



Product Discovery

Project Planning

Product Definition

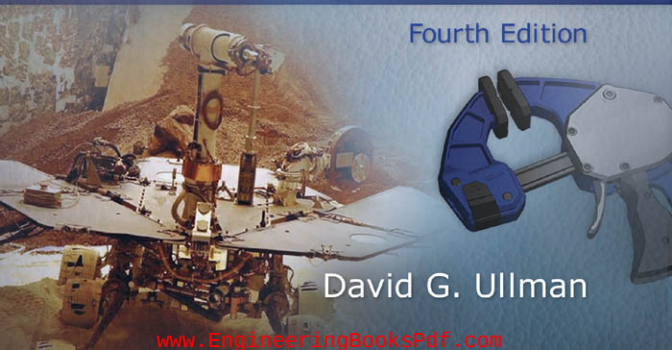
Conceptual Design

Product Development

Product Support

The Mechanical Design Process

Fourth Edition



David G. Ullman

www.EngineeringBooksPdf.com

The Mechanical Design Process

McGraw-Hill Series in Mechanical Engineering

Alciatore/Histand

Introduction to Mechatronics and Measurement System

Anderson

Fundamentals of Aerodynamics

Anderson

Introduction to Flight

Anderson

Modern Compressible Flow

Barber

Intermediate Mechanics of Materials

Beer/Johnston

Vector Mechanics for Engineers

Beer/Johnston

Mechanics of Materials

Budynas

Advanced Strength and Applied Stress Analysis

Budynas/Nisbett

Shigley's Mechanical Engineering Design

Cengel

Heat Transfer: A Practical Approach

Cengel

Introduction to Thermodynamics & Heat Transfer

Cengel/Boles

Thermodynamics: An Engineering Approach

Cengel/Cimbala

Fluid Mechanics: Fundamentals and Applications

Cengel/Turner

Fundamentals of Thermal-Fluid Sciences

Dieter

Engineering Design: A Materials & Processing Approach

Doebelin

Measurement Systems: Application & Design

Dorl/Byers

Technology Ventures: From Idea to Enterprise

Dunn

Measurement & Data Analysis for Engineering and Science

Fianemore/Franzial

Fluid Mechanics with Engineering Applications

Hamrock/Schmid/Jacobson

Fundamentals of Machine Elements

Heywood

Internal Combustion Engine Fundamentals

Holman

Experimental Methods for Engineers

Holman

Heat Transfer

Hutton

Fundamental of Finite Element Analysis

Kays/Crawford/Welgand

Convective Heat and Mass Transfer

Meirovioeh

Fundamentals of Vibrations

Norton

Design of Machinery

Palm

System Dynamics

Reddy

An Introduction to Finite Element Method

Schey

Introduction to Manufacturing Processes

Shames

Mechanics of Fluids

Smith/Hashemi

Foundations of Materials Science & Engineering

Turns

An Introduction to Combustion: Concepts and Applications

Ugural

Mechanical Design: An Integrated Approach

Ullman

The Mechanical Design Process

White

Fluid Mechanics

White

Viscous Fluid Flow

Zeid

CAD/CAM Theory and Practice

Zeid

Mastering CAD/CAM

The Mechanical Design Process

Fourth Edition

David G. Ullman

Professor Emeritus, Oregon State University



Higher Education

Boston Burr Ridge, IL Dubuque, IA New York San Francisco St. Louis
Bangkok Bogotá Caracas Kuala Lumpur Lisbon London Madrid Mexico City
Milan Montreal New Delhi Santiago Seoul Singapore Sydney Taipei Toronto



Higher Education

THE MECHANICAL DESIGN PROCESS, FOURTH EDITION

Published by McGraw-Hill, a business unit of The McGraw-Hill Companies, Inc., 1221 Avenue of the Americas, New York, NY 10020. Copyright © 2010 by The McGraw-Hill Companies, Inc. All rights reserved. Previous editions © 2003, 1997, and 1992. No part of this publication may be reproduced or distributed in any form or by any means, or stored in a database or retrieval system, without the prior written consent of The McGraw-Hill Companies, Inc., including, but not limited to, in any network or other electronic storage or transmission, or broadcast for distance learning.

Some ancillaries, including electronic and print components, may not be available to customers outside the United States.

This book is printed on acid-free paper.

1 2 3 4 5 6 7 8 9 0 DOC/DOC 0 9

ISBN 978-0-07-297574-1

MHID 0-07-297574-1

Global Publisher: *Raghothaman Srinivasan*

Senior Sponsoring Editor: *Bill Stenquist*

Director of Development: *Kristine Tibbetts*

Senior Marketing Manager: *Curt Reynolds*

Senior Project Manager: *Kay J. Brimeyer*

Senior Production Supervisor: *Sherry L. Kane*

Lead Media Project Manager: *Stacy A. Patch*

Associate Design Coordinator: *Brenda A. Rolwes*

Cover Designer: *Studio Montage, St. Louis, Missouri*

Cover Image: *Irwin clamp: © Irwin Industrial Tools; Marin bike: © Marin Bicycles; MER: © NASA/JPL.*

Senior Photo Research Coordinator: *John C. Leland*

Compositor: *S4Carlisle Publishing Services*

Typeface: *10.5/12 Times Roman*

Printer: *R. R. Donnelley Crawfordsville, IN*

Library of Congress Cataloging-in-Publication Data

Ullman, David G., 1944-

The mechanical design process / David G. Ullman.—4th ed.

p. cm.—(McGraw-Hill series in mechanical engineering)

Includes index.

ISBN 978-0-07-297574-1—ISBN 0-07-297574-1 (alk. paper)

1. Machine design. I. Title.

TJ230.U54 2010

621.8'15—dc22

2008049434

www.mhhe.com

ABOUT THE AUTHOR

David G. Ullman is an active product designer who has taught, researched, and written about design for over thirty years. He is president of Robust Decisions, Inc., a supplier of software products and training for product development and decision support. He is Emeritus Professor of Mechanical Design at Oregon State University. He has professionally designed fluid/thermal, control, and transportation systems. He has published over twenty papers focused on understanding the mechanical product design process and the development of tools to support it. He is founder of the American Society Mechanical Engineers (ASME)—Design Theory and Methodology Committee and is a Fellow in the ASME. He holds a Ph.D. in Mechanical Engineering from the Ohio State University.

CONTENTS

Preface xi

CHAPTER 1

Why Study the Design Process? 1

- 1.1 Introduction 1
- 1.2 Measuring the Design Process with Product Cost, Quality, and Time to Market 3
- 1.3 The History of the Design Process 8
- 1.4 The Life of a Product 10
- 1.5 The Many Solutions for Design Problems 15
- 1.6 The Basic Actions of Problem Solving 17
- 1.7 Knowledge and Learning During Design 19
- 1.8 Design for Sustainability 20
- 1.9 Summary 21
- 1.10 Sources 22
- 1.11 Exercises 22

CHAPTER 2

Understanding Mechanical Design 25

- 2.1 Introduction 25
- 2.2 Importance of Product Function, Behavior, and Performance 28
- 2.3 Mechanical Design Languages and Abstraction 30
- 2.4 Different Types of Mechanical Design Problems 33
- 2.5 Constraints, Goals, and Design Decisions 40
- 2.6 Product Decomposition 41
- 2.7 Summary 44

2.8 Sources 44

2.9 Exercises 45

2.10 On the Web 45

CHAPTER 3

Designers and Design Teams 47

- 3.1 Introduction 47
- 3.2 The Individual Designer: A Model of Human Information Processing 48
- 3.3 Mental Processes That Occur During Design 56
- 3.4 Characteristics of Creators 64
- 3.5 The Structure of Design Teams 66
- 3.6 Building Design Team Performance 72
- 3.7 Summary 78
- 3.8 Sources 78
- 3.9 Exercises 79
- 3.10 On the Web 80

CHAPTER 4

The Design Process and Product Discovery 81

- 4.1 Introduction 81
- 4.2 Overview of the Design Process 81
- 4.3 Designing Quality into Products 92
- 4.4 Product Discovery 95
- 4.5 Choosing a Project 101
- 4.6 Summary 109
- 4.7 Sources 110
- 4.8 Exercises 110
- 4.9 On the Web 110

CHAPTER 5**Planning for Design 111**

- 5.1 Introduction 111
- 5.2 Types of Project Plans 113
- 5.3 Planning for Deliverables—
The Development of Information 117
- 5.4 Building a Plan 126
- 5.5 Design Plan Examples 134
- 5.6 Communication During the
Design Process 137
- 5.7 Summary 141
- 5.8 Sources 141
- 5.9 Exercises 142
- 5.10 On the Web 142

CHAPTER 6**Understanding the Problem and
the Development of Engineering
Specifications 143**

- 6.1 Introduction 143
- 6.2 Step 1: Identify the Customers:
Who Are They? 151
- 6.3 Step 2: Determine the Customers'
Requirements: *What* Do the Customers
Want? 151
- 6.4 Step 3: Determine Relative Importance of the
Requirements: *Who Versus What* 155
- 6.5 Step 4: Identify and Evaluate the Competition:
How Satisfied Are the Customers *Now*? 157
- 6.6 Step 5: Generate Engineering
Specifications: *How* Will the Customers'
Requirement Be Met? 158
- 6.7 Step 6: Relate Customers' Requirements to
Engineering Specifications: *How* to Measure
What? 163
- 6.8 Step 7: Set Engineering Specification Targets
and Importance: *How* Much Is Good
Enough? 164

- 6.9 Step 8: Identify Relationships Between
Engineering Specifications: How Are the
Hows Dependent on Each Other? 166

- 6.10 Further Comments on QFD 168

- 6.11 Summary 169

- 6.12 Sources 169

- 6.13 Exercises 169

- 6.14 On the Web 170

CHAPTER 7**Concept Generation 171**

- 7.1 Introduction 171

- 7.2 Understanding the Function of Existing
Devices 176

- 7.3 A Technique for Designing with Function 181

- 7.4 Basic Methods of Generating Concepts 189

- 7.5 Patents as a Source of Ideas 194

- 7.6 Using Contradictions to Generate Ideas 197

- 7.7 The Theory of Inventive Machines, TRIZ 201

- 7.8 Building a Morphology 204

- 7.9 Other Important Concerns During Concept
Generation 208

- 7.10 Summary 209

- 7.11 Sources 209

- 7.12 Exercises 211

- 7.13 On the Web 211

CHAPTER 8**Concept Evaluation and
Selection 213**

- 8.1 Introduction 213

- 8.2 Concept Evaluation Information 215

- 8.3 Feasibility Evaluations 218

- 8.4 Technology Readiness 219

- 8.5 The Decision Matrix—Pugh's Method 221

- 8.6 Product, Project, and Decision Risk 226

- 8.7 Robust Decision Making 233
- 8.8 Summary 239
- 8.9 Sources 239
- 8.10 Exercises 240
- 8.11 On the Web 240

CHAPTER 9

Product Generation 241

- 9.1 Introduction 241
- 9.2 BOMs 245
- 9.3 Form Generation 246
- 9.4 Materials and Process Selection 264
- 9.5 Vendor Development 266
- 9.6 Generating a Suspension Design for the Marin 2008 Mount Vision Pro Bicycle 269
- 9.7 Summary 276
- 9.8 Sources 276
- 9.9 Exercises 277
- 9.10 On the Web 278

CHAPTER 10

Product Evaluation for Performance and the Effects of Variation 279

- 10.1 Introduction 279
- 10.2 Monitoring Functional Change 280
- 10.3 The Goals of Performance Evaluation 281
- 10.4 Trade-Off Management 284
- 10.5 Accuracy, Variation, and Noise 286
- 10.6 Modeling for Performance Evaluation 292
- 10.7 Tolerance Analysis 296
- 10.8 Sensitivity Analysis 302
- 10.9 Robust Design by Analysis 305
- 10.10 Robust Design Through Testing 308
- 10.11 Summary 313

- 10.12 Sources 313
- 10.13 Exercises 314

CHAPTER 11

Product Evaluation: Design For Cost, Manufacture, Assembly, and Other Measures 315

- 11.1 Introduction 315
- 11.2 DFC—Design For Cost 315
- 11.3 DFV—Design For Value 325
- 11.4 DFM—Design For Manufacture 328
- 11.5 DFA—Design-For-Assembly Evaluation 329
- 11.6 DFR—Design For Reliability 350
- 11.7 DFT and DFM—Design For Test and Maintenance 357
- 11.8 DFE—Design For the Environment 358
- 11.9 Summary 360
- 11.10 Sources 361
- 11.11 Exercises 361
- 11.12 On the Web 362

CHAPTER 12

Wrapping Up the Design Process and Supporting the Product 363

- 12.1 Introduction 363
- 12.2 Design Documentation and Communication 366
- 12.3 Support 368
- 12.4 Engineering Changes 370
- 12.5 Patent Applications 371
- 12.6 Design for End of Life 375
- 12.7 Sources 378
- 12.8 On the Web 378

APPENDIX A**Properties of 25 Materials Most Commonly Used in Mechanical Design** 379

- A.1 Introduction 379
- A.2 Properties of the Most Commonly Used Materials 380
- A.3 Materials Used in Common Items 393
- A.4 Sources 394

APPENDIX B**Normal Probability** 397

- B.1 Introduction 397
- B.2 Other Measures 401

APPENDIX C**The Factor of Safety as a Design Variable** 403

- C.1 Introduction 403

- C.2 The Classical Rule-of-Thumb Factor of Safety 405

- C.3 The Statistical, Reliability-Based, Factor of Safety 406

- C.4 Sources 414

APPENDIX D**Human Factors in Design** 415

- D.1 Introduction 415
- D.2 The Human in the Workspace 416
- D.3 The Human as Source of Power 419
- D.4 The Human as Sensor and Controller 419
- D.5 Sources 426

Index 427

PREFACE

I have been a designer all my life. I have designed bicycles, medical equipment, furniture, and sculpture, both static and dynamic. Designing objects has come easy for me. I have been fortunate in having whatever talents are necessary to be a successful designer. However, after a number of years of teaching mechanical design courses, I came to the realization that I didn't know how to teach what I knew so well. I could show students examples of good-quality design and poor-quality design. I could give them case histories of designers in action. I could suggest design ideas. But I could not tell them what to do to solve a design problem. Additionally, I realized from talking with other mechanical design teachers that I was not alone.

