary and sufficient condition for a point (x_0, λ_0, μ_0) to be a saddle point for the Lagrangian L defined in (8) on the domain

$$D = \{x, \lambda, g : x \in \mathbb{R}^n, \lambda \in \mathbb{R}^m, \lambda \geq \theta, g \in \mathbb{R}^m\}$$

is that

$$\begin{aligned} (i) \quad & L(x_{ij}, \lambda_{ij}, \mu_{ij}) = \min\{L(x, \lambda_{ij}, \mu_{ij}) : x \in R^n\} \\ (ii) \quad & g(x_{ij}) \leq 0, \quad h(x_{ij}) = 0 \\ (iii) \quad & \langle \lambda_{ij}, g(x_{ij}) \rangle = 0. \end{aligned} \quad (18)$$

Moreover, if (18i) and (18ii) hold, then

$$f(x_{ij}) = L(x_{ij}, \lambda_{ij}, \mu_{ij}). \quad (19)$$

Proof. It follows from the definition of L , that if (18i) and (18ii) hold, then (19) holds.

Let $(x_{ij}, \lambda_{ij}, \mu_{ij})$ be a saddle point for L . Then (18) follows from the definition of saddle point. The inequality in (18i) and the equation (18ii) were established in the course of proving Theorem 2.6. The equality in (18i) also follows from Theorem 2.6, since there are converted the equality constraints to two inequality constraints, $h(x) \leq 0$ and $-h(x) \leq 0$, and then $h(x_{ij}) \leq 0$ and $-h(x_{ij}) \leq 0$. Now suppose that (18) holds. Condition (18ii) states that x_{ij} is feasible. Conditions (18i) and (18) imply that (5) holds. It now follows from Remark 2.4 that $(x_{ij}, \lambda_{ij}, \mu_{ij})$ is a saddle point for $L(x, \lambda, \mu)$.

5. PERTURBATION THEORY

In problems arising in applications the values of various parameters are usually not known precisely but are known to lie in certain ranges of values. Thus one might assign a value to the parameters and then study the effect of perturbing the parameters. We shall assume that this is modeled by perturbing the constraints and shall study the effect of perturbations on the value of the infimum of the problem. Given the uncertainty in the parameter values, one would hope for a certain stability of the problem in the sense that small perturbations in the constraints will not lead to large changes in the value of the infimum. We shall show that if the unperturbed convex programming problem is strongly convex and has a solution, then the change in the value of the infimum is bounded below by a linear function of the perturbations. In the preceding section we showed that under reasonable assumptions the set $\mathcal{M}$ of multipliers and the set $\mathcal{H}$ of Kuhn-Tucker vectors are nonempty and equal. In this section we shall show that the vectors in $\mathcal{H}$ and $\mathcal{M}$ are the coefficient vectors of the aforementioned linear function of the perturbations. Moreover, the sets $\mathcal{H}$ and $\mathcal{M}$ are the negative subdifferential of the infimum as a function of the perturbation, evaluated at the unperturbed value. Thus, in a sense, the multipliers measure the rate of change of the value of the infimum as we perturb the constraints. We shall also give an economic interpretation of the multipliers. These statements will be made precise in this section.

Throughout this section we assume that the problem CPH is consistent, that is, it has a nonempty feasible set. For each $x = (x_0, x_1)$ with $x_0 \in \mathbb{R}^n$ and $x_1 \in \mathbb{R}^m$ we define a perturbed programming problem CPH(x) as follows.

Problem CPH(x)

Minimize $f(x)$ subject to x in $X(x)$, where

$$X(x) = \{x : g(x) \leq x_1, h(x) = x_0\},$$

the function f is a convex function from $\mathbb{R}^n$ to $\mathbb{R}$, the function g is a convex function from $\mathbb{R}^n$ to $\mathbb{R}^m$, and the function h is an affine function from $\mathbb{R}^n$ to $\mathbb{R}^m$.

The constraint set $X(x)$ can also be described by

$$X(x) = \{x : g(x) - x_1 \leq 0, h(x) - x_0 = 0\}.$$

For each fixed x the function $g(x) - x_1$ is convex and the function $h(x) - x_0$ is affine. Thus, each problem CPH(x) is a convex programming problem. Note that the problem CPH is the problem CPH(θ).

For each x in $\mathbb{R}^{n+m}$ the feasible set $X(x)$ is either empty or convex. For each x such that $X(x) \neq \emptyset$, let

$$v(x) = \inf\{f(x) : x \in X(x)\}.$$

Thus we have a function v whose domain is the set of x such that $X(x)$ is not empty and which can assume the value $-\infty$. We shall designate the domain of v by $\text{dom}(v)$ and shall call v the value function of the family of perturbed problems CPH(x), or simply the value function. For x in $\text{dom}(v)$ we shall call $v(x)$ the value of CPH(x). We shall denote the infimum of the unperturbed problem CPH by v . Thus

$$v = \inf\{f(x) : g(x) \leq 0, h(x) = 0\},$$

$$\text{or } v = \inf \emptyset.$$

Since $v(x) = -\infty$ is possible, we shall adjoin $-\infty$ to the real $\mathbb{R}$. Although we shall not need to do so in this section, we shall also adjoin $+\infty$ to $\mathbb{R}$. To work with $-\infty$ and $+\infty$ in a consistent way, we define the following operations:

$$\begin{aligned} (+\infty) + (+\infty) &= +\infty, \\ (-\infty) + (-\infty) &= -\infty, \\ (-\infty) + x &= -\infty, & x \text{ real}, \\ x + (-\infty) &= -\infty, & x \text{ real } x > 0, \\ x + (-\infty) &= +\infty, & x \text{ real } x < 0, \\ -\infty < x < +\infty, & x \text{ real}. \end{aligned}$$

The operation $(-\infty) + (+\infty)$ is not defined.

For functions f defined on a convex set C and taking on either real values or the value of $-\infty$, the definition of convexity is unchanged. Namely, f is said to be convex on C if for every x_0, x_1 in C and every (α, β) in P_1

$$\beta f(x_1) + \beta f(x_2) \leq \alpha f(x_0) + \beta f(x_1),$$

where the values of $-\infty$ are to be operated on as in (1). As in Chapter 1, an equivalent definition of f convex is that the set

$$\operatorname{epi}(f) = \{(x, y) : x \in C, y \in \mathbb{R}, y \geq f(x)\}$$

is convex.

Lemma 3.1. *The domain of the value function ω is a convex set and ω is a convex function on $\operatorname{dom}(\omega)$. If the problem CFI is strongly consistent, then Φ is an interior point of $\operatorname{dom}(\omega)$.*

Proof. We first show that $\operatorname{dom}(\omega)$ is a convex set. We shall suppose that the affine equality constraints have been written as pairs of inequality constraints. Since we assume that CFI is consistent, Φ is in $\operatorname{dom}(\omega)$ and thus $\operatorname{dom}(\omega) \neq \emptyset$. We must show that if w_1 and w_2 are two points in $\operatorname{dom}(\omega)$, then so is $\alpha w_1 + \beta w_2$ for all (α, β) in P_1 . Since w_1 and w_2 are in $\operatorname{dom}(\omega)$, there exist points x_1 and x_2 in P' such that

$$g(x_1) \leq w_1 \quad \text{and} \quad g(x_2) \leq w_2.$$

Since g is convex, for any (α, β) in P_1

$$\alpha g(x_1) + \beta g(x_2) \leq \alpha g(x_1) + \beta g(x_2) \leq \alpha w_1 + \beta w_2.$$

Thus, $\alpha w_1 + \beta w_2$ is in $N(w_1 + \beta w_2)$ and $\alpha w_1 + \beta w_2$ is in $\operatorname{dom}(\omega)$.

To show that ω is convex, again let w_1 and w_2 be two points in $\operatorname{dom}(\omega)$ and let (α, β) be in P_1 . Let $x_1 \in N(w_1)$ and let $x_2 \in N(w_2)$. Then as in the preceding paragraph $\alpha x_1 + \beta x_2$ is in $N(\alpha w_1 + \beta w_2)$. Hence

$$\begin{aligned} \omega(\alpha w_1 + \beta w_2) &= \inf\{f(x) : g(x) \leq \alpha w_1 + \beta w_2\} \\ &\leq \inf\{f(x) : x \in \alpha x_1 + \beta x_2, g(x) \leq w_1, g(x) \leq w_2\} \\ &= \inf\{f(x) : x_1 \in N(w_1), x_2 \in N(w_2)\} \\ &\leq \inf\{\alpha f(x_1) + \beta f(x_2) : x_1 \in N(w_1), x_2 \in N(w_2)\}. \end{aligned} \quad (1)$$

If $\alpha(w_1)$ and $\alpha(w_2)$ are both finite, then by Exercise 1.6.4 the rightmost expression in (1) is equal to

$$\inf\{\alpha f(x_1) : x_1 \in N(w_1)\} + \inf\{\beta f(x_2) : x_2 \in N(w_2)\} = \alpha\omega(w_1) + \beta\omega(w_2).$$

If either $\alpha(\mathbf{x}_1) = -\infty$ or $\alpha(\mathbf{x}_2) = -\infty$, then the rightmost expression in (1) is equal to $-\infty$. Thus,

$$\alpha(\mathbf{x}_1 + \beta \mathbf{x}_2) \leq \alpha(\mathbf{x}_1) + \beta \alpha(\mathbf{x}_2), \quad (2)$$

and the convexity of α is established.

We now suppose that the problem CPH is strongly consistent and shall show that $\mathbf{0}$ is an interior point of $\mathcal{A}$. We shall not incorporate the affine equality constraints into the inequality constraints. Recall that we are writing the vector $\mathbf{z}$ which is the perturbation of the constraints as $\mathbf{z} = (\mathbf{z}_1, \mathbf{z}_2)$, where $\mathbf{z}_1 \in \mathbb{R}^p$ is the perturbation of the inequality constraints and $\mathbf{z}_2 \in \mathbb{R}^q$ is the perturbation of the equality constraints.

Since CPH is strongly consistent, there exists an α_0 such that

$$g(\mathbf{x}_0) < 0 \quad \text{and} \quad \mathbf{h}(\mathbf{x}_0) = \mathbf{0}. \quad (3)$$

Let $\mathbf{x}_0 = (x_{01}, \dots, x_{0n})$. For each $i = 1, \dots, m$, the function G_i defined on $\mathbb{R}^n \times \mathbb{R}$ by

$$G_i(\mathbf{x}, x_{0i}) = g_i(\mathbf{x}) - x_{0i}$$

is continuous. This holds because G_i is convex and a convex function is continuous on the interior of its domain of definition, which in this case is $\mathbb{R}^n \times \mathbb{R}$. From (3) we get that $G_i(\mathbf{x}_0, 0) < 0$. Hence, by the continuity of G_i , there exist $\alpha_i > 0$ and $\delta_i > 0$ such that whenever $|x_{0i}| < \delta_i$ and $\|\mathbf{x} - \mathbf{x}_0\| < \alpha_i$

$$g_i(\mathbf{x}) - x_{0i} < 0.$$

Let $\alpha_0 = \min\{\alpha_i : i = 1, \dots, m\}$ and let $\delta_0 = \min\{\delta_i : i = 1, \dots, m\}$. Then

$$g(\mathbf{x}) < \alpha_0 \quad (4)$$

when $|x_{0i}| < \delta_0$ and $\|\mathbf{x} - \mathbf{x}_0\| < \alpha_0$.

We now consider the perturbed equality constraints, which in accordance with (3) we write as

$$\mathbf{Ax} = \mathbf{b} + \mathbf{z}_2, \quad (5)$$

Recall that in Remark 1.1 we assumed that A has rank k . We assume that the square matrix A_1 consisting of the first k columns of A is nonsingular. There is no loss of generality in this assumption since we can relabel the components to achieve this. Let $A_{2,k}$ be the matrix consisting of the last $n - k$ columns of A . Let $\mathbf{x} = (x_1, x_{k+1})$, where x_1 is the vector consisting of the first k components of $\mathbf{x}$ and x_{k+1} is the vector consisting of the last $n - k$ components of $\mathbf{x}$. Then

we may write (8) as

$$A_0x_0 + A_{0-1}x_{0-1} = b + r_0.$$

From this equation and the equation $A_0x_0 = b$ we get that

$$(x_0 - x_{0-1}) = A_0^{-1}[-d_{0-1}(x_{0-1} - x_{0-1-1}) + r_1], \quad (9)$$

where $x_{0-1} = (x_{0-1,1}, \dots, x_{0-1,p_0})$ and $x_{0-1-1} = (x_{0-1-1,1}, \dots, x_{0-1-1,p_0})$. It follows from (6) that there exist a $\beta > 0$ and a $0 < \rho < \rho_0$ such that if $\|x_{0-1} - x_{0-1-1}\| < \frac{1}{2}\rho$ and $\|r_1\| < \beta$, then $\|x_0 - x_{0-1}\| < \frac{1}{2}\rho_0$. Since for any x we have $\|x\|^2 = \|x_0\|^2 + \|x_{0-1}\|^2$, it follows that for every vector x_{0-1} satisfying $\|x_{0-1} - x_{0-1-1}\| < \frac{1}{2}\rho$ and for every vector x_1 satisfying $\|x_1\| < \beta$ the vector $x = (x_0, x_{0-1})$, where x_0 is given by (9), is such that $\|x - x_0\| < \rho_0/\sqrt{2}$ and (5) is satisfied. The last statement together with (4) shows that for every $x = (x_0, x_0)$ with $\|x\| < \min(A_0, \beta)$ there exists a vector x_0 that is feasible. Thus Φ is an interior point of $\text{dom}(\phi)$.

Lemma 3.1 guarantees that $\text{dom}(\phi)$ is a convex set and ϕ is a convex function. It also guarantees that if CFI is strongly consistent, $\text{dom}(\phi)$ has nonempty interior. There is no guarantee, however, that in general $\text{dom}(\phi)$ will have a nonempty interior. To state two important corollaries to Lemma 3.1, it will therefore be necessary to consider the notions of relative interior and relative boundary, which were introduced in Section 7 of Chapter II.

COROLLARY 1. If $\phi(x_0) = -\infty$ at some point x_0 in $\text{dom}(\phi)$, then $\phi(x) = -\infty$ at all points of $\text{dom}(\phi)$ with the possible exception of relative boundary points.

Proof. It follows from (7) that if $\phi(x_0) = -\infty$ at some point x_0 in $\text{dom}(\phi)$, then $\phi(x) = -\infty$ at all points of the half open line segment $[x_0, x_0]$, where $x_0 \neq x_0$ is another point of $\text{dom}(\phi)$.

Let x_0 be a point of $\text{dom}(\phi)$ that is not a relative boundary point. Then x_0 is a relative interior point of $\text{dom}(\phi)$. Hence the line segment $[x_0, x_0]$, which lies in $\text{dom}(\phi)$, can be extended to a line segment $[x_0, x_0]$ that lies in $\text{dom}(\phi)$. The point x_0 is interior to $[x_0, x_0]$, and so $\phi(x_0) = -\infty$.

COROLLARY 2. If ϕ is finite at a relative interior point of its domain, then ϕ is finite over its entire domain.

This follows from Corollary 1, for if ϕ were not finite at some point, it would have to be infinite at all relative interior points of its domain.

Remark 3.1. The only property of ϕ that was used in the proof of Corollary 1 was the convexity of ϕ . Thus Corollaries 1 and 2 are valid for all convex functions admitting $-\infty$ as a value.

Lemma 3.2. *Let problem CPH be strongly consistent and let $v = \alpha(\theta)$ be finite. Then α is finite on $\text{dom}(\alpha)$, the subdifferential $\text{sub}(\theta) \neq \emptyset$ and each η in $-\text{sub}(\theta)$ can be written as $\eta = \bar{\lambda} \cdot \bar{p}$, where $\bar{\lambda} \in \mathbb{R}^n$, $\bar{\lambda} \geq 0$, $\bar{p} \in \mathbb{R}^n$ and such that for all $x = (x_1, x_2)$ in $\text{dom}(\alpha)$*

$$-\alpha(x) \leq -\alpha(\theta) = \langle \eta, x \rangle = \alpha(\theta) = \langle \bar{\lambda}, x_2 \rangle = \langle \bar{p}, x_2 \rangle. \quad (7)$$

Proof. It follows from Lemma 3.1 that θ is an interior point of $\text{dom}(\alpha)$ and that α is a convex function on the convex set $\text{dom}(\alpha)$. Since $\alpha(\theta)$ is finite, it follows from Corollary 2 to Lemma 3.1 that α is finite at all points of $\text{dom}(\alpha)$. Therefore, by Theorem III.2.1, the set of subgradients of α at θ is not empty. Relation (7) then follows by taking $\eta = -\bar{\xi}$, where $\bar{\xi}$ is any element in $\text{sub}(\theta)$.

We now show that $\bar{\lambda} \geq 0$. Since the problem is consistent, $\{x: g(x) \leq 0, h(x) = 0\}$ is nonempty. Hence the set $X(\lambda)$ is nonempty for any vector $x = (x_1, x_2)$ with $x_1 \geq 0$ and $x_2 = 0$. Thus, all such vectors are in $\text{dom}(\alpha)$.

If $\bar{\lambda}$ were not nonnegative, then at least one component of $\bar{\lambda}$, say $\bar{\lambda}_i$, would be negative. Let $\bar{x}_i = (x_{1i}, x_{2i})$ where $x_{1i} = \bar{x}_i$ the vector in $\mathbb{R}^n$ all of whose components except the i th are zero and whose i th component is 1. Then $\bar{x}_i$ is in $\text{dom}(\alpha)$ and

$$-\alpha(\bar{x}_i) \geq -\alpha(\theta) = \langle \bar{\lambda}, \bar{x}_i \rangle = \langle \bar{p}, \bar{x}_i \rangle = \alpha(\theta) = \bar{\lambda}_i = -\alpha(\bar{x}_i), \quad (8)$$

the last inequality holding because $\bar{\lambda}_i < 0$.

Since $\bar{x}_{1i} \geq 0$ and $\bar{x}_{2i} = 0$, we have $X(\bar{\lambda}) \subseteq X(\bar{x}_i)$. Therefore,

$$-\alpha(\bar{x}_i) = \inf\{\bar{\lambda}(x): x \in X(\bar{x}_i)\} \leq \inf\{\bar{\lambda}(x): x \in X(\bar{\theta})\} = -\alpha(\theta),$$

which contradicts (8). Thus $\bar{\lambda} \geq 0$.

Remark 3.2. The vector $-\eta$ is any vector in $\text{sub}(\theta)$ and we need not be unique. By Theorem III.3.1, the vector $-\eta$ is unique if and only if α is differentiable at θ , in which case $-\eta = \nabla \alpha(\theta)$. Thus, $-\eta$ gives the direction of maximum rate of increase at θ of the value function α .

If α is not differentiable, then by Exercise III.2.3,

$$\alpha(\theta, v) \geq \langle -\eta, v \rangle \quad \text{for all } v \text{ in } \mathbb{R}^n,$$

where $\alpha(\theta, v)$ is the directional derivative of α at θ in the direction v . Thus, $-\eta$ again in a sense measures the changes in α as we move away from the origin. The vector η is therefore called a *sensitivity vector*.

We noted in Remark 2.5 that if we know the multipliers $(\bar{\lambda}, \bar{\mu})$ of Theorem 2.2, then we may replace the constrained problem by an unconstrained problem. In other words, the multipliers $(\bar{\lambda}, \bar{\mu})$ of Theorem 2.2 are Kuhn-Tucker vectors. By Corollary 1 to Theorem 2.5, if CPH has a solution, then the set X^* of Kuhn-Tucker vectors is equal to the set $\mathcal{M}$ of multiplier vectors.

In the next lemma we shall show that if $\omega(\mathbf{B})$ is finite, then every vector in $-\text{lin}(\mathbf{B})$ is a Kuhn-Tucker vector.

Lemma 3.1. Let

$$\omega(\mathbf{B}) = \inf\{f(x) : x \in X\},$$

where $X = \{x : g(x) \leq \lambda, h(x) = 0\}$, be finite and let $\omega(\mathbf{B}) \neq \emptyset$. Then for each $q = (\lambda, \mu)$ such that $-q \in \text{lin}(\mathbf{B})$ the vector λ is nonnegative and

$$\omega(\mathbf{B}) = \inf\{f(x) + \langle \lambda, g(x) \rangle + \langle \mu, h(x) \rangle : x \in \mathbb{R}^n\}.$$

Proof. Since $-q \in \text{lin}(\mathbf{B})$, for each x in $\text{dom}(g)$

$$\omega(x) \geq \omega(\mathbf{B}) + \langle q, x \rangle. \quad (9)$$

From the proof of Lemma 3.2 we get that $\lambda \geq 0$. For all x in $\mathbb{R}^n$, $g(x) \leq g(x)$ and $h(x) = h(x)$. Therefore $x = (g(x), h(x))$ is in $\text{dom}(q)$. Taking $x = (g(x), h(x))$ in (9) gives

$$\omega(g(x), h(x)) + \langle \lambda, g(x) \rangle + \langle \mu, h(x) \rangle \geq \omega(\mathbf{B}) \quad (10)$$

for all x in $\mathbb{R}^n$. Now

$$\omega(g(x), h(x)) = \inf\{f(y) : y \in N(g(x), h(x))\},$$

where $N(g(x), h(x)) = \{y : g(y) \leq g(x), h(y) = h(x)\}$. But $x \in N(g(x), h(x))$. Hence

$$f(x) \geq \omega(g(x), h(x)).$$

Substituting this into (10) gives

$$f(x) + \langle \lambda, g(x) \rangle + \langle \mu, h(x) \rangle \geq \omega(\mathbf{B})$$

for all x . Hence

$$\omega(\mathbf{B}) \leq \inf\{f(x) + \langle \lambda, g(x) \rangle + \langle \mu, h(x) \rangle\}.$$

We now establish the reverse inequality and thus prove the lemma. We have

$$\begin{aligned} & \inf\{f(x) + \langle \lambda, g(x) \rangle + \langle \mu, h(x) \rangle : x \in \mathbb{R}^n\} \\ & \leq \inf\{f(x) + \langle \lambda, g(x) \rangle + \langle \mu, h(x) \rangle : g(x) \leq \lambda, h(x) = 0\}. \end{aligned}$$

Since $\lambda \geq 0$, $g(x) \leq 0$, and $h(x) = 0$, the right-hand side of this inequality is less than or equal to

$$\inf_{x \in N} \{g(x) \leq 0, h(x) = 0\} = \omega(0).$$

Remark 3.2. Because of relationships (7) and (9), which bound the change in ω from below by a linear function, problems with $f(x) \neq 0$ are said to be *stable*.

The next theorem is an immediate consequence of Lemmas 3.2 and 3.3.

Theorem 3.1. Let the programming problem CPM be strongly consistent and let $\omega(0)$ be finite. Then $\text{dom}(0)$ is nonempty and for every $q = (q_1, q_2)$ such that $-q \in \text{dom}(0)$ the vector q is nonnegative and

$$\omega(0) = \inf_{x \in N} \{g(x) + (q_1, q_2)x : x \in N\}.$$

The next theorem summarizes the relationships among the set $\text{dom}(0)$, the multiplier vectors, and the Kuhn-Tucker vectors.

Theorem 3.2. Let the programming problem CPM be strongly consistent and let it have a solution. Then the set M of multiplier vectors, the set K of Kuhn-Tucker vectors, and the set $-\text{dom}(0)$ are all nonempty and equal.

Proof. That M and K are nonempty and equal was noted in Corollary 1 of Theorem 2.3 and Remark 2.6. Theorem 3.1 states that if CPM is strongly consistent and $\omega(0)$ is finite, then $\text{dom}(0)$ is nonempty and every vector in $-\text{dom}(0)$ is a Kuhn-Tucker vector.

To complete the proof of the theorem, we must show that each Kuhn-Tucker vector is an element of $-\text{dom}(0)$. Let $q_0 = (q_0, q_1)$ be a Kuhn-Tucker vector. Since CPM is consistent, $0 \in \text{dom}(0)$ and hence $\text{dom}(0) \neq \emptyset$. For any x in $\text{dom}(0)$

$$\begin{aligned} \omega(0) &= \inf_{x \in N} \{g(x) + (q_0, q_1)x\} = (q_0, q_1)x \leq \inf_{x \in N} \{f(x) \\ &\quad + (q_0, q_1)x\} = (q_0, q_1)x \leq \omega(x). \end{aligned}$$

Since $g(x) \leq 0$ and $h(x) = 0$ for x in N and since $\lambda_0 \geq 0$, we get that

$$\begin{aligned} \omega(0) &\leq \inf_{x \in N} \{f(x) + (q_0, q_1)x\} = (q_0, q_1)x \leq \omega(x) \\ &= \omega(0) + (q_0, q_1)x = (q_0, q_1)x. \end{aligned}$$

Hence

$$\omega(x) - \omega(0) = (q_0, q_1)x = (q_0, q_1)x$$

for all x in $\text{dom}(0)$. Thus $q_0 \in -\text{dom}(0)$.

The assumption of strong consistency in Theorem 3.2, as already noted, guarantees that the sets $\mathcal{W}$ and $-\text{int}(\mathcal{W})$ are nonempty. In any specific problem, the validity of the strong consistency assumption is easily verified a priori from the data of the problem. In the next theorem we show that if we drop the assumption of strong consistency but assume that one of the sets $\mathcal{X}$, $\mathcal{W}$, or $-\text{int}(\mathcal{W})$ is nonempty, then all of the sets are nonempty and are equal. This assumption, however, is not verifiable a priori.