This situation reminded me of an experience I had once had on ice skates. As a novice skater I could stand up and go forward, lamely. A friend (a teacher by trade) could easily skate forward and backward as well. He had been skating since he was a young boy, and it was second nature to him. One day while we were skating together, I asked him to teach me how to skate backward. He said it was easy, told me to watch, and skated off backward. But when I tried to do what he did, I immediately fell down. As he helped me up, I asked him to tell me exactly what to do, not just show me. After a moment's thought, he concluded that he couldn't actually describe the feat to me. I still can't skate backward, and I suppose he still can't explain the skills involved in skating backward. The frustration that I felt falling down as my friend skated with ease must have been the same emotion felt by my design students when I failed to tell them exactly what to do to solve a design problem.

This realization led me to study the process of mechanical design, and it eventually led to this book. Part has been original research, part studying U.S. industry, part studying foreign design techniques, and part trying different teaching approaches on design classes. I came to four basic conclusions about mechanical design as a result of these studies:

1. The only way to learn about design is to do design.
2. In engineering design, the designer uses three types of knowledge: knowledge to generate ideas, knowledge to evaluate ideas and make decisions, and knowledge to structure the design process. Idea generation comes from experience and natural ability. Idea evaluation comes partially from experience and partially from formal training, and is the focus of most engineering education. Generative and evaluative knowledge are forms of domain-specific knowledge. Knowledge about the design process and decision making is largely independent of domain-specific knowledge.
3. A design process that results in a quality product can be learned, provided there is enough ability and experience to generate ideas and enough experience and training to evaluate them.

4. A design process should be learned in a dual setting: in an academic environment and, at the same time, in an environment that simulates industrial realities.

I have incorporated these concepts into this book, which is organized so that readers can learn about the design process at the same time they are developing a product. Chaps. 1–3 present background on mechanical design, define the terms that are basic to the study of the design process, and discuss the human element of product design. Chaps. 4–12, the body of the book, present a step-by-step development of a design method that leads the reader from the realization that there is a design problem to a solution ready for manufacture and assembly. This material is presented in a manner independent of the exact problem being solved. The techniques discussed are used in industry, and their names have become buzzwords in mechanical design: quality function deployment, decision-making methods, concurrent engineering, design for assembly, and Taguchi's method for robust design. These techniques have all been brought together in this book. Although they are presented sequentially as step-by-step methods, the overall process is highly iterative, and the steps are merely a guide to be used when needed.

As mentioned earlier, domain knowledge is somewhat distinct from process knowledge. Because of this independence, a successful product can result from the design process regardless of the knowledge of the designer or the type of design problem. Even students at the freshman level could take a course using this text and learn most of the process. However, to produce any reasonably realistic design, substantial domain knowledge is required, and it is assumed throughout the book that the reader has a background in basic engineering science, material science, manufacturing processes, and engineering economics. Thus, this book is intended for upper-level undergraduate students, graduate students, and professional engineers who have never had a formal course in the mechanical design process.

ADDITIONS TO THE FOURTH EDITION

Knowledge about the design process is increasing rapidly. A goal in writing the fourth edition was to incorporate this knowledge into the unified structure—one of the strong points of the first three editions. Throughout the new edition, topics have been updated and integrated with other best practices in the book. Some specific additions to the new edition include:

1. Improved material to ensure team success.
2. Over twenty blank templates are available for download from the book's website (www.mhhe.com/ullman4e) to support activities throughout the design process. The text includes many of them filled out for student reference.
3. Improved material on project planning.

4. Improved sections on Design for the Environment and Design for Sustainability.
5. Improved material on making design decisions.
6. A new section on using contradictions to generate ideas.
7. New examples from the industry, with new photos and diagrams to illustrate the examples throughout.

Beyond these, many small changes have been made to keep the book current and useful.

ELECTRONIC TEXTBOOK

CourseSmart is a new way for faculty to find and review eTextbooks. It's also a great option for students who are interested in accessing their course materials digitally and saving money. CourseSmart offers thousands of the most commonly adopted textbooks across hundreds of courses from a wide variety of higher education publishers. It is the only place for faculty to review and compare the full text of a textbook online, providing immediate access without the environmental impact of requesting a print exam copy. At CourseSmart, students can save up to 50% off the cost of a print book, reduce their impact on the environment, and gain access to powerful Web tools for learning including full text search, notes and highlighting, and email tools for sharing notes between classmates. www.CourseSmart.com

ACKNOWLEDGMENTS

I would like to thank these reviewers for their helpful comments:

Patricia Brackin, *Rose-Hulman Institute of Technology*

William Callen, *Georgia Institute of Technology*

Xiaoping Du, *University of Missouri-Rolla*

Ian Grosse, *University of Massachusetts–Amherst*

Karl-Heinrich Grote, *Otto-von-Guericke University, Magdeburg, Germany*

Mica Grujicic, *Clemson University*

John Halloran, *University of Michigan*

Peter Jones, *Auburn University*

Mary Kasarda, *Virginia Technical College*

Jesa Kreiner, *California State University–Fullerton*

Yuyi Lin, *University of Missouri–Columbia*

Ron Lumia, *University of New Mexico*

Spencer Magleby, *Brigham Young University*

Lorin Maletsky, *University of Kansas*

Make McDermott, *Texas A&M University*
Joel Ness, *University of North Dakota*
Charles Pezeshki, *Washington State University*
John Renaud, *University of Notre Dame*
Keith Rouch, *University of Kentucky*
Ali Sadegh, *The City College of The City University of New York*
Shin-Min Song, *Northern Illinois University*
Mark Steiner, *Rensselaer Polytechnic Institute*
Joshua Summers, *Clemson University*
Meenakshi Sundaram, *Tennessee Technical University*
Shih-Hsi Tong, *University of California–Los Angeles*
Kristin Wood, *University of Texas*

Additionally, I would like to thank Bill Stenquist, senior sponsoring editor for mechanical engineering of McGraw-Hill, Robin Reed, developmental editor, Kay Brimeyer, project manager, and Lynn Steines, project editor, for their interest and encouragement in this project. Also, thanks to the following who helped with examples in the book:

Wayne Collier, *UGS*
Jason Faircloth, *Marin Bicycles*
Marci Lackovic, *Autodesk*
Samir Mesihovic, *Volvo Trucks*
Professor Bob Paasch, *Oregon State University*
Matt Popik, *Irwin Tools*
Cary Rogers, *GE Medical*
Professor Tim Simpson, *Penn State University*
Ralf Strauss, *Irwin Tools*
Christopher Voorhees, *Jet Propulsion Laboratory*
Professor Joe Zaworski, *Oregon State University*

Last and most important my thanks to my wife, Adele, for her never questioning confidence that I could finish this project.

1

C H A P T E R

Why Study the Design Process?

KEY QUESTIONS

- What can be done to design quality mechanical products on time and within budget?
- What are the ten key features of design best practice that will lead to better products?
- What are the phases of a product's life cycle?
- How are design problems different from analysis problems?
- Why is it during design, the more you know, the less design freedom you have?
- What are the Hanover Principles?

1.1 INTRODUCTION

Beginning with the simple potter's wheel and evolving to complex consumer products and transportation systems, humans have been designing mechanical objects for nearly five thousand years. Each of these objects is the end result of a long and often difficult design process. This book is about that process. Regardless of whether we are designing gearboxes, heat exchangers, satellites, or doorknobs, there are certain techniques that can be used during the design process to help ensure successful results. Since this book is about the *process* of mechanical design, it focuses not on the design of any one type of object but on techniques that apply to the design of all types of mechanical objects.

If people have been designing for five thousand years and there are literally millions of mechanical objects that work and work well, why study the design process? The answer, simply put, is that there is a continuous need for new, cost-effective, high-quality products. Today's products have become so complex that most require a team of people from diverse areas of expertise to develop an idea into hardware. The more people involved in a project, the greater is the need for assistance in communication and structure to ensure nothing important

is overlooked and customers will be satisfied. In addition, the global marketplace has fostered the need to develop new products at a very rapid and accelerating pace. To compete in this market, a company must be very efficient in the design of its products. It is the process that will be studied here that determines the efficiency of new product development. Finally, it has been estimated that 85% of the problems with new products not working as they should, taking too long to bring to market, or costing too much are the result of a poor design process.

The goal of this book is to give you the tools to develop an efficient design process regardless of the product being developed. In this chapter the important features of design problems and the processes for solving them will be introduced. These features apply to any type of design problem, whether for mechanical, electrical, software, or construction projects. Subsequent chapters will focus more on mechanical design, but even these can be applied to a broader range of problems.

Consider the important factors that determine the success or failure of a product (Fig. 1.1). These factors are organized into three ovals representing those factors important to product design, business, and production.

Product design factors focus on the product's function, which is a description of what the object does. The importance of function to the designer is a major topic of this book. Related to the function are the product's form, materials, and manufacturing processes. Form includes the product's architecture, its shape, its color, its texture, and other factors relating to its structure. Of equal importance to form are the materials and manufacturing processes used to produce the product. These four variables—function, form, materials, and manufacturing processes—

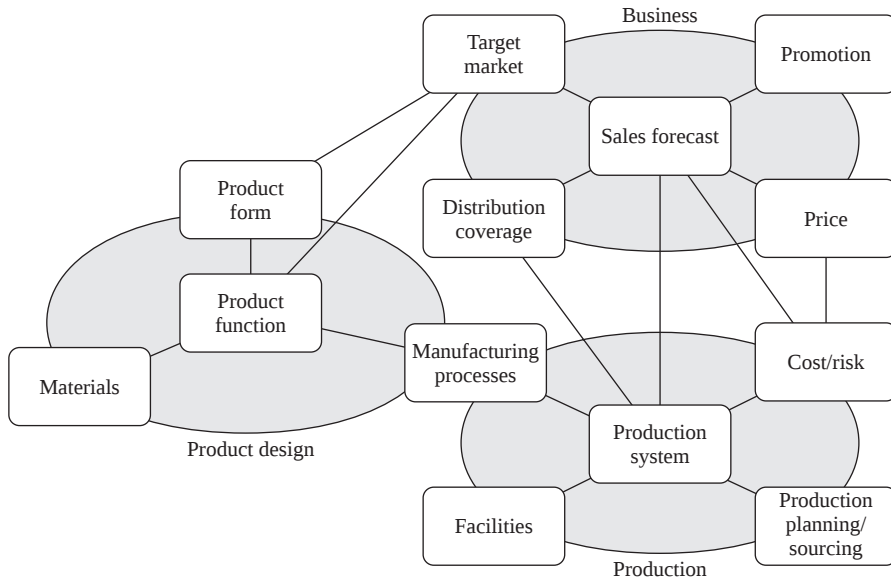


Figure 1.1 Controllable variables in product development.

are of major concern to the designer. This product design oval is further refined in Fig. 9.3.

The product form and function is also important to the business because the customers in the target market judge a product primarily on what it does (its function) and how it looks (its form). The target market is one factor important to the business, as shown in Fig. 1.1. The goal of a business is to make money—to meet its sales forecasts. Sales are also affected by the company's ability to promote the product, distribute the product, and price the product, as shown in Fig. 1.1.

The business is dependent not only on the product form and function, but also on the company's ability to produce the product. As shown in the production oval in Fig. 1.1, the production system is the central factor. Notice how product design and production are both concerned with manufacturing processes. The choice of form and materials that give the product function affects the manufacturing processes that can be used. These processes, in turn, affect the cost and hence the price of the product. This is just one example of how intertwined product design, production, and businesses truly are. In this book we focus on the product design oval. But, we will also pay much attention to the business and production variables that are related to design. As shown in the upcoming sections, the design process has a great effect on product cost, quality, and time to market.

1.2 MEASURING THE DESIGN PROCESS WITH PRODUCT COST, QUALITY, AND TIME TO MARKET

The three measures of the effectiveness of the design process are product cost, quality, and time to market. Regardless of the product being designed—whether it is an entire system, some small subpart of a larger product, or just a small change in an existing product—the customer and management always want it cheaper (lower cost), better (higher quality), and faster (less time).

The actual cost of designing a product is usually a small part of the manufacturing cost of a product, as can be seen in Fig. 1.2, which is based on data from Ford Motor Company. The data show that only 5% of the manufacturing cost of a car (the cost to produce the car but not to distribute or sell it) is for design activities that were needed to develop it. This number varies with industry and product, but for most products the cost of design is a small part of the manufacturing cost.

However, the effect of the quality of the design on the manufacturing cost is much greater than 5%. This is most accurately shown from the results of a detailed study of 18 different automatic coffeemakers. Each coffeemaker had the same function—to make coffee. The results of this study are shown in Fig. 1.3. Here the effects of changes in manufacturing efficiency, such as material cost, labor wages, and cost of equipment, have been separated from the effects of the design process. Note that manufacturing efficiency and design have about the same influence on the cost of manufacturing a product. The figure shows that

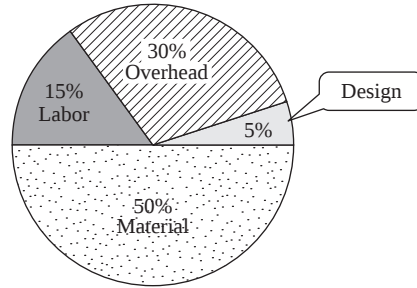


Figure 1.2 Design cost as fraction of manufacturing cost.

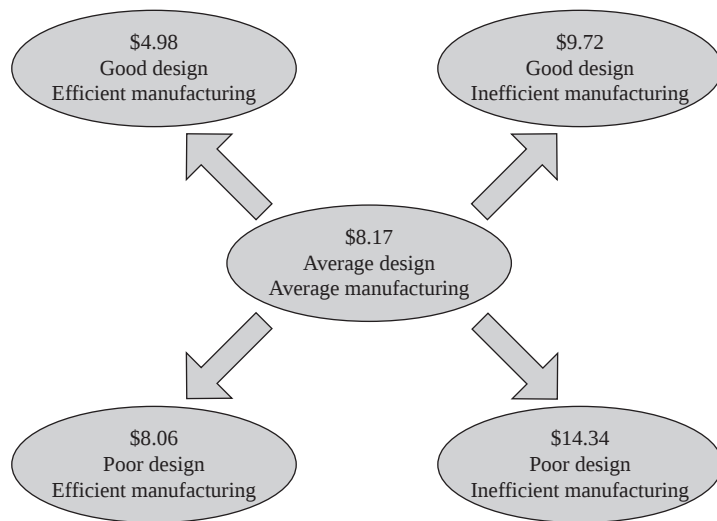


Figure 1.3 The effect of design on manufacturing cost.