THEOREM 3.3. *Let the programming problem (P1) have a solution. Let one of the sets $-\text{int}(\mathcal{W})$, the set $\mathcal{X}$ of Kuhn-Tucker multipliers, or the set $\mathcal{W}$ of multipliers as in Theorem 3.2 be nonempty. Then all of the sets $\mathcal{W}$, $\mathcal{X}$, and $-\text{int}(\mathcal{W})$ are nonempty and*

$$\mathcal{W} = \mathcal{X} = -\text{int}(\mathcal{W}).$$

Proof. By Corollary 1 to Theorem 3.2, the sets $\mathcal{W}$ and $\mathcal{X}$ are equal. By Lemma 3.1, if $\text{int}(\mathcal{W})$ is not empty, then $-\text{int}(\mathcal{W}) \subseteq \mathcal{X}$. Of course, if $\text{int}(\mathcal{W}) = \emptyset$, this is also true. In proving Theorem 3.2, we showed that $\mathcal{X} \subseteq -\text{int}(\mathcal{W})$. Thus, $\mathcal{X} = -\text{int}(\mathcal{W})$, and the theorem is proved. $\square$

Remark 3.4. The relationship between perturbations and multiplier vectors has an interesting economic interpretation. We can consider the function f as giving the cost of an activity described by a vector x , where x is subject to constraints $g(x) \leq b$ and $h(x) = 0$. Our objective is to minimize f subject to the constraints. We perturb the constraints by a vector $a = (a_1, a_2)$ in the hope of reducing the optimal cost $\omega(b)$. If we must pay a price $\lambda = (\lambda_1, \lambda_2)$ per unit change in the constraints, then the total cost of the perturbed problem with constraints $g(x) \leq a_1$ and $h(x) = a_2$ is

$$\omega(a) + \langle \lambda, a \rangle.$$

Thus it will be advantageous to purchase a change in the constraints if and only if

$$\omega(a) + \langle \lambda, a \rangle < \omega(b).$$

But

$$\omega(a) \geq \omega(b) - \langle \lambda, a_1 \rangle - \langle \mu, a_2 \rangle, \quad (\lambda, \mu) \in \mathcal{W}, \quad (a_1, a_2) \in \mathbb{R}^{m+p},$$

and so

$$\omega(a) + \langle \lambda, a_1 \rangle + \langle \mu, a_2 \rangle \geq \omega(b).$$

Also, any (λ, μ) for which the last inequality holds is in $-\text{int}(\text{epi} \mathcal{W})$, and so is in

18. Thus we can consider a multiplier vector (λ, p) to be a price such that the purchase of a perturbation at that price will not result in a lowering of the cost. Since at these prices no perturbation is worth buying, We may consider the cost to be at equilibrium at these prices. Therefore, such prices are called equilibrium prices.

We now illustrate the results of this section by several examples.

Example 1.1. Minimize

$$f(x) = e^{-2x_1 + x_2}$$

subject to

$$x_1^2 + x_2^2 - 20 \leq 0.$$

This problem is a perturbation of the problem in Example IV.5.5. Since f and the constraint function $g(x) = x_1^2 + x_2^2 - 20$ are convex, the unperturbed problem is a convex programming problem, as are the perturbed problems. The problem is strongly consistent since $g(0) < 0$.

Let

$$\hat{L}(x, \lambda, \eta) = e^{-2x_1 + x_2} + \lambda(g(x) + \eta) + \eta(x_1^2 + x_2^2 - 20 - \eta).$$

It follows from Theorem 2.5 that a solution $\hat{x} = (\hat{x}_1, \hat{x}_2)$ is characterized by the existence of a positive number $\hat{\lambda}$ such that $\hat{\lambda}(g(\hat{x}) - \eta) = 0$ and $\partial \hat{L}(\hat{x}, \hat{\lambda}, \eta) = 0$, where the partial derivatives are evaluated at $(\hat{x}, \hat{\lambda})$. Thus

$$-e^{-2\hat{x}_1 + \hat{x}_2} + \hat{\lambda}\hat{x}_1 = 0, \quad -e^{-2\hat{x}_1 + \hat{x}_2} + \hat{\lambda}\hat{x}_2 = 0. \quad (11)$$

From this it follows that $\hat{\lambda} \neq 0$ and that

$$\hat{\lambda}\hat{x}_1 = \hat{\lambda}\hat{x}_2.$$

Hence $\hat{x}_1 = \hat{x}_2$. From $\hat{\lambda} \neq 0$ and $\hat{\lambda}(g(\hat{x}) - \eta) = 0$ it follows that

$$\hat{x}_1^2 + \hat{x}_2^2 = 20 + \eta.$$

Let $\hat{\xi}$ denote the common value of $\hat{\xi}_1$ and $\hat{\xi}_2$. Then

$$2\hat{\xi}^2 = 20 + \eta. \quad (12)$$

and so $\hat{\xi} = \sqrt{10 + \frac{1}{2}\eta}$. It follows from (11) and (12) that

$$\hat{\lambda} = e^{-2\hat{\xi}} = (10 + \frac{1}{2}\eta)^{-2}. \quad (13)$$

It also follows from (12) and from the relation $f(\bar{x}) = z^*$ that

$$u(x) = z^{-1/2} = (1 + \frac{1}{2}x)^{-1/2}.$$

Hence

$$u'(x) = -(1/2)(1 + \frac{1}{2}x)^{-3/2}.$$

From (13) we get that the value of the multiplier λ_0 for the unperturbed problem is $10^{-1/2}$. Also, $u'(0) = -10^{-1/2}$ and $u - u'(0)x = -\lambda_0x$, in agreement with the theory.

We now show that the solution of the constrained problem can be obtained by solving the unconstrained problem of minimizing $F(x)$, where

$$F(x) = z^{-1/2}(1+x) + \lambda_0 x^2 + x^2 - 20x.$$

Substituting the value of λ_0 in the preceding expression gives

$$F(x) = z^{-1/2}(1+x) + 10^{-1/2}x^2 + x^2 - 20x.$$

Since F is convex, any critical point furnishes a minimum. Setting DF/dx , and DF/dx_0 equal to zero gives

$$z^{-1/2}(1+x) + 2z^{-1/2}x = 0, \quad -z^{-1/2}(1+x) + (1+2x)z = 0.$$

Hence $x_1 = x_0$. Solving $x = x_1 = x_0$ gives

$$x^* = 10^2 z^{-1/2},$$

whence $x = 10$. This is in agreement with the solution of the constrained problem given by (12) with $z = 0$.

Example 5.3. Find the point in the half space defined by $x_1 + x_2 \leq 1$ that is closest to the origin. This problem has a solution. (Why?)

We may formulate this problem analytically as follows: Minimize $\|x\|^2$ subject to $x_1 + x_2 \leq 1$. Thus the problem is a problem of type CPH(x). If we take $z = 0$, we see that CPH is strongly consistent. On sketching the feasible sets for $z > 0$, it becomes clear that for $z > 0$ the solution is the origin and that $\alpha(z) = 0$. For $z = 0$ the minimum is the point of intersection of the line $x_1 = x_2$ and the line $x_1 + x_2 = 1$. Thus the minimum is achieved at $(\frac{1}{2}, \frac{1}{2})$ and $\alpha(z) = \frac{1}{2}z^2$ for $z < 0$. The function α is differentiable at all points, and $u'(0) = 0$. Therefore 0 is the unique Kuhn-Tucker vector, and the constrained problem CPH can be replaced by the problem of minimizing $\|x\|^2$. The conditions of Theorem 2.5 are clearly satisfied if we take $x_0 = (0, 0)$ and $\lambda_0 = 0$.

We may also formulate the problem analytically as follows: Minimize $\phi(x)$ subject to $x_1 + x_2 \leq \alpha$. Again, the solution is the origin if $\alpha \geq 0$ and the point $(\frac{1}{2}\alpha, \frac{1}{2}\alpha)$ if $\alpha < 0$. The function ϕ is now given by $\phi(x) = 0$ for $x \geq 0$ and $\phi(x) = -\alpha/\sqrt{2}$ for $x < 0$. Thus, ϕ is not differentiable at the origin. It is a straightforward calculation to show that $\text{dom}(\nabla)$ is the interval $[-1/\sqrt{2}, 0]$. Hence by Theorem 3.2 the Kuhn-Tucker vectors and multiplier vectors are all points in the interval $[0, 1/\sqrt{2}]$. We now verify this assertion independently.

For λ to be a Kuhn-Tucker vector, we must have

$$\text{inf}\{[x] + \lambda(x_1 + x_2)\} = 0.$$

When $\alpha = 0$, the quantity in brackets is clearly zero. Therefore, for λ to be a Kuhn-Tucker vector, we require that $\lambda \geq 0$ and

$$\sqrt{x_1^2 + x_2^2} + \lambda(x_1 + x_2) \geq 0 \quad \text{for all } x \in \mathbb{R}^2.$$

To see what requirement this places on λ , we introduce polar coordinates. The preceding inequality is then satisfied at the origin and whenever

$$\lambda(\cos \theta + \sin \theta) \geq -1, \quad 0 \leq \theta < 2\pi. \quad (14)$$

The function $\mu(\theta) = \cos \theta + \sin \theta$ achieves its minimum on the interval $[0, 2\pi]$ at $\theta = 3\pi/4$ and the value of the minimum is $-\sqrt{2}$. Hence we must have $\lambda \leq 1/\sqrt{2}$ as well as $\lambda \geq 0$.

We can replace the constrained problem by the unconstrained problem

$$\text{min}\{[x] + \lambda(x_1 + x_2)\}$$

where λ is any point in $[0, 1/\sqrt{2}]$. The preferred choice of λ is clearly $\lambda = 0$. We leave it to the reader to verify that any λ in $[0, 1/\sqrt{2}]$ will serve as a multiplier λ_0 as in Theorem 3.2.

Example 3.2 furnishes another illustration of the advantage of minimizing the square of the distance rather than the distance itself.

Example 3.3. Minimize $f(x) = x_1$ subject to

$$g_1(x) = x_1 \leq 1, \quad g_2(x) = x_1^2 - x_2 \leq 0.$$

When $\alpha = 0$, the problem reduces to the problem considered in Example 2.1. We saw there that when $\alpha = 0$, there exists no Kuhn-Tucker vector. By drawing an appropriate figure, the reader can see that $\text{dom}(\phi) = \{x: x_1 + x_2 \geq 0\}$ and that $\phi(x) = -\sqrt{x_1 + x_2}$. We leave the proof of this assertion as an exercise. Thus neither Φ nor any point on the line $x_1 + x_2 = 0$ are interior to $\text{dom}(\phi)$. It is easy to verify that $\text{deg}(\phi) = \frac{1}{2}$ at each x on the boundary of $\text{dom}(\phi)$.

The next example has discontinuous value function at the origin.

Example 3.4. Minimize $x^{*T}x$ subject to $\sqrt{x_1^2 + x_2^2} - x_3 \leq 1$. For $z = 0$, the feasible set is $\{x = (x_1, x_2, x_3) : x_3 \geq 0, x_1 = 0\}$. The minimum is attained at each point of the feasible set since $x^{*T}x = x^{*T}x = 1$. Hence $\inf(z) = 1$. If $z \neq 0$, the feasible set $K(z)$ is empty. If $z > 0$, points x of the feasible set satisfy

$$x_3^2 \leq 2x_3z + z^2.$$

Thus the feasible set is the set of points on and "inside" the parabola $x_3^2 = 2zx_3 + z^2$. This parabola is symmetric with respect to the x_1 -axis, has vertex at $(-\frac{1}{2}z, 0)$, and opens to the right. Hence $\inf\{x^{*T}x : x \in K(z)\} = 0$, and $\inf(z) = 0$ for $z < 0$. Thus, $\text{dom}(v) = [0, \infty)$, and v is discontinuous at $z = 0$. Also, $\partial\inf(z) = \emptyset$. Note that if $z > 0$, then the infimum is never attained at any point of $K(z)$. Thus the problem CPl(z) has no solution if $z > 0$.

Exercise 3.1. Prove the assertions in Example 3.3 that

- (i) $\text{dom}(v) = \{z : x_1 + x_2 \geq 0\}$,
- (ii) $\inf(z) = -\sqrt{z_1 + z_2}$, and
- (iii) $\partial\inf(z) = \emptyset$ at each z on the boundary of $\text{dom}(v)$.

Exercise 3.2. In each of the following problems find v , $\text{dom}(v)$, and $\partial\inf(z)$ if it is not empty:

- (i) Minimize $x^{*T}x$ subject to $-x \leq 0$.
- (ii) Minimize x_1 subject to $x_1^2 + x_2^2 \leq 1$.
- (iii) Minimize $x_1^2 + x_2^2 + 2x_1 + 2x_2$ subject to $x_1 + x_2 \geq 2$.
- (iv) Minimize x_1 subject to $x^{*T}x = x_1 \leq 0$.

Exercise 3.3. Show directly that there are no Kuhn-Tucker vectors for the problem with $z = 0$ in Example 3.4.

4. LAGRANGIAN DUALITY

If CPl has a solution and is strongly consistent, then Theorem 2.4 suggests a method for computing the value of the minimum. Let

$$S_0 = \{\eta : \eta = (\lambda, \mu) (\lambda \in \mathbb{R}^m, \lambda \geq 0, \mu \in \mathbb{R}^p)\}. \quad (1)$$

and for each $\mathbf{q} \in \mathbf{F}_q$ let

$$\theta(\mathbf{q}) = \inf \{ f(\mathbf{x}) + \langle \mathbf{A}_1, \mathbf{g}(\mathbf{x}) \rangle + \langle \mathbf{q}, \mathbf{h}(\mathbf{x}) \rangle : \mathbf{x} \in \mathbb{R}^n \}. \quad (2)$$

Let

$$\delta = \sup \{ \theta(\mathbf{q}) : \mathbf{q} \in \mathbf{F}_q \}. \quad (3)$$

Then by Theorem 2.4, $\delta = v$, where

$$v = \inf \{ f(\mathbf{x}) : \mathbf{q}(\mathbf{x}) \leq \mathbf{0}, \mathbf{h}(\mathbf{x}) = \mathbf{0} \}. \quad (4)$$

A possible procedure for finding $\mathbf{x}_q$, a point at which the minimum is achieved, is to first find all solutions to the equation $f(\mathbf{x}_q) = v$ and then retain those solutions that satisfy the constraints.

This discussion suggests that it might be fruitful to consider the following problem, which we call the problem dual to CFI, or simply the dual problem, and denote by DCFI.

Problem DCFI

Let

$$\mathbf{F} = \{ \mathbf{q} : \mathbf{q} \in \mathbf{F}_q, \theta(\mathbf{q}) > -\infty \}, \quad (5)$$

where $\mathbf{F}_q$ is defined in (1) and $\theta(\mathbf{q})$ in (2). Then

$$\text{Maximize } \{ \theta(\mathbf{q}) : \mathbf{q} \in \mathbf{F} \}.$$

Note that in formulating the dual problem it is not assumed that CFI has a solution or is strongly consistent. The set $\mathbf{F}$ is the feasible set for DCFI. If $\mathbf{F} \neq \emptyset$, then the dual problem is said to be consistent. The original problem CFI is often referred to as the primal problem.

The next theorem indicates how the dual problem gives information about the primal problem and vice-versa.

THEOREM 4.1 (LAGRANGIAN DUALITY THEOREM FOR CONVEX PROGRAMMING).

- (i) If $\mathbf{X} \neq \emptyset$ and $\mathbf{F} \neq \emptyset$, then for each $\mathbf{x} \in \mathbf{X}$ and $\mathbf{q} \in \mathbf{F}$

$$\theta(\mathbf{q}) \leq f(\mathbf{x}). \quad (6)$$

Moreover, if δ and v are as in (3) and (4), then both are finite and $\delta \leq v$.

- (ii) Let $\mathbf{F} \neq \emptyset$. If $\theta(\mathbf{q})$ is unbounded above on $\mathbf{F}$, then $\mathbf{X} = \emptyset$.
 (iii) Let $\mathbf{X} \neq \emptyset$. If $f(\mathbf{x})$ is unbounded below on $\mathbf{X}$, then $\mathbf{F} = \emptyset$.
 (iv) If there exists an $\mathbf{q}_0 \in \mathbf{F}$ and an $\mathbf{x}_0 \in \mathbf{X}$ such that $f(\mathbf{x}_0) = \theta(\mathbf{q}_0)$, then

$$\delta = \theta(\mathbf{q}_0) = f(\mathbf{x}_0) = v.$$

Thus $\mathbf{x}_0$ is a solution of CFI and $\mathbf{q}_0$ is a solution of DCFI.

Proof. The key to the proof of the theorem is the observation that since $\lambda \geq 0$ for $\mathbf{q}_0 \in F$ and since $g(\mathbf{x}) \leq 0$ and $h(\mathbf{x}) = 0$ for $\mathbf{x} \in X$, it follows that for $\mathbf{x} \in X$ and $\mathbf{q}_0 \in F$

$$f(\mathbf{x}) + \langle \mathbf{q}_0, g(\mathbf{x}) \rangle + \langle \mathbf{q}_0, h(\mathbf{x}) \rangle \leq f(\mathbf{x}). \quad (7)$$

In (7) we take the infimum over all $\mathbf{x}$ in $\bar{X}$. This infimum is less than the infimum in (2) taken over all $\mathbf{x} \in X$. From these observations and from (7) the inequalities (i) and (i') follow. Statement (ii) follows from (i), for if X were not empty, then $\#(g)$ would be bounded above. Also, (iii) follows from (ii) by a similar argument. Statement (iv) follows from the chain

$$\#(\mathbf{q}_0) \leq \beta \leq \gamma \leq f(\mathbf{x}_0) = \#(\mathbf{q}_0).$$

We now see how the dual problem gives useful information about the primal problem and vice-versa. Statement (i) shows how a feasible point for the dual problem gives a lower bound for the value of the primal problem and how a feasible point for the primal problem gives an upper bound for the value of the dual problem. Statement (ii) gives a way of checking whether $X = \emptyset$ in the absence of more direct checks. Statement (iii) gives a condition under which the dual problem cannot be formulated.

Definition 4.1. Problems CPE and DCPH are said to exhibit a *duality gap* if $\beta < \gamma$.

The next theorem gives a sufficient condition for the absence of a duality gap and relates the multipliers $\mathbf{q}_0 = (\mathbf{q}_0, \mathbf{p}_0)$ to the solutions of the dual problem.

Theorem 4.2. Let the primal problem CPH be strongly consistent and have a solution. Then:

- (i) There is no duality gap.
- (ii) If $\mathbf{q}_0 = (\mathbf{q}_0, \mathbf{p}_0)$ is a multiplier for CPH, then $\mathbf{q}_0$ is a solution of the dual problem.
- (iii) If $\mathbf{q}_0 = (\mathbf{q}_0, \mathbf{p}_0)$ is a solution of the dual problem, then $\mathbf{q}_0$ is a multiplier vector for the primal problem.

Proof. Let $\mathbf{x}_0$ denote the point at which CPH achieves its minimum. Since the problem is strongly consistent by Theorem 3.2, there exists a multiplier vector $\mathbf{q}_0 = (\mathbf{q}_0, \mathbf{p}_0)$ for CPH. Thus,

$$f(\mathbf{x}_0) \leq f(\mathbf{x}) + \langle \mathbf{q}_0, g(\mathbf{x}) \rangle + \langle \mathbf{q}_0, h(\mathbf{x}) \rangle \quad (8)$$

for all $\mathbf{x}$ in $\bar{X}$. Hence, from (2), $\#(\mathbf{q}_0) \geq f(\mathbf{x}_0)$. By Theorem 4.1, $\#(\mathbf{q}_0) \leq f(\mathbf{x}_0)$, so $\#(\mathbf{q}_0) = f(\mathbf{x}_0)$. Conclusions (i) and (ii) now follow from (iv) of Theorem 4.1.

Now let $\mathbf{q}_0 = (\lambda_0, \mathbf{p}_0)$ be a solution of the dual problem. From (2) and from the absence of a duality gap we get that

$$\inf\{f(\mathbf{x}) + (\lambda_0, \mathbf{p}_0) \cdot \mathbf{b}(\mathbf{x}) : \mathbf{x} \in \mathbb{R}^n\} = \theta(\mathbf{q}_0) = f(\mathbf{x}_{\mathbf{q}_0}).$$

Thus $(\lambda_0, \mathbf{p}_0)$ is a Kuhn-Tucker vector, and so by Theorem 3.3, $(\lambda_0, \mathbf{p}_0)$ is a multiplier vector for the primal problem.

We now present several examples to illustrate the theory developed in this section.

Example 4.1. We take the primal problem to be the problem in Exercise IV.5.3 or the problem in Example 3.1 with $r = 0$. The primal problem is

$$\begin{aligned} \text{Minimize } f(\mathbf{x}) &= e^{-2x_1 + 3x_2} \\ \text{subject to } x_1^2 + x_2^2 &\leq 10. \end{aligned}$$

Let

$$L(\mathbf{x}, \lambda) = e^{-2x_1 + 3x_2} + \lambda(x_1^2 + x_2^2 - 10).$$

Then for each $\lambda \geq 0$ define

$$\theta(\lambda) = \inf\{L(\mathbf{x}, \lambda) : \mathbf{x} \in \mathbb{R}^2\}.$$

The dual problem is to maximize $\theta(\lambda)$.

We saw in Example 3.1 that the primal problem is strongly consistent and that it has a solution at $\mathbf{x}_0 = (\frac{1}{2}, \frac{1}{2})$, where $\frac{1}{2} = \ln 10$. The value of the minimum is 10^{-1} and the value of the multiplier λ_0 is 10^{-1} . Thus we would expect the value of the maximum of θ to be 10^{-1} and for θ to assume this value at $\lambda_0 = 10^{-1}$. We now show that this is indeed so.

We first determine the function θ . For each $\lambda \geq 0$, the function $L(\cdot, \lambda)$ is readily seen to be a strictly convex function of $\mathbf{x}$. Hence a necessary and sufficient condition that $L(\cdot, \lambda)$ attain a unique minimum at $\hat{\mathbf{x}} = (\hat{x}_1, \hat{x}_2)$ is that $\partial L/\partial x_1$ and $\partial L/\partial x_2$ both equal zero at $\hat{\mathbf{x}}$. As in Example 3.1 we see that for $\lambda \neq 0$ this implies that $\hat{x}_1 = \hat{x}_2$. Denoting this common value by $\hat{x}$, we get that

$$\theta(\lambda) = \inf\{\phi(\hat{x}, \lambda) : \hat{x} \in \mathbb{R}\}, \quad (4)$$

where

$$\phi(\hat{x}, \lambda) = e^{-2\hat{x} + 3\hat{x}} + 2\lambda(\hat{x}^2 - 10).$$

For fixed positive λ , $\phi(\cdot, \lambda)$ is a strictly convex function of $\hat{x}$. A necessary and sufficient condition that $\phi(\cdot, \lambda)$ attain a unique minimum at a point $\hat{x}_0$ is that

$\lambda g_j(\lambda)$ vanish at λ_0 . Now

$$\frac{d\theta}{d\lambda} = -2e^{-2\lambda} + 2\lambda e^{\lambda}.$$

Setting $d\theta/d\lambda = 0$, we get that $\exp(-2\lambda) = \lambda^2 e^{\lambda}$. Substituting this into (9) gives

$$\theta(\lambda) = 1\lambda^{1/3} - 2\lambda, \quad \lambda > 0.$$

Since $\theta(\lambda) = 0$ when $\lambda = 0$, the preceding formula is valid for all $\lambda \geq 0$. An elementary calculus argument shows that θ is maximized at the unique point at which $\theta(\lambda) = 0$. This value of λ is readily seen to be 10^{-10} and the corresponding value of θ is 10^{-10} . These values are as required.

Example 4.2. Let the primal problem be as follows: Minimize $\|x\|$ subject to $x_1 + x_2 \leq 0$. This problem was analyzed in Example 3.2 with $\tau = 0$. This problem is strongly consistent and has $x = 0$ as its solution. The value of the minimum is clearly zero. Any λ in the interval $[0, 1/\sqrt{2}]$ is a multiplier vector.

The dual problem is as follows: Maximize $\{\theta(\lambda) : \lambda \geq 0, \theta(\lambda) > -\infty\}$, where $\theta(\lambda) = \inf\{\|x\| + \lambda(x_1 + x_2) : x \in \mathbb{R}^2\}$. We introduce polar coordinates (r, ϕ) and obtain

$$\theta(\lambda) = \inf\{r\} + \lambda(\cos \phi + \sin \phi)r : r \geq 0, 0 \leq \phi < 2\pi\}.$$

Let

$$\gamma(\phi, \lambda) = 1 + \lambda(\cos \phi + \sin \phi).$$

Then $\theta(\lambda) = 0$ if $\gamma(\phi, \lambda) \geq 0$ for all $0 \leq \phi < 2\pi$, and $\theta(\lambda) = -\infty$ if this is not so and $\tau > 0$. The argument used in Example 3.2 shows that $\gamma(\phi, \lambda) \geq 0$ for all ϕ in $[0, 2\pi]$ if and only if $0 \leq \lambda \leq 1/\sqrt{2}$. Thus $\theta(\lambda) = 0$ for λ in $[0, 1/\sqrt{2}]$ and $\theta(\lambda) = -\infty$ otherwise. Hence $\max\{\theta(\lambda) : \lambda \in T\} = 0$ and the maximum is attained at each $\lambda \in [0, 1/\sqrt{2}]$. Comparison with the solution to the primal problem shows that there is no duality gap and the points at which the dual problem obtains its maximum are the multipliers for the primal problem.

Example 4.3. Let the primal problem be the problem treated in Example 2.1 and in Example 3.1 with $\tau = 0$. The primal problem has the origin as the only feasible point. This point furnishes the solution and the value of the minimum is zero. This problem is clearly not strongly consistent.

The dual problem is as follows: Maximize $\{\theta(\lambda) : \lambda \in T\}$, where

$$T = \{\lambda : \lambda \geq 0, \theta(\lambda) > -\infty\} \quad \text{and} \quad \theta(\lambda) = \inf\{x_1 + \lambda_1 x_1 + \lambda_2(x_1^2 - x_2) : x \in \mathbb{R}^2\}.$$

If we write

$$x_1 + \lambda_1 x_2 + \lambda_2 (x_1^2 - x_2) = x_1 + \lambda_1 x_1^2 + (\lambda_1 - \lambda_2) x_2,$$

we see that if $\lambda_1 \neq \lambda_2$ then $\theta(\lambda) = -\infty$. For $\lambda \in \mathbb{R}$ with $\lambda_1 = \lambda_2$ we have $\lambda_2 \neq 0$ and

$$\theta(\lambda) = \inf \{x_1 + \lambda_1 x_1^2 : x_1 \in \mathbb{R}\}.$$

The function inside the curly brackets attains its minimum at $x_1 = -1/(2\lambda_2)$ and

$$\theta(\lambda) = -1/(4\lambda_2). \quad (F)$$

For $\lambda = 0$ we have

$$\theta(0) = -\infty.$$

Thus $F = \{\lambda = (\lambda_1, \lambda_2) : \lambda_1 = \lambda_2, \lambda_2 > 0\}$. Since $\lambda_2 > 0$, it follows from (F) that $\sup\{\theta(\lambda) : \lambda \in F\} = 0$. The supremum is not attained in the dual problem, even though there is no duality gap. Thus we have an example in which there is no duality gap, the primal problem has a solution but the dual problem has no solution. We also recall from Example 3.3 that $\text{Att}(\theta) = \emptyset$ and there are no Kuhn-Tucker vectors.

Example 4.4. In this example we will exhibit a duality gap. Let the primal problem be the problem in Example 3.4 with $\varepsilon = 0$. We saw there that the minimum is attained at every point of the feasible set and that the value of the minimum is 1.

The dual problem is to maximize $\{\theta(\lambda) : \lambda \in \mathbb{R}, \theta(\lambda) > -\infty\}$, where

$$\theta(\lambda) = \inf \{e^{-x_1} + \lambda(x_1/\sqrt{x_1^2 + x_2^2} - x_2) : x_1 \in \mathbb{R}\}.$$

Let

$$L(x_1, x_2) = e^{-x_1} + \lambda(x_1/\sqrt{x_1^2 + x_2^2} - x_2).$$

Then $L(x_1, x_2) \geq 0$ for all x_1 in $\mathbb{R}^2$. We shall show that $\theta(\lambda) = 0$.

Let $x_1 > 0$ and let $x_2 > 0$. Then

$$\begin{aligned} L(x_1, x_2) &= e^{-x_1} + \frac{\lambda x_1^2}{\sqrt{x_1^2 + x_2^2} + x_2} \\ &= e^{-x_1} + \lambda \frac{x_1^2}{x_2 \left[\sqrt{1 + \left(\frac{x_1}{x_2}\right)^2} + 1 \right]}. \end{aligned}$$

Now let $x_1 = (x_2)^2$ and then let $x_1 \rightarrow \infty$. We see that $f(x, \lambda) \rightarrow 0$ along this curve. Thus $\theta(\lambda) = 0$ for all $\lambda \geq 0$ and $\max\{\theta(\lambda) : \lambda \geq 0\} = 0$. Thus $d = 0$, and since $s = 1$, we have a duality gap.

Theorem 4.2 states that if a solution to the primal problem exists, then strong consistency is a sufficient condition for the absence of a duality gap and the existence of a solution to the dual problem. The following example shows that strong consistency is not a necessary condition for the absence of a duality gap and for the existence of a solution to the dual problem.

Example 4.3. Minimize $f(x) = x_1$ subject to

$$g_1(x) = (x_1 + 1)^2 + (x_2)^2 - 1 \leq 0, \quad g_2(x) = -x_1 \leq 0.$$

The problem is clearly a convex programming problem. The feasible set consists of the origin $\bar{x}$ in $\mathbb{R}^2$, and the problem is not strongly consistent. The solution to the primal problem is $x_p = \bar{x}$ and $v = 0$.

The dual objective function is given by

$$\theta(\lambda) = \inf\{x_1 + \lambda_1[(x_1 + 1)^2 + (x_2)^2 - 1] - \lambda_2 x_1 : (x, \lambda) \in \mathbb{R}^2\},$$

where $\lambda \geq 0$. If we take $\lambda = \lambda_0 = (0, 1)$ we get that $\theta(\lambda_0) = 0$. It then follows from (v) of Theorem 4.1 that there is no duality gap and that λ_0 is a solution of the dual problem.

We will now obtain two necessary and sufficient conditions for the absence of a duality gap and the existence of a solution to the dual problem.

Theorem 4.3. Let problem CPM have a solution x_p . Then a necessary and sufficient condition for the absence of a duality gap and for $\eta_0 = (\lambda_0, \mu_0)$, $\lambda_0 \geq 0$, to be a solution of the dual problem is that (x_p, λ_0, μ_0) is a saddle point for L .

Proof. Let (x_p, λ_0, μ_0) be a saddle point for L . Then (2.16) and (2.18) of Theorem 2.7 hold. Thus (i) of Theorem 4.2 holds. That there is no duality gap and that (λ_0, μ_0) is a solution to the dual problem follows by the same arguments used to establish (ii) and (iii) in Theorem 4.2.

Now let $\eta_0 = (\lambda_0, \mu_0)$ be a solution of the dual problem and let there be no duality gap. Then, as in the proof of statement (iii) in Theorem 4.2 with $\lambda_0 = \lambda_0$ and $\mu_0 = \mu_0$, we get that (λ_0, μ_0) is a multiplier as in Theorem 2.2 for the primal problem. But (7) of Theorem 2.2 states that (x_p, λ_0, μ_0) is a saddle point for L .

Theorem 4.4. Let CPM have a solution x_p . Then a necessary and sufficient condition for the absence of a duality gap and the existence of a solution to the dual problem is that $\text{int} W \neq \emptyset$. In this case every vector η in $-\text{int} W$ is a solution of the dual problem.

Proof. We first prove the sufficiency of the condition $\text{dual}(\mathbf{B}) \neq \emptyset$. Since the primal problem has a solution, $\text{val}(\mathbf{P})$ is finite. By hypothesis, $\text{dual}(\mathbf{B}) \neq \emptyset$. Hence by Lemma 3.3 each vector $\mathbf{q} = (\mathbf{Q}_0, \mathbf{q})$ such that $-\mathbf{q} \in \text{dual}(\mathbf{B})$ is a Kuhn-Tucker vector, and such vectors exist. By Theorem 3.3 each such vector is in $\mathbf{R}^n$, the set of multiplier vectors, as in Theorem 2.3. Thus $(\mathbf{x}_0, \mathbf{q}, \mathbf{p})$ is a saddle point for L , and so by Theorem 4.3 there is no duality gap and $(\mathbf{Q}_0, \mathbf{q})$ is a solution to the dual problem. Now suppose that there is no duality gap and that the dual problem has a solution $\mathbf{u}_0 = (\mathbf{Q}_0, \mathbf{u}_0)$. Then by Theorem 4.3 the point $(\mathbf{x}_0, \mathbf{Q}_0, \mathbf{u}_0)$ is a saddle point for L . By Theorem 3.3, $(\mathbf{Q}_0, \mathbf{u}_0)$ is a Kuhn-Tucker vector. It then follows from Theorem 3.3 that $(\mathbf{Q}_0, \mathbf{u}_0) \in -\text{dual}(\mathbf{B})$.

COROLLARY 1. Let CPH have a solution $\mathbf{x}_0$. Then a necessary and sufficient condition for the absence of a duality gap and the existence of a solution to the dual problem is that one of the sets $\text{dual}(\mathbf{B})$, $\mathbf{X}$, or $\mathbf{R}^n$ be empty.

The corollary is an immediate consequence of Theorems 4.4 and 3.3.

Remark 4.1. In Remark 3.4 we indicated an economic interpretation of the multiplier vectors as equilibrium prices. In view of the preceding corollary, we see that the solution of the dual problem gives the equilibrium price.

Exercise 4.1. Formulate, and solve whenever possible, the problems dual to the problems in Exercise 3.2. Discuss your results as they relate to Theorem 4.1.

Exercise 4.2. The problem posed in Exercise IF.5.8 is a convex programming problem. Solve this problem by formulating the dual and solving the dual. Then, to simplify the algebra involved, first solve the problem for $\mathbf{u}_0 = \mathbf{0}$. Then use a translation of coordinates to obtain the result for general $\mathbf{u}_0$.

Exercise 4.3. Show that the function f defined by (1.2) is concave.

5. GEOMETRIC INTERPRETATION

In this section we present a geometric interpretation of the duality and perturbation results. We shall assume that all the constraints are inequality constraints and that we are considering CPH: Minimize $f(\mathbf{x})$ subject to $g(\mathbf{x}) \leq \mathbf{0}$.

The mapping

$$\zeta = f(\mathbf{x}), \quad \mathbf{z} = g(\mathbf{x}),$$

is a mapping from $\mathbb{R}^n$ to $\mathbb{R}^{n+1}$. We have indicated two image sets,

$$I = \{(\mathbf{0}, \zeta), \quad \zeta = f(\mathbf{x}), \quad \mathbf{z} = g(\mathbf{x}), \mathbf{x} \in \mathbb{R}_+^n\}, \quad (1)$$

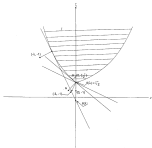


Figure 3.1

Thus, for all (ξ, η) in F

$$(\xi - x_0, -\eta) \text{ in } (C) \quad \xi - g_0(\xi)$$

and on the hyperplane

$$(\xi - x_0, -\eta) \text{ in } (C) \quad \xi - g_0(\xi) \quad (2)$$

is a supporting hyperplane for the set F . The equation of this hyperplane can also be written as

$$\xi = (x_0, y_0) + g_0(\xi) \quad (3)$$

Hence $g_0(\xi)$ is the ξ -intercept of the hyperplane (2). See Figures 3.1 and 3.2. The

geometric interpretation of the dual problem is to find, among supporting hyperplanes (7) to f with $\lambda \geq 0$, those whose λ -intercept is maximized provided such hyperplanes exist.

In Figure 5.1 the supporting hyperplane to f at any point P with x -coordinate greater than x^* has $\lambda = 0$. Thus these hyperplanes do not enter into the competition. The unique hyperplane through M_1 with normal $(0, -1)$ is the required hyperplane. The λ -intercept of this horizontal hyperplane is $\frac{1}{2}$. Thus there is no duality gap and the dual problem has a solution. The point M_1 corresponds to a point x_1 such that $g(x_1) = \frac{1}{2}$. Therefore, from the complementary slackness condition $(\lambda_{x_1}, g(x_{x_1})) = 0$, we have that $\lambda = 0$. This is in agreement with the geometry.

In Figure 5.2 any hyperplane through M_2 whose normal lies in the angle formed by u and the vector $(0, -1)$ is a supporting hyperplane with $\lambda \geq 0$ having maximal λ -intercept equal to $\frac{1}{2}$. Thus, again there is no duality gap and the dual problem has a solution.

From Figures 5.1 and the definition of the function ω we see that the graph of ω is the lower boundary curve of J for $x \leq x_1^*$ and then the straight line $\zeta = \frac{1}{2}$ for $x \geq x_1^*$. At M_1 the function ω is differentiable and $\omega'(0) = 0$. The dual problem has the unique solution $\lambda_{x_1} = 0$, and so $\omega'(0) = \lambda_{x_1}$.

In Figure 5.2 the graph of ω is the lower boundary curve of f for $x \leq 0$ and then the straight line $\zeta = \frac{1}{2}$ for $x \geq 0$. At M_2 the function ω is not differentiable. The set of vectors $\lambda_{x_2} \geq 0$ such that $(-\lambda_{x_2}, -1)$ lies in the angle formed by u and $(0, -1)$ is the set $\text{con}(0)$. The dual problem attains its solution at any vector in $-\text{con}(0)$.

Exercise 5.1. In Exercises 3.1(i), (ii), and (iv) sketch the sets J and carry out the geometric arguments of this section in these specific problems, whenever possible. For problems in which the arguments made in connection with Figures 5.1 and 5.2 seem to fail, explain why.

6. QUADRATIC PROGRAMMING

Quadratic Programming Problem (QP)

Let Q be a positive semidefinite $n \times n$ matrix, let k be a vector in $\mathbb{R}^n$, let A be an $m \times n$ matrix, and let c be a vector in $\mathbb{R}^m$. Minimize

$$f(x) = \frac{1}{2}(x, Qx) + (k, x)$$

subject to

$$Ax \leq c.$$

The formulation just given includes equality constraints of the form $Ax = d$. This is achieved by the device already used of writing each equality constraint $(k_i, x) = d_i$ where k_i is the i th row of A as two equality constraints.

Because $\tilde{Q}$ is positive semidefinite, the function f is convex. The constraints are affine. Hence the problem is a convex programming problem. We leave it as an exercise to show that if the matrix $\tilde{Q}$ is positive definite, the problem does have a solution, while the problem need not have a solution if $\tilde{Q}$ is only positive semidefinite. If $\tilde{Q}$ is positive definite, then f is strictly convex. Since in this case we are minimizing a strictly convex function over a convex set, the minimum is achieved at a unique point x_0 .

For the quadratic programming problem the Lagrangian L is defined by

$$L(x, \lambda) = \frac{1}{2}\langle x, \tilde{Q}x \rangle - \langle b, x \rangle + \langle \lambda, Ax - a \rangle. \quad (1)$$

THEOREM 5.1. *A necessary and sufficient condition for a point x_0 to be a solution of QP is that*

$$Ax_0 \leq a \quad (2)$$

and there exist a vector $\lambda_0 \geq 0$ in $\mathbb{R}^m$ such that

$$\tilde{Q}x_0 - b + A^T\lambda_0 = 0, \quad (3)$$

$$\langle \lambda_0, a - Ax_0 \rangle = 0. \quad (4)$$

Proof. We first prove the necessity of the conditions. If x_0 is a solution of QP, then x_0 is feasible, so (2) holds. By Theorem IV.5.1 or by Remark 2.3, there exists a real number $\lambda_0 \geq 0$ and a vector $\lambda_0 \geq 0$ in $\mathbb{R}^m$ such that $(\lambda_0, b_0) \in F(\lambda, b)$.

$$\lambda_0(\lambda_0 x_0 - b) + A^T\lambda_0 = 0 \quad (5)$$

and (4) holds. We shall show that $\lambda_0 = 0$ is not possible. Hence we may take $\lambda_0 = 1$ and get (3).

If $\lambda_0 = 0$, then $\lambda_0 \neq 0$ and from (5) we get that $A^T\lambda_0 = 0$. It then follows from (4) that $\langle \lambda_0, a \rangle = 0$. Thus the system

$$\begin{pmatrix} A^T \\ a^T \end{pmatrix} y = 0, \quad y \geq 0, \quad y \neq 0, \quad y \in \mathbb{R}^n,$$

has a solution $y = \lambda_0 \neq 0$. It then follows from Gordan's theorem (Theorem II.3.2) that

$$(A, -a) \begin{pmatrix} x \\ 1 \end{pmatrix} \geq 0$$

for all $x \in \mathbb{R}^n$, $1 \in \mathbb{R}$. Hence $Ax \geq a$ for all x in $\mathbb{R}^n$ and $a \in \mathbb{R}$. In particular,

this inequality holds for $\lambda = 0$. Hence $Ax \geq 0$ and $A(-x) \geq 0$ for all x in $\mathbb{R}^n$. Thus $Ax = 0$ for all x in $\mathbb{R}^n$, which is not possible since A is not the zero matrix.

The sufficiency of the conditions follows from Theorem IV.3.1 with $f(x) = \frac{1}{2}(x, Qx)$ and $g(x) = Ax - a$.

Remark 6.1. In Theorem 6.1 we showed that in the quadratic programming problem the KKT condition is a necessary condition for a minimum without assuming a priori that a constraint qualification holds. The linearity of the constraints, however, played a crucial role.

Theorem 6.1 suggests the following procedure for solving QP. Solve the linear system (1) and then reject all solutions of (1) that do not satisfy (2) and (4). In actual practice, however, various computational algorithms related to the simplex method have been devised to solve quadratic programming problems. For a discussion of these see Rosen, Solov'yev, and Shorah [1991].

An alternate, and possibly more fruitful, approach is through the dual problem, which is

$$\max\{k\lambda : k \geq 0, k\lambda \geq -\alpha\}, \quad (5)$$

where

$$k\lambda = \inf\{L(x, k) : x \in \mathbb{R}^n\} \quad (6)$$

and L is defined in (1). A justification for this procedure is provided by the following theorem.

THEOREM 6.2. Let Problem QP have a solution x_0 . Then there is no duality gap between QP and the dual problem, the dual problem has a solution, and the solution is given by the multiplier λ_0 of Theorem 6.1.

Proof. For fixed $k \geq 0$, the function $L(\cdot, k)$ defined in (1) is convex on $\mathbb{R}^n$. Therefore, by the corollary to Theorem IV.3.2, a necessary and sufficient condition for a point x_0 to minimize L is that $\nabla_x L(x_0, k) = 0$, where ∇_x denotes the gradient with respect to x . Since x_0 is a solution to QP, it follows from (3) of Theorem 6.1 that $\nabla_x L(x_0, \lambda_0) = 0$. Hence,

$$L(x_0, \lambda_0) = \min\{L(x, \lambda_0) : x \in \mathbb{R}^n\}.$$

From this and from (2) and (4) we see that (2.16) of Theorem 2.7 holds, where k is the zero function. Thus, (x_0, λ_0) is a saddle point for L . The conclusion of the theorem now follows from Theorem 4.1.

Remark 6.2. Note that we have not assumed strong consistency, but we have assumed that QP has a solution. We have already remarked that a sufficient condition for the existence of a solution is that Q be positive definite.

In determining the explicit form of the dual problem, it is useful to consider the unconstrained quadratic programming problem: Minimize

$$f(\mathbf{x}) = \frac{1}{2}(\mathbf{x}, \mathbf{Q}\mathbf{x}) - (\mathbf{b}, \mathbf{x}), \quad \mathbf{x} \in \mathbb{R}^n. \quad (8)$$

Since $\mathbf{Q}$ is positive semidefinite, there exists an orthogonal matrix P such that PQP^T is a diagonal matrix. Under the transformation of coordinates $\mathbf{y} = P\mathbf{x}$, (8) becomes

$$F(\mathbf{y}) = f(P^T\mathbf{y}) = \frac{1}{2}\mathbf{y}^T P\mathbf{Q}P^T\mathbf{y} - \mathbf{b}^T P^T\mathbf{y} = \frac{1}{2}(\alpha_1 y_1^2 + \cdots + \alpha_k y_k^2) + \sum_{i=k+1}^n d_i y_i,$$

where the α_i are the positive eigenvalues of $\mathbf{Q}$. If k is positive definite, then $k = n$; if k is positive semidefinite, then $k < n$. If a coefficient d_i with $i > k$ is different from zero, then by considering vectors $p\mathbf{e}_i$, where p is a scalar and $\mathbf{e}_i$ is the vector whose i th component is 1 and whose other components are zero, we see that $\inf\{F(\mathbf{y}) : \mathbf{y} \in \mathbb{R}^n\} = -\infty$. If $d_i = 0$ for all $i > k$, then $\nabla F(\mathbf{y}_0) = \mathbf{0}$ at any point $\mathbf{y}_0 = (y_{01}, \dots, y_{0n})$ with

$$y_{0i} = \begin{cases} \frac{d_i}{\alpha_i}, & i = 1, \dots, k, \\ \text{arbitrary}, & i = k+1, \dots, n. \end{cases}$$

Since F is convex, it follows from the corollary to Theorem IV.1.1 that F attains a minimum at any such point $\mathbf{y}_0$. If $\mathbf{Q}$ is positive definite, then $k = n$ and the point $\mathbf{y}_0$ is unique. We summarize this discussion in the following lemma.

LEMMA 6.1. *The unconstrained programming problem either has infimum equal to $-\infty$ or attains a minimum at a finite point.*

We return to the determination of the explicit form of the dual problem. For fixed $\lambda \in \mathbb{R}$, the problem of minimizing $L(\mathbf{x}, \lambda)$ over all $\mathbf{x}$ in $\mathbb{R}^n$ is an unconstrained quadratic programming problem, since

$$\begin{aligned} L(\mathbf{x}, \lambda) &= \frac{1}{2}(\mathbf{x}, \mathbf{Q}\mathbf{x}) - (\mathbf{b}, \mathbf{x}) + (\mathbf{b}, A\mathbf{x} - \mathbf{c}) \\ &= \frac{1}{2}(\mathbf{x}, \mathbf{Q}\mathbf{x}) - (\mathbf{b} - \lambda A\mathbf{x}, \mathbf{x}) - (\mathbf{b}, \mathbf{c}) \\ &= \frac{1}{2}(\mathbf{x}, \mathbf{Q}\mathbf{x}) - (\mathbf{d}, \mathbf{x}) - (\mathbf{b}, \mathbf{c}), \end{aligned}$$

where $\mathbf{d} = \mathbf{b} - \lambda A\mathbf{c}$. It then follows from Lemma 6.1 that if $\mathbf{Q}$ is finite, then we may replace the infimum in (7) by a minimum. Thus, for fixed λ such that

$\theta(\lambda) \geq -\infty$, if we denote a point at which min $L(\cdot, \lambda)$ is attained by $\mathbf{x}_\lambda$, then

$$\theta(\lambda) = L(\mathbf{x}_\lambda, \lambda) = \frac{1}{2}(\mathbf{x}_\lambda, Q\mathbf{x}_\lambda) - (\mathbf{b}, \mathbf{x}_\lambda) + (\lambda, A\mathbf{x}_\lambda - \mathbf{r}). \quad (9)$$

Since for fixed λ , the function $L(\cdot, \lambda)$ is convex on E^n , it follows from the corollary to Theorem IV.3.2 that the minimum is attained at a point $\mathbf{x}_\lambda$ if and only if $\nabla_\mathbf{x} L(\mathbf{x}_\lambda, \lambda) = \mathbf{0}$. A straightforward calculation shows that $\nabla_\mathbf{x} L(\mathbf{x}_\lambda, \lambda) = \mathbf{0}$ if and only if

$$Q\mathbf{x}_\lambda = \mathbf{b} + A'\lambda. \quad (10)$$

If we take the inner product of each side of (10) with $\mathbf{x}_\lambda$, we get

$$(\mathbf{x}_\lambda, Q\mathbf{x}_\lambda) = (\mathbf{b}, \mathbf{x}_\lambda) + (\lambda, A\mathbf{x}_\lambda) = 0.$$

Substitution of this equation into (9) gives

$$\theta(\lambda) = -\frac{1}{2}(\mathbf{x}_\lambda, Q\mathbf{x}_\lambda) = (\lambda, \mathbf{r}).$$

Thus the explicit formulation of the dual problem is as follows.

Dual Quadratic Programming Problem (DQPP)

Maximize

$$-\frac{1}{2}(\mathbf{x}, Q\mathbf{x}) + (\lambda, \mathbf{r}) \quad (11)$$

subject to

$$Q\mathbf{x} + A'\lambda = \mathbf{b}, \quad \lambda \geq \mathbf{0}. \quad (12)$$

If Q is positive definite, then the dual problem can be simplified. From (12) we get that

$$\mathbf{x} = Q^{-1}(\mathbf{b} - A'\lambda). \quad (13)$$

If we substitute this into (11), then a straightforward calculation, using the fact that Q^{-1} is symmetric, gives

$$\theta(\lambda) = -\frac{1}{2}(\lambda, R\lambda) + (\mathbf{d}, \lambda) + \alpha, \quad (14)$$

where

$$R = AQ^{-1}A', \quad \mathbf{d} = AQ^{-1}\mathbf{b} - \mathbf{r}, \quad \alpha = -\frac{1}{2}(\mathbf{b}, Q^{-1}\mathbf{b}).$$

Thus the dual problem becomes the following: Maximize $\theta(k)$ subject to $k \geq 0$, where $\theta(k)$ is given by (14).

This problem is a quadratic programming problem whose constraint set has a simple structure, namely $\{k: k \in \mathbb{R}^n, k \geq 0\}$. By Theorem 6.1 there is no duality gap and this problem has a solution k_0 , where k_0 is a multiplier as in Theorem 5.1. Thus, from (2) we get that the solution of QP is

$$x_0 = Q^{-1}(\bar{b} - A^T k_0).$$

Exercise 6.1. Show that if Q is positive definite, then the quadratic programming problem has a unique solution. Does your proof apply to the unconstrained problem?

Exercise 6.2. The problems posed in Exercise IV.5.8 and Exercise 3.2(B) are quadratic programming problems. Use the duality developed in this section to solve these problems. Compare with the duality in Theorem 4.1 and 4.2.

7. DUALITY IN LINEAR PROGRAMMING

The linear programming problem LPI formulated in Section 1 in Chapter IV can be viewed as a convex programming problem with $f(x) = c - b, x$ and $g(x) = Ax - c$. Thus all of the results of this chapter hold for the linear programming problem. The linear nature of LPI, however, results in a much stronger duality than that of Theorem 4.1. We proceed to determine this duality.

To maintain a symmetry in the duality, we shall not incorporate the constraint $x \geq 0$ in the matrix A and the vector c , as was done in Chapter IV, but shall write our primal problem (P) as

$$\text{Minimize } c - b, x \quad \text{subject to } \begin{pmatrix} A \\ -I_n \end{pmatrix} x \leq \begin{pmatrix} c \\ b_0 \end{pmatrix}, \quad (P)$$

where I_n is the $n \times n$ identity matrix and b_0 is the zero vector in $\mathbb{R}^n$.

We now formulate the dual problem. Let $q = (q, \mu)$, where $q \in \mathbb{R}^n$, $\mu \in \mathbb{R}^n$, $\mu \geq 0$. Let

$$\begin{aligned} \theta(q) &= \inf \{ c - b, x + (q, Ax - c) - (q, \mu) : x \in \mathbb{R}^n \} \\ &= \inf \{ (c - b - \mu + A^T q), x - (q, c) : x \in \mathbb{R}^n \}. \end{aligned}$$

The expression inside the curly braces is linear in x . Therefore, for it to have a finite minimum over all x in $\mathbb{R}^n$, the coefficient of each component of x must be zero, and so

$$A^T q - b - \mu = 0. \quad (1)$$

Since $\mathbf{p} \geq \mathbf{0}$, equation (1) is equivalent to

$$\mathcal{A}^T \mathbf{x} - \mathbf{b} \geq \mathbf{0}.$$

If (1) holds, then $\theta(\mathbf{q}) = \langle -\mathbf{q}, \mathbf{x} \rangle$. Thus the problem dual to P is

$$\text{Maximize } \langle -\mathbf{q}, \mathbf{k} \rangle \quad \text{subject to } \begin{pmatrix} \mathcal{A}^T \\ \mathbf{I}_m \end{pmatrix} \mathbf{x} \geq \begin{pmatrix} \mathbf{b} \\ \mathbf{q}_m \end{pmatrix}, \quad (2)$$

where $\mathcal{A}_m$ is the $m \times m$ identity matrix, and $\mathbf{q}_m$ is the zero vector in $\mathbb{R}^m$.

The dual problem (2) can be written as a linear programming problem:

$$\text{Minimize } \langle -(-\mathbf{q}), \mathbf{k} \rangle \quad \text{subject to } \begin{pmatrix} -\mathcal{A}^T \\ -\mathbf{I}_m \end{pmatrix} \mathbf{k} \leq \begin{pmatrix} -\mathbf{b}^T \\ \mathbf{q}_m \end{pmatrix}.$$

The dual of this problem is

$$\text{Maximize } \langle \mathbf{b}, \mathbf{y} \rangle \quad \text{subject to } \begin{pmatrix} -\mathcal{A} \\ \mathcal{I}_n \end{pmatrix} \mathbf{y} \geq \begin{pmatrix} -\mathbf{q} \\ \mathbf{0}_n \end{pmatrix}.$$

If we write this problem as a minimization problem and relabel the vector $\mathbf{y} \in \mathbb{R}^n$ as $\mathbf{x}$, we see that the dual of the dual problem is the primal problem. We summarize this discussion in the following lemma.