(Source: Data reduced from “Assessing the Importance of Design through Product Archaeology,” *Management Science*, Vol. 44, No. 3, pp. 352–369, March 1998, by K. Ulrich and S. A. Pearson.)

Designers cost little, their impact on product cost, great.

good design, regardless of manufacturing efficiency, cuts the cost by about 35%. In some industries this effect is as high as 75%.

Thus, comparing Fig. 1.2 to Fig. 1.3, we can conclude that *the decisions made during the design process have a great effect on the cost of a product but cost very little*. Design decisions directly determine the materials used, the goods

Product cost is committed early in the design process and spent late in the process.

purchased, the parts, the shape of those parts, the product sold, the price of the product, and the sales.

Another example of the relationship of the design process to cost comes from Xerox. In the 1960s and early 1970s, Xerox controlled the copier market. However, by 1980 there were over 40 different manufacturers of copiers in the marketplace and Xerox's share of the market had fallen significantly. Part of the problem was the cost of Xerox's products. In fact, in 1980 Xerox realized that some producers were able to sell a copier for less than Xerox was able to manufacture one of similar functionality. In one study of the problem, Xerox focused on the cost of individual parts. Comparing plastic parts from their machines and ones that performed a similar function in Japanese and European machines, they found that Japanese firms could produce a part for 50% less than American or European firms. Xerox attributed the cost difference to three factors: materials costs were 10% less in Japan, tooling and processing costs were 15% less, and the remaining 25% (half of the difference) was attributable to how the parts were designed.

Not only is much of the product cost committed during the design process, it is committed early in the design process. As shown in Fig. 1.4, about 75% of the manufacturing cost of a typical product is committed by the end of the conceptual phase process. This means that decisions made after this time can influence only 25% of the product's manufacturing cost. Also shown in the figure is the amount of cost incurred, which is the amount of money spent on the design of the product.

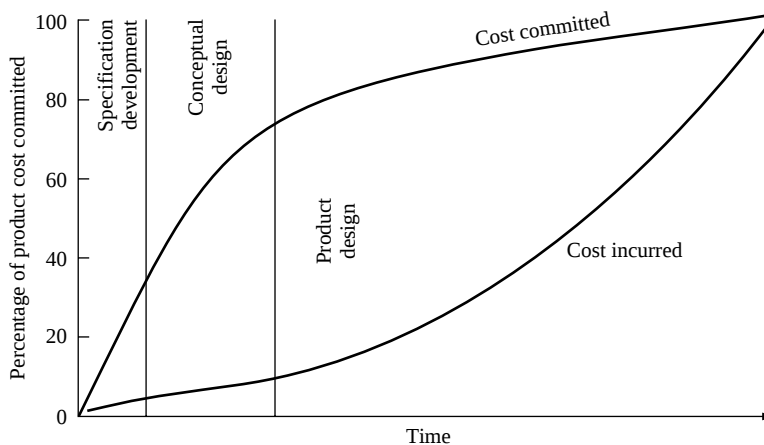


Figure 1.4 Manufacturing cost commitment during design.

Table 1.1 What determines quality

	1989	2002
Works as it should	4.99 (1)	4.58 (1)
Lasts a long time	4.75 (2)	3.93 (5)
Is easy to maintain	4.65 (3)	3.29 (5)
Looks attractive	2.95 (4–5)	3.58 (3–4)
Incorporates latest technology/features	2.95 (4–5)	3.58 (3–4)

Scale: 5 = very important, 1 = not important at all, brackets denote rank.

Sources: Based on a survey of consumers published in *Time*, Nov. 13, 1989, and a survey based on quality professional, R. Sebastianelli and N. Tamimi, “How Product Quality Dimensions Relate to Defining Quality,” *International Journal of Quality and Reliability Management*, Vol. 19, No. 4, pp. 442–453, 2002.

It is not until money is committed for production that large amounts of capital are spent.

The results of the design process also have a great effect on product quality. In a survey taken in 1989, American consumers were asked, “What determines quality?” Their responses, shown in Table 1.1, indicate that “quality” is a composite of factors that are the responsibility of the design engineer. In a 2002 survey of engineers responsible for quality, what is important to “quality” is little changed. Although the surveys were of different groups, it is interesting to note that in the thirteen years between surveys, the importance of being easy to maintain has dropped, but the main measures of quality have remained unchanged.

Note that the most important quality measure is “works as it should.” This, and “incorporates latest technology/features,” are both measures of product function. “Lasts a long time” and most of the other quality measures are dependent on the form designed and on the materials and the manufacturing process selected. What is evident is that the decisions made during the design process determine the product’s quality.

Besides affecting cost and quality, the design process also affects the time it takes to produce a new product. Consider Fig. 1.5, which shows the number of design changes made by two automobile companies with different design philosophies. The data points for Company B are actual for a U.S. automobile manufacturer, and the dashed line for Company A is what is typical for Toyota. Iteration, or change, is an essential part of the design process. However, changes occurring late in the design process are more expensive than those occurring earlier, as prior work is scrapped. The curve for Company B shows that the company was still making changes after the design had been released for production. In fact, over 35% of the cost of the product occurred after it was in production. In essence, Company B was still designing the automobile as it was being sold as a product. This causes tooling and assembly-line changes during production and the possibility of recalling cars for retrofit, both of which would necessitate significant expense, to say nothing about the loss of customer confidence. Company A, on the other hand, made many changes early in the design process and finished the design of the car before it went into production. Early design changes require more engineering time and effort but do not require changes in hardware or documentation. A change that would cost \$1000 in engineering time if made

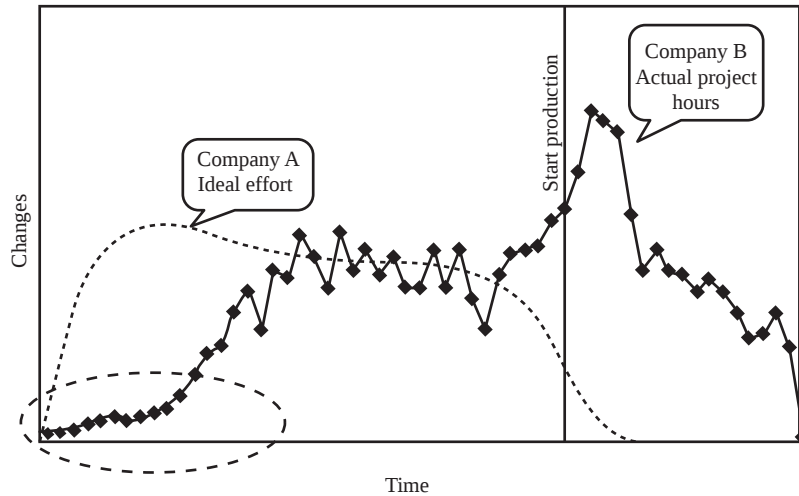


Figure 1.5 Engineering changes during automobile development.
 (Source: Data from Tom Judd, Cognition Corp., “Taking DFSS to the Next Level,”
 WCBF, Design for Six Sigma Conference, Las Vegas, June 2005.)

Fail early; fail often.

early in the design process may cost \$10,000 later during product refinement and \$1,000,000 or more in tooling, sales, and goodwill expenses if made after production has begun.

Figure 1.5 also indicates that Company A took less time to design the automobile than Company B. This is due to differences in the design philosophies of the companies. Company A assigns a large engineering staff to the project early in product development and encourages these engineers to utilize the latest in design techniques and to explore all the options early to preclude the need for changes later on. Company B, on the other hand, assigns a small staff and pressures them for quick results, in the form of hardware, discouraging the engineers from exploring all options (the region in the oval in the figure). The design axiom, *fail early, fail often*, applies to this example. Changes are required in order to find a good design, and early changes are easier and less expensive than changes made later. The engineers in Company B spend much time “firefighting” after the product is in production. In fact, many engineers spend as much as 50% of their time firefighting for companies similar to Company B.

An additional way that the design process affects product development time is in how long it takes to bring a product to market. Prior to the 1980s there was little emphasis on the length of time to develop new products. Since then competition has forced new products to be introduced at a faster and faster rate. During the 1990s development time in most industries was cut by half. This trend

has continued into the twenty-first century. More on how the design process has played a major role in this reduction is in Chap. 4.

Finally, for many years it was believed that there was a trade-off between high-quality products and low costs or time—namely, that it costs more and takes more time to develop and produce high-quality products. However, recent experience has shown that increasing quality and lowering costs and time can go hand in hand. Some of the examples we have discussed and ones throughout the rest of the book reinforce this point.

1.3 THE HISTORY OF THE DESIGN PROCESS

During design activities, ideas are developed into hardware that is usable as a product. Whether this piece of hardware is a bookshelf or a space station, it is the result of a process that combines people and their knowledge, tools, and skills to develop a new creation. This task requires their time and costs money, and if the people are good at what they do and the environment they work in is well structured, they can do it efficiently. Further, if they are skilled, the final product will be well liked by those who use it and work with it—the customers will see it as a quality product. *The design process, then, is the organization and management of people and the information they develop in the evolution of a product.*

In simpler times, one person could design and manufacture an entire product. Even for a large project such as the design of a ship or a bridge, one person had sufficient knowledge of the physics, materials, and manufacturing processes to manage all aspects of the design and construction of the project.

By the middle of the twentieth century, products and manufacturing processes had become so complex that one person no longer had sufficient knowledge or time to focus on all the aspects of the evolving product. Different groups of people became responsible for marketing, design, manufacturing, and overall management. This evolution led to what is commonly known as the “over-the-wall” design process (Fig. 1.6).

In the structure shown in Fig. 1.6, the engineering design process is walled off from the other product development functions. Basically, people in marketing communicate a perceived market need to engineering either as a simple, written request or, in many instances, orally. This is effectively a one-way communication and is thus represented as information that is “thrown over the wall.”

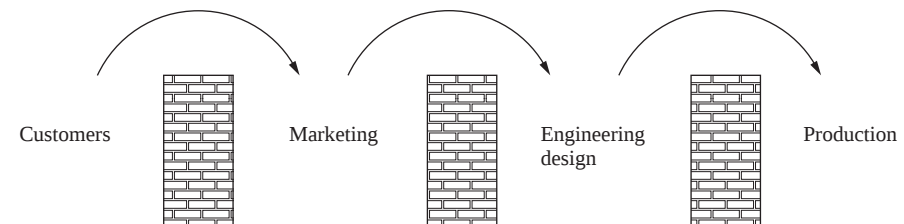


Figure 1.6 The over-the-wall design method.

Engineering interprets the request, develops concepts, and refines the best concept into manufacturing specifications (i.e., drawings, bills of materials, and assembly instructions). These manufacturing specifications are thrown over the wall to be produced. Manufacturing then interprets the information passed to it and builds what it thinks engineering wanted.

Unfortunately, often what is manufactured by a company using the over-the-wall process is not what the customer had in mind. This is because of the many weaknesses in this product development process. First, marketing may not be able to communicate to engineering a clear picture of what the customers want. Since the design engineers have no contact with the customers and limited communication with marketing, there is much room for poor understanding of the design problem. Second, design engineers do not know as much about the manufacturing processes as manufacturing specialists, and therefore some parts may not be able to be manufactured as drawn or manufactured on existing equipment. Further, manufacturing experts may know less-expensive methods to produce the product. Thus, this single-direction over-the-wall approach is inefficient and costly and may result in poor-quality products. Although many companies still use this method, most are realizing its weaknesses and are moving away from its use.

In the late 1970s and early 1980s, the concept of *simultaneous engineering* began to break down the walls. This philosophy emphasized the simultaneous development of the manufacturing process with the evolution of the product. Simultaneous engineering was accomplished by assigning manufacturing representatives to be members of design teams so that they could interact with the design engineers throughout the design process. The goal was the simultaneous development of the product and the manufacturing process.

In the 1980s the simultaneous design philosophy was broadened and called *concurrent engineering*, which, in the 1990s, became *Integrated Product and Process Design (IPPD)*. Although the terms *simultaneous*, *concurrent*, and *integrated* are basically synonymous, the change in terms implies a greater refinement in thought about what it takes to efficiently develop a product. Throughout the rest of this text, the term *concurrent engineering* will be used to express this refinement.

In the 1990s the concepts of *Lean* and *Six Sigma* became popular in manufacturing and began to have an influence on design. Lean manufacturing concepts were based on studies of the Toyota manufacturing system and introduced in the United States in the early 1990s. Lean manufacturing seeks to eliminate waste in all parts of the system, principally through teamwork. This means eliminating products nobody wants, unneeded steps, many different materials, and people waiting downstream because upstream activities haven't been delivered on time. In design and manufacturing, the term "lean" has become synonymous with minimizing the time to do a task and the material to make a product. The Lean philosophy will be refined in later chapters.

Where Lean focuses on time, Six Sigma focuses on quality. Six Sigma, sometimes written as (6σ) was developed at Motorola in the 1980s and popularized in the 1990s as a way to help ensure that products were manufactured to the highest

Table 1.2 The ten key features of design best practice

-
1. Focus on the entire product life (Chap. 1)
 2. Use and support of design teams (Chap. 3)
 3. Realization that the processes are as important as the product (Chaps. 1 and 4)
 4. Attention to planning for information-centered tasks (Chap. 4)
 5. Careful product requirements development (Chap. 5)
 6. Encouragement of multiple concept generation and evaluation (Chaps. 6 and 7)
 7. Awareness of the decision-making process (Chap. 8)
 8. Attention to designing in quality during every phase of the design process (throughout)
 9. Concurrent development of product and manufacturing process (Chaps. 9–12)
 10. Emphasis on communication of the right information to the right people at the right time (throughout and in Section 1.4.)
-

standards of quality. Six Sigma uses statistical methods to account for and manage product manufacturing uncertainty and variation. Key to Six Sigma methodology is the five-step DMAIC process (Define, Measure, Analyze, Improve, and Control). Six Sigma brought improved quality to manufactured products. However, quality begins in the design of products, and processes, not in their manufacture. Recognizing this, the Six Sigma community began to emphasize quality earlier in the product development cycle, evolving DFSS (Design for Six Sigma) in the late 1990s.

Essentially DFSS is a collection of design best practices similar to those introduced in this book. DFSS is still an emerging discipline.