Lemma 3.1. *The dual of the linear programming problem*

$$\text{Minimize } \langle -\mathbf{b}, \mathbf{x} \rangle \quad \text{subject to } \begin{pmatrix} \mathcal{A} \\ -\mathbf{I}_n \end{pmatrix} \mathbf{x} \leq \begin{pmatrix} \mathbf{a} \\ \mathbf{q}_n \end{pmatrix} \quad (3)$$

is

$$\text{Maximize } \langle -\mathbf{a}, \mathbf{k} \rangle \quad \text{subject to } \begin{pmatrix} \mathcal{A}^T \\ \mathcal{I}_n \end{pmatrix} \mathbf{k} \geq \begin{pmatrix} \mathbf{b} \\ \mathbf{q}_n \end{pmatrix}. \quad (4)$$

The dual of the dual problem (4) is the primal problem (3).

THEOREM 3.1 (DUALITY THEOREM OF LINEAR PROGRAMMING). *Let X denote the feasible set for the primal problem and let Y denote the feasible set for the dual problem. Then:*

- (i) For each $\mathbf{x}$ in X and $\mathbf{k}$ in Y

$$\langle -\mathbf{a}, \mathbf{k} \rangle \leq \langle -\mathbf{b}, \mathbf{x} \rangle \quad (\text{weak duality})$$

- (ii) Let T be nonempty. Then X is empty if and only if $\zeta(-a, b)$ is unbounded above on X .
- (iii) Let X be nonempty. Then T is empty if and only if $\zeta(-b, a)$ is bounded below on X .
- (iv) (Strong duality). Let the primal problem (P) have a solution x_0 . Then the dual problem (D) has a solution y_0 and $\zeta(x, y_0) = (b, x_0)$.

Remark 7.1. Since problem P is also a convex programming problem, all of the conclusions of Theorem 4.1 apply to the linear programming problem. Thus, (i) is a special case of (i) of Theorem 4.1. In (ii) the statement that if $\zeta(-a, b)$ is unbounded above then X is empty is a special case of (ii) of Theorem 4.1. In the linear programming problem the stronger "if and only if" statement holds. Analogous comments hold for (iii).

Also note that the phenomenon exhibited in Example 4.3 in which there is no duality gap, yet the dual problem has no solution, cannot occur.

We now prove the theorem. We pointed out in Remark 7.1 that (i) follows from (i) of Theorem 4.1.

From the discussion of (ii) in Remark 7.1 it follows that to establish (ii) we must show that if X is empty then $\zeta(-a, b)$ is unbounded above. We first note that if X is empty, then we must have $a \neq 0$. For if $a = 0$, then $x = 0$ would be feasible and X would be nonempty.

If X is empty, then the system

$$\begin{pmatrix} A \\ -I_n \end{pmatrix} x \leq \begin{pmatrix} b \\ 0 \end{pmatrix}$$

has no solution. It then follows from Theorem II.5.1 that if X is empty, then there exists a vector $f = (y_1, y_2)$, where $y_1 \in \mathbb{R}^m$, $y_1 \neq 0$, $y_2 \in \mathbb{R}^n$, $y_2 \neq 0$, such that

$$\zeta(a, 0), \zeta(y_1, y_2) = -1, \quad \zeta(a', -I_n) \begin{pmatrix} y_1 \\ y_2 \end{pmatrix} = (b_0, 0). \quad (2)$$

Hence (2) becomes

$$(a, y_2) = -1, \quad A'y_1 = b_0, \quad y_1 \geq 0.$$

Since $T \neq \emptyset$, there exists a $k_0 \geq 0$ such that $A'k_0 \geq b$. Then for any real $\alpha > 0$, the vector $k_0 + \alpha y_1$ is in T and

$$\zeta(-a, k_0 + \alpha y_1) = \zeta(-a, k_0) + \alpha.$$

Thus, since α can be taken to be arbitrarily large, $\zeta(-a, k)$ is unbounded above.

We view the primal problem as the dual to the dual and view the dual problem as primal. Then from (ii) we conclude that if X is nonempty, then F is empty if and only if $\langle \mathbf{k}, \mathbf{x} \rangle$ is unbounded above, or equivalently that $\langle -\mathbf{k}, \mathbf{x} \rangle$ is unbounded below. Thus, (iii) is proved.

We now take up the proof of (iv). Since the primal problem has a solution, there exists a vector $\mathbf{x}$ in E^n as in Theorem IV.4.1. It is readily seen that this vector is a multiplier as in the KKT conditions of Theorem IV.5.2. The set $\mathcal{B}$ is thus nonempty. The absence of a duality gap and the existence of a solution follow from Corollary 1 to Theorem 4.4. Since there is no duality gap and both the primal and dual have a solution, $\langle -\mathbf{c}, \mathbf{x}_0 \rangle = \langle -\mathbf{k}, \mathbf{x}_0 \rangle$.

COROLLARY 1. *A necessary and sufficient condition that a point $\mathbf{x}_0$ in X and a point $\mathbf{x}_0$ in F be solutions of the primal and dual problems respectively is that*

$$\langle \mathbf{k}, \mathbf{x}_0 \rangle = \langle \mathbf{c}, \mathbf{x}_0 \rangle.$$

Proof. The necessity of the condition follows from (iv) of Theorem 3.1 and the sufficiency from (iv) of Theorem 4.1.

COROLLARY 2. *Necessary and sufficient conditions that a point $\mathbf{x}_0$ in X and a point $\mathbf{x}_0$ in F be solutions of the primal and dual problems respectively are*

- (i) $\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{x}_0^* \rangle = 0$ and
- (ii) $\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{k}, \mathbf{x}_0^* \rangle = 0$

Proof. Let $\mathbf{x}_0$ be optimal for the primal problem and let $\mathbf{x}_0$ be optimal for the dual problem. Then $\mathbf{c}, \mathbf{x}_0 \in \mathcal{C}$ and $\mathbf{k}, \mathbf{x}_0 \in \mathcal{C}$. Hence

$$\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{x}_0^* \rangle \leq 0. \quad (1)$$

Since $\mathbf{x}_0$ is feasible for the dual problem, $\mathbf{c}, \mathbf{x}_0 \geq \mathbf{k}$. Since $\mathbf{x}_0 \geq 0$, we have

$$\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{k}, \mathbf{x}_0^* \rangle \geq 0. \quad (2)$$

By Corollary 1, $\langle \mathbf{k}, \mathbf{x}_0 \rangle = \langle \mathbf{c}, \mathbf{x}_0 \rangle$, so $\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{k}, \mathbf{x}_0^* \rangle = \langle \mathbf{c}, \mathbf{x}_0 \rangle \geq 0$. Therefore $\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{x}_0^* \rangle \geq 0$. Comparing this with (1) gives (i). Also,

$$\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{x}_0^* \rangle = \langle \mathbf{c}, \mathbf{x}_0, \mathbf{x}_0^* \rangle - \langle \mathbf{c}, \mathbf{x}_0^* \rangle = \langle \mathbf{c}, \mathbf{x}_0 - \mathbf{k}, \mathbf{x}_0^* \rangle$$

and so from (2) we get that $\langle \mathbf{c}, \mathbf{x}_0 - \mathbf{k}, \mathbf{x}_0^* \rangle \leq 0$. This inequality and (1) give (ii).

To prove the sufficiency, note that (i) and (ii) give

$$\langle \mathbf{c}, \mathbf{x}_0^* \rangle = \langle \mathbf{c}, \mathbf{x}_0 - \mathbf{x}_0^* \rangle = \langle \mathbf{c}, \mathbf{x}_0 - \mathbf{k}, \mathbf{x}_0^* \rangle = \langle \mathbf{k}, \mathbf{x}_0^* \rangle.$$

It then follows from the equality of the extremum quantities and Corollary 1 that x_p is a solution of the primal problem and that λ_p is a solution of the dual.

The next theorem is a stronger version of the strong duality statement (iv) of Theorem 7.1 in that we obtain strong duality by merely assuming the feasible sets X and Y to be nonempty.

THEOREM 7.2. *The feasible sets X and Y for the primal and dual problems respectively are nonempty if and only if there exists a solution x_p to the primal problem and a solution λ_p to the dual problem. In this event*

$$\langle c, \lambda_p \rangle = \langle b, x_p \rangle. \quad (5)$$

Proof. If the primal and dual problems have solutions, then the sets X and Y are clearly nonempty.

To show that $X \neq \emptyset$ and $Y \neq \emptyset$ imply that both the dual and primal problems have solutions and that (5) holds, it suffices to show that the primal problem has a solution. For then by (iv) of Theorem 7.1, the dual problem has a solution and (5) holds.

In Remark IV.4.1 we noted that relation (2v) of Theorem IV.4.1 is a consequence of (i), (ii), (iii), and (c). Thus in Corollary IV.4.2 the same observation holds. Hence a necessary and sufficient condition for a vector x_p to be a solution of P is that there exist a vector λ_p in $\mathbb{R}^m$ such that

$$\begin{pmatrix} A \\ -I_n \end{pmatrix} x_p \leq \begin{pmatrix} c \\ b \end{pmatrix}, \quad \begin{pmatrix} A' \\ I_m \end{pmatrix} \lambda_p \geq \begin{pmatrix} b \\ 0 \end{pmatrix}, \quad (6)$$

$$\langle b, \lambda_p \rangle = \langle c, x_p \rangle.$$

Since the existence of vectors x_p in $\mathbb{R}^n$ and vector λ_p in $\mathbb{R}^m$ that satisfy (6) is a necessary and sufficient condition that x_p be a solution of P, it follows that if P has no solution, then the system

$$\begin{aligned} \text{(i)} \quad & \begin{pmatrix} A \\ -I_n \end{pmatrix} x \leq \begin{pmatrix} c \\ b \end{pmatrix}, \\ \text{(ii)} \quad & \begin{pmatrix} -A' \\ -I_m \end{pmatrix} u \leq \begin{pmatrix} -b \\ 0 \end{pmatrix}, \\ \text{(iii)} \quad & \langle -b, x \rangle + \langle u, u \rangle = 0 \end{aligned} \quad (7)$$

has no solution (x, u) , $x \in \mathbb{R}^n$, $u \in \mathbb{R}^m$. Condition (i) is the condition that x be feasible for P and condition (ii) is the condition that u be feasible for the dual problem. By assumption, there exist feasible vectors $x \in X$ and $u \in Y$. By (i) of Theorem 7.1, for such vectors we have $\langle -b, x \rangle \geq \langle -c, u \rangle$, and thus $\langle -b, x \rangle + \langle u, u \rangle \geq 0$. Hence for a vector (x, u) that satisfies (i) and (ii) we

fail to satisfy (7ii), it must fail to satisfy

$$c(-h_0, u) + (q_0, u) \leq 0, \quad (7iii)$$

so we may replace (7ii) by (7iii) in (7).

Now consider the system

$$\begin{aligned} & \lambda > 0, \quad \lambda \in \mathbb{R}, \\ & \begin{pmatrix} A \\ -I_m \end{pmatrix} u \leq \begin{pmatrix} b \\ \alpha \end{pmatrix}, \quad \begin{pmatrix} -A' \\ -I_m \end{pmatrix} u \leq \begin{pmatrix} -b' \\ \alpha \end{pmatrix}, \\ & c(-h_0, u) + (q_0, u) \leq 0. \end{aligned} \quad (8)$$

If (7) has a solution (u_0, u_0) , then (8) has a solution $(\lambda_0, u_0, u_0) = (1, u_0, u_0)$. Conversely, if (8) has a solution (λ_0, u_0, u_0) , then $(u_0, u_0) = (u_0/\lambda_0, u_0/\lambda_0)$ is a solution of (7). Therefore (7) has no solution if and only if (8) has no solution.

System (8) can be written as

$$\begin{aligned} & c(u_0, u_0, 1), (u_0, u_0, 1) \geq 0, \\ & \begin{pmatrix} A & 0 & -b' \\ -I_m & 0 & b_1 \\ 0 & -A' & b \\ 0 & -I_m & b_0 \\ -b' & b' & 0 \end{pmatrix} \begin{pmatrix} u \\ u \\ v \end{pmatrix} \leq \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix} \end{aligned} \quad (9)$$

where 0 represents a zero matrix of appropriate dimension and b_1 is the zero vector in $\mathbb{R}^l$. Since this system has no solution, it follows from Theorem B.5.1 (Farkas's lemma) that there exists a vector $(p_1, p_2, q_1, q_2, v) \geq 0_{m+l+m+l}$ with $p_1 \in \mathbb{R}^m, p_2 \in \mathbb{R}^l, q_1 \in \mathbb{R}^l, q_2 \in \mathbb{R}^m$, and $v \in \mathbb{R}$ such that

$$\begin{pmatrix} A' & -I_m & 0 & 0 & -b' \\ 0 & 0 & -A & -I_m & b \\ -b' & b_1 & b' & b_0 & 0 \end{pmatrix} \begin{pmatrix} p_1 \\ p_2 \\ q_1 \\ q_2 \\ v \end{pmatrix} = \begin{pmatrix} 0_1 \\ 0_2 \\ 1 \end{pmatrix}.$$

Thus

$$\begin{aligned} & A'p_1 - p_2 - vA = b_1, \\ & -Ap_2 - q_1 + vA = b_0, \\ & c(-p_2, A) + (q_2, b) = 1. \end{aligned} \quad (10)$$

We shall show that this system cannot have a solution $(p_1, p_2, q_1, q_2) \in \mathcal{D}_{\mathcal{A}_1, \dots, \mathcal{A}_m, \mathcal{B}_1, \mathcal{B}_2}$. This contradiction will prove that the primal problem has a solution, since we arrived at (10) by assuming that it did not.

Let $x_0 \in X$ and let $y_0 \in Y$. Such points exist since $X \neq \emptyset$ and $Y \neq \emptyset$. Suppose first that $\alpha = 0$. Then $\lambda p_1 = p_1$ and $\lambda q_1 = -q_1$. Since $x_0 \in X$, we have $\lambda x_0 \in X$. Thus, since $p_1 \geq 0$,

$$\langle -p_1, x \rangle \leq \langle -p_1, \lambda x_0 \rangle = -\langle \lambda p_1, x_0 \rangle = -\langle p_1, x_0 \rangle \leq 0, \quad (11)$$

where the last inequality follows from $p_1 \geq 0$, $x_0 \geq 0$. Similarly,

$$\langle q_1, b \rangle \leq \langle q_1, \lambda p_2 \rangle = \langle \lambda q_1, p_2 \rangle = -\langle q_1, p_2 \rangle \leq 0. \quad (12)$$

From (11) and (12) we get

$$\langle -p_1, x \rangle + \langle q_1, b \rangle \leq 0, \quad (13)$$

which contradicts the last equation in (10).

Suppose now that $\alpha > 0$. The first equation in (10) implies that p_1/α is feasible for the dual problem. The second equation in (10) implies that q_1/α is feasible for the primal problem. By the weak duality, $\langle -b, q_1/\alpha \rangle \geq \langle -x, p_1/\alpha \rangle$. Hence we again arrive at (13), which contradicts the last equation in (10).

VI

SIMPLEX METHOD

1. INTRODUCTION

In Section 4 of Chapter IV we formulated the linear programming problem and derived a necessary and sufficient condition that a point x_0 be a solution to a linear programming problem. In Section 7 of Chapter V we presented the duality theorems for linear programming. We remarked in Section 4 of Chapter IV that in practice the dimension of the space $\mathbb{R}^n$ precluded the use of the necessary and sufficient conditions to solve linear programming problems and that instead a method known as the simplex method is used. In this chapter we shall present the essentials of the simplex method, which is an algorithm for solving linear programming problems.

To motivate the rationale of the simplex methods, we consider several examples in $\mathbb{R}^2$.

Example 1.1. Minimize $f(x) = x_1 - 2x_2$ subject to

$$\begin{aligned} -x_1 - x_2 &\leq -1, \\ -x_1 + x_2 &\leq 0, \\ x_1 + 2x_2 &\leq 4, \\ x_1 &\geq 0, \quad x_2 \geq 0. \end{aligned} \tag{1}$$

The feasible set is indicated by the hatched area in Figure 6.1. The level curves $f(x) = 0$ and $f(x) = -1$ are indicated by dashed lines; the arrows indicate the directions of the gradient vectors of f . Since the gradient vectors indicate the direction of increasing f , we see from the figure that f attains a minimum at the extreme point $(\frac{1}{3}, \frac{1}{3})$ of the feasible set.

Example 1.2. This example differs from Example 1.1 in that we take the function f to be $f(x) = x_1 + x_2$. The level curves are lines parallel to the line $x_1 + x_2 = 1$, a segment of which is part of the boundary of the feasible set. Since the gradient vector of f is $(1, 1)$ and this is the direction of increasing f ,

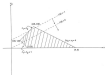


Figure 4.1

it follows that Z attains its minimum on all boundary points of the feasible set that lie on the line $x_1 + x_2 = 1$. In particular, the minimum is attained at the two extreme points $(\frac{1}{2}, \frac{1}{2})$ and $(1, 0)$.

Example 7.1. Minimize $F(x) = -x_1 - x_2$ subject to

$$-x_1 + x_2 \leq 1/2$$

$$x_1 \leq 2x_2 + 1 \quad (1)$$

$$x_1 \geq 0, \quad x_2 \geq 0.$$

The feasible set is the shaded area in Figure 4.2. The level curve $F(x) = -1/2$ is indicated by a dashed line. Since the problem of F^* for $j = 1, \dots, n_1$ it follows from the figure that this problem has no solution. The feasible set in this example is unbounded, the unbounded feasible set, however, does not imply that a linear programming problem has no solution. One can illustrate this by taking $F(x) = x_1 + x_2$. Then the problem has a solution at the extreme point $(1, 1)$.

In the examples that we considered either the problem had no solution or the problem had a solution and there was an extreme point of the feasible set at which the solution was attained. There is, of course, a third possibility.

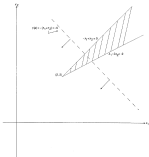


Figure 5.1

surely that the feasible set is empty. In the next section we shall show that in a linear programming problem there are the only possibilities.

The simplex method is applied to problems in canonical form, CLP (see Section 4 of Chapter IV for definitions). To put the problem in Example 1.1 in canonical form we introduce additional slack variables x_3, x_4, x_5 and write the problem as

$$\begin{aligned}
 &\text{Minimize } f(x) = x_1 - 2x_2 \\
 &\text{subject to } -x_3 - x_4 + x_5 = -4, \\
 &\quad -x_1 + x_3 + x_4 = 0, \\
 &\quad x_5 + 2x_1 + x_2 = 4,
 \end{aligned} \tag{2}$$

Since the variables x_{10} , x_{11} , x_{12} do not appear in the formula for f , the problem with feasible set defined by (2) has its minimum attained at

$$x_1 = \frac{1}{2}, \quad x_2 = \frac{2}{3}, \quad x_3 = \frac{1}{2}, \quad x_4 = x_5 = x_6 = 0.$$

It is readily verified that this point is an extreme point of the set defined by (2). We shall leave it to the reader to show that the extreme points of the constraint set in problem SLP corresponds to the extreme points in the problem CLP obtained by introducing slack variables.

Exercise 1.1. Solve the following linear programming problem graphically. Maximize $f(x) = x_1 + 2x_2$ subject to

$$\begin{aligned} x_1 + x_2 &\leq 5, \\ -x_1 + x_2 &\leq 1, \\ x_2 - x_1 &\leq 0, \\ -2x_1 - x_2 &\leq -1, \\ x_1 &\geq 0, \quad x_2 \geq 0. \end{aligned}$$

Exercise 1.2. Let x_0 be a feasible point for problem SLP and let $(x_0, \bar{z}_0)$ be the corresponding feasible point for the related problem CLP; then

$$\begin{aligned} Ax_0 &\leq b, & Ax_0 + \bar{z}_0 &= b, \\ x_0 &\geq 0, & \bar{z}_0 &\geq 0. \end{aligned}$$

Show that x_0 is an extreme point of the feasible set for problem SLP if and only if $(x_0, \bar{z}_0)$ is an extreme point for problem CLP.

Exercise 1.3. In problem LP, which is Minimize $\langle -b, x \rangle$ subject to $Ax \leq b$, show that if the feasible set is compact, then the problem has a solution and there exists an extreme point of the feasible set at which the minimum is attained.

2. EXTREME POINTS OF FEASIBLE SET

In this section we shall show that if a linear programming problem has a solution, then there exists an extreme point at which the minimum is attained. We shall prove this for problems in canonical form without assuming that the feasible set is compact, as was done in Exercise 1.3. By virtue of Exercise 1.2 this also holds for problem SLP. We shall also give an analytic characterization of the extreme points of the feasible set for CLP problems. The significance of

these results is that in looking for a solution to a linear programming problem we can confine ourselves to the extreme points of the feasible set, points for which we have an analytic characterization. The simplex method is an algorithm for efficiently determining whether a solution exists and, if so, for finding an extreme point which solves the problem.

We now consider the linear programming problem in canonical form:

$$\begin{aligned} & \text{Maximize } \langle \mathbf{b}, \mathbf{x} \rangle \\ & \text{subject to } A\mathbf{x} = \mathbf{a}, \quad \mathbf{x} \geq \mathbf{0}. \end{aligned} \quad (1)$$

In keeping with much of the linear programming literature we have replaced the objective of minimizing $\langle -\mathbf{b}, \mathbf{x} \rangle$ by the equivalent objective of maximizing $\langle \mathbf{b}, \mathbf{x} \rangle$. We assume that $\mathbf{x}$ and $\mathbf{0}$ are in $\mathbb{R}^n$, that $\mathbf{a}$ is in $\mathbb{R}^m$, that A is $m \times n$, and that

$$(i) \quad m < n, \quad (ii) \quad \text{rank } A = m. \quad (2)$$

The rationale for the assumption (2) is precisely the same as given in Remark V.1.2 for the assumption that the $k \times n$ matrix A defining the affine constraints in the convex programming problem CPH has rank k and that $k < n$.

Before presenting the principal results of this section, we emphasize that the reader should keep in mind that the constraint $\mathbf{x} \geq \mathbf{0}$ implies that the components x_j of $\mathbf{x}$ are either positive or zero.

THEOREM 2.1. *A feasible point $\mathbf{x}_0$ is an extreme point of the feasible set if and only if the columns $\mathbf{A}_j$ of the matrix A that correspond to positive components of $\mathbf{x}_0$ are linearly independent.*

Proof. Suppose that the columns $\mathbf{A}_j$ of A that correspond to positive components of $\mathbf{x}_0$ are linearly independent. We shall show that $\mathbf{x}_0$ is an extreme point of the feasible set.

If $\mathbf{x}_0$ were not an extreme point of the feasible set, there would exist two distinct feasible points $\mathbf{x}_1$ and $\mathbf{x}_2$ and a scalar λ such that $0 < \lambda < 1$ and

$$\mathbf{x}_0 = \lambda \mathbf{x}_1 + (1 - \lambda) \mathbf{x}_2. \quad (3)$$

Let B denote the submatrix of A formed by the columns of A that correspond to positive components of $\mathbf{x}$ and let N denote the submatrix of A formed by the columns of A corresponding to zero components of $\mathbf{x}$. To simplify notation, let us assume that B consists of the first k columns of A and that N consist of the remaining $n - k$ columns. There is no loss of generality in making this assumption since it can be achieved by a relabeling of the coordinates. Since the columns of B are linearly independent, it follows from (2) that $k \leq m$. It then further follows from (2) that $n - k > 0$. For $\mathbf{x}$ in $\mathbb{R}^n$, let

$$\mathbf{x}_B = (x_1, \dots, x_k), \quad \mathbf{x}_N = (x_{k+1}, \dots, x_n).$$

Then

$$\begin{aligned}x_{ij} &= (x_{i0}, x_{i1}) = (x_{i0}, \theta_j) \\ \mathbf{x}_i &= (x_{i0}, x_{i1}), \quad i = 1, 2,\end{aligned}$$

where x_{i0}, x_{i1} , and $\mathbf{x}_i$ are as in (1) and θ_j is the zero vector in $\mathbb{R}^{n-k}$.

Since $k > 0$ and $1 - k > 0$, it follows from (1) that a component of $\mathbf{x}_0$ is equal to zero if and only if the corresponding components of $\mathbf{x}_1$ and $\mathbf{x}_2$ are zero. Hence $x_{10} = x_{20} = x_{30} = \theta_0$ and we may write

$$\mathbf{x}_i = (x_{i0}, \theta_i), \quad i = 0, 1, 2.$$

Since $\mathbf{x}_1$ is feasible,

$$\mathbf{c} \cdot \mathbf{x}_1 = (\mathbf{B}, \mathbf{F}) \begin{pmatrix} x_{10} \\ \theta_1 \end{pmatrix} = Bx_{10}.$$

Similarly, we get that $\mathbf{c} = Bx_{20}$. Hence

$$B(x_{10} - x_{20}) = \mathbf{0}. \quad (6)$$

Since $\mathbf{x}_1 \neq \mathbf{x}_2$ and $x_{10} = x_{20}$, we have that $x_{10} - x_{20} \neq \mathbf{0}$. But this cannot be, since the columns of B are linearly independent. This contradiction shows that $\mathbf{x}_0$ is an extreme point.

Now let us suppose that $\mathbf{x}_0$ is an extreme point of the feasible set. To simplify notation, we suppose that the first k components of $\mathbf{x}_0$ are positive and the remaining components are zero. We shall show that the first k columns $\mathbf{A}_1, \dots, \mathbf{A}_k$ of A are linearly independent. If these columns were not linearly independent, there would exist constants $\alpha_1, \dots, \alpha_k$ not all zero such that

$$\alpha_1 \mathbf{A}_1 + \dots + \alpha_k \mathbf{A}_k = \mathbf{0}.$$

Let the $\mathbf{u}$ -vector $\mathbf{u}$ be defined by

$$\mathbf{u} = (\alpha_1, \dots, \alpha_k, 0, \dots, 0).$$

Then for any real ϵ ,

$$\mathbf{u}\mathbf{u} = \epsilon \sum_{j=1}^k \alpha_j \mathbf{A}_j = \mathbf{0}. \quad (7)$$

For $\epsilon > 0$ and sufficiently small the vectors $\mathbf{x}_0 + \epsilon\mathbf{u}$ and $\mathbf{x}_0 - \epsilon\mathbf{u}$ have their first k components positive and the remaining $n - k$ components equal to zero. Also

$$B(\mathbf{x}_0 \pm \epsilon\mathbf{u}) = B\mathbf{x}_0 \pm \epsilon\mathbf{u}\mathbf{u} = \mathbf{c}.$$

where the last equality follows from (5) and the feasibility of x_0 . Thus for sufficiently small $\epsilon > 0$ the vectors $x_0 + \epsilon u$ and $x_0 - \epsilon u$ are distinct and feasible. But then the relation

$$x_0 = \frac{1}{2}(x_0 + \epsilon u) + \frac{1}{2}(x_0 - \epsilon u)$$

contradicts the hypothesis that x_0 is an extreme point. Hence the vectors $A_1, \dots, A_k$ are linearly independent.