Beyond these formal methodologies, during the 1980s and 1990s many design process techniques were introduced and became popular. They are essential building blocks of the design philosophy introduced throughout the book.

All of these methodologies and best practices are built around a concern for the ten key features listed in Table 1.2. These ten features are covered in the chapters shown and are integrated into the philosophy covered in this book. The primary focus is on the integration of teams of people, design tools and techniques, and information about the product and the processes used to develop and manufacture it.

The use of teams, including all the “stakeholders” (people who have a concern for the product), eliminates many of the problems with the over-the-wall method. During each phase in the development of a product, different people will be important and will be included in the product development team. This mix of people with different views will also help the team address the entire life cycle of the product.

Tools and techniques connect the teams with the information. Although many of the tools are computer-based, much design work is still done with pencil and paper. Thus, the emphasis in this book is not on computer-aided design but on the techniques that affect the culture of design and the tools used to support them.

1.4 THE LIFE OF A PRODUCT

Regardless of the design process followed, every product has a life history, as described in Fig. 1.7. Here, each box represents a phase in the product’s life.

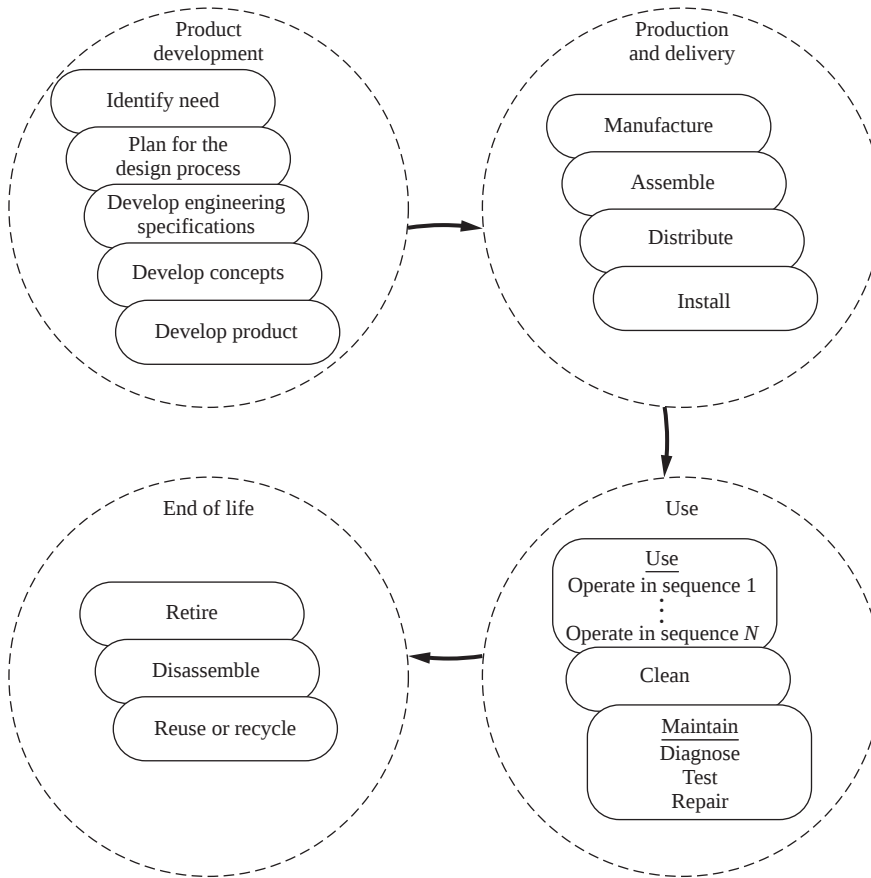


Figure 1.7 The life of a product.

These phases are grouped into four broad areas. The first area concerns the development of the product, the focus of this book. The second group of phases includes the production and delivery of the product. The third group contains all the considerations important to the product's use. And the final group focuses on what happens to the product after it is no longer useful. Each phase will be introduced in this section, and all are detailed later in the book. Note that designers, responsible for the first five phases, must fully understand all the subsequent phases if they are to develop a quality product.

The design phases are:

Identify need. Design projects are initiated either by a market requirement, the development of a new technology, or the desire to improve an existing product.

The design process not only gives birth to a product but is also responsible for its life and death.

Plan for the design process. Efficient product development requires planning for the process to be followed. Planning for the design process is the topic of Chap. 4.

Develop engineering requirements. The importance of developing a good set of specifications has become one of the key points in concurrent engineering. It has recently been realized that the time spent evolving complete specifications prior to developing concepts saves time and money and improves quality. A technique to help in developing specifications is covered in Chap. 6.

Develop concepts. Chapters 7 and 8 focus on techniques for generating and evaluating new concepts. This is an important phase in the development of a product, as decisions made here affect all the downstream phases.

Develop product. Turning a concept into a manufacturable product is a major engineering challenge. Chapters 9–12 present techniques to make this a more reliable process. This phase ends with manufacturing specifications and release to production.

These first five phases all must take into account what will happen to the product in the remainder of its lifetime. When the design work is completed, the product is released for production, and except for engineering changes, the design engineers will have no further involvement with it.

The production and delivery phases include:

Manufacture. Some products are just assemblies of existing components. For most products, unique components need to be formed from raw materials and thus require some manufacturing. In the over-the-wall design philosophy, design engineers sometimes consider manufacturing issues, but since they are not experts, they sometimes do not make good decisions. Concurrent engineering encourages having manufacturing experts on the design team to ensure that the product can be produced and can meet cost requirements. The specific consideration of *design for manufacturing* and product cost estimation is covered in Chap. 11.

Assemble. How a product is to be assembled is a major consideration during the product design phase. Part of Chap. 11. is devoted to a technique called *design for assembly*, which focuses on making a product easy to assemble.

Distribute. Although distribution may not seem like a concern for the design engineer, each product must be delivered to the customer in a safe and cost-effective manner. Design requirements may include the need for the product to be shipped in a prespecified container or on a standard pallet. Thus, the

design engineers may need to alter their product just to satisfy distribution needs.

Install. Some products require installation before the customer can use them. This is especially true for manufacturing equipment and building industry products. Additionally, concern for installation can also mean concern for how customers will react to the statement, “Some assembly required.”

The goal of product development, production, and delivery is the use of the product. The “Use” phases are:

Operate. Most design requirements are aimed at specifying the use of the product. Products may have many different operating sequences that describe their use. Consider as an example a common hammer that can be used to put in nails or take them out. Each use involves a different sequence of operations, and both must be considered during the design of a hammer.

Clean. Another aspect of a product’s use is keeping it clean. This can range from frequent need (e.g., public bathroom fixtures) to never. Every consumer has experienced the frustration of not being able to clean a product. This inability is seldom designed into the product on purpose; rather, it is usually simply the result of poor design.

Maintain. As shown in Fig. 1.7, to *maintain* a product requires that problems must be *diagnosed*, the diagnosis may require *tests*, and the product must be *repaired*.

Finally, every product has a finite life. End-of-life concerns have become increasingly important.

Retire. The final phase in a product’s life is its retirement. In past years designers did not worry about a product beyond its use. However, during the 1980s increased concern for the environment forced designers to begin considering the entire life of their products. In the 1990s the European Union enacted legislation that makes the original manufacturer responsible for collecting and reusing or recycling its products when their usefulness is finished. This topic will be further discussed in Section 12.8.

Disassemble. Before the 1970s, consumer products could be easily disassembled for repair, but now we live in a “throwaway” society, where disassembly of consumer goods is difficult and often impossible. However, due to legislation requiring us to recycle or reuse products, the need to design for disassembling a product is returning.

Reuse or recycle. After a product has been disassembled, its parts can either be reused in other products or recycled—reduced to a more basic form and used again (e.g., metals can be melted, paper reduced to pulp again).

This emphasis on the life of a product has resulted in the concept of Product Life-cycle Management (PLM). The term PLM was coined in the fall of 2001 as a blanket term for computer systems that support both the definition or authoring of product information from cradle to grave. PLM enables management

of this information in forms and languages understandable by each constituency in the product life cycle—namely, the words and representations that the engineers understand are not the same as what manufacturing or service people understand.

A predecessor to PLM was Product Data Management (PDM), which evolved in the 1980s to help control and share the product data. The change from “data” in PDM to life cycle in PLM reflects the realization that there is more to a product than the description of its geometry and function—the processes are also important.

As shown in Fig. 1.8, PLM integrates six different major types of information. In the past these were separate, and communications between the communities

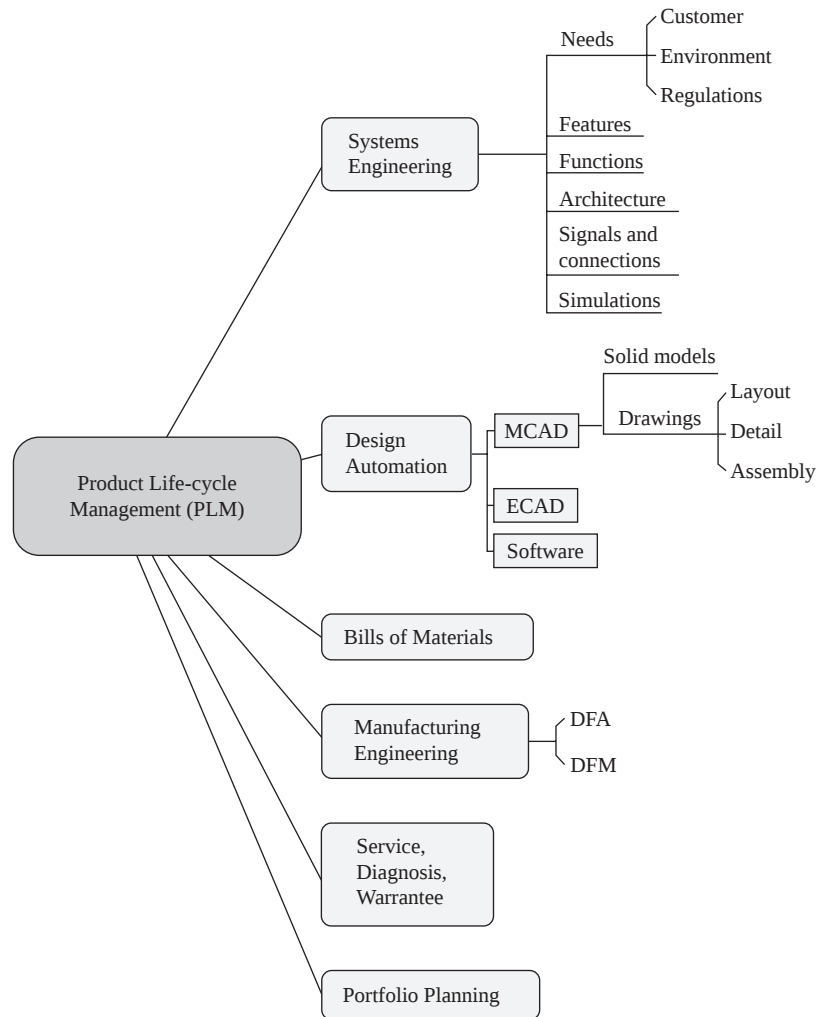


Figure 1.8 Product Life-cycle Management.

was poor (think of the over-the-wall method, Fig. 1.6). Whereas Fig. 1.7 focuses on the activities that happen during a product's life, PLM, Fig. 1.8 focuses on the information that must be managed to support that life. What PLM calls "Systems Engineering" is support for the technical development of the function of the product. The topics listed under Systems Engineering are all covered in this book.

What historically was called CAD (Computer-Aided Design) is now often referred to as MCAD for Mechanical CAD to differentiate it from Electronic CAD (ECAD). These two, along with software are all part of *design automation*. Like most of PLM, this structure grew from the twigs to the root of the tree. Traditional drawings included layout and detailed and assembly drawings. The advent of solid models made them a part of an MCAD system.

Bills Of Materials (BOMs) are effectively parts lists. BOMs are fundamental documents for manufacturing. However, as product is evolving in systems engineering so does the BOM; early on there may be no parts to list. In manufacturing, PLM manages information about Design For Manufacturing (DFM) and Assembly (DFA).

Once the product is launched and in use, there is a need to maintain it, or as shown in Fig. 1.7, diagnose, test, and repair it. These activities are supported by service, diagnosis, and warrantee information in a PLM system. Finally, there is need to manage the product portfolio—namely, of the products that could be offered, which ones are chosen to be offered (the organization's portfolio). Portfolio decisions are the part of doing business that determines which products will be developed and sold.

This description of the life of a product and systems to manage it, gives a good basic understanding of the issues that will be addressed in this book. The rest of this chapter details the unique features of design problems and their solution processes.

1.5 THE MANY SOLUTIONS FOR DESIGN PROBLEMS

Consider this problem from a textbook on the design of machine components (see Fig. 1.9):

What size SAE grade 5 bolt should be used to fasten together two pieces of 1045 sheet steel, each 4 mm thick and 6 cm wide, which are lapped over each other and loaded with 100 N?

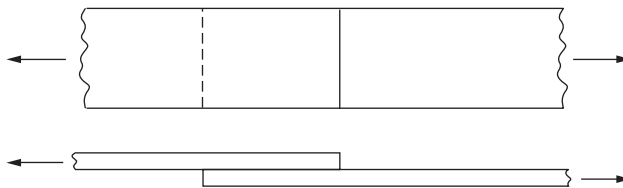


Figure 1.9 A simple lap joint.

Design problems have many satisfactory solutions but no clear best solution.

In this problem the need is very clear, and if we know the methods for analyzing shear stress in bolts, the problem is easily understood. There is no necessity to design the joint because a design solution is already given, namely, a grade 5 bolt, with one parameter to be determined—its diameter. The product evaluation is straight from textbook formulas, and the only decision made is in determining whether we did the problem correctly.

In comparison, consider this, only slightly different, problem:

Design a joint to fasten together two pieces of 1045 sheet steel, each 4 mm thick and 6 cm wide, which are lapped over each other and loaded with 100 N.

The only difference between these problems is in their opening clauses (shown in italics) and a period replacing the question mark (you might want to think about this change in punctuation). The second problem is even easier to understand than the first; we do not need to know how to design for shear failure in bolted joints. However, there is much more latitude in generating ideas for potential concepts here. It may be possible to use a bolted joint, a glued joint, a joint in which the two pieces are folded over each other, a welded joint, a joint held by magnets, a Velcro joint, or a bubble-gum joint. Which one is best depends on other, unstated factors. This problem is not as well defined as the first one. To evaluate proposed concepts, more information about the joint will be needed. In other words, the problem is not really understood at all. Some questions still need to be answered: Will the joint require disassembly? Will it be used at high temperatures? What tools are available to make the joint? What skill levels do the joint manufacturers have?

The first problem statement describes an analysis problem. To solve it we need to find the correct formula and plug in the right values. The second statement describes a design problem, which is ill-defined in that the problem statement does not give all the information needed to find the solution. The potential solutions are not given and the constraints on the solution are incomplete. This problem requires us to fill in missing information in order to understand it fully.

Another difference between the two problems is in the number of potential solutions. For the first problem there is only one correct answer. For the second there is no correct answer. In fact, there may be many good solutions to this problem, and it may be difficult if not impossible to define what is meant by the “best solution.” Just consider all the different cars, televisions, and other products that compete in the same market. In each case, all the different models solve essentially the same problem, yet there are many different solutions. The goal in design is to find a good solution that leads to a quality product with the least commitment of time and other resources. *All design problems have a multitude of satisfactory solutions and no clear best solution.* This is shown graphically

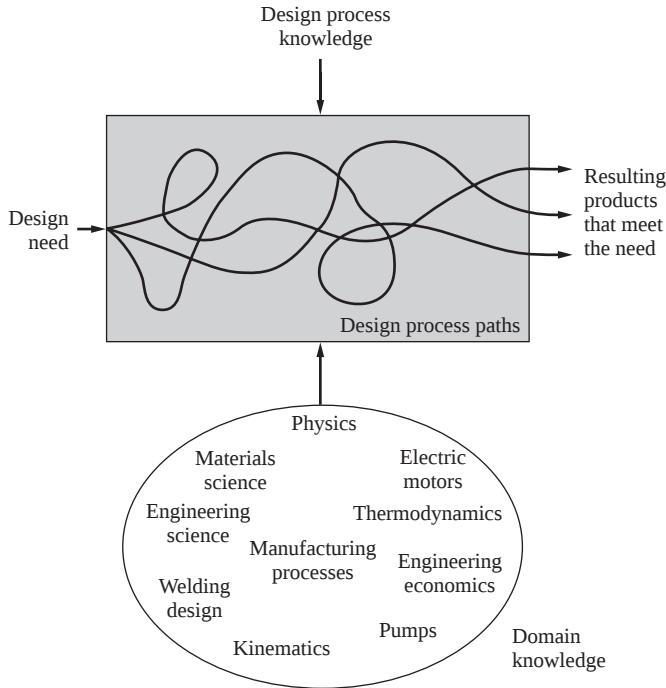


Figure 1.10 The many results of the design process.

in Fig. 1.10 where the factors that affect exactly what solution is developed are noted. Domain knowledge is developed through the study of engineering physics and other technical areas and through the observation of existing products. It is the study of science and engineering science that provides the basis on which the design process is based. Design process knowledge is the subject of this book.

For mechanical design problems in particular, there is an additional characteristic: the solution must be a piece of working hardware—a product. Thus, mechanical design problems begin with an ill-defined need and result in an object that behaves in a certain way, a way that the designers feel meets this need. This creates a paradox. *A designer must develop a device that, by definition, has the capabilities to meet some need that is not fully defined.*

1.6 THE BASIC ACTIONS OF PROBLEM SOLVING

Regardless of what design problem we are solving, we always, consciously or unconsciously, take six basic actions:

1. *Establish* the need or realize that there is a problem to be solved.
2. *Plan* how to solve the problem.

3. *Understand* the problem by developing requirements and uncovering existing solutions for similar problems.
4. *Generate* alternative solutions.
5. *Evaluate* the alternatives by comparing them to the design requirements and to each other.
6. *Decide* on acceptable solutions.

This model fits design whether we are looking at the entire product (see the product life-cycle diagram, Fig. 1.7) or the smallest detail of it.

These actions are not necessarily taken in 1-2-3 order. In fact they are often intermingled with solution generation and evaluation improving the understanding of the problem, enabling new, improved solutions to be generated. This iterative nature of design is another feature that separates it from analysis.

The list of actions is not complete. If we want anyone else on the design team to make use of our results, a seventh action is also needed:

7. *Communicate* the results.

The need that initiates the process may be very clearly defined or ill-defined. Consider the problem statements for the design of the simple lap joint of two pieces of metal given earlier (Fig. 1.9). The need was given by the problem statement in both cases. In the first statement, understanding is the knowledge of what parameters are needed to characterize a problem of this type and the equations that relate the parameters to each other (a model of the joint). There is no need to generate potential solutions, evaluate them, or make any decision, because this is an analysis problem. The second problem statement needs work to understand. The requirements for an acceptable solution must be developed, and then alternative solutions can be generated and evaluated. Some of the evaluation may be the same as the analysis problem, if one of the concepts is a bolt.

Some important observations:

- New needs are established throughout the design effort because new design problems arise as the product evolves. Details not addressed early in the process must be dealt with as they arise; thus, the design of these details poses new subproblems.
- Planning occurs mainly at the beginning of a project. Plans are always updated because understanding is improved as the process progresses.
- Formal efforts to understand new design problems continue throughout the process. Each new subproblem requires new understanding.
- There are two distinct modes of generation: concept generation and product generation. The techniques used in these two actions differ.
- Evaluation techniques also depend on the design phase; there are differences between the evaluation techniques used for concepts and those used for products.
- It is difficult to make decisions, as each decision requires a commitment based on incomplete evaluation. Additionally, since most design problems

are solved by teams, a decision requires consensus, which is often difficult to obtain.

- Communication of the information developed to others on the design team and to management is an essential part of concurrent engineering.

We will return to these observations as the design process is developed through this text.

1.7 KNOWLEDGE AND LEARNING DURING DESIGN

When a new design problem is begun, very little may be known about the solution, especially if the problem is a new one for the designer. As work on the project progresses, the designer's knowledge about the technologies involved and the alternative solutions increases, as shown in Fig. 1.11. Therefore, after completing a project, most designers want a chance to start all over in order to do the project properly now that they fully understand it. Unfortunately, few designers get the opportunity to redo their projects.

Throughout the solution process knowledge about the problem and its potential solutions is gained and, conversely, design freedom is lost. This can also be seen in Fig. 1.11, where the time into the design process is equivalent to exposure to the problem. The curve representing knowledge about the problem is a learning curve; the steeper the slope, the more knowledge is gained per unit time. Throughout most of the design process the learning rate is high. The second curve in Fig. 1.11 illustrates the degree of design freedom. As design decisions are made, the ability to change the product becomes increasingly limited. At the beginning the designer has great freedom because few decisions have been made and little capital has been committed. But by the time the product is in production,

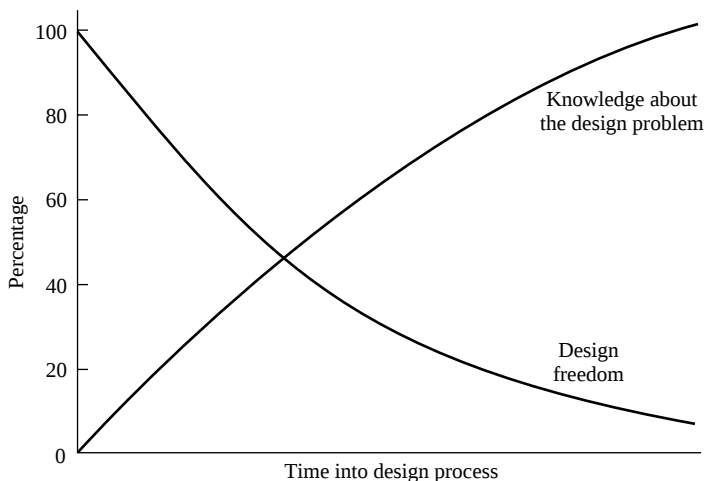


Figure 1.11 The design process paradox.

A design paradox: The more you learn the less freedom you have to use what you know.

any change requires great expense, which limits freedom to make changes. Thus, *the goal during the design process is to learn as much about the evolving product as early as possible in the design process because during the early phases changes are least expensive.*

1.8 DESIGN FOR SUSTAINABILITY

It is important to realize that design engineers have much control over what products are designed and how they interact with the earth over their lifetime. The responsibility that goes with designing is well summarized in the Hannover Principles. These were developed for EXPO 2000, The World's Fair in Hannover, Germany. These principles define the basics of Designing For Sustainability (DFS) or Design For the Environment (DFE). DFS requires awareness of the short- and long-term consequences of your design decisions.

The Hannover Principles aim to provide a platform on which designers can consider how to adapt their work toward sustainable ends. According to the World Commission on Environment and Development, the high-level goal is “Meeting the needs of the present without compromising the ability of future generations to meet their own needs.”

The Hannover Principles are:

1. **Insist on rights of humanity and nature to coexist** in a healthy, supportive, diverse, and sustainable condition.
2. **Recognize interdependence.** The elements of human design interact with and depend on the natural world, with broad and diverse implications at every scale. Expand design considerations to recognizing even distant effects.
3. **Accept responsibility for the consequences of design** decisions on human well-being, the viability of natural systems and their right to coexist.
4. **Create safe objects of long-term value.** Do not burden future generations with requirements for maintenance or vigilant administration of potential danger due to the careless creation of products, processes, or standards.
5. **Eliminate the concept of waste.** Evaluate and optimize the full life cycle of products and processes to approach the state of natural systems in which there is no waste.
6. **Rely on natural energy flows.** Human designs should, like the living world, derive their creative forces from perpetual solar income. Incorporate this energy efficiently and safely for responsible use.
7. **Understand the limitations of design.** No human creation lasts forever and design does not solve all problems. Those who create and plan should practice

You are responsible for the impact of your products on others.

humility in the face of nature. Treat nature as a model and mentor, not as an inconvenience to be evaded or controlled.

8. **Seek constant improvement by the sharing of knowledge.** Encourage direct and open communication between colleagues, patrons, manufacturers, and users to link long-term sustainable considerations with ethical responsibility, and reestablish the integral relationship between natural processes and human activity.
9. **Respect relationships between spirit and matter.** Consider all aspects of human settlement including community, dwelling, industry, and trade in terms of existing and evolving connections between spiritual and material consciousness.

We will work to respect these principles in the chapters that follow. We introduced the concept of “lean” earlier in this chapter as the effort to reduce waste (Principle 5). We will revisit this and the other principles throughout the book. In Chap. 11, we will specifically revisit DFS as part of Design for the Environment. In Chap. 12, we focus on product retirement. Many products are retired to landfills, but in keeping with the first three principles, and focusing on the fifth principle, it is best to design products that can be reused and recycled.

1.9 SUMMARY

The design process is the organization and management of people and the information they develop in the evolution of a product.

- The success of the design process can be measured in the cost of the design effort, the cost of the final product, the quality of the final product, and the time needed to develop the product.
- Cost is committed early in the design process, so it is important to pay particular attention to early phases.
- The process described in this book integrates all the stakeholders from the beginning of the design process and emphasizes both the design of the product and concern for all processes—the design process, the manufacturing process, the assembly process, and the distribution process.
- All products have a life cycle beginning with establishing a need and ending with retirement. Although this book is primarily concerned with planning for the design process, engineering requirements development, conceptual design, and product design phases, attention to all the other phases is important. PLM systems are designed to support life-cycle information and communication.

- The mechanical design process is a problem-solving process that transforms an ill-defined problem into a final product.
- Design problems have more than one satisfactory solution.
- Design for Sustainability embodied in the Hannover Principles is becoming an increasingly important part of the design process.

1.10 SOURCES

Creveling, C. M., Dave Antis, and Jeffrey Lee Slutsky: *Design for Six Sigma in Technology and Product Development*, Prentice Hall PTR, 2002. A good book on DFSS.

Ginn, D., and E. Varner: *The Design for Six Sigma Memory Jogger*, Goal/QPC, 2004. A quick introduction to DFSS

The Hannover Principles, Design for Sustainability. Prepared for EXPO 2000, Hannover, Germany, <http://www.mcdonough.com/principles.pdf>

Product life-cycle management (PLM) description based on work at Siemens PLM supplied by Wayne Embry their PLM Functional Architect.

http://www.plm.automation.siemens.com/en_us/products/teamcenter/index.shtml

<http://www.johnstark.com/epwl4.html> PLM listing of over 100 vendors.

Ulrich, K. T., and S. A. Pearson: “Assessing the Importance of Design through Product Archaeology,” *Management Science*, Vol. 44, No. 3, pp. 352–369, March 1998, or “Does Product Design Really Determine 80% of Manufacturing Cost?” working paper 3601–93, Sloan School of Management, MIT, Cambridge, Mass., 1993. In the first edition of *The Mechanical Design Process* it was stated that design determined 80% of the cost of a product. To confirm or deny that statement, researchers at MIT performed a study of automatic coffeemakers and wrote this paper. The results show that the number is closer to 50% on the average (see Fig. 1.3) but can range as high as 75%.

Womack, James P., and Daniel T. Jones: *Lean Thinking: Banish Waste and Create Wealth in Your Corporation*, Simon and Schuster, New York, 1996.

1.11 EXERCISES

- 1.1 Change a problem from one of your engineering science classes into a design problem. Try changing as few words as possible.
- 1.2 Identify the basic problem-solving actions for
 - a. Selecting a new car
 - b. Finding an item in a grocery store
 - c. Installing a wall-mounted bookshelf
 - d. Placing a piece in a puzzle
- 1.3 Find examples of products that are very different yet solve exactly the same design problem. Different brands of automobiles, bikes, CD players, cheese slicers, wine bottle openers, and personal computers are examples. For each, list its features, cost, and perceived quality.
- 1.4 How well do the products in Exercise 1.3 meet the Hannover Principles?
- 1.5 To experience the limitations of the over-the-wall design method try this. With a group of four to six people, have one person write down the description of some object that is

not familiar to the others. This description should contain at least six different nouns that describe different features of the object. Without showing the description to the others, describe the object to one other person in such a manner that the others can't hear. This can be done by whispering or leaving the room. Limit the description to what was written down. The second person now conveys the information to the third person, and so on until the last person redescribes the object to the whole group and compares it to the original written description. The modification that occurs is magnified with more complex objects and poorer communication. (Professor Mark Costello of Georgia Institute of Technology originated this problem.)