The following corollary is an immediate consequence of the theorem and the fact that A has rank m .

COROLLARY 1. *An extreme point of the feasible set has at most m positive components.*

COROLLARY 2. *The number of extreme points of the feasible set is less than or equal to*

$$\binom{n}{m} = \frac{n!}{m!(n-m)!}.$$

Proof. If x is an extreme point of the feasible set, then x has $k \leq m$ positive components. These k positive components correspond to k linearly independent columns of A . We assert that if x' is another extreme point with the same nonzero components as x , then $x = x'$. To see this, let B denote the submatrix of A corresponding to the nonzero components of x and x' , let $x = (x_B, 0)$, and let $x' = (x'_B, 0)$. Then by the argument used to establish (4), we get that $Bx_B = x'_B = b$. Since the columns of B are linearly independent, $x_B = x'_B$, and so $x = x'$.

If $k < m$, then since $\text{rank } A = m$, we can adjoin an additional $m - k$ columns of A to obtain a set of m linearly independent columns associated with x . If $k = m$, then this certainly holds. Hence the number of extreme points cannot exceed the number of ways in which m columns can be selected from among n columns.

THEOREM 2.1. *Let the canonical linear programming problem as stated in (1) have a solution x_0 . Then there exists an extreme point x of the feasible set that is also a solution.*

Proof. By Theorem 2.1, if the columns of A that correspond to positive components of x_0 are linearly independent, then x_0 is an extreme point of the feasible set and we let $x = x_0$.

Let us now suppose that the columns of A that correspond to the positive components of x_0 are linearly dependent. To simplify notation, we suppose that the last p components of x_0 are positive and that the remaining $n - p$ com-

ponents are equal to zero. Then the first p columns of A are linearly dependent and there exist scalars $\gamma_1, \dots, \gamma_p$ not all zero such that

$$\sum_{j=1}^p \gamma_j A_j = 0. \quad (6)$$

For each real number σ define a vector $x(\sigma)$ as follows:

$$x_j(\sigma) = \begin{cases} x_{0j} - \sigma \gamma_j & j = 1, \dots, p \\ x_j & j = p+1, \dots, n \end{cases} \quad (7)$$

Then, from the fact that $x_{0j} = 0$ for $j = p+1, \dots, n$, from the feasibility of x_0 and from (6) we get that

$$Ax(\sigma) = Ax_0 - \sigma \sum_{j=1}^p \gamma_j A_j = 0.$$

Let

$$\sigma_0 = \min \left\{ \frac{x_{0j}}{|\gamma_j|} : j = 1, \dots, p \text{ and } \gamma_j \neq 0 \right\}. \quad (8)$$

It follows from (7) that $x(\sigma) \geq 0$ whenever $|\sigma| \leq \sigma_0$. Hence $x(\sigma)$ is feasible whenever $|\sigma| \leq \sigma_0$. Thus,

$$\langle k, x_0 \rangle \geq \langle k, x(\sigma) \rangle = \langle k, x_0 \rangle - \sigma \sum_{j=1}^p k_j \gamma_j$$

for all σ such that $|\sigma| \leq \sigma_0$. Hence

$$\sum_{j=1}^p k_j \gamma_j = 0 \quad \text{and} \quad \langle k, x_0 \rangle = \langle k, x(\sigma) \rangle.$$

Thus, $x(\sigma)$ is a solution whenever $|\sigma| \leq \sigma_0$. Let the minimum in (8) be attained at the index r . Let $\sigma_1 = x_{0r}/\gamma_r$. Then $\sigma_1 = \sigma_0$ if $\gamma_r > 0$ and $\sigma_1 = -\sigma_0$ if $\gamma_r < 0$. Hence $x(\sigma_1)$ is feasible and is a solution. Moreover $x(\sigma_1)$ has fewer than p positive components. If the corresponding columns of A are linearly independent, then $x(\sigma_1)$ is an extreme point. If the columns are not linearly independent, we repeat the process, taking x_0 to be $x(\sigma_1)$. In a finite number of iterations we shall obtain a solution such that the column vectors corresponding to the positive components are linearly independent. This solution will be an extreme point of the feasible set.

2. PRELIMINARIES TO SIMPLEX METHOD

In the last section we showed that if a solution to the canonical linear programming problem exists, then it exists at an extreme point of the feasible set. We also showed that there are at most $n(n-1)/2$ such points. For a moderate system of 6 constraint equations and 30 variables, there can be as many as

$$\binom{30}{2} = \frac{30}{2} \cdot 29 = 90 \cdot 29$$

extreme points. Thus the simple-minded scheme of evaluating the objective function at all possible extreme points and then selecting the point at which the largest value of the objective function is attained is not practical. The simplex method is an algorithm for finding a solution by going from one extreme point of the feasible set to another in such a way that the objective function is not decreased.

We now introduce the notation and definitions that will be used. Consider the system

$$Ax = r \quad (1)$$

that defines the feasible set for CLP, the canonical linear programming problem. Thus, A is $m \times n$ and of rank m . Let B denote any $m \times m$ submatrix of A obtained by taking m linearly independent columns of A and let N denote the submatrix of A consisting of the remaining $n - m$ columns. We can write a vector x in $\mathbb{R}^n$ as

$$x = (x_B, x_N),$$

where x_B consists of those components of x corresponding to the columns in B and x_N consists of those components of x corresponding to the columns in N .

Definition 1.1. The columns of B are called a *basis*. The components of x_B are called *basis variables*. The components of x_N are called *nonbasic variables* or *free variables*.

We may write the system (1) as

$$(B, N) \begin{pmatrix} x_B \\ x_N \end{pmatrix} = r. \quad (2)$$

Any vector of the form (x_B, x_N) where x_N is arbitrary and

$$x_B = B^{-1}r - B^{-1}Nx_N$$

will be a solution of (2). (This is why the components of $\mathbf{u}_B$ are called *free variables*.) There is a unique solution of (2) with $\mathbf{u}_B = \mathbf{b}_B$, namely

$$\mathbf{b}_N = (\mathbf{b}_B, \mathbf{b}_N), \quad \mathbf{b}_N = \mathbf{B}^{-1}\mathbf{b}. \quad (3)$$

Definition 2.1. (i) A solution of (1) of the form (3) is called a *basic solution*. (ii) If $\mathbf{b}_B \geq \mathbf{0}_B$, then $\mathbf{b}$ is *feasible* and is called a *basic feasible solution*. (iii) If some of the components of $\mathbf{b}_B$ are zero, then $\mathbf{b}$ is called a *degenerate basic solution*.

We emphasize the following aspects of our notation. By $\mathbf{u}_B$ we mean the vector consisting of those components of $\mathbf{u}$ that correspond to the columns of $\mathbf{B}$. By $\mathbf{u}_N$ we mean the vector consisting of those components of $\mathbf{u}$ that correspond to the columns of $\mathbf{N}$. By $\mathbf{b}_B$ we mean the value of the basic variables in a basic solution. By $\mathbf{b}$ we mean the value of a solution of (1) and by $\mathbf{b}_B$ we mean the value of a particular $\mathbf{u}_B$.

We also call attention to a change in terminology. Up to now we have used the term *solution* to mean a solution of the linear programming problem. We are now using it to mean a solution of (1). Henceforth we shall use the term *optimal solution* to designate a solution of the linear programming problem, that is, a *feasible point* at which the maximum is attained.

Remark 2.1. It follows from the preceding and from Theorem 2.1 that a basic feasible solution is an extreme point and conversely. It follows from Theorem 2.2 that if an optimal solution exists, then there exists a basic feasible solution that is an optimal solution. Therefore, we need only consider basic feasible solutions in our search for the optimal solution.

To apply the simplex method, we reformulate the problem stated in (i) of Section 2 by introducing a scalar variable x and writing the problem as

$$\begin{aligned} &\text{Maximize:} \\ &\text{subject to } A\mathbf{x} = \mathbf{b}, \\ &\quad (\mathbf{b}, \mathbf{x}) = x = 0, \\ &\quad x \geq 0. \end{aligned} \quad (4)$$

The simplex method consists of two phases. In *phase I* the method either determines a basic feasible solution or determines that the feasible set is empty. In *phase II* the method starts with a basic feasible solution and either determines that no optimal solution exists or determines an optimal solution. In the latter case it does so by moving from one basic feasible solution to another in such a way that the value of x is not decreased.

We conclude this section with a review of some aspects of linear algebra involving the solution of systems of linear equations and the calculations of inverses of matrices. These aspects of linear algebra arise in the simplex method. Consider a system of linear equations

$$M\mathbf{y} = \mathbf{d}, \quad (5)$$

where M is a $k \times n$ matrix of rank k , $\mathbf{y}$ is a vector in $\mathbb{R}^n$, and $\mathbf{d}$ is a vector in $\mathbb{R}^k$. A system of linear equations $M'\mathbf{y} = \mathbf{d}'$ is said to be equivalent to (5) if the solution sets of both equations are equal. We define an elementary operation on the system (5) to be any one of the following operations.

- (i) Interchange two equations.
- (ii) Multiply an equation by a nonzero constant.
- (iii) Add a nonzero multiple of an equation to another equation.

Any sequence of elementary operations applied to (5) will yield a system equivalent to (5).

Let R denote a $k \times k$ submatrix of M that has rank k and let F denote the submatrix of M formed by the remaining $n - k$ columns. Then we may write (5) as

$$(R, F) \begin{pmatrix} \mathbf{y}_k \\ \mathbf{y}_n \end{pmatrix} = \mathbf{d}, \quad (6)$$

where $\mathbf{y}_k$ is the vector obtained from $\mathbf{y}$ by taking the components of $\mathbf{y}$ that correspond to the columns in R and $\mathbf{y}_n$ has an analogous meaning. By an appropriate sequence of elementary operations, we can obtain a system equivalent to (6) having the form

$$(U, F) \begin{pmatrix} \mathbf{y}_k \\ \mathbf{y}_n \end{pmatrix} = \mathbf{d}, \quad (7)$$

where U is a $k \times k$ upper triangular matrix with diagonal entries equal to 1 and F is a $k \times (n - k)$ matrix. To simplify notation, assume that R consists of the first k columns of M . Then (7) in component form reads

$$\begin{aligned} x_1 + a_{12}x_2 + \cdots + a_{1n}x_n + a_{1,k+1}x_{k+1} + \cdots + a_{1n}x_n &= d'_1, \\ x_2 + \cdots + a_{2n}x_n + a_{2,k+1}x_{k+1} + \cdots + a_{2n}x_n &= d'_2, \\ &\vdots \\ x_k + a_{k,k+1}x_{k+1} + \cdots + a_{kn}x_n &= d'_k. \end{aligned} \quad (8)$$

The last equation in (8) expresses x_k in terms of $x_{k+1}, \dots, x_n$. Proceeding

upward through the system (3) by backward substitution, we can express $x_1, \dots, x_n$ in terms of $x_{n+1}, \dots, x_r$, and thus obtain the general solution of (7) and hence of (5). The method of solving (5) just described is called *Gauss elimination*.

In solving the system (5) or (6) we need not write down the equations and the variables. We need only to write down the augmented matrix (A, b) or $(R, \bar{b})$ and apply the elementary operations to the rows of the augmented matrix. In so doing we speak of elementary row operations. The reader will recall, or can verify, that an elementary row operation on a $k \times n$ matrix A can be effected in the following fashion. Perform the elementary operation on the $k \times k$ identity matrix to obtain a matrix E . Then premultiply the matrix A by E . The matrices E are called elementary matrices. In computer programs for solving systems of linear equations by Gauss elimination, elementary row operations are executed by premultiplications by elementary matrices. Such computer programs also incorporate rules for choosing the components p_i of $\bar{p}$ so as to increase accuracy and speed of computation. We shall not take up these questions here.

By an additional sequence of elementary row operations, we can obtain a system equivalent to (7) having the form

$$(I_k, 0) \begin{pmatrix} x_1 \\ x_2 \end{pmatrix} = d, \quad (9)$$

where I_k is the $k \times k$ identity matrix. We can rewrite (9) as

$$I_k x_1 = d - 0x_2, \quad (10)$$

and thus obtain x_1 in terms of x_2 directly, without backward substitution. The method of solving (5) by reducing it to the equivalent system (9) is known as *Gauss-Jordan elimination*. For most large systems, computer implementations of Gauss elimination are superior to those of Gauss-Jordan elimination.

We next recall an effective way of calculating the inverse of a nonsingular matrix R . We first note that an elementary matrix E is nonsingular. Since R is nonsingular, we can reduce R to the identity matrix by applying a sequence of elementary row operations. But each elementary row operation is effected by a premultiplication by an elementary matrix. Thus, we have a sequence $E_p, E_q, \dots, E_1$ of elementary matrices such that

$$E_p \dots E_q E_1 R = I.$$

Multiplying this equation through on the right by R^{-1} gives

$$E_p \dots E_q I = R^{-1}. \quad (11)$$

This is the justification for the following method of calculating R^{-1} for matrices of low dimensions. Reduce the matrix (R, I) to the form (I, T) by a sequence of elementary operations. Then $T = R^{-1}$. The so-called product form of the inverse will also be used in the simplex method.

Remark 1.1. In reducing (6) to (8), we performed a sequence of elementary operations, which in effect was a sequence of premultiplications of (6) by elementary matrices. Thus,

$$E_p \cdots E_2(R, I, d) = (I_p, R(d')). \quad (12)$$

Since the sequence of elementary operations reduced R to I_p , we have that $R^{-1} = E_p \cdots E_2 I$. Hence, from (12) we have

$$W = R^{-1}d \quad \text{and} \quad d' = R^{-1}d.$$

Substituting this into (10) gives

$$x_i = R^{-1}d - R^{-1}E_2 x_p.$$

4. PHASE II OF SIMPLEX METHOD

Our aim in this and the next section is to explain the essentials of the simplex method; it is not to present the most current computer implementations of the simplex method. In Section 5, where we discuss the revised simplex method, we shall go into greater detail concerning the computations in current computer codes. For more complete discussions of computational questions see Dantzig [1963], Chvátal [1983], and Gass [1985].

We suppose that the problem is formulated as in relation (4) of Section 3. Phase I of the simplex, which we will describe in Section 5, either determines that the feasible set is empty or determines a basic feasible solution. We begin our discussion of phase II with the assumption that by means of phase I, or otherwise, we have obtained a basic feasible solution $\bar{q} = (\bar{q}_0, \bar{q}_p)$, where $\bar{q}_0$ is the zero vector in $\mathbb{R}^{m-p}$. We denote the value of the objective function at $\bar{q}$ by $\bar{z}$. Thus, $\bar{z} = (\bar{q}_0, \bar{q}_p)$. We call $\bar{q} = (\bar{q}_0, \bar{q}_p)$ the current basic feasible solution, the matrix B associated with $\bar{q}_p$ the current basis, and $\bar{z}$ the current value of the objective function. We abuse the phrase for good terminology and also call the components of $\bar{q}_p$ the current basis. The simplex method consists of a sequence of iterations. In a given iteration we determine either that the current basic feasible solution is optimal or that there is no optimal solution or we determine a new basic feasible solution that does not decrease the value of the objective function. In the last event, we repeat the iteration, with the new basic feasible solution as the current one. Each iteration consists of four steps. We shall illustrate the process using Example 1.1.

Step 1. Express z and the basic variables x_B in terms of the nonbasic variables x_N .

The equation in (1) of Section 3 can be written as

$$\begin{pmatrix} B & N & 0_n \\ B_0 & N_0 & -1 \end{pmatrix} \begin{pmatrix} x_B \\ x_N \\ z \end{pmatrix} = \begin{pmatrix} b \\ c \\ 0 \end{pmatrix}, \quad (1)$$

where 0_n is the zero vector in $\mathbb{R}^n$. If we multiply the augmented matrix corresponding to the first m equations in (1) by B^{-1} on the left, we see that (1) is equivalent to

$$\begin{pmatrix} I_m & B^{-1}N & 0_n \\ 0 & N_0 & -1 \end{pmatrix} \begin{pmatrix} x_B \\ x_N \\ z \end{pmatrix} = \begin{pmatrix} B^{-1}b \\ c \\ 0 \end{pmatrix}, \quad (2)$$

where I_m is the $m \times m$ identity. Setting $x_B = \bar{x}_B$ and $x_N = \bar{x}_N$ in (2) gives

$$\bar{x}_B = B^{-1}b. \quad (3)$$

From (2) and (3) we get that

$$x_B = \bar{x}_B - B^{-1}N x_N. \quad (4)$$

From (2) and (4) we get that

$$z = \langle \bar{b}, \bar{x}_B \rangle + \langle \bar{b}_0 - N_0^T B^{-1}N, \bar{x}_N \rangle.$$

Setting $x_N = \bar{x}_N$ gives

$$\bar{z} = \langle \bar{b}_0, \bar{x}_B \rangle \quad (5)$$

and so

$$z = \bar{z} + \langle \bar{b}_0 - N_0^T B^{-1}N, x_N \rangle. \quad (6)$$

Equations (4) and (6) express x_B and z in terms of x_N .

Let

$$\bar{d}_j = \bar{b}_0 - N_0^T B^{-1}N = (d_{j_{m+1}}, \dots, d_{j_n}), \quad (7)$$

where $j_{m+1}, \dots, j_n$ are the indices of the components of $\bar{x}$ that constitute $\bar{x}_B$. Then we may write (6) as

$$z = \bar{z} + \langle \bar{d}_0, x_N \rangle = \bar{z} + \sum_{j=m+1}^n d_{j_{m+1}} x_{j_{m+1}} \quad (8)$$

Remark 2.2. Note that by virtue of (3) the right-hand side of (2) can be written as $\bar{b}_0 \bar{c}'$.

The difficult calculation in (4) and (6) is the calculation of B^{-1} . The other calculations involve straightforward matrix and vector-matrix multiplications. We now present a second, seemingly different, method of obtaining x_0 and z in terms of x_0 . Computer implementation of the second method turns out in essence to be the same as the computer implementation of the previous method. Our use of the second method is didactic. For small-scale problems solved by hand this method is advantageous and will be the one used in our illustrative example. The authors method is a method that systematically records the sequence of steps and calculations that we shall describe.

We again start with the augmented matrix of (1). By a sequence of elementary operations applied to the first m rows of (1), we can transform this matrix to an equivalent matrix

$$\begin{pmatrix} I_m & B & \bar{b}_0 & \bar{c}' \\ \bar{b}_0' & \bar{b}_0' & -1 & 0 \end{pmatrix}. \quad (7)$$

Each of these operations can be effected by a premultiplication by an elementary matrix E_i . Since the sequence of operations transforms B to I , the product of the elementary matrices associated with these transformations is equal to B^{-1} , as was shown in the preceding section. Thus, the matrix (7) can also be obtained by a premultiplication of the first m rows of the augmented coefficient matrix (1) by the matrix B^{-1} . Hence, $\bar{B} = B^{-1}B$ and $\bar{c}' = B^{-1}c$, and we have, as before, that (1) is equivalent to a system with augmented matrix

$$\begin{pmatrix} I_m & B^{-1}B & \bar{b}_0 & \bar{c}' \\ \bar{b}_0' & \bar{b}_0' & -1 & 0 \end{pmatrix}, \quad (8)$$

where $\bar{c}' = B^{-1}c$. From (8) we again get (3) and (4). Let us now suppose that the columns of B are the columns $j_1, \dots, j_m$ of A . If we successively multiply row 1 of (8) by $-b_{1j_1}$ and add the result to the last row, multiply row 2 of (8) by $-b_{2j_1}$ and add the result to the last row, and so on, then we transform (8) into

$$\begin{pmatrix} I_m & B^{-1}B & \bar{b}_0 & \bar{c}' \\ \bar{b}_0' & \bar{b}_0' - \bar{b}_0' B^{-1}B & -1 & -(\bar{b}_0' \bar{c}') \end{pmatrix}. \quad (9)$$

From the last row of this matrix and from the relation $\bar{c}' = B^{-1}c$, we immediately obtain

$$z = (\bar{b}_0' B^{-1} \bar{c}') + (\bar{b}_0' - \bar{b}_0' B^{-1} B) x_0.$$

Remark 4.1. From (11) and $x' = B^{-1}b$ we again get (3) and (5). Consequently, the right-hand side of (11) can be written as $b_0 = -25$. From this relation and from (3) and (5) we again get (8).

We illustrate the preceding description of Step 1 using the second method.

Example 4.1. We first introduce slack variables x_4, x_5, x_6 and write the problem as in (2) of Section 1. We then write the problem as a maximization problem in the format of (4) of Section 3.

Maximize z subject to

$$\begin{aligned} -x_2 - x_1 + x_3 &= -1, \\ -x_2 + x_1 &+ x_4 = 0, \\ x_1 + 2x_2 &+ x_5 = 4, \\ -x_6 + 2x_1 &- z = 0. \end{aligned} \quad (12)$$

It is readily verified that $\bar{x} = (-1, 0, 0, 1, 0)$, which corresponds to the point (1, 0) in Figure 3.1, is a basic feasible solution. The value of z at this point is -1 . Thus, x_4, x_5, x_6 are the basic variables and x_2, x_1 are the nonbasic variables. The augmented coefficient matrix of (12) is

$$\begin{array}{c} x_4 \\ x_5 \\ x_6 \end{array} \left\{ \begin{array}{cccccc|c} x_1 & x_2 & x_3 & x_4 & x_5 & z & c \\ -1 & -1 & 1 & 0 & 0 & 0 & -1 \\ -1 & 1 & 0 & 1 & 0 & 0 & 0 \\ 1 & 2 & 0 & 0 & 1 & 0 & 4 \\ -1 & 2 & 0 & 0 & 0 & -1 & 0 \end{array} \right\} \quad (13)$$

Above each column of the matrix we have indicated the variable corresponding to the column. At the left of the matrix we have listed the current basic variables. The last column is the column of right-hand sides of (12). In the tabular description of the simplex method the matrix (13) is the initial tableau. The matrix B in (13) is the submatrix formed by the first three rows and columns 1, 4, and 5. Thus

$$B = \begin{pmatrix} -1 & 0 & 0 \\ -1 & 1 & 0 \\ 1 & 0 & 1 \end{pmatrix}$$

To obtain the matrix corresponding to (13), we must first transform B to the identity matrix by a sequence of elementary row operations on (13). From the form of B it is clear that the most effective sequence of operations is the following. (i) Multiply row 1 by -1 . (ii) Add the transformed row 1 to row 2.

(iii) Add the negative of the transformed row 1 to row 3. To obtain the row corresponding to the last row of (11), add the transformed row 1 to row 4. We thus obtain

$$\begin{array}{l} \left(\begin{array}{cccccc|c} x_1 & x_2 & x_3 & x_4 & x_5 & z & 1 \\ x_1 & 1 & -1 & 0 & 0 & 0 & 1 \\ x_4 & 0 & 2 & -1 & 1 & 0 & 1 \\ x_3 & 0 & 1 & 1 & 0 & 1 & 2 \\ 0 & 3 & -1 & 0 & 0 & -1 & 1 \end{array} \right). \end{array} \quad (14)$$

From this we read off

$$\begin{aligned} x_1 &= 1 - x_3 + x_{5r} \\ x_4 &= 1 - 2x_{5r} + x_{3r} \\ x_3 &= 2 - x_1 - x_{5r} \end{aligned} \quad (15)$$

and

$$z = -1 + 3x_{5r} - x_{3r} \quad (16)$$

which are (4) and (6) for this example. Note that the last column of (14) gives the current values of $\bar{z}_r$, $\bar{c}_{5r}$, $\bar{c}_3 = \bar{c}$ in agreement with (4) and (6) and Remark 4.2.

Step 2. Check the current basic feasible solution for optimality.

If every component of the vector $\mathbf{d}_r$ defined in (7) is nonpositive (i.e., ≤ 0), then the current basic feasible solution is optimal and the process terminates.

To see that this is so, recall that by Remark 3.1 we need only consider basic feasible solutions in our search for the optimal solution. The current basic feasible solution is the unique basic feasible solution with $\mathbf{x}_B = \mathbf{b}_B$. Thus any other basic feasible solution will have at least one of the components of $\mathbf{x}_B$ positive. But since $\mathbf{d}_r \leq \mathbf{0}$, we have from (8) that the value of the objective function will either decrease or be unchanged. Thus the current basic feasible solution is optimal.

In our example, we see from (16) that the current solution is not optimal.

Step 3. If the current basic feasible solution is not optimal, either determine a new basis and corresponding basic feasible solution that does not decrease the current value of z or determine that z is unbounded and that therefore the problem has no solution.

We first motivate Step 3 using Example 4.1.

From (14) we see that if we keep the value of x_3 at zero and increase the value of x_6 , then we shall increase the value of z . The variables x_1, x_2, x_3, x_4, x_5 must satisfy (15) and the condition $x_i \geq 0, i = 1, \dots, 5$, for feasibility. Setting $x_6 = 1$ in (15) and taking into account the requirement that x_1, x_2 , and x_5 must be nonnegative give the following inequalities that x_3 must satisfy:

$$1 - x_3 \geq 0, \quad 1 - 2x_3 \geq 0, \quad 3 - x_3 \geq 0.$$

For all three of these inequalities to hold we can only increase the value of x_3 to $\frac{1}{2}$. If $x_3 = \frac{1}{2}$ and $x_6 = 0$, then from (15) we get that $x_1 = \frac{1}{2}, x_2 = 0$, and $x_5 = \frac{3}{2}$. We now have another feasible solution $(\frac{1}{2}, \frac{1}{2}, 0, 0, \frac{3}{2})$. This corresponds to the point $(\frac{1}{2}, \frac{1}{2})$ in Figure 6.1. The value of z is increased from -1 to $\frac{1}{2}$.

We now show that this solution is basic. The matrix corresponding to the system (15) is the submatrix of (14) consisting of the first three rows of (14). The matrix B corresponding to $(\frac{1}{2}, \frac{1}{2}, 0, 0, \frac{3}{2})$ is the matrix formed by taking columns 1, 2, and 5 of this submatrix. Thus

$$B = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 2 & 0 \\ 0 & 1 & 1 \end{pmatrix}.$$

Since the columns of B are linearly independent, the point $(\frac{1}{2}, \frac{1}{2}, 0, 0, \frac{3}{2})$ is a basic feasible solution.

In summary, the original basic variable x_4 whose value was 1 is now a nonbasic variable and the original nonbasic variable x_6 is now a basic variable with value $\frac{1}{2}$. The variable x_3 is said to be the leaving variable or to leave the basis and the variable x_6 is said to be the entering variable or to enter the basis.