2

CHAPTER

Understanding Mechanical Design

KEY QUESTIONS

- What is the difference between function, behavior, and performance?
- Why does mechanical design flow from function to form?
- What are the languages of mechanical design?
- Are all design problems the same?
- What can you learn from dissecting products?

2.1 INTRODUCTION

For most of history, the discipline of mechanical design required knowledge of only mechanical parts and assemblies. But early in the twentieth century, electrical components were introduced in mechanical devices. Then, during World War II, in the 1940s, electronic control systems became part of the mix. Since this change, designers have often had to choose between purely mechanical systems and systems that were a mix of mechanical and electronic components and systems. These electronic systems have matured from very simple functions and logic to the incorporation of computers and complex logic. Many electromechanical products now include microprocessors. Consider, for example, cameras, office copiers, cars, and just about everything else. Systems that have mechanical, electronic, and software components are often called *mechatronic* devices. What makes the design of these devices difficult is the necessity for domain and design process knowledge in three overlapping but clearly different disciplines. But, no matter how electronic or computer-centric devices become, nearly all products require mechanical functions and a mechanical interface with humans. Additionally, all products require mechanical machinery for manufacture and assembly

and mechanical components for housing. Thus, no matter how “smart” products become, there will always be the need for mechanical design.

To explore systems that have significant mechanical components consider two examples that will be used throughout the book, the Irwin Quick-Grip clamp (Fig. 2.1) and the drive wheel assembly for the NASA Mars Exploration Rover (MER) developed by Cal Tech’s Jet Propulsion Laboratory (JPL) (Fig. 2.2).



Figure 2.1 Irwin Quick-Grip clamp.
(Reprinted with permission of Irwin Industrial Tools.)

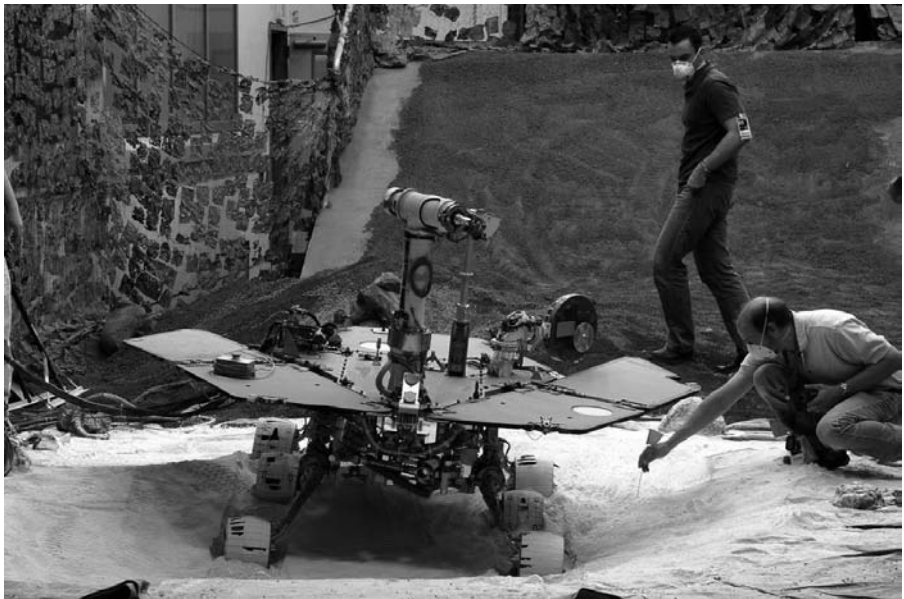


Figure 2.2 The Mars Exploration Rover being tested by JPL engineers. (Reprinted with permission of NASA/JPL.)

Irwin is one of the largest manufacturers of one-handed bar clamps. What makes the model shown in Fig 2.1 unique is that it can generate over 550 lb (250 kg) of force with the strength of only one hand. Irwin introduced this product in 2006 and sells many tens of thousands of them a month. In contrast to the purely mechanical, high-production-volume Quick-Grip, only two MERs were made and they are highly mechatronic.

The two MERs were launched toward Mars on June 10 and July 7, 2003, in search of answers about the history of water on Mars. They landed on Mars January 3 and January 24, 2004. They were designed for 90 Sol (Martian days, about 40 min longer than an Earth day) and were still operating in 2008, over 1300 Sols (over 3.5 years) past their design life. One of the Rovers, *Opportunity*, had traveled over 11 km (7.1 mi) during its five years of life.

Each Rover is a six-wheeled, solar-powered robot that stands 1.5 m (4.9 ft) high and is 2.3 m (7.5 ft) wide and 1.6 m (5.2 ft) long. They weigh 180 kg (400 lb) on Earth, 35 kg (80 lb) of which is the wheel and suspension system. Mars has only 38% the gravitational pull of Earth. So they weigh 68.4 kg (152 lb) on Mars. As shown in Fig. 2.3, a very simplified diagram of the MER's systems, propulsion and steering are two of the subsystems. Later in this chapter, we delve further into the MER, and in later chapters we will detail the wheels.

In general, during the design process the function of the system and its decomposition are considered first. After the function has been decomposed into the finest subsystems possible, assemblies and components are developed to provide these functions. For mechanical devices, the general decomposition is system–subsystem–assembly–component. Figure 2.3 shows the MER propulsion system, within which the motor and transmission are two subsystems. The wheel is a component. Systems, subsystems, and components all have *features*, specific attributes that are important, such as dimensions, material properties, shapes, or functional details. For the MER propulsion system, an important feature is that it can propel the MER at 5 cm/sec. For the transmission, a feature is that it has a 1500:1 reduction ratio. For the MER wheel, some of the important features are its diameter, tread pattern, and flexibility.

We must also note that many systems have both electrical and mechanical subsystems and components. Electrical systems generally provide energy, sensing, and control functions. The function of these electrical systems is fulfilled by circuits (electrical assemblies) that can be decomposed into electrical components (e.g., switches, transistors, and ICs), much as with mechanical objects. Finally, some of the control functions are filled by microprocessors. Physically, these are electric circuits, but the actual control function is provided by software programs in the processor. These programs are assemblies of coding modules composed of individual coding statements. Note that the function of the microprocessor could be filled by an electrical or possibly even a purely mechanical system. During the early phases of the design process, when developing systems is the focus of the effort, it is often unclear whether the actual function will be met by mechanical assemblies, electrical circuits, software programs, or a mix of these elements.

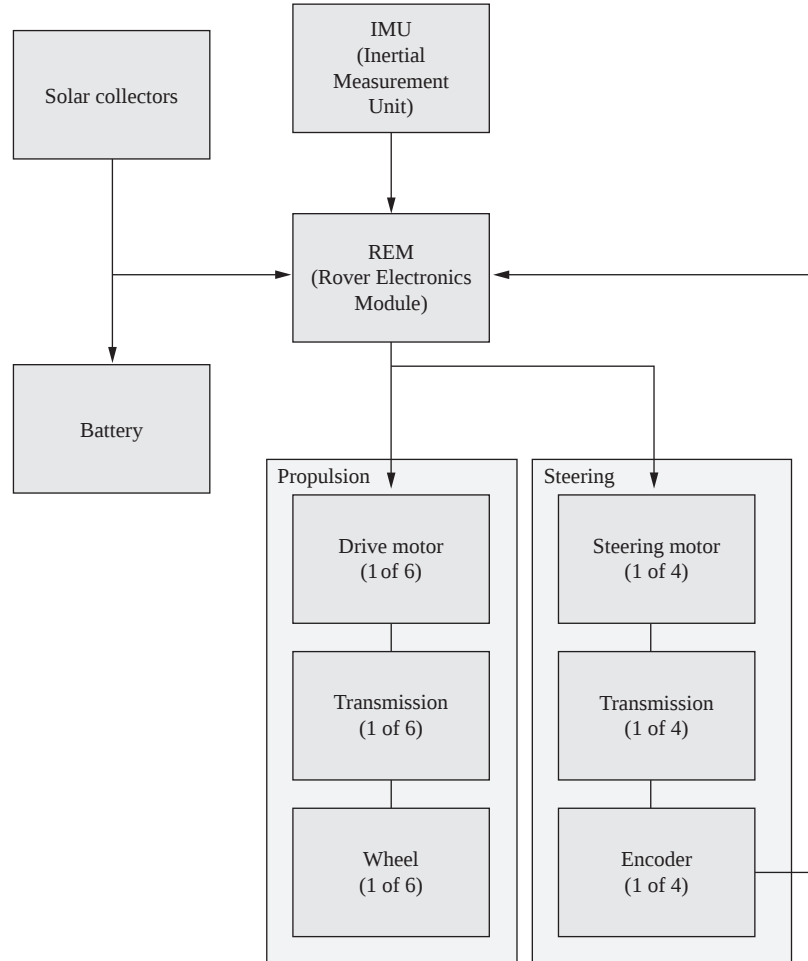


Figure 2.3 The MER Propulsion System showing some of the sub-systems and components.

2.2 IMPORTANCE OF PRODUCT FUNCTION, BEHAVIOR, AND PERFORMANCE

What is the function of the Irwin clamp? How does it behave? Does it have good performance? These three questions revolve around the terms “function,” “behavior,” and “performance”—similar, but different attributes of the clamp.

There are many synonyms for the word *function*. In mechanical engineering, we commonly use the terms *function*, *operation*, and *purpose* to describe *what* a device does. A common way of classifying mechanical devices is by their *function*. In fact, some devices having only one main function are named for

Function determines form and form, in turn, enables function.

that function. For example, a screwdriver has the function of enabling a person to insert or remove a screw. The terms *drive*, *insert*, and *remove* are all verbs that tell what the screwdriver does. In telling what the screwdriver does, we have given no indication of how the screwdriver accomplishes its function. To discover how, we must have some information on the form of the device. The term *form* relates to any aspect of physical shape, geometry, construction, material, or size. As we shall see in Chap. 3, one of the main ways engineers mentally index their knowledge about the mechanical world is by function. Now reread this paragraph and replace the screwdriver example with the Quick-Grip clamp.

In Fig. 2.3, we physically decomposed the Mars Rover propulsion and steering systems into subsystems and components at its physical boundaries. Functional decomposition is often much more difficult than physical decomposition, as each function may use part of many components and each component may serve many functions. Consider the handlebar of a bicycle. The handlebar is a bent piece of tubing, a single component that serves many functions. It enables the rider to “steer the bicycle” (“steer” is a verb that tells what the device does), and the handlebar “supports the rider” (again, a function telling what the handlebar does). Further, it not only “supports the brake levers” but also “transforms (another function) the gripping force” to a pull on the brake cable. The shape of the handlebar and its relationship with other components determine how it provides all these different functions. The handlebar, however, is not the only component needed to steer the bike. Additional components necessary to perform this function are the front fork, the bearings between the fork and the frame, the front wheel, and miscellaneous fasteners. Actually, it can be argued that all the components on a bike contribute to steering, since a bike without a seat or rear wheel would be hard to steer. In any case, the handlebar performs many different functions, but in fulfilling these functions, the handlebar is only a part of various assemblies. Similarly, the steering on the MER cannot actually steer it without the wheels in the propulsion system. The coupling between form and function makes mechanical design challenging.

Many common devices are cataloged by their function. If we want to specify a bearing, for example, we can search a bearing catalog and find many different styles of bearings (plain, ball, or tapered roller, for example). Each “style” has a different geometry—a different form—though all have the same primary function, namely, to reduce friction between a shaft and another object. Cataloging is possible in mechanical design as long as the primary function is clearly defined by a single piece of hardware, either a single component or an assembly. In other words, the form and function are decomposed along the same boundaries. This is true of many mechanical devices, such as pumps, valves, heat exchangers, gearboxes, and fan blades, and is especially true of many electrical circuits and components, such as resistors, capacitors, and amplifier circuits.

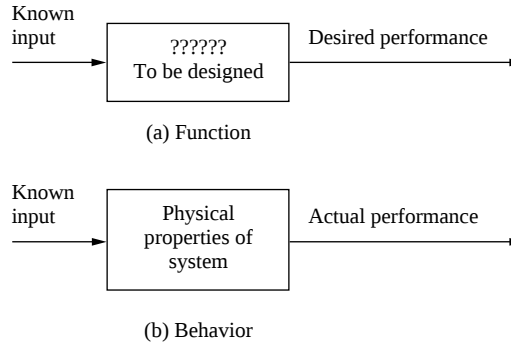


Figure 2.4 Function and behavior.

Two other terms often related to function are *behavior* and *performance*. *Function* and *behavior* are often used synonymously. However, there is a subtle difference, as shown in Fig. 2.4. In this figure there are two standard system blocks with an input represented by an arrow into the box, the system acted on by the input represented by the box, and the reaction of the system to the input represented by the arrow out of the box. The box in the upper part of the figure shows that *function is the desired output* from a system that is yet to be designed. When we begin to design a device, the device itself is unknown, but what we want it to do is known. If the system is known, as in the second part of the figure, then the behavior of the system can be found. *Behavior is the actual output*, the response of the system's physical properties to the input energy or control. Thus, the behavior can be simulated or measured, whereas function is only a desire.

Performance is the measure of function and behavior—how well the device does what it is designed to do. When we say that one function of the handlebar is to steer the bicycle, we say nothing about how well it serves this purpose. Before designing a handlebar, we must develop a clear picture of its desired performance. For example, one design functional goal is that the handlebar must “support 50 kg,” a measurable desired performance for the handlebar. The development of clear performance measures is the focus of Chap. 6. Further, after designing the handlebar we can simulate its strength analytically or measure the strength of a prototype to find the actual performance for comparison to that desired. This comparison is a major focus of Chap. 10.

2.3 MECHANICAL DESIGN LANGUAGES AND ABSTRACTION

Many “languages” or representations can be used to describe a mechanical object. Consider for a moment the difference between a detailed drawing of a component and the actual hardware that *is* the component. Both the drawing and the hardware represent the same object; however, they each represent it in a different language.

A skilled designer speaks many languages.

Extending this example further, if the component we are discussing is a bolt, then the word *bolt* is a textual (semantic or word) description of the component, a third language. Additionally, the bolt can be represented through equations (the final language) that describe its functionality and possibly its form. For example, the ability of the bolt to “carry shear stress” (a function) is described by the equation $\tau = F/A$; the shear stress τ is equal to the shear force F on the bolt divided by the stress area A of the bolt.