Before presenting Step 3 in general, we call attention to another possible situation. Suppose that we have an example such that instead of (15) we arrived at a system

$$x_1 = 1 + \alpha x_6 + x_3,$$

$$x_2 = 1 + \beta x_6 + x_3,$$

$$x_5 = 3 + \gamma x_6 + x_3,$$

with α as in (16) and where α, β , and γ are all nonnegative. If we set $x_6 = 0$ and let $x_3 = z$, then we obtain feasible solutions for all $z \geq 0$. If we let $z \rightarrow \infty$, we get that $z \rightarrow \infty$. Thus, the objective function is unbounded and the problem has no solution.

We now describe Step 3, which we break up into two substeps, for the general problem.

Step 3(i). Choose the nonbasic variable that will enter the new basis.

Let $\bar{d}_q$ be as in (7). Since the current basic feasible solution is not optimal, at least one of the components of $\bar{d}_q$ is positive. Choose the component of $\bar{d}_q$ that is the largest. If there is more than one such component, choose the one with the smallest index. Suppose that the component of $\bar{d}_q$ that is chosen is the q th component of $\bar{x} = (\bar{x}_1, \dots, \bar{x}_n)$. The variable x_q will enter the basis. The column of A corresponding to x_q is $\bar{a}_q$.

Step 3(i). Determine that the problem has no solution or choose the basic variable that will leave the basis.

Let x_i be the column of $B^{-1}A$ corresponding to x_q . Then $x_q = B^{-1}\bar{a}_q$. Let

$$x_q = (x_{q1}, \dots, x_{qm}) \quad \bar{a}_q = (\bar{a}_{q1}, \dots, \bar{a}_{qm}) \quad (17)$$

and let $\bar{x}_q = (x_{q1}, \dots, x_{qm})$. Set $x_q = t$, where $t \geq 0$ and let the other components of x_q remain at zero. For the resulting vector to be a feasible solution, x_q must satisfy the conditions (4) and $x_q \geq 0$. Equation (4) becomes

$$x_p = \bar{a}_{pq} - t\bar{a}_{qj} \quad (18)$$

Thus, the condition $x_p \geq 0$ imposes the restrictions

$$0 \leq \bar{a}_{pq} - t\bar{a}_{qj}, \quad j = 1, \dots, m. \quad (19)$$

From (8) we get that

$$z = \bar{z} + d_q t. \quad (20)$$

Since $\bar{z}$ is feasible, $\bar{a}_{pq} \geq 0$ and so $\bar{a}_{qj} \geq 0$ for all $j = 1, \dots, m$. Hence if $\bar{a}_{qj} < 0$ for all $j = 1, \dots, m$, then (18) and (19) hold for all $t \geq 0$. Since $d_q > 0$, it follows that z is unbounded above and the problem has no solution.

If not all of the $\bar{a}_{qj}$ are nonpositive, let

$$\rho = \min \left\{ \frac{\bar{a}_{pq}}{\bar{a}_{qj}}; j = 1, \dots, m \text{ and } \bar{a}_{qj} > 0 \right\}. \quad (21)$$

Let ρ denote the smallest index j at which the minimum is achieved. Let $\bar{a}_{q\rho} = \rho$ and let

$$\begin{aligned} \bar{z}'_1 &= \bar{z}_1 - \rho \bar{a}_{q1} & \bar{z}'_j &= \bar{z}_j \\ \bar{z}'_p &= \rho \\ \bar{z}'_q &= 0 \\ \bar{z}'_j &= 0 & \text{all other indices.} \end{aligned} \quad (22)$$

Let $\bar{q}^*$ denote the vector whose coordinates are given by (22). Then

$$\bar{q}^* = (q_{1,p}^*, \dots, q_{i_{p-1},p}^*, 0, q_{i_{p+1},p}^*, \dots, q_{i_p,p}^*, 0, \dots, 0, r^*, 0, \dots, 0),$$

where r^* is the q th coordinate of $\bar{q}^*$. It follows from the way $\bar{q}^*$ was constructed that $\bar{q}^*$ is a feasible solution. We shall show that $\bar{q}^*$ is also basic.

We can write $\bar{q}^*$ as $\bar{q}^* = (q_{1,p}^*, \bar{q}_q^*)$ with

$$q_{1,p}^* = (q_{1,p}^*, \dots, q_{i_{p-1},p}^*, r^*, q_{i_{p+1},p}^*, \dots, q_{i_p,p}^*), \quad \bar{q}_q^* = 0.$$

We shall show that the coordinate variables of

$$s_p = (s_{1,p}, \dots, s_{i_{p-1},p}, s_q, s_{i_{p+1},p}, \dots, s_{i_p,p}) \quad (23)$$

are a basis. Since s_q is obtained from s_p by replacing s_{i_p} by s_q , this will also justify our saying that s_q enters the basis and s_{i_p} leaves the basis.

From (17) we get that the matrix I_p in (2) and (5) is

$$(s_1, \dots, s_{i_{p-1}}, s_i, s_{i_{p+1}}, \dots, s_{i_p}),$$

where s_i is the vector whose i th component is 1 and whose other components are zero. Let

$$R = (s_1, \dots, s_{i_{p-1}}, s_q, s_{i_{p+1}}, \dots, s_{i_p}).$$

That is, R is the matrix obtained from I_p by replacing the i_p th column by s_q . Let R' be the matrix obtained from $R^{-1}R$ by replacing column s_q by s_{i_p} . Let s_{i_p} be as in (21) and let s_q be the vector obtained from s_{i_p} by replacing s_{i_p} by s_q . Let d_q denote the q th coordinate of the vector d_q defined in (7) and let $\bar{d}_q$ denote the vector obtained from d_q by replacing the q th coordinate of d_q by 0. The system of equations

$$\begin{pmatrix} s_1 & \cdots & s_{i_{p-1}} & s_q & s_{i_{p+1}} & \cdots & s_{i_p} & R' & d_q \\ 0 & \cdots & 0 & d_q & 0 & \cdots & 0 & \bar{d}_q & -1 \end{pmatrix} \begin{pmatrix} s_1 \\ \vdots \\ s_{i_{p-1}} \\ s_q \\ \vdots \\ s_{i_p} \end{pmatrix} = \begin{pmatrix} R^{-1}r \\ -1 \end{pmatrix} \quad (24)$$

is equivalent to the system whose augmented matrix is (11), since (24) is obtained from (11) by a reordering of the variables.

We assert that the matrix R is a basis and hence that the variables s_p are basic. To prove this, we must show that the columns of R are linearly independent. By (21) and the definition of the index p , the i_p th component of s_q is positive. The i_p th component of each of the other columns of R is zero. Hence s_q cannot be written as a linear combination of the other columns of R . Since the remaining columns of R are linearly independent, it follows that all the columns of R are linearly independent. Hence the coordinate variables of

$\mathbf{u}_q$ are a basis, and the vector $\bar{\mathbf{z}}$ is a basic feasible solution. From (28) we get that

$$\bar{\mathbf{z}} = \bar{\mathbf{z}} + d_q \mathbf{e}^q.$$

Since $d_q > 0$ and $\epsilon > 0$, it follows that the value of the objective function is not decreased.

We say that in Step 3 we have updated the unprimal quantities to the primal quantities. The primal, or updated, quantities are the "new" or "updated current" quantities.

Remark 4.1. In Step 3(i) we choose the entering variable x_q to be the variable corresponding to the largest component of $\mathbf{d}_q$. Upon starting Step 3, it should be clear that the choice of entering variable could be any variable corresponding to a positive component of $\mathbf{d}_q$. The choice of largest component was motivated by the fact that at a given iteration this choice would appear to result in the largest increase in the value of z .

Step 4. Return to Step 1 using the updated current quantities and start the next iteration.

We now illustrate Steps 3 and 4 using Example 4.1.

Step 3(i). From (14) we see that $\mathbf{d}_q = (d_1, d_2)$ and that $d_1 > 0$ and $d_2 < 0$. Thus, we take x_1 as the entering variable. The entering column $\mathbf{v}_1$ is the second column of (14), and $d_q = d_1 = 3$.

Step 3(ii). By Remark 4.2, the ratios in the last column of (14) give

$$-z_1 = \bar{z}_2 - \bar{z}_3 = (\bar{z}_2 - \bar{z}_3 - \bar{z}_4) = (1, 1, 3), \quad \zeta = -1.$$

From this and the fact that the first three entries in column 2 are positive, we get that

$$\frac{\bar{z}_1}{v_1} = \frac{\bar{z}_2}{v_1} = 1, \quad \frac{\bar{z}_3}{v_1} = \frac{\bar{z}_4}{v_1} = \frac{1}{2}, \quad \frac{\bar{z}_5}{v_1} = \frac{\bar{z}_6}{v_1} = 3.$$

Hence $\theta = \frac{1}{2}$ and x_4 is the leaving variable. The new, or updated, basic variables are (x_1, x_2, x_3) and the updated nonbasic variables are x_4 and x_6 . From (22) we get the updated basic feasible solution

$$\begin{aligned} \bar{z}'_1 = \bar{z}'_2 = 1 - \frac{1}{2}(1) = \frac{1}{2}, \quad \bar{z}'_3 = \bar{z}'_4 = \frac{1}{2} \\ \bar{z}'_5 = 0, \quad \bar{z}'_6 = \bar{z}'_4 = 3 - \frac{1}{2}(1) = \frac{5}{2}. \end{aligned}$$

From (20) we get

$$z^* = (-1) + 3\left(\frac{1}{2}\right) = \frac{1}{2}.$$

These calculations, however, are unnecessary as we will get them automatically in Step 1 of the next iteration.

Step 4. We go to Step 1 and start the next iteration using the updated data. We do so dropping the primes that indicated updated data. The variables $x_1, \dots, x_6$ satisfy the system of equations represented by (24). We rewrite (14) with the labelling of the new basic variables:

$$\begin{matrix} x_1 & x_2 & x_3 & x_4 & x_5 & z & \\ x_1 & \left\{ \begin{array}{cccccc} 1 & 1 & -1 & 0 & 0 & 0 & 1 \\ 0 & 2 & -1 & 1 & 0 & 0 & 1 \\ 0 & 1 & 1 & 0 & 1 & 0 & 2 \\ 0 & 3 & -1 & 0 & 0 & -1 & 1 \end{array} \right. \end{matrix} \quad (25)$$

Step 1 requires us to transform (25) into a form corresponding to (11). The matrix B is now

$$B = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 2 & 0 \\ 0 & 1 & 1 \end{pmatrix}.$$

Thus, to achieve the desired result, we apply the following sequence of transformations to (25).

- (i) Multiply row 2 by $\frac{1}{2}$.
- (ii) Add the negative of the new row 2 to row 1.
- (iii) Add the negative of the new row 2 to row 3.
- (iv) Multiply the new row 2 by -3 and add the resulting row to row 4.

We get

$$\begin{matrix} x_1 & x_2 & x_3 & x_4 & x_5 & z & \\ x_1 & \left\{ \begin{array}{cccccc} 1 & 0 & -\frac{1}{2} & -\frac{1}{2} & 0 & 0 & \frac{1}{2} \\ 0 & 1 & -\frac{1}{2} & \frac{1}{2} & 0 & 0 & \frac{1}{2} \\ 0 & 0 & \frac{1}{2} & -\frac{1}{2} & 1 & 0 & \frac{3}{2} \\ 0 & 0 & \frac{1}{2} & -\frac{1}{2} & 0 & -1 & -\frac{1}{2} \end{array} \right. \end{matrix} \quad (26)$$

From (26) we can read off, if we desire, the basic variables $(x_1, x_2, x_3) = (x_5, x_4, x_6)$ and z in terms of the nonbasic variables (x_3, x_4) . By Remark 4.3 the last column of (26) gives

$$\hat{z}_1 = \frac{1}{2} \quad \hat{z}_2 = \frac{1}{2} \quad \hat{z}_3 = \frac{3}{2} \quad \hat{z} = \frac{1}{2}$$

for $\mathbf{x}_p$, the current basic feasible solution, and ζ , the current value of the objective function. Note that we have agreement with the previous calculation of these quantities. The current basic feasible solution corresponds to the point $(\frac{1}{2}, \frac{1}{2})$ in Figure 6.1. The current value $\zeta = \frac{1}{2}$ is greater than the previous current value of -1 .

Before proceeding to Step 2, we introduce some terminology and comment on the carrying-out of Step 1 in general. The entry a_{ij} in row i , column j in this example is called the *pivot* and the position in the matrix is called the *pivot position*. In performing the sequence of elementary operations used to obtain (26) from (25), we say that we have *pivoted* on the element a_{ij} .

In our example for pivot was the entry in the row and column of the entering variable. We now show that in each iteration after the initial iteration we always pivot on the element in the row and column of the entering variable. The system (26) gives the relationship among the variables x_1, x_2 , and x at the end of Step 4 of each iteration. It expresses x and each component of

$$\mathbf{x}_p = (x_1, \dots, x_{j_0}, x, x_{j_0+1}, \dots, x_{m_1})$$

except x_{j_0} in terms of the components of $\mathbf{x}_p$ and the variable x_{j_0} . Step 1 of the current iteration requires that $\mathbf{x}_p$ and x be expressed in terms of the components of $\mathbf{x}_p$. This will be achieved if we express x_{j_0} in terms of the components of $\mathbf{x}_p$ and do so in a way that does not destroy the already established expressions of x and the x_{j_0} , $j = 1, \dots, m_1$, $j \neq j_0$, in terms of x_1 and x_{j_0} . We assert that pivoting on row i_p , the entering row, and on column j_0 (column $\mathbf{x}_{j_0}$, the entering column), achieves this.

By the definition of j_0 and (21), we have that $a_{i_p j_0} > 0$. Hence in (24) we may multiply row i_p by $1/a_{i_p j_0}$ and obtain a row with entry 1 in the (i_p, j_0) position. If we now multiply row i_p by $-a_{1j_0}$ and add to row 1, the resulting first row will have a zero in the $(1, j_0)$ position; that is, the column vector $\mathbf{x}_{j_0}$ now has a 1 in the i_p entry and a zero in the first position. Since each of the vectors $\mathbf{x}_{j_0}$, $j = 1, \dots, m_1$, $j \neq j_0$, has a zero in the (i_p, j_0) entry, the first entries of these vectors are not changed. Similarly, if for each $i = 2, \dots, m_1$, $i \neq j_0$, we multiply row i_p by $-a_{ij_0}$ and add the resulting row to row i , then the resulting i th row will have a zero in the (i, j_0) position. Each of the vectors $\mathbf{x}_{j_0}$, $j = 1, \dots, m_1$, $j \neq j_0$, will have j th entries unchanged. Finally, if we multiply row i_p by $-d_{i_p}$ and add the resulting row to the last row of (26), then we shall get a system

$$\begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} x_1 \\ x_2 \\ x \\ x_{j_0} \end{pmatrix} = \begin{pmatrix} c_1' \\ c_2' \\ c_3' \\ -c_4' \end{pmatrix}.$$

This system expresses $\mathbf{x}_p$ and x in terms of $\mathbf{x}_{j_0}$.

We now return to Step 2 and (26) of the current iteration. Since $d_{i_p} = d_{j_0} - \hat{d}_{j_0}$, the current basic feasible solution is not optimal. The current value $\zeta = \frac{1}{2}$ is an improvement, however.

Step 3(f). Since $d_1 > 0$ and $d_2 < 0$, we take x_2 to be the entering variable.

Step 3(g). The column corresponding to x_2 is the third column of (26). Since the only positive entry in this column is the entry in row 3 corresponding to x_3 , the variable x_2 will be leaving variable. We rewrite the matrix (26) with the updated basic variables indicated on the left:

$$\begin{array}{l} x_1 \\ x_2 \\ x_3 \end{array} \left\{ \begin{array}{cccccc} x_1 & x_2 & x_3 & x_4 & x_5 & r \\ 1 & 0 & -\frac{1}{2} & -\frac{1}{2} & 0 & 0 & \frac{3}{2} \\ 0 & 1 & -\frac{1}{2} & \frac{1}{2} & 0 & 0 & \frac{5}{2} \\ 0 & 0 & \frac{1}{2} & -\frac{1}{2} & 1 & 0 & \frac{1}{2} \\ 0 & 0 & \frac{1}{2} & -\frac{1}{2} & 0 & -1 & -\frac{1}{2} \end{array} \right\}$$

We now follow Step 4 and begin another iteration with Step 1. We follow our pivot rule and pivot on $\frac{1}{2}$ in row 3, column 3. This will result in a matrix

$$\begin{array}{l} x_1 \\ x_2 \\ x_3 \end{array} \left\{ \begin{array}{cccccc} x_1 & x_2 & x_3 & x_4 & x_5 & r \\ 1 & 0 & 0 & -\frac{1}{2} & \frac{1}{2} & 0 & \frac{5}{2} \\ 0 & 1 & 0 & \frac{1}{2} & \frac{1}{2} & 0 & \frac{7}{2} \\ 0 & 0 & 1 & -\frac{1}{2} & \frac{1}{2} & 0 & \frac{3}{2} \\ 0 & 0 & 0 & -\frac{1}{2} & -\frac{1}{2} & -1 & -\frac{5}{2} \end{array} \right\}$$

Step 2. $\bar{b}_3 = (-\frac{1}{2}, -\frac{1}{2})$. All components of $\bar{b}_3$ are negative, so the current basic feasible solution is optimal. From the last column we get that

$$z_{\text{opt}} = z_{x_1, x_2, x_3} = \frac{1}{2} \cdot \frac{5}{2}$$

and that the value of the objective function is $z = \frac{5}{4}$. This solution corresponds to the point $(\frac{5}{4}, \frac{1}{4})$ in Figure 5.1.

8. TERMINATION AND CYCLING

In Section 4 we showed that each iteration of the simplex method results in a value of the objective function that is greater than or equal to the value of the objective function at the preceding iteration. If we could show that each iteration produces a value of the objective function that is strictly greater than the value at the preceding iteration, then the simplex method would terminate at an optimal solution in a finite number of steps, for there are finite number of basic feasible solutions. Unfortunately, it is not true that each iteration

produces a value of the objective function that is strictly greater than the preceding value, as the following simple example shows.

Let us suppose that at the end of Step 2 of an iteration we had arrived at a system given by

$$\begin{array}{l} x_1 \\ x_2 \\ x_3 \end{array} \begin{pmatrix} x_1 & x_2 & x_3 & x_4 & x_5 & b & \theta \\ 1 & 1 & 1 & 0 & 0 & 0 & 1 \\ 0 & 2 & -1 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 & 1 & 0 & 2 \\ 0 & 2 & -1 & 0 & 0 & -1 & 0 \end{pmatrix}.$$

Then the entering variable is x_1 and the entering column $\mathbf{v}_1$ is the second column. Recalling that the last column gives $\mathbf{b}_{\text{new}} = \mathbf{b}$, we get that

$$\frac{b_1}{v_{11}} = \frac{0}{1} = 0, \quad \frac{b_2}{v_{21}} = \frac{0}{2} = 0, \quad \frac{b_3}{v_{31}} = \frac{2}{1} = 2. \quad (1)$$

Thus $r^* = 0$ and the leaving variable is x_2 . It then follows from (4.26) that $\bar{c}^* = c^*$ and the updated value of a is not changed. From (4.22) we get that the updated basic feasible solution is

$$\mathbf{b}_{\text{new}} = \mathbf{b}_{\text{old}} - \mathbf{v}_1 \mathbf{v}_1^T \mathbf{b}_{\text{old}} = \mathbf{b}_{\text{old}} - \mathbf{v}_1 \mathbf{v}_1^T \mathbf{b} = (1, 0, 2).$$

Looking at (1), we see that $r^* = 0$ because $\bar{b}_1 = 0$. In other words, $r^* = 0$ is a consequence of the current basic feasible solution being degenerate. In this example, the updated basic solution is also degenerate. In the next example we show that a degenerate basic solution need not lead to another degenerate solution with no increase in the value of a .

Let us now suppose that at the end of Step 2 we arrived at the system described by

$$\begin{array}{l} x_1 \\ x_2 \\ x_3 \end{array} \begin{pmatrix} x_1 & x_2 & x_3 & x_4 & x_5 & b & \theta \\ 1 & 1 & 0 & 0 & 0 & 0 & 1 \\ 0 & -2 & -1 & 1 & 0 & 0 & 0 \\ 0 & 1 & 1 & 0 & 1 & 0 & 2 \\ 0 & 2 & -1 & 0 & 0 & -1 & 0 \end{pmatrix}.$$

The current basic feasible solution $(\mathbf{b}_1, \mathbf{b}_2, \mathbf{b}_3) = (1, 0, 2)$ is again degenerate. The entering variable is again x_1 . Now, however, the ratios used to determine r^* are

$$\frac{b_1}{v_{11}} = \frac{1}{1} = 1, \quad \frac{b_2}{v_{21}} = \frac{0}{-2} = 0,$$

Thus $\bar{r}^* = 1$ and the leaving variable is x_1 . Since $\bar{d}_1 = 1$, $\bar{r}^* = -1$, and $\bar{r}^* = 1$, from (4.20) we get

$$\bar{r} = -1 + M(1) = 2.$$

Thus the value of s is increased. From (4.22) we calculate the new basic feasible solution to be

$$\bar{b}_0 = (b_1, b_2, b_3) = (b_2, b_1, b_3) = (3, 1, 1),$$

which is nondegenerate.

We now consider the general situation. Let $\bar{b}_0$ be a basic degenerate feasible solution with zero component $\bar{b}_i$. If the corresponding component $\bar{d}_i$ of the leaving column $\bar{r}$ is positive, then $\bar{r}^* = 0$; the updated basic feasible solution will be degenerate, and the updated value of the objective function will be unchanged. This follows from (4.11), (4.12), and (4.20).

The following question now arises: Can there exist a sequence of iterations $\bar{b}_0, \bar{b}_1, \dots, \bar{b}_{p-1}$ such that the basic feasible solutions $\bar{b}_{k+1}, j = 0, 1, \dots, p$, are all degenerate and such that $\bar{b}_{k+1, j} = \bar{b}_{k, j}$? If this occurs, then all current values $\bar{c}_0, \bar{c}_1, \dots, \bar{c}_{p-1}$ of the objective will be equal and the sequence of degenerate basic solutions $\bar{b}_{k+1}, \bar{b}_{k+2}, \dots, \bar{b}_{k+p}$ will repeat infinitely often. Thus, the simplex method will not terminate. The phenomenon just described is called *cycling* and the simplex method is said to *cycle*.

The bad news is that cycling can occur. The good news is that cycling appears to be extremely rare in problems arising in applications, and there are modifications to the simplex method that prevent cycling.

The following example due Brink [1973] cycles in six iterations.

Example 5.1. Minimize $z = 22x_1 + 90x_2 + 21x_3 + 24x_4$ subject to

$$\begin{aligned} -6x_1 + 8x_2 + 28x_3 + x_4 + 9x_5 &= 0, \\ \frac{1}{2}x_1 + \frac{1}{2}x_2 + \frac{1}{2}x_3 + x_4 + x_5 &= 0, \\ x_2 &+ x_3 = 1, \\ x_j &\geq 0, \quad j = 1, \dots, 5. \end{aligned}$$

From the preceding discussion and from Section 4 we have the following:

Lemma 5.1. Phase II of the simplex method does one of the following:

- (i) determines an optimal solution in a finite number of steps,
- (ii) determines that no solution exists in a finite number of steps, or
- (iii) cycles.

We shall present a very simple rule for the selection of entering and leaving variables that makes cycling impossible. This rule is the smallest index rule or Bland's rule, named after its originator R. G. Bland [1977]. Other methods for preventing cycling use the perturbation method and the related lexicographic entering rule. For a discussion of these methods see Chvátal [1983], Dantzig [1963], and Quey [1985].

Before stating the smallest index rule, we recall that we pointed out in Remark 4.3 that the choice of entering variable could be any variable corresponding to a positive component of $\bar{b}_k$. Also, it is possible to choose any variable x_i at which r^* in (21) of Section 4 is attained.

Definition 5.1 (Smallest Index Rule). From among all possible entering variables, choose the one with smallest index. From among all possible leaving variables, choose the one with smallest index.

Remark 5.1. According to the smallest index rule, the choice of entering variable will, in general, be different from the choice of index such that $\bar{b}_k$ is a maximum. The choice of leaving variable is the same as that previously used.

Theorem 5.1. If at each iteration of the simplex method the smallest index rule is used, then cycling will not occur.

Proof. We shall prove the theorem by showing that the use of the smallest index rule and the assumption of cycling lead to a contradiction. The argument that we shall give follows the one by Chvátal [1983].

Let us assume that cycling occurs and that $\bar{B}_0, \bar{B}_{i_1}, \dots, \bar{B}_{i_p}$ is a sequence of bases corresponding to iterations $i_0, i_1, \dots, i_p$, and such that $\bar{B}_{i_{p+1}} = \bar{B}_{i_0}$. Since cycling occurs, all of these bases are degenerate.

Let J denote the set of variables that are basic in some iterations and nonbasic in others. Let x_i denote the variable in J with the greatest index. If x_i is basic in $\bar{B}_{i_0}$, then there must exist an iteration i_{i_1} with $0 < i < p$ such that x_i leaves $\bar{B}_{i_{i_1}}$ and is nonbasic at the next iteration. If x_i is nonbasic in $\bar{B}_{i_0}$, then there must exist a basis $\bar{B}_{i_{i_1}}$ with $0 < i < p$ such that x_i enters $\bar{B}_{i_{i_1}}$. Since $\bar{B}_{i_{p+1}} = \bar{B}_{i_0}$, there exists a basis $\bar{B}_{i_{i_2}}$ with $i < j < p$ such that x_i leaves $\bar{B}_{i_{i_2}}$ and is nonbasic at the next iteration. In conclusion, we have shown that there is a basis $\bar{B}$ in the sequence $\bar{B}_0, \dots, \bar{B}_{i_p}$ such that x_i leaves $\bar{B}$ at the iteration J corresponding to $\bar{B}$ and becomes nonbasic in the next basis. Let x_j denote the entering variable at the iteration J . Thus, $x_j \notin \bar{B}$ and x_j is in the next basis. Since cycling occurs, there is basis $\bar{B}^*$ at an iteration i^* beyond J and in the sequence $\bar{B}_0, \bar{B}_{i_1}, \dots, \bar{B}_{i_p}, \bar{B}_{i_{p+1}}, \dots, \bar{B}_{i_{i_p}}$ such that x_j enters the basis at i^* . That is, x_j is not in $\bar{B}^*$ but is in the succeeding basis.