Based on this, we can use four different representations or languages to describe the bolt. These four can be used to describe any mechanical object:

Semantic. The verbal or textual representation of the object—for example, the word *bolt*, or the sentence, “The shear stress on the bolt is the shear force divided by the stress area.”

Graphical. The drawings of the object—for example, scale representations such as solid models, orthogonal drawings, sketches, or artistic renderings.

Analytical. The equations, rules, or procedures representing the form or function of the object—for example, $\tau = F/A$.

Physical. The hardware or a physical model of the object.

In most mechanical design problems, the initial need is expressed in a semantic language as a written specification or a verbal request by a customer or supervisor. The result of the design process is a physical object. Although the designer produces a graphical representation of the product, not the hardware itself, all the languages will be used as the product is refined from its initial, abstract semantic representation to its final physical form.

Further complicating how we refer to objects being designed, consider two drawings for a MER wheel, as shown in Fig. 2.5. Figure 2.5a is a rough sketch, which gives only abstract information about the component. It centers on the

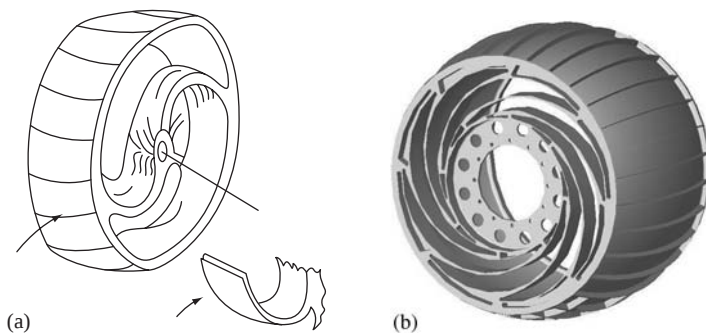


Figure 2.5 Abstract sketch and solid model of a MER wheel.

function of the wheel’s spokes to act like springs. Figure 2.5*b* is a solid model of the same component, focused on the final form of the wheel. In progressing from the sketch to the solid model, the *level of abstraction* of the device is *refined*.

Some design process techniques are better suited for abstract levels and others for levels that are more concrete. There are no true levels of abstractions, but rather a continuum on which the form or function can be represented. Descriptions of three levels of abstraction in each of the four languages are given in Table 2.1. The object we call a bolt is used as an example in Table 2.2.

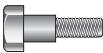
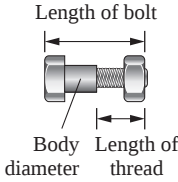
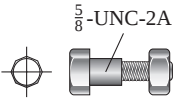
Another term that is often used in describing the analytical row in Table 2.1 is *simulation fidelity*. As analytical models or simulations increase in fidelity, their representation of the actual object or system becomes a more accurate representation of reality. Simulation fidelity will be further refined in Chap. 10.

The process of making an object less abstract (or more concrete) is called *refinement*. Mechanical design is a continuous process of refining the given needs

Table 2.1 Levels of abstraction in different languages

Language	Level of abstraction		
	Abstract	→	Concrete
Semantic	Qualitative words (e.g., <i>long, fast, lightest</i>)	Reference to specific parameters or components	Reference to the values of the specific parameters or components
Graphical	Rough sketches	Scale drawings	Solid models with tolerances
Analytical	Qualitative relations (e.g., <i>left of</i>)	Back-of-the-envelope calculations	Detailed analysis
Physical	None	Models of the product	Final hardware

Table 2.2 Levels of abstraction in describing a bolt

Language	Level of abstraction		
	Abstract	→	Concrete
Semantic	A bolt	A short bolt	A 1" 1/4–20 UNC Grade 5 bolt
Graphical			
Analytical	Right-hand rule	$\tau = F/A$	$\tau = F/A$
Physical	—	—	

to the final hardware. The refinement of the bolt in Table 2.4 is illustrated on a left-to-right continuum. In most design situations, the beginning of the problem appears in the upper left corner and the final product in the lower right. The path connecting these is a mix of the other representations and levels of abstractions.

2.4 DIFFERENT TYPES OF MECHANICAL DESIGN PROBLEMS

Traditionally, we decompose mechanical engineering by discipline: fluids, thermodynamics, mechanics, and so on. In categorizing the types of mechanical design problems, this discipline-oriented approach is not appropriate. Consider, for example, the simplest kind of design problem, a selection design problem. Selection design means picking one (maybe more) item from a list such that the chosen item meets certain requirements. Common examples are selecting the correct bearing from a bearings catalog, selecting the correct lenses for an optical device, selecting the proper fan for cooling equipment, or selecting the proper heat exchanger for a heating or cooling process. The design process for each of these problems is essentially the same, even though the disciplines are very different. The goal of this section is to describe different types of design problems independently of the discipline.

Before beginning, we must realize that most design situations are a mix of various types of problems. For example, we might be designing a new type of consumer product that will accept a whole raw egg, break it, fry it, and deliver it on a plate. Since this is a new product, there will be a lot of *original design* work to be done. As the design process proceeds, we will *configure* the various parts. To determine the thickness of the frying surface we will analyze the heat conduction of the frying component, which is *parametric design*. And we will *select* a heating element and various fasteners to hold the components together. Further, if we are clever, we may be able to *redesign* an existing product to meet some or all of the requirements. Each of the italicized terms is a different type of design problem. It is rare to find a problem that is purely one type.

2.4.1 Selection Design

Selection design involves choosing one item (or maybe more) from a list of similar items. We do this type of design every time we choose an item from a catalog. It may sound simple, but if the catalog contains more than a few items and there are many different features to the items, the decision can be quite complex.

To solve a selection problem we must start with a clear need. The catalog or the list of choices then effectively generates potential solutions for the problem. We must evaluate the potential solutions with respect to our specific requirements to make the right choice. Consider the following example. During the process of designing a product, an engineer must select a bearing to support a shaft. The known information is given in Fig. 2.6. The shaft has a diameter of 20 mm (0.787 in.). There is a radial force of 6675 N (1500 lb) on the shaft at the bearing,

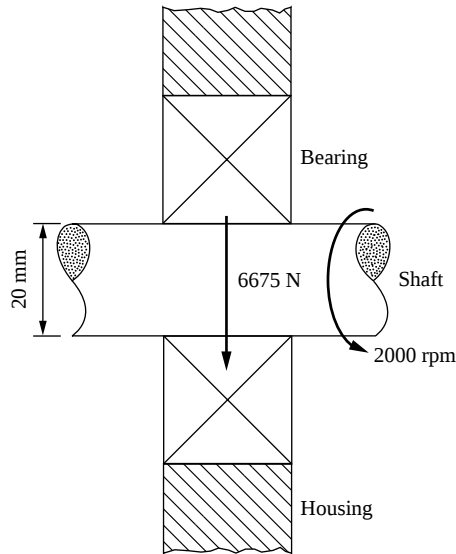


Figure 2.6 Load on a shaft.

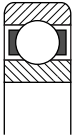
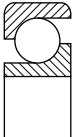
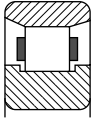
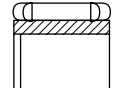
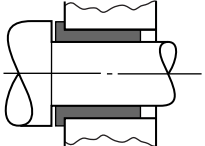
and the shaft rotates at a maximum of 2000 rpm. The housing to support the bearing is still to be designed. All we need to do is select a bearing to meet the needs. The information on shaft size, maximum radial force, and maximum rpm given in bearing catalogs enables us to quickly develop a list of potential bearings (Table 2.3). This is the simplest type of design problem we could have, but it is still incompletely defined. We do not have enough information to select among the five possible choices. Even if a short list is developed—the most likely candidates being the 42-mm-deep groove ball bearing and the 24-mm needle bearing—there is no way to make a good decision without more knowledge of the function of the bearing and of the engineering requirements on it.

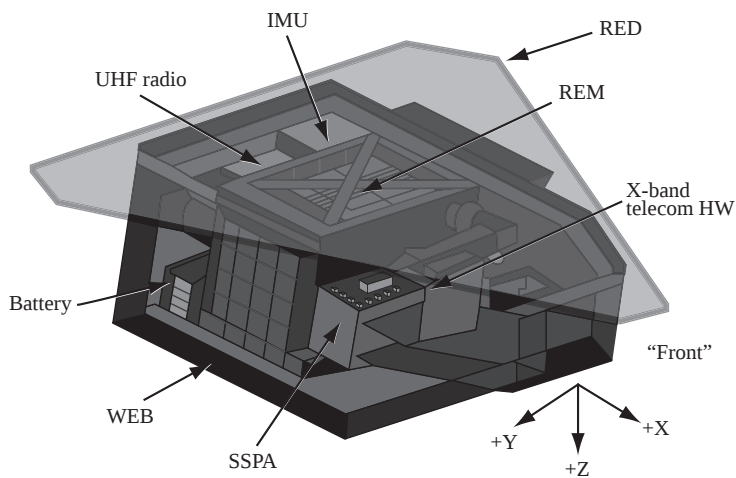
2.4.2 Configuration Design

A slightly more complex type of design is called configuration or packaging design. In this type of problem, all the components have been designed and the problem is how to assemble them into the completed product. Essentially, this type of design is similar to playing with an Erector set or other construction toy, or arranging living-room furniture.

Consider packaging of the assemblies in the MER. The body of the MER is made up of a Rover Equipment Deck (RED) where all the experiments are mounted, a Rover Electronics Module (REM), an Inertial Measurement Unit (IMU), a Warm Electronics Box (WEB), a battery, a UHF radio, an X-band telecom HW, and a Solid-State Power Amplifier (SSPA), as shown in Fig. 2.7. Each of these assemblies is of known size and has certain constraints on its position. For example, the RED must be on top and the WEB on the bottom, but

Table 2.3 Potential bearings for a shaft

Type		Outside diameter (mm)	Width (mm)	Load rating (lb)	Speed limit (rpm)	Catalog number
Deep-groove ball bearing		42	8	1560	18,000	6000
		47	14	2900	15,000	6204
		52	15	3900	9000	6304
Angular-contact ball bearing		47	14	3000	13,000	7204
		37	9	1960	34,000	71,904
Roller bearing		47	14	6200	13,000	204
		52	15	7350	13,000	220
Needle bearing		24	20	1930	13,000	206
		26	12	2800	13,000	208
Nylon bushing		23	Variable	290 ⋮ 8	10 ⋮ 500	4930

**Figure 2.7** The major assemblies in the MER.

many of the other major assemblies can be anywhere inside the envelop defined by these two.

Configuration design answers the question, How do we fit all the assemblies in an envelop? or Where do we put what? One methodology for solving this type of problem is to randomly select one component from the list and position it so that all the constraints on that assembly are met. We could start with the REM in the middle, then we select and place a second component. This procedure is continued until either we run into a conflict or all the components are in the MER. If a conflict arises, we back up and try again. For many configuration problems, some of the components to be fit into the assembly can be altered in size, shape, or function, giving the designer more latitude to determine potential configurations and making the problem solution more difficult. There are other methods to configure assemblies. They will be covered in Chap. 11.

2.4.3 Parametric Design

Parametric design involves finding values for the features that characterize the object being studied. This may seem easy enough—just find some values that meet the requirements. However, consider a very simple example. We want to design a cylindrical storage tank that must hold 4 m^3 of liquid. This tank is described by the parameters r , its radius, and l , its length and its volume is determined by

$$V = \pi r^2 l$$

Given a volume equal to 4 m^3 , then

$$r^2 l = 1.273$$

We can see that an infinite number of values for the radius and length will satisfy this equation. To what values should the parameters be set? The answer is not obvious, nor even completely defined with the information given. (This problem will be readdressed in Chap. 10, where the accuracy to which the radius and the length can be manufactured will be used to help find the best values for the parameters.)

Let us extend the concept further. It may be that instead of a simple equation, a whole set of equations and rules govern the design. Consider the instance in which a major manufacturer of copying machines had to design paper-feed mechanisms for each new copier. (A paper feed is a set of rollers, drive wheels, and baffles that move a piece of paper from one location to another in the machine.) Many parameters—the number of rollers, their positions, the shape of the baffles, and the like—characterize this particular design problem, but obviously there are certain similarities in paper feeders, regardless of the relative positions of the beginning and end points of the paper, the obstructions (other components in the machine) that must be cleared, and the size and weight of the paper. The company developed a set of equations and rules to aid designers in developing workable paper paths, and using this information, the designers could generate values for parameters in new products.

2.4.4 Original Design

Any time the design problem requires the development of a process, assembly, or component not previously in existence it calls for an original design. (It can be said that if we have never seen a wheel and we design one, then we have an original design.) Though most selection, configuration, and parametric problems are represented by equations, rules, or some other logical scheme, original design problems usually cannot be reduced to any algorithm. Each one represents something new and unique.

In many ways the other types of design problems—selection, configuration, and parametric—are simply constrained subsets of an original design. The potential solutions are limited to a list, an arrangement of components, or a set of related characterizing values. Thus, if we have a clear methodology for performing original design, we should be able to solve any design problem with a more limited set of potential solutions.

2.4.5 Redesign

Most design problems solved in industry are for the redesign of an existing product. Suppose a manufacturer of hydraulic cylinders makes a product that is 0.25 m long. If the customer needs a cylinder 0.3 m long, the manufacturer might lengthen the outer cylinder and the piston rod to meet this special need. These changes may require only parameter changes, or they may require something more extensive. What if the materials are not available in the needed length, or cylinder fill time becomes too slow with the added length? Then the redesign effort may require much more than parameter changes. Regardless of the change, this is an example of *redesign*, the modification of an existing product to meet new requirements.

Many redesign problems are *routine*; the design domain is so well understood that the method used can be put in a handbook as a series of formulas or rules. The parameter changes in the example of the hydraulic cylinder are probably routine for the manufacturer.

The hydraulic cylinder can also be used as an example of a *mature design*, in that it has remained virtually unchanged over many years. There are many examples of mature designs in our everyday lives: pencil sharpeners, hole punches, and staplers are a few found on the average desk. For these products, knowledge about the design problem is high. There is little more to learn.