We shall write $i \in \bar{B}$ to mean that x_i is a basic variable in the basis $\bar{B}$ and $i \notin \bar{B}$ to mean that x_i is nonbasic. If we denote the elements of the matrix $B^{-1}b$

in relation (4) of Section 4 by x_j , then we can write (4) and (5) of Section 4 as

$$\begin{aligned}x_i &= \zeta_i + \sum_{j \in B} a_{ij} x_j \quad i \in B_1 \\x &= \zeta + \sum_{j \in B} d_j x_j.\end{aligned}\quad (2)$$

Since x_i enters at B_1 of B , a solution of (2) is then given by

$$\begin{aligned}x_i &= \beta, \quad \beta \text{ arbitrary,} \\x_i &= 0, \quad i \in B, \quad i \neq i_1 \\x_i &= \zeta_i - a_{i_1 i} \beta \quad i \in B_1 \\x &= \zeta + d_{i_1} \beta.\end{aligned}\quad (3)$$

For the basic B^* equations (4) and (5) in Section 4 become

$$\begin{aligned}x_{B^*} &= \zeta_{B^*} + (B^*)^{-1} B^* x_B^*, \\x &= \zeta^* + \sum_{j \in B^*} d_j^* x_j,\end{aligned}\quad (4)$$

where $d_j^* = 0$ for $j \in B^*$.

The systems (2) and (4) are equivalent; that is, they have the same solution sets. In particular, (3) is a solution of (4). Since all of the bases from B to B^* are degenerate, the value of the objective function does not change from iteration to iteration. Hence, $\zeta^* = \zeta$. Equating the expressions for x in (2) and (4) and using (3) and $\zeta^* = \zeta$ give

$$\zeta + d_i x = \zeta + d_{i_1}^* \beta + \sum_{j \in B^*} d_j^* (\zeta_j - a_{i_1 j} \beta) - a_{i_1 i} \beta.$$

Hence

$$\left(d_i - d_{i_1}^* + \sum_{j \in B^*} d_j^* a_{i_1 j} \right) \beta = \sum_{j \in B^*} d_j^* \zeta_j \quad \text{for all } \beta.$$

Since the right-hand side is constant and the left-hand side holds for all β , we must have

$$d_i - d_{i_1}^* + \sum_{j \in B^*} d_j^* a_{i_1 j} = 0. \quad (5)$$

Since x_i enters B at iteration k , $d_i = 0$. Also, since x_i is leaving, $i \neq i_1$. Hence $i < i_1$. Since x_{i_1} is not entering at B^* and $i < i_1$, it follows that $d_{i_1}^* < 0$. For

otherwise, by the smallest index rule, x_i would not be the entering variable at P^r . It then follows from (3) that there exists an index r in B such that

$$d_i^r x_{i_0} < 0. \quad (5)$$

Since $d_i^r \neq 0$, the variable x_i is not in the basis B^r . But $r \in B$, so $x_r \in B$. Hence $r \leq i$.

It follows from (11) in Section 4 that since x_i is leaving in B and x_i is entering, $x_{i_0} > 0$. Since x_i is entering at B^r , it follows that $d_i^r > 0$. Hence $d_i^r x_{i_0} > 0$. Comparing this with (5) shows that $r \neq i$. Hence $r < i$. Since $r < i$, since the variable x_i is not in B^r , and since x_i enters B^r , we conclude from the smallest index rule that $d_i^r < 0$. From (5) we conclude that $d_i^r < 0$ and

$$x_{i_0} > 0.$$

Since $x_i \notin B^r$, for basic feasible solution Q^r at B^r has component $d_i^r = 0$. Let ζ_r denote the i th component of the basic feasible solution ξ_r corresponding to the basis B . We assert that $\zeta_r = 0$. To prove this assertion, we note that because all the bases $B^1, \dots, B^r$ are degenerate, it follows from (21) of Section 4 that $d^r = 0$ at all the iterations $1, \dots, P^r$. Hence if $\zeta_r > 0$, then x_i would not leave any of the bases $B, \dots, B^r$. This contradicts the fact that $x_i \notin B^r$, and this proves the assertion.

Since $x_{i_0} > 0$ and $\zeta_r = 0$, it follows from (21) of Section 4 that x_i is a candidate for leaving the basis B . Yet, we chose x_r with $r > i$ to leave the basis, in contradiction to the smallest index rule.

Remark 5.2. Computational experience indicates that using the smallest subscript rule or the lexicographic rule instead of the maxi rule considerably increases the time required for the simplex method or the revised simplex method to find an optimal solution. Since very few instances of cycling have been reported, some computer codes do not incorporate any anticycling routines. Other computer codes introduce an anticycling routine only after a sequence of specified length of degenerate bases occurs. The anticycling routine will then lead to a nondegenerate basis, at which point the use of the largest coefficient rule can be resumed.

Exercise 5.1. Carry out the iterations of the simplex method for Example 5.1 and show that cycling occurs at the sixth iteration. Since $h_0 = (h_0, f_0, e_1) = (0, 0, 1)$ is a basic feasible solution.

Exercise 5.2. Solve the problem in Example 5.1 using the simplex method and Bland's rule.

4. PHASE I OF SIMPLEX METHOD

As noted in Section 3, phase I of the simplex method either determines a basic feasible solution or determines that none exists, in which case the problem has no optimal solution.

We begin with the linear programming problem in standard form (SLP):

$$\begin{aligned} & \text{Maximize } (b, x) \\ & \text{subject to } Ax \leq a, \quad x \geq 0. \end{aligned} \quad (1)$$

Let us first suppose that the vector a is nonnegative. We then put the problem in canonical form (CLP) by adjoining a vector of slack variables,

$$x_p = (x_{n+1}, \dots, x_{n+m})', \quad x_{n+i} \geq 0, \quad i = 1, \dots, m,$$

and write the problem as

$$\begin{aligned} & \text{Maximize } (b, x) \\ & \text{subject to } (A, I) \begin{pmatrix} x \\ x_p \end{pmatrix} = c, \quad x \geq 0, \quad x_p \geq 0. \end{aligned} \quad (2)$$

Since $x \geq 0$, the vector $(0, x_p) = (0_p, x_p)$, where 0_p is the zero vector in $\mathbb{R}^n$, is a basic feasible solution for the problem in (2). We then immediately pass to phase II with initial basic feasible solution $(0_p, x_p) = (0_p, c)$.

If some of the components of a are negative, the vector $(0_p, c)$ is still a basic solution of the system of equations in (2), but it is not feasible. We therefore resort to another device and introduce an auxiliary problem.

Auxiliary Problem

Minimize x_0 subject to

$$\begin{aligned} & (A, I, -e) \begin{pmatrix} x \\ x_p \\ x_0 \end{pmatrix} = 0 \\ & x \geq 0, \quad x_p \geq 0, \quad x_0 \geq 0, \end{aligned} \quad (3)$$

where $e = (1, \dots, 1)$ and I is the $m \times m$ identity matrix.

In component form the system in (3) is

$$\begin{aligned} & \sum_{j=1}^n a_{ij}x_j + x_{n+i} - x_0 = c_i, \quad i = 1, \dots, m, \\ & x_j \geq 0, \quad j = 0, 1, \dots, n+m. \end{aligned}$$

A vector $(\bar{b}_1, \bar{b}_2)$ is feasible for relation (2) if and only if $(\bar{b}_1, \bar{b}_2, 0)$ is feasible for the auxiliary problem. Therefore, $(\bar{b}_1, \bar{b}_2)$ is feasible for relation (2) if and only if $(\bar{b}_1, \bar{b}_2, 0)$ is optimal for the auxiliary problem.

If we write the objective function of the auxiliary problem as "Minimize $w = -x_6$ " then the auxiliary problem becomes a linear programming problem in canonical form to which we can apply phase II of the simplex method, provided we can get an initial basic feasible solution. A basic solution to the system (1) is $(\bar{b}_1, \bar{b}_2, \bar{b}_3) = (0, 0, 0)$. This solution, however, is not feasible since some components of $\bar{x}$ are negative. We can obtain a basic feasible solution from this solution by replacing an appropriate basic variable by x_6 . Before we do this in general, we shall illustrate using Example 4.1.

In solving Example 4.1, we obtained an initial feasible solution from Figure 6.1. We now ignore the figure and write the auxiliary problem in canonical form: Maximize $w = -x_6$ subject to

$$\begin{aligned} -x_6 - x_2 + x_4 & & -x_6 &= -1 \\ -x_6 + x_2 & + x_4 & -x_6 &= 0 \\ x_6 + 2x_2 & + x_4 + x_5 &= 4 \\ x_i &\geq 0, \quad i = 0, 1, \dots, 5. \end{aligned} \quad (3)$$

The vector $(\bar{b}_6, \bar{b}_2, \bar{b}_4, \bar{b}_5, \bar{b}_3, \bar{b}_1) = (0, 0, 0, -1, 0, 0)$ is a basic solution but is not feasible. The only negative entry on the right-hand side is -1 and occurs in the first equation. Therefore, if in the augmented coefficient matrix of (4) we pivot on x_6 in row 1, then the right-hand side of the resulting system should be positive. Performing the indicated pivot gives

$$\left\{ \begin{array}{ccccccc} x_1 & x_2 & x_3 & x_4 & x_5 & x_6 & w \\ 1 & 1 & -1 & 0 & 0 & 1 & 1 \\ 0 & 2 & -1 & 1 & 0 & 0 & 1 \\ 2 & 2 & -1 & 0 & 1 & 0 & 5 \end{array} \right\} \quad (5)$$

Thus, $(\bar{b}_6, \bar{b}_2, \bar{b}_3, \bar{b}_4, \bar{b}_5, \bar{b}_1) = (1, 0, 0, 0, 1, 5)$ is a basic feasible solution. From (5) we also get that

$$w = -x_6 = -1 + x_2 + x_5 - x_4. \quad (6)$$

We can now start phase II for the problem of minimizing (6) subject to the constraining equations given by (5). We leave this to the reader and return to the general situation.

We now obtain a system equivalent to (4) and having a basic feasible solution. Let x_p denote the most negative component of $\bar{x}$. Then,

$$x_p = \min\{x_i \mid i = 1, \dots, m\}.$$

If there are more than one such components, choose the one with smallest index. Next, pivot on a_{ip} in the p th row of the system (7). The result will be an equivalent system

$$(x_1', x_2, \dots, x_{p-1}, -x_p, x_{p+1}, \dots, x_m, x_p) \begin{pmatrix} b_1 \\ b_2 \\ \vdots \\ b_p \\ b_{p+1} \\ \vdots \\ b_m \end{pmatrix} = c', \quad (7')$$

where $a = (1, \dots, 1)$ and

$$c'_p = -c_p \geq 0, \quad c'_i = c_i - c_p \geq 0, \quad i = 1, \dots, p-1, p+1, \dots, m.$$

The system (7') is component form is

$$\begin{aligned} \sum_{j=1}^n a'_{ij} x_j - x_{m+p} + x_{m+1} &= c'_i, & i=1, \\ \sum_{j=1}^n a'_{mj} x_j - x_{m+p} + x_m &= c'_p. \end{aligned}$$

Thus a basic feasible solution to (7') is

$$\begin{aligned} x_m &= c'_p, & x_{m+1} &= c'_i, & i=1, \dots, p-1, p+1, \dots, m, \\ x_j &= 0, & j=1, \dots, n, m+p, \end{aligned} \quad (8)$$

and

$$w = -x_m = -c'_p = -x_{m+1} + \sum_{j=1}^n a'_{mj} x_j. \quad (9)$$

The basic variables are

$$x_{m+p-1}, \dots, x_{m+2}, x_p, x_{m+1}, x_{m+2}, \dots, x_{m+p-1}. \quad (10)$$

We have transformed the auxiliary problem (4) to the following problem. Maximize w as given in (9) subject to (8) and $x_i \geq 0, i=1, \dots, n+m$. For this problem we have determined a basic feasible solution given by (8) with basic variables given by (10). We can now apply phase II to the problem in this form. If we incorporate an anticycling sub routine in our implementation of Phase II, then according to Lemma 5.1, we shall, in a finite number of steps, determine that the auxiliary problem either has no solution or determine an optimal solution to the auxiliary problem. If the auxiliary problem has no solution, then the original linear programming problem (1) has no feasible solution. For, as we saw, a vector $(\bar{x}, \bar{z})$ is feasible for the original problem if and only if $(\bar{x}, \bar{z}, 0)$ is optimal for the auxiliary problem.

The remaining alternative is that the auxiliary problem has a solution. We now explore the possibilities for this alternative. The variable x_m is basic initially. In applying phase II of the simplex method to the auxiliary problem,

we note that our rule for choosing the leaving variable is such that if at some iteration x_k is a candidate for leaving the basis, then x_k will be chosen to be the leaving variable. If at some iteration k , the variable x_k leaves the basis, then at the iteration k_{+1} , the variable x_k will have the value zero. The solution $(\bar{b}_k, \bar{b}_{k+1}, \bar{b}_{k+2})$ of the constraining equations at k_{+1} will also be a solution of (2). Hence if $\bar{b}_k = 0$, then $w = 0$ and $(\bar{b}_k, \bar{b}_{k+1})$ is optimal for relation (1). Thus, if x_k ever leaves the basis at an iteration k , at the next iteration k_{+1} , we will have a solution to the auxiliary problem and hence a basic feasible solution for relation (2). We can then begin phase II with this basic feasible solution.

An optimal solution could also conceivably occur under the following circumstances:

- (i) x_k basic, $w = 0$, and
- (ii) x_k basic, $w \neq 0$.

If (i) occurs, then the original problem has no feasible solution. We now show that (ii) is impossible. If (i) did occur, then $x_k = -w = 0$. Since the previous iteration did not end in an optimal solution, $x_k \neq 0$ at the start of the previous iteration and became zero in the previous iteration. If this happened, then x_k was a candidate for leaving the basis but did not do so. This contradicts the rule for determining whether x_k leaves the basis.

The next lemma follows from the discussion in this section.

LEMMA 5.1. *If an anticycling routine is included, then phase I of the simplex method either determines that no feasible solution exists or determines a basic feasible solution.*

LEMMA 5.1 and 5.2 imply the following.

THEOREM 5.1. *The simplex method incorporating an anticycling routine applied to a linear programming problem determines one of the following in a finite number of iterations:*

- (i) No feasible solutions exist.
- (ii) The problem does not have an optimal solution.
- (iii) An optimal solution.

EXERCISE 5.1. Complete the solution of the illustrative example.

EXERCISE 5.2. Let the linear programming problem be given originally in canonical form

$$\begin{aligned} &\text{Maximize } (\mathbf{b}, \mathbf{x}) \\ &\text{subject to } A\mathbf{x} = \mathbf{d}, \quad \mathbf{x} \geq \mathbf{0}, \end{aligned}$$

Show that an appropriate auxiliary problem for phase I is

$$\begin{aligned} \text{Maximize } w &= -\sum_{i=1}^m b_i \\ \text{subject to } (A, I) \begin{pmatrix} x \\ y \end{pmatrix} &= b, \quad x \geq 0, \quad y \geq 0. \end{aligned}$$

Show here one obtains an initial basic feasible solution for the auxiliary problem.

Exercise 4.3. Solve the following linear programming problems using the simplex method:

(a) Minimize $x_1 + 3x_2$ subject to

$$\begin{aligned} x_1 + 2x_2 + 3x_3 &\leq 4, \\ x_1 + 3x_2 + x_3 &\leq 5, \\ x_i &\geq 0, \quad i = 1, 2, 3. \end{aligned}$$

(b) Minimize $3x_1 - x_2$ subject to

$$\begin{aligned} x_1 + 2x_2 &\geq 2, \\ 3x_1 + x_2 &\leq 5, \\ x_i &\leq b_i, \quad x_i \geq 0, \quad i = 1, 2, 3, 4. \end{aligned}$$

(c) Minimize $-2x_1 - x_2 + x_3 + x_4$ subject to

$$\begin{aligned} x_1 - x_2 + 2x_3 - x_4 &= 2, \\ 2x_1 + x_2 - 3x_3 + x_4 &= 8, \\ x_1 + x_2 + x_3 + x_4 &= 7. \end{aligned}$$

7. REVISED-SIMPLEX METHOD

The revised simplex method does not differ from the simplex method in principle, that is, in going from one basic solution to another in such a way that the objective function is not decreased but differs in the manner in which some of the computations are carried out. In describing the revised simplex method, we assume, as we did in our description of the simplex method in Section 4, that the problem has been formulated as in (4) of Section 3. The reader will recall from Section 4 that at each iteration a basis B is determined

and the calculation of $B^{-1}N$ is used to determine the entering and leaving variable. The revised simplex method differs from the version of the simplex method given in Section 4 in that a system (1) in Section 4 with current B is not reduced to (1.1) in Section 4 and the current basis not given by (2.2) in Section 4. In the revised simplex method the entering variable x_q is chosen as before. If A_q is the column of A corresponding to x_q , then the updated basis $\bar{B}$ is obtained from B by replacing the column A_p of the leaving variable x_p by A_q . This procedure enables us to calculate $(\bar{B})^{-1}$ from B^{-1} very easily and expedites other calculations, as we shall show.

Before we present the steps in an iteration of the revised simplex method, we take up the calculation of $(B')^{-1}$, the updated inverse. At the current iteration l let the basis B be given by

$$B = (A_{i_1} \dots A_{i_m}), \quad (1)$$

where the A_{i_k} are columns of B . We then write

$$N = (A_{j_1} \dots A_{j_n}).$$

Let x_q be the entering variable at l , determined as in Section 4, and let A_q denote the column of N corresponding to x_q . Let the leaving variable x_p be determined as in Section 4. Let $\bar{B}$ denote the matrix obtained from B by replacing the column A_{i_p} by A_q and let $\bar{N}$ denote the matrix obtained from N by replacing column A_{j_p} by A_q . Then

$$\begin{aligned} \bar{B} &= (A_{i_1} \dots A_{i_{p-1}} A_q A_{i_{p+1}} \dots A_{i_m}), \\ \bar{N} &= (A_{j_1} \dots A_{j_{p-1}} A_q A_{j_{p+1}} \dots A_{j_n}). \end{aligned} \quad (2)$$

Hence

$$\begin{aligned} \bar{B}^{-1}\bar{B} &= B^{-1}(A_{i_1} \dots A_{i_{p-1}} A_q A_{i_{p+1}} \dots A_{i_m}) \\ &= \begin{pmatrix} 1 & 0 & \dots & 0 & v_{1q} & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 & v_{2q} & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 & v_{pq} & 0 & \dots & 1 \end{pmatrix} \\ &= (v_1 \dots v_{p-1} v_p v_{p+1} \dots v_m), \end{aligned} \quad (3)$$

where $v_g = B^{-1}A_q$. Note that v_g is the g th column of $B^{-1}N$ associated with x_q and this is as in Step III of an iteration in the simplex method. From the way p is determined we have that $v_{pq} > 0$, and hence the matrix on the right in (3) is nonsingular. If we multiply both sides of (3) by $\bar{B}$ on the left, we get that

$$\bar{B} = \bar{B}(v_1 \dots v_{p-1} v_p v_{p+1} \dots v_m).$$

Since B is the product of nonsingular matrices, it is nonsingular and so is a basis. We also get

$$(B)^{-1} = (A_1 \dots A_{j-1} A_{j+1} \dots A_m)^{-1} B^{-1}. \quad (8)$$

It is readily verified that multiplying the matrix on the right-hand side of the following by the matrix on the right-hand side of (8) gives the $m \times m$ identity matrix. Hence

$$\begin{pmatrix} 1 & 0 & \dots & 0 & r_{1j} & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 & r_{2j} & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 & r_{mj} & 0 & \dots & 0 \end{pmatrix} = \begin{pmatrix} 1 & 0 & \dots & 0 & q_{1j} & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 & q_{2j} & 0 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \dots & 0 & q_{mj} & 0 & \dots & 1 \end{pmatrix} \quad (9)$$

where

$$q_{ij} = -\frac{a_{ij}}{r_{ij}}, \quad i = 1, \dots, m, \quad i \neq j, \quad q_{jj} = \frac{1}{r_{jj}}.$$

The matrix on the right in (9) is called an *eta* matrix and will be denoted by E_j . Thus we may write (9) as

$$(B)^{-1} = E_j B^{-1}, \quad (10)$$

which updates B^{-1} for the next iteration.

We summarize this discussion in the following lemma.

Lemma 7.1. Let B be a basis matrix as in (1) and let B^{-1} be a matrix as in (2). Then B^{-1} is basic and (9) holds, where E_j is the matrix on the right in (9).

Remark 7.1. In the various calculations in the simplex method the essential matrix is B^{-1} , not B . The matrix B is needed, in principle, only to calculate B^{-1} . Equation (9) enables us to calculate $(B')^{-1}$ from B^{-1} without forming B or B' . It is this fact that underlies the revised simplex method. Also note that an *eta* matrix differs from an identity matrix in one column.

We now present a generic revised-simplex method. Most computer codes for solving linear programming problems use the revised simplex method. Codes used in practice incorporate various modifications to improve accuracy, reduce the number of iterations needed, decrease the running time, and so on. For details the reader is referred to Charnal [1963], Dantzig [1963], Gass [1965], and Marquag [1961].

As in our presentation of the simplex method we assume that the problem is in canonical form (Maximize z subject to

$$Ax = b, \quad \langle b, e \rangle = z = 0, \quad x \geq 0. \quad (7)$$

Let the initial stage or iteration of phase III be called the zeroth iteration. At the k th iteration, $k = 0, 1, 2, \dots$, let $B(k)$ denote the basis, let $N(k)$ denote the submatrix of A complementary to $B(k)$, let $x_{B(k)}$ denote the vector of basic variables, let $x_{N(k)}$ denote the vector of nonbasic variables, let $z(k) = \langle b_{B(k)}, x_{B(k)} \rangle$ be the corresponding basic feasible solution, and let $z(k)$ be the current value of the objective function. Equations (4)–(6) from Section 4 yield, for $k = 0, 1, 2, \dots$,

$$\begin{aligned} x_{B(k)} &= x_{B(k)} - [B(k)]^{-1} N(k) x_{N(k)}, \\ z(k) &= \langle b_{B(k)}, x_{B(k)} \rangle, \\ z &= z(k) - \langle b_{B(k)} - b_{B(k)} [B(k)]^{-1} N(k), x_{N(k)} \rangle. \end{aligned} \quad (8)$$

From Lemma 7.1 we get that

$$[B(k+1)]^{-1} = E_{k+1} [B(k)]^{-1}, \quad k = 0, 1, 2, \dots,$$

where E_{k+1} is an approximate eta matrix. Proceeding inductively, we get

$$[B(k+1)]^{-1} = E_{k+1} E_k \cdots E_1 [B(0)]^{-1} = E_k. \quad (9)$$

A sequential file in the computer storing the matrices $E_1, \dots, E_k, \dots, E_l$ is called an *eta file*. The calculation as indicated in (9) (from right to left) is called a *forward transformation*. The matrix $[B(k)]^{-1}$ is calculated using elementary matrices as in (11) from Section 3.

We now present the steps in the k th iteration, $k \geq 1$, of our generic revised simplex algorithm. In essence, the steps mirror the steps of the simplex method as presented in Section 4. Some of the calculations are different and Step 2 of the simplex method is deleted. For typographical convenience we will replace the designations $B(k)$, $N(k)$, $x_{B(k)}$, $x_{N(k)}$, ... of the current quantities by B , N , x_B , x_N , ... The updated quantities $B(k+1)$, $N(k+1)$, $x_{B(k+1)}$, ... will be designated as B' , N' , $x_{B'}$, $x_{N'}$, ... As in (7) from Section 4 let

$$\bar{x}_k = b_k - b_k B^{-1} x = (x_{k+1}, \dots, x_n), \quad (10)$$

where $x_{k+1}, \dots, x_n$ are the components of x that constitute x_N . We may then write the last equation in (8) as

$$z = z_k + \langle \bar{b}_k, x_k \rangle = z_k + \sum_{j=k+1}^n \bar{c}_j x_j,$$

Step 1. Calculate $\bar{b}_j$. This is done in two steps.

(i) Calculate

$$w = B_0^T B_1 B_2 \cdots B_k [B(k)]^{-1}$$

from left to right, that is, as

$$1 \cdots (B_0^T B_1 B_2 \cdots B_k) \cdots (B_k [B(k)]^{-1}).$$

This backward computation involves repeated multiplications of a vector by a matrix, whereas the forward calculation (from right to left) involves repeated multiplications of matrices. Clearly, the backward transformation involves fewer operations.

(ii) Calculate

$$\bar{b}_j = b_{ij} - (c_j, \bar{A}_j), \quad j = m + 1, \dots, n,$$

where the i_j are the indices of the nonzero variables at the i -th stage and the $\bar{A}_j$ are the corresponding columns of N .

Step 2. If all the $\bar{b}_j$ are less than or equal to zero, the current basic feasible solution is optimal and the algorithm terminates. If the algorithm terminates, then we use (8) to calculate z_0 , the optimal value.

If the algorithm does not terminate at Step 2, we go to Step 3.

Step 3. Choose an entering variable.

Choose an index j_0 such that $\bar{b}_{j_0} > 0$. This may be done according to the smallest index rule, the ratio $\bar{b}_{j_0}$ rule, or any other rule, as long as $\bar{b}_{j_0} > 0$. We denote the index j_0 by q and take x_q to be the entering variable. The column $\bar{A}_q$ will be the entering column for the basis at the $(k + 1)$ -st iteration.

Step 4. Either determine that the problem is unbounded or determine the leaving variable.

(i) Calculate

$$v_j = B_1 B_2 \cdots B_k [B(k)]^{-1} \bar{A}_q$$

from right to left. Note that v_j is as in Section 4.

(ii) If all components of v_j are nonpositive, then the problem is unbounded and no solution exists. The justification for this conclusion is the same as the one given in Step 3(ii) of Section 4.

- (iii) If $s = (s_{x_1}, \dots, s_{x_m})$ has at least one positive component, then as in (21) of Section 4 let p denote the smallest index r of which

$$r^* = \min \left\{ \frac{s_r}{a_{rp}} : a_{rp} > 0 \right\}$$

is attained. Choose x_p to be the leaving variable.

Step 5. Calculate E_{k+1} as given by the matrix on the right in (7).

Step 6. Update the data for the $(k + 1)$ st iteration.

- (i) Update x_B to $x_{B'}$ by replacing x_{x_p} by x_p .
- (ii) Update x_N to $x_{N'}$ by replacing x_{x_p} by x_{x_p} .
- (iii) Update B to B' by replacing B_{x_p} by B_p .
- (iv) Update N to N' by replacing B_p by B_{x_p} .
- (v) Update $(\bar{b}_B, \bar{b}_N)$ to $(\bar{b}_{B'}, \bar{b}_{N'})$ using (22) from Section 4.
- (vi) Update B^{-1} to $(B')^{-1}$ using (8).
- (vii) Update $\bar{b}_B$ and $\bar{b}_N$ to $\bar{b}_{B'}$ and $\bar{b}_{N'}$.