However, consider the bicycle. The basic configuration of the bicycle—the two tensioned, spoked wheels of equal diameter, the diamond-shaped frame, and the chain drive—was fairly refined late in the nineteenth century. While the 1890 Humber shown in Fig. 2.8 looks much like a modern bicycle, not all bicycles of this era were of this configuration. The Otto dicycle, shown in Fig. 2.9, had two spoked wheels and a chain; stopping and steering this machine must have been a challenge. In fact, the technology of bicycle design was so well developed by the end of the nineteenth century that a major book on the subject, *Bicycles and Tricycles: An Elementary Treatise on Their Design and Construction*, was

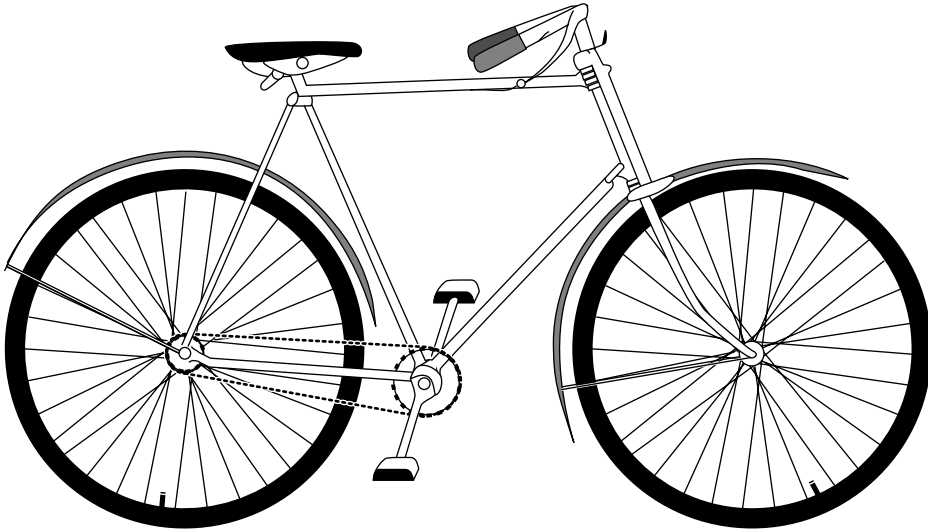


Figure 2.8 1890 Humber bicycle.

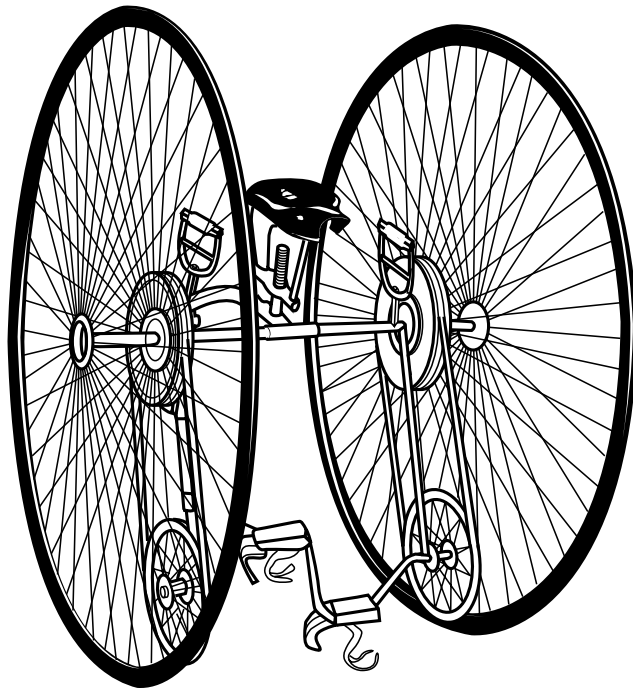


Figure 2.9 The Otto dicycle.



Figure 2.10 The Marin Mount Vision. (Reprinted with permission of Marin Bicycles.)

Most design problems are redesign problems since they are based on prior, similar solutions. Conversely, most design problems are original as they contain something new that makes prior solutions inadequate.

published in 1896.¹ The only major change in bicycle design since the publication of that book was the introduction of the derailleur in the 1930s.

However, in the 1980s the traditional bicycle design began to change again. For example, the mountain bike shown in Fig. 2.10 no longer has a diamond-shaped frame. Why did a mature design like a bicycle begin evolving again? First, customers are always looking for improved performance. Bicycles of the style shown in Fig. 2.10 are better able to handle rough terrain than traditional bikes. Second, there is improved understanding of human comfort, ergonomics, and suspensions. Third, customers are always looking for something new and exciting even if performance is not greatly improved. Fourth, materials and components have improved.

The point is that even mature designs change to meet new needs, to attract new customers, or to take advantage of new materials. Part of the design of a new bicycle like the Marin Mount Vision is routine, and part is original. Additionally,

¹The book, written by Archibald Sharp, has recently (1977) been reprinted by the MIT Press, Cambridge, Mass.

many subproblems were parametric problems, selection problems, and configuration problems. Thus, the redesign of a product, even a mature one, may require a wide range of design activity.

2.4.6 Variant Design

Sometimes companies will produce a large number of variants as their products. A variant is a customized product designed to meet the needs of the customer. For example, when you order a new computer from companies such as Dell, you can specify one of three graphics cards, two battery configurations, three communication options, and two levels of memory. Any combination of these is a variant that is specifically tuned to your needs. Also, Volvo trucks estimates that of the 50,000 parts it has in its inventory it annually supplies over 5000 variants, different truck models specifically assembled to meet the needs of the customer.

2.4.7 Conceptual Design and Product Design

Two other terms that will be used throughout the book are *conceptual design* and *product design*. These are catchall terms for two parts of the product development process. First, you must develop a concept and then refine the concept into a product. The activities during the conceptual and product development phases may make use of original, parametric, and selection design and redesign as needed.

2.5 CONSTRAINTS, GOALS, AND DESIGN DECISIONS

The progression from the initial need (the design problem) to the final product is made in increments punctuated by *design decisions*. Each design decision changes the *design state*. The state of a product is a snapshot of all the information known about it at any given time during the process. In the beginning, the design state is just the problem statement. During the process, the design state is a collection of all the knowledge, drawings, models, analyses, and notes thus far generated.

Two different views can be taken of how the design process progresses from one design state to the next. One view is that products evolve by a continuous comparison between the design state and the *goal*, that is, the requirements for the product given in the problem statement. This philosophy implies that all the requirements are known at the beginning of the design problem and that the difference between them and the current design state can be easily found. This difference controls the process. This philosophy is the basis for the methods in Chap. 6.

Another view of the design process is that when a new problem is begun, the design requirements effectively constrain the possible solutions to a subset of all possible product designs. As the design process continues, other *constraints* are added to further reduce the potential solutions to the problem, and potential solutions are continually eliminated until there is only one final design. In other

Constraints are often opportunities in disguise.

words, design is the successive development and application of constraints until only one unique product remains.

Beyond the constraints in the original problem specifications, constraints added during the design process come from two sources. The first is from the designer's knowledge of mechanical devices and the specific problem being solved. If a designer says, "I know bolted joints are good for fastening together sheet metal," this piece of knowledge constrains the solution to bolted joints only. Since every designer has different knowledge, the constraints introduced into the design process make each designer's solution to a given problem unique. The second type of constraint added during the design process is the result of design decisions. If a designer says, "I will use 1-cm-diameter bolts to fasten these two pieces of sheet metal together," the solution is constrained to 1-cm-diameter bolts, a constraint that may affect many other decisions—clearance for tools to tighten the bolt, thickness of materials used, and the like. During the design process, a majority of the constraints are based on the results of design decisions. Thus, the individual designer's ability to make well-informed decisions throughout the design process is essential. Decision-making techniques are emphasized in Chap. 8.

2.6 PRODUCT DECOMPOSITION

We will conclude this chapter with a method that can be the basis for understanding existing products. As such, it can serve as a starting point whether doing redesign, original design, or some other type of design, whether at the system or subsystem level. This *product decomposition* or "benchmarking" method helps us understand how a product is built, its parts, its assembly, and its function. It cannot be overemphasized how important it is to do decomposition and how it is the starting place for all design. In this chapter, we will decompose to understand the parts and assembly. In Chap. 7, the decomposition begun here will be extended to understand function.

Figure 2.11 shows a template that can be used to organize the decomposition. It is partially filled in for a pre-2003 version of the Irwin Quick-Grip. This version is the starting point for the redesign effort that resulted in the product shown in Fig. 2.1

The template begins with a brief description of the product and how it works—its function. This follows with a section showing each part. Only a selection of the parts is shown for the clamp in Fig. 2.11. Each part is given a name, the number required, its material, and the manufacturing process.

Often it can be hard to determine the material and manufacturing process. For plastics, there is a set of simple experiments for rough identification. Over the last few years, handheld devices have been developed that can identify materials



Product Decomposition

Design Organization: Example for the Mechanical Design Process

Date: Aug. 14, 2007

Product Decomposed: Irwin Quick Grip—pre 2007

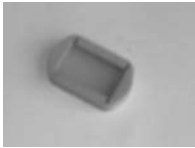
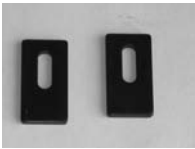
Description: This is the Quick-Grip Product that has been on the market for many years



How it works: Squeeze the pistol grip repeatedly to move the jaws closer together and increase the clamping force. Squeeze the release trigger to release the clamping force. The foot (the part on the left in the picture that holds the face that is clamped against) is reversible so the clamping force can be made to push apart rather than squeeze together.

Parts:

Part #	Part Name	# Req'd.	Material	Mfg. Process	Image
1	Main body	1	PPO or PVC	Injection molded	
2	Trigger	1	PVC	Injection molded	
4	Face plate, left	1	Polyethylene	Injection molded	

Part #	Part Name	# Req'd.	Material	Mfg. Process	Image
8	Pad	2	??	Injection molded	
13	Power spring	1	Steel	Wound wire	
14	Jam plates	2	Steel	Stamped sheet	

Disassembly:

Step #	Procedure	Part #s removed	Image
1	Take off left face plate	4	
12	Remove jam plates and power spring from main body assembly	13, 14, 1	
13	Remove trigger from main body assembly	2	
14	Pry off pad from main body assembly	8	

The Mechanical Design Process
Copyright 2008, McGraw-Hill

Designed by Professor David G. Ullman
Form # 1.0

Figure 2.11 Product decomposition samples for an older version of the Irwin Quick-Grip.
(Photos reprinted with permission of Irwin Industrial Tools.)

just by pointing the device at a sample of the material. While the main market for these devices is recycling, they are very useful when decomposing a product. Details on these are given in the Sources section at the end of this chapter.

The final section of the template is for the disassembly of the product. To build this section of the Product Decomposition report, remove one part at a time. Document the procedure needed to remove the part and the part numbers for those parts removed. Document what was done with a photograph. Figure 2.11 shows only a couple of the steps. Usually disassembly and part naming occur at the same time. Disassembly step 1 shows the left face plate, Part #4, was removed from the product. The internal parts of the clamp can now be seen in the photo. As this is a digital image in the actual template, it can easily be rescaled and studied as needed. Steps 12–14 are shown using a single image. The first one shows the removal of two parts, #13 and #14, at the same time as they come out together. Note how each procedure begins with a verb or verb phrase to tell what has to be done to remove the parts. Make these as descriptive as possible.

2.7 SUMMARY

- A product can be divided into functionally oriented operating *systems*. These are made-up of mechanical *assemblies*, electronic circuits, and computer programs. Mechanical assemblies are built of various *components*.
- The important form and function aspects of mechanical devices are called *features*.
- Function and behavior tell *what* a device does; form describes *how* it is accomplished.
- Mechanical design moves from function to form.
- One component may play a role in many functions, and a single function may require many different components.
- There are many different types of mechanical design problems: selection, configuration, parametric, original, redesign, routine, and mature.
- Mechanical objects can be described semantically, graphically, analytically, or physically.
- The design process is a continuous constraining of the potential product designs until one final product evolves. This constraining of the design space is made through repeated decisions based on comparison of design alternatives with design requirements.
- Mechanical design is the refinement from abstract representations to a final physical artifact.
- Product dissection is a useful way to understand the structure of a product.

2.8 SOURCES

Good books on designing new products

Clausing, Don, and Victor Fey: *Effective Innovation: Development of Winning Technologies*, ASME Press, 2004.

Cooper, Robert G.: *Winning at New Products*, 3rd ed., Perseus Publishing, 2001.

Vogel, C.M., J. Cagan, and P. Boatwright: *The Design of Things to Come*, Wharton School Publishing, 2005.

Plastics identification

The PHAZIR is a handheld, battery-powered, point-and-shoot plastic identifier. It weighs only 4 lb (1.8 kg) and takes 1–2 sec to determine the makeup of the sample.
www.polychromix.com

Metals identification

The iSort is a handheld, battery-powered, point-and-shoot spectrometer for on-site identification and analysis of all common metal alloys. Metal identification just requires pointing the gun-shaped iSort at a clean metal sample. The iSort is fairly expensive.
<http://www.spectro.com/pages/e/p010101.htm>

An inexpensive method uses the color of a chemical deposition to identify the metal. The process requires putting a drop of solution on the sample, then using a battery-powered electric charge through the solution to cause a chemical deposition on a piece of blotter paper. The color of the resulting deposit identifies the metal. <http://www.alloyid.com>

2.9 EXERCISES

- 2.1 Decompose a simple system such as a home appliance, bicycle, or toy into its assemblies, components, electrical circuits, and the like. Figures 2.3 and 2.11 will help.
- 2.2 For the device decomposed, list all the important features of one component.
- 2.3 Select a fastener from a catalog that meets these requirements:
 - Can attach two pieces of 14-gauge sheet steel (0.075 in., 1.9 mm) together
 - Is easy to fasten with a standard tool
 - Can only be removed with special tools
 - Can be removed without destroying either base materials or fastener
- 2.4 Sketch at least five ways to configure two passengers in a new four-wheeled commuter vehicle that you are designing.
- 2.5 You are a designer of diving boards. A simple model of your product is a cantilever beam. You want to design a new board so that a 150-lb (67-kg) woman deflects the board 3 in. (7.6 cm) when standing on the end. Parametrically vary the length, material, and thickness of the board to find five configurations that will meet the deflection criterion.
- 2.6 Find five examples of mature designs. Also, find one mature design that has been recently redesigned. What pressures or new developments led to the change?
- 2.7 Describe your chair in each of the four languages at the three levels of abstraction, as was done with the bolt in Table 2.2.

2.10 ON THE WEB

A template for the following document is available on the book's website:
www.mhhe.com/Ullman4e

- Product Decomposition



Designers and Design Teams

KEY QUESTIONS

- Why is it important to know how people do design?
- How