Step 7. Begin iteration $k + 1$ using the updated quantities.

Remark 7.2. If the initial basis consists of the slack variables, then $B(0) = I$ and (7) becomes

$$[B(k + 1)]^{-1} = E_{k+1}B_1 \dots B_k B.$$

Remark 7.3. The via file can become very long. In some computer codes, once the via file reaches a predetermined length, say $k + 1$, then $B(k + 1)$ is set as the initial basis $B(0)$, the matrix $[B(k + 1)]^{-1}$ is set as $[B(0)]^{-1}$, and a new via file is constructed.

Exercise 7A. Solve the linear programming problem in Example 1.1 using the revised simplex method. Start with phase I.

Exercise 7B. Solve the linear programming problem in Exercises 5.5(a), (b), and (c) using the revised simplex method.

BIBLIOGRAPHY

1. Barto, R. G. and D. B. Stewart, *The Elements of Real Analysis*, 3rd ed., John Wiley & Sons, New York, 1979.
2. Bazaraa, M. S., H. D. Sherali, and C. M. Shetty, *Nonlinear Programming, Theory and Algorithms*, 2nd ed., John Wiley & Sons, New York, 1993.
3. Beale, E. M. L., "Cycling in the Dual Simplex Algorithm," *Rand. Research Logistics Quarterly* 1 (1955), 204-205.
4. Beale, E. M. L., "New Finite Stopping Rules for the Simplex Method," *Mathematics of Operations Research* 2 (1977), 108-107.
5. Charnoi, T., *Linear Programming*, W. H. Freeman, New York, 1964.
6. Dantzig, G. B., *Linear Programming and Extensions*, Princeton University Press, Princeton, N.J., 1963.
7. Dantzig, M., *Games of Strategy: Theory and Applications*, Prentice-Hall, Englewood Cliffs, N.J., 1965.
8. Fedotkin, I., "Über die Theorie der reellen Ungleichungen," *Journal für die reine und angewandte Mathematik* 129 (1961), 1-27.
9. Fiacco, A. V. and G. P. McCormick, *Nonlinear Programming: Sequential Unconstrained Minimization Techniques*, John Wiley & Sons, New York, 1968.
10. Fleming, W., *Functions of Several Variables*, 2nd ed., Springer, New York, 1977.
11. Gale, D., *The Theory of Linear Economic Models*, McGraw-Hill, New York, 1960.
12. Gill, S. L., *Linear Programming: Methods and Applications*, 3d ed., McGraw-Hill, New York, 1993.
13. Golub, G. H. and C. F. Van Loan, *Matrix Computations*, Johns Hopkins University Press, Baltimore, 1980.
14. Gordan, P., "Über die Auflösungen linearer Gleichungen mit reellen Koeffizienten," *Mathematische Annalen* 6 (1873), 23-28.
15. John, N. P. and B. L. Womersley, "Trust-Region Functions in Nonlinear Programming," *Mathematical Programming* 17 (1979), 231-267.
16. John, F., *Extremum Problems with Inequalities as Subsidiary Conditions*, Studies and Essays: Courant Anniversary Volume (R. G. Feferick, G. B. Neugebauer, and I. I. Stekol, Eds.), Interscience, New York, 1948, pp. 87-104.
17. Karush, W., "Minima of Functions of Several Variables with Inequalities as Side Conditions," Master's Dissertation, University of Chicago, Chicago, IL, Dec. 1939.
18. Kuhn, H. W. and A. W. Tucker, *Nonlinear Programming: Proceedings of the Second Berkeley Symposium on Mathematical Statistics and Probability* (J. Neyman, Ed.), University of California Press, Berkeley, CA, 1961, pp. 481-482.

- Marquardt, D. L. *Nonlinear Programming*, McGraw-Hill, New York, 1969.
- Moffatt, E. J., "The Lagrange Multiplier Rule," *American Mathematical Monthly* 66 (1959), 922-924.
- Mostafa, T. K., "Störungs- und Störungs-Grenzen Ungleichungen," *Integration* (Frankfurt), 1976.
- Marshall, R. A., *Advanced Linear Programming: Computation and Practice*, McGraw-Hill, New York, 1969.
- Pontryagin, L. S., "An Indirect Variational Proof for the Problem of Lagrange with Differential Inequalities on Side Conditions," *Transactions of the American Mathematical Society* 74 (1953), 171-176.
- Rockafellar, R. T., *Convex analysis*, Princeton University Press, Princeton, N.J., 1970.
- Rudin, W., *Principles of Mathematical analysis*, 3rd ed., McGraw-Hill, New York, 1966.

INDEX

- Affine hull, 74
- Affine transformations, 78–114
- Angles of sets, 4–5
- Alternative theorems
 - convex functions, 113–117
 - convex sets, linear inequalities, 81–88
 - Farkas's lemma, 43–45, 87–88
 - Gale's theorem, 88
 - Gordan's theorem, 85–87
 - Motzkin's theorem, 89–97
- Anticipating cycles, simplex method
 - phase I procedure, 124
 - termination and cycling, 130, 134
- Backward transformation, revised simplex method, 126
- Barzilai's algorithm, 74
- Basic solution, simplex method, 181
- Basic variables
 - phase I procedure, 131–140
 - simplex method, 120
- Bland's rule, simplex method, termination and cycling, 146–150
- Bauer–Wurman property, 12
- Boundary points
 - convex function optimization, 147
 - supporting hyperplanes, 81–87
- Convex linear programming problem, 124–126
 - optimization, 144
 - simplex method, 126, 130
 - extreme points, feasible sets, 125–129
 - phase I, 141, 146
 - revised simplex method, 128–129
- Convexity theorems, 41–45
- Convex products, 61–68
- Convex–Schwarz inequality, 3–4
- Closed ball space, 31–36
- Closed line segments, 31
- Closed sets 4–8
- Closed point characterization, 58–59
- Convex functions, 35
- Colonial Botic game, 122, 127
- Compositions, 14–15
 - convex sets, separation theorems, 14–15
 - supporting hyperplanes, extreme points, 58–61
- Complementary slackness condition, linear programming, 148–149
- Completeness (property), real numbers, 12–13
- Concave functions
 - defined, 89
 - game theory, 127
- Consistency
 - convex programming (subadditive theory), 182, 188
 - convex programming (problem 11 CPTs), 180–181
 - strong in CPTs, 184–187
- Convexity qualifications
 - convex programming
 - feasibility conditions, 183–188
 - quadratic programming, 181
 - differentiable convex programming, 189, 193
 - Kuhn–Tucker constraint qualifications, 182–183
 - unoptimal constraints, 185–178
- Continuity, 2–11
- Convex cones, 44, 62
- Convex functions
 - alternative theorems, 113–117
 - definitions and properties, 87–110
 - differentiable functions, 109–111
 - differentiable convex programming, 109–111

- Convex functions (Continued)
 - unit theory applications, 117–123
 - linear's inequality, 98–99
 - optimization problems, 137–150
 - partial derivatives, 95, 108–109
 - subgradients, 102–104
- Convex hulls, 40–49
 - extreme points, 41–48
- Convex programming, duality
 - primal's interpretation, 205–210
 - Lagrangian duality, 200–203
 - linear programming, 210–223
 - quadratic programming, 232
- Convex programming problem I (CPI), problem statement, 179–180
- Convex programming problem II (CPII), problem dual to (CPII), 208–207
 - problem statement, 180–184
- Convex sets
 - affine geometry, 65–68
 - hyperplane, coefficients, 74
 - dependence, 71–72
 - half-directions, 74
 - independence, 72–75
 - linear manifold, 69–71
 - linear transformations, 70
 - non-void interiors, 76–77
 - parallel subspaces, 70
 - simplex directions, 74
 - extreme points, 76–80
 - hyperplane support, 80–81
 - linear' dependence, systems of, 41–48
 - Farkas's lemma, 42–43, 47
 - Gale's lemma, 48
 - Gordan's lemma, 44–47
 - Minkowski's lemma, 46, 47
 - nonempty relative interiors, 69–92
 - properties of, 35–40
 - Catastrophic theorem, 41–42
 - convex sets, 44
 - convex hull, 40–41
 - half spaces, 55–56
 - planes, convex combinations, 45–48
 - vector combinations, 37
 - relative interior points, 69–81
 - separation theorems, 42–54
 - disjoint convex sets, 53–54, 51
 - proper separation, 46, 50
 - strict separation, 46
 - strong separation, 47–49, 51
- Convex form, simplex method, 234–240
- Convex value simplex method, 241–245
- Cycling, simplex method, 246–248
- Disjunctive hull solution, simplex method, 245, 248–247
- Derivatives
 - convex functions, 94–96, 99
 - Hessian's response (H),
 - differentiation, 91–93
 - linear transformations, 20–22
 - second derivative rule, differentiable
 - unconstrained problems, 100–104
- Determinability
 - convex functions, 95, 106–111
 - Hessian's response (H), linear transformations, 34
- Differentiable nonlinear programming
 - first-order conditions, 145–148
 - second-order conditions, 143–179
- Differentiable unconstrained problems, (CPI–14)
 - linear' separation, 134–135
 - local minimizers, 131
 - refined non-optimal approximations, 130, 134
 - positive semidefinite functions, 131–132
 - second derivative rule, 131
 - strict minimizers, 130–134
- Differentiation, Hessian's response (H), 91–93
- Disjunctive derivation, convex functions, 44–46
- Disjoint convex sets, hyperplane separation, 41
- Duality theorems
 - convex programming
 - dual quadratic programming (DQP), 210–223
 - primal's interpretation, 207–210
 - Lagrangian duality, 200–207
 - linear programming, 211–221
 - zero, convex programming, 202–203
 - quadratic programming, 232
- Elementary operations, simplex method, 232, 236
- Empty set, 4–5
- Feasible variable
 - vertex simplex method, 234–240
 - simplex method, 239–241
- Epigraphs, convex functions, 68–69
- Equality constraints
 - convex programming problem (CPI), 180
 - differentiable nonlinear programming
 - necessary conditions, 146–148, 149–171
 - sufficient conditions, 159–168, 171–177
- Epistemic plane, 170

- Equivalent norms, Euclidean versus (ℓ^p), 34–35
- For the revised simplex method, 298, 299
- The norm, revised simplex method, 297–300
- Euclidean norm, 2–3
- Euclidean versus (ℓ^p)
 - defined, 1
 - matrix topology, 1, 3
 - vectors, 1–4
- Extreme points
 - convex sets, 56–62, 65
 - simplex method, feasible sets, 223–228
- Farkas's lemma
 - differentiable nonlinear programming, 162–163
 - linear inequalities, 34–36, 67–68
 - linear programming, 142
 - linear programming duality, 225
- Feasible sets
 - optimization problems, 129
 - simplex method
 - initially problem, 131–135
 - interim points, 225–228
- Finite-iteration property, 15
- First-order conditions, differentiable nonlinear programming, 145–152
- Forward transformation, revised simplex method, 298, 299
- Frechet differentiable, 34
- Free variables, simplex method, 299
- Fritz-Kuhn theorem, differentiable nonlinear programming, 149–163
 - second-order conditions, 171–174
- Functional extension theorem, 18–20
- Gale's theorem, linear inequalities, 68
- Gauss-Jordan, 127–127
 - optimal pivot strategies, 126–127
 - optimal pivot strategies, 125–126
- Gauge functions, 66
- Gauss elimination, 202–204
- Gauss-Jordan elimination, 201–202
- Gordan's interpretation, convex programming duality, 205–206
- Gordan's theorem
 - generalization to convex functions, 210
 - linear inequalities, 64–67
 - quadratic programming, 211–213
- Half-space first stepsize, 38
- Half space, 34–36
 - intersection theorem, 35–36
- Hessian matrix
 - differentiable convex functions, 110–111
 - differentiable unconstrained problems, 110–112
- Hyperplane
 - convex functions subgradients, 102–106
 - convex programming (duality, geometric interpretation), 208–210
 - convex set properties, 36, 38
 - Farkas's lemma, 67
 - definition, 33–34
 - supporting hyperplanes, 34–36, 62, 65
- Implicit function theorem, 163–164
- Infimum, 12–14
- Inner product, 2
- Interior points
 - continuity of convex functions, 99–100
 - convex sets
 - α -dimensional, 94
 - relative interior points, 90–93
- Jordan matrix, 20, 201
- Jordan's inequality, convex functions, 99–100
- Karush-Kuhn-Tucker theorem, differentiable nonlinear programming, 159–161
 - constraint qualifications, 154–160
 - second-order conditions, 169–172
 - slackness conditions, 159–162
- Krein-Milman theorem, 66
- Kuhn-Tucker matrix
 - definition, 164
 - examples, 177–180
 - is multiplicity, 164–165
 - is subgradients of value function, 164–165
- Lagrange multipliers
 - convex programming
 - in Kuhn-Tucker matrix, 164–165
 - necessary conditions, 153–164
 - optimality conditions, 167–169
 - is subgradients of value function, 164–165
 - slackness conditions, 156–160
 - differentiable nonlinear programming, first-order conditions, 149–152
 - Lagrangian duality, convex programming, 209–211
 - quadratic programming, 210–211
- Lifting method
 - revised simplex method, 298–299
 - simplex method, 298–300
- Left-hand limit, real numbers, 14

- Limit points
 - continuity, 5–11
 - microtopology, 4
- Limits, 5–11
- Linear functionals
 - Euclidean n -space ($\mathbb{R}^n$), 31–34
 - hyperplane definition, 33–34
- Linear inequalities
 - convex sets, functions of the alternative, 41–48
 - Farkas's lemma, 43–45, 47
 - Gale's theorem, 48
 - Gordan's theorem, 45–47
 - Motzkin's theorem, 46–47
- Linear inequalities, *see also* Nonlinear programming
 - convex linear programming problem, 144
 - complementary slackness condition, 140–143
 - duality in, 141–143
 - formulation, 135–140
 - necessary and sufficient conditions, 140–143
 - simplex method
 - artificial problem, 151–155
 - examples, 227–229
 - extreme points, feasible set, 226, 228
 - phase I procedure, 254–255
 - phase II procedure, 254–258
 - primal-dual, 250–254
 - primal method, 241–248
 - termination and cycling, 241–248
 - standard vs. canonical form forms, 149–149
- Linear regression, 154–155
- Linear transformation
 - as derivative, 22–26
 - Euclidean n -space ($\mathbb{R}^n$), 21
- Line segments, 8
- Linear-convex-concave functions, 98
- Local extrema
 - convex functions, 137
 - differentiable nonlinear programming, second-order sufficiency theorems, 175–178
 - differentiable unconstrained problems, 175–178
- Logarithmic convexity, 114
- Lower bounds, real numbers, 12–13
- Matrix games, 122–127
- Maximization of convex functions, 173–180
 - Minisquare approximation by orthogonal functions, 175
 - Motzkin topology: Euclidean n -space ($\mathbb{R}^n$), 5–8
- Minimizers
 - convex function optimization, 173–180
 - convex programming, necessary and sufficient conditions, 180, 188
 - quadratic programming, 213–215
 - differentiable nonlinear programming, 145–147
 - fundamental extremal theorems, 39–39
 - linear programming, 140–147
- Minkowski distance function, 97
- Mixed strategies, 126–127
- Morgan's test, 4–5
- Motzkin's theorem, linear inequalities, 46–47
- Multiplication, *see* Kronecker product; Laplace multipliers
- Necessary conditions
 - convex programming, 180–188
 - differentiable nonlinear programming
 - first-order conditions, 145–145
 - second-order conditions, 165–178
 - linear programming duality, 213–221
 - quadratic programming, 213–215
- Negative half space, convex set properties, 31–34
- Nonbasic variables, simplex method, 238
- Nonlinear programming, differentiable problems
 - first-order conditions, 145–145
 - second-order conditions, 165–178
- Nonstrict supporting hyperplane, 36–37, 39
- Null set, 4–5
- Open sets, 14–16
- Open line segments, 8
- Open sets, matrix topology, 4–8
- Optimal solution, simplex method, 241, 246–248
- Optimal strategies, *see* Game theory
- Orthogonality, 4
- Parallel hyperplanes, 34
- Parallel subspaces, 70
- Payoff matrix, 123–125
- "Pseudo factor" point, differentiable nonlinear programming, 171–171
- Perturbation theory, 189–190
- First problem, simplex method, 240–248
- unitary problems, 252–253

- Positive definite matrix, 111, 112, 115, 126
- Positive semi-definite, 11–14
- Primal versus programming problem (PVP), 180–188
 - Lagrangian duality, 200–207
 - Primal separation, correct sets, 46
- Quadratic programming, 204–210
 - least quadratic programs (LQP), 204–211
- Real numbers: homomorphism properties, 11–14
- Relative homomorphisms, correct sets, support and separation theorems, 81–86
- Relative interior, correct sets, support and separation theorems, 86–88
- Revised simplex method, 231–260
- Right-hand side, 11–14
- Saddle point
 - convex programming
 - Lagrangian duality, 206–207
 - necessary and sufficient conditions, 182–188
 - game theory applications, 187–197
- Second derivative test, differentiable unconstrained problems, 100–111
- Second-order conditions, differentiable nonlinear programming, 161, 176
- Second-order test vector, differential nonlinear programming, 167–173
- Sensitivity vector, convex programming, get objective theory, 193
- Separation theorems, 41–54
 - alternating theorems, 44–48
 - best possible in SP, 83–85
 - disjoint correct sets, 14–16, 92
 - point and correct sets, 48–53
 - primal separation, 46, 14–16
 - weak separation, 46–47
 - strong separation, 47–49, 54
- Spectral criterion, Hessian regular (H²) basis, 10, 11
- Simplex method
 - examples, 222–226
 - correct points feasible set, 125–129
 - phase I procedure, 231–240
 - phase II procedure, 244–245
 - polynomiality, 248–254
 - revised method, 255, 260
 - termination and cycling, 245–250
- Stack variables
 - linear programming, 144
 - simplex method, phase I, 241–243
- Statistical index, *also* simplex method-I, termination and cycling, 245–248
- Strict local minimum
 - defined, 128
 - differentiable nonlinear programming, second-order sufficient conditions, 170, 176
 - differentiable unconstrained problems, 136–138
- Strictly convex function
 - defined, 87–88
 - differential criteria, 106–111
- Strong separation, correct sets, 46–47
- Strong separation, correct sets, 47–48
- Subdifferentiable, correct functions, 102–110
- Subgradients, correct functions, 102–106
 - differentiable functions, 105–107
- Submodular-concave game, 129–137
- Subsequences, continuity, 5–11
- Sufficient conditions
 - convex programming, 181–187
 - for absence of duality gap, 206–207
 - quadratic programming, 211
 - differentiable nonlinear programming, 169–180
 - second-order conditions, 170–176
 - linear programming theory, 218–219
- Support function, 161
- Support function properties, 54–55, 95
- Supremum, 12–18
- Tallness method, 126
- Uniqueness criterion, differentiable nonlinear programming, 149–170
 - second-order sufficient theorems, 170–176
- Termination, simplex method, 245–250
- Theorems of the alternative: linear inequalities, 60–66
 - Farkas's lemma, 60–65, 47–48
 - Gale's theorem, 66
 - Gordan's theorem, 64–67
 - Minkowski's theorem, 60, 67
- Three-sheet property, correct functions, 93–94
- Triangle inequality, 1
- Two-person zero-sum games, 120–127
- Unbounded sets
 - revised simplex method, 260
- Simplex method, 225, 248–249
- Unconstrained problems
 - differentiable problems, 129–136

Unconstrained problems (continued)

- least squares, 134–135
- orthogonal mean square approximations, 135–136
- normal distribution test, 139–141
- order minimizers, 136, 139

Uniform continuity, 117

Upper bounds, real numbers, 71–74

Value function, perturbation theory, 189–200

Vector-valued set-valued functions, alternative theorems, 105–117

Van Neumann minimax theorem, 117–121

PURE AND APPLIED MATHEMATICS

A Wiley-Interscience Series of Texts, Monographs, and Tracts

Founded by **RICHARD COURANT**

Editors: **MYRON B. ALLEN**, **IL. DAVID A. COS**, **PETER LAX**

Editors Emeriti: **PETER HILTON**, **HARRY HOCHSTADT**, **JOHN TOLAND**

A complete list of the titles in this series appears at the end of this volume.

PURE AND APPLIED MATHEMATICS

A Wiley-Interscience Series of Texts, Monographs, and Tracts

Founded by RICHARD COURANT

Editors: MYRON H. ALLEN II, DAVID S. COX, PETER LAX

Editors Emeriti: PETER HILTON, HARRY HOFSTADT, JOHN TOLAND

ALDAMEX, FIEBELICH, and STRICKER—Abstract and Concrete Categories

ANDERSON and KIRKWOOD—Logic of Mathematics

ANDREWS and COHEN—A Posteriori Error Estimation in Finite Element Analysis

ARONIS and GOLDBERG—Continuous Differential Geometry and Its Generalizations

BARON and HALLIDAY—Statistical Analysis for Applied Science

*BATES—Geometric Algebra

BENSON—Applied Functional Analysis, Second Edition

BLOOM and GOREVSKY—Linear Operators in Spaces with an Indefinite Metric

BROOK—The Fourier-analytic theory of Quadratic Reciprocity

BURMAN, NEUMANN, and STERN—Nonnegative Matrices in Dynamic Systems

BURQUETTE—Convexity and Continuation in R^n

BYVANDTROP—Methods of Solving Singular Systems of Ordinary Differential Equations

BYRNE—Lebesgue Measure and Integration: An Introduction

*CARTER—Finite Groups of Lie Type

CASTELLO, CORIO, JURETSKY, and POLONELLA—Orthogonal Sets and Finite Methods in Linear Algebra: Applications to Matrix Calculations, Systems of Equations, Inequalities, and Linear Programming

CASTELLO, CORIO, FRECHETTI, GARCIA, and ALONSO—Modeling and Solving Mathematical Programming Models in Engineering and Science

CHATELIN—Representation of Matrices

CLARK—Mathematical Bioeconomics: The Optimal Management of Renewable Resources, Second Edition

COX—Primes of the Form $x^2 + ny^2$, Fermat, Class Field Theory, and Complex Multiplication

*CURTIS and REEDER—Representation Theory of Finite Groups and Associative Algebras

*CURTIS and REEDER—Methods of Representation Theory: With Applications to Finite Groups and Orders, Volume I

CURTIS and REEDER—Methods of Representation Theory: With Applications to Finite Groups and Orders, Volume II

DINTELEMAN—Vector Integration and Stochastic Integration in Banach Spaces

*DUNFORD and SCHWARTZ—Linear Operators

Part I.—General Theory

Part II.—Spectral Theory, Self-Adjoint Operators in Hilbert Space

Part III.—Spectral Operators

FABRIS and KEMALID—Positive Linear Systems: Theory and Applications

FOLKLAND—Real Analysis: Modern Techniques and Their Applications

FRIEDLICH and KREIDL—Linear Spaces and Information Theory

GILBESON—Tribonacci Theory and Quadratic Differentials

GRINE and KILANTZ—Function Theory of One Complex Variable

*Now available in a lower priced paperback edition in the Wiley Classics Library

†Now available in paperback.

- KRÖPFTER and HARBES—*Principles of Algebraic Geometry*
 KRULL—*Algebra*
 KROGH—*Groups and Characters*
 KRÜTZER, KRÜSS and OLGER—*Time-Dependent Problems and Difference Methods*
 KUNO and KOTLAND—*Fourier Series, Transform, and Boundary Value Problems, Second Edition*
 KUNZE—*Applied and Computational Complex Analysis*
 Volume 1, Power Series—Integration—Conformal Mapping—Location of Zeros
 Volume 2, Special Functions—Integral Transforms—Asymptotics—Continued Fractions
 Volume 3, Discrete Fourier Analysis, Cauchy Integrals, Construction of Conformal Maps, Uniform Functions
 KUTLER and WU—*A Course in Modern Algebra*
 KUTZNETZ—*Integral Equations*
 KURT—*Two-Dimensional Geometry, Variational Problems*
 KYAMIS and KIKI—*An Introduction to Metric Spaces and Fixed Point Theory*
 KYRIAKOPOULOS and KOSHEL—*Fundamentals of Differential Geometry, Volume I*
 KYRIAKOPOULOS and KOSHEL—*Fundamentals of Differential Geometry, Volume II*
 KÖRNY—*Fibonacci and Lucas Numbers with Applications*
 L&S—*Linear Algebra*
 LIONS—*An Introduction to Nonlinear Partial Differential Equations*
 MCCONNELL and MONROE—*Noncommutative Noetherian Rings*
 MEYERSON—*Functional Analysis: An Introduction to Banach Space Theory*
 NATHAN—*Perturbation Methods*
 NATHAN and NICK—*Nonlinear Oscillations*
 NIKKIN—*The Hilbert Transform of Schwartz Distributions and Applications*
 PETROV—*Geometry of Reflecting Rays and Inverse Spectral Problems*
 PRINTEP—*Spinors and Variational Methods*
 RADJ—*Measure Theory and Integration*
 RAJAGAN and SENG—*Finite Group Decompositions in Mathematical Analysis*
 RECHT—*Metric Systems and Quasiconformal Mappings*
 REYER—*Calculus of Polynomials: From Approximation Theory to Algebra and Number Theory, Second Edition*
 REICHEL—*Network Flows and Monotropic Optimization*
 ROTHMAN—*Introduction to Modern Set Theory*
 RUDIN—*Functional Analysis on Groups*
 SENG—*The Averaged Method of Iteration: Applications in Numerical Methods and Approximation*
 SENG and SENG—*The Averaged Method of Iteration*
 SIEGEL—*Topics in Complex Function Theory*
 Volume 1—Elliptic Functions and Uniformization Theory
 Volume 2—Automorphic Functions and Abelian Integrals
 Volume 3—Abelian Functions and Modular Functions of Several Variables
 SMITH and BROMANOWICZ—*Post-Modern Algebra*
 STARK—*Linear's Functions and Boundary Value Problems, Second Edition*
 STERN—*Differential Geometry*
 STOKER—*Nonlinear Vibrations in Mechanical and Electrical Systems*
 STOKER—*Water Waves: The Mathematical Theory with Applications*
 WHEELING—*An Introduction to Multigrid Methods*
 WHELAN—*Linear and Nonlinear Waves*
 YAUZBES—*Partial Differential Equations of Applied Mathematics, Second Edition*

*Now available in a lower priced paperback edition in the Wiley Classics Library.
 †Now available in paperback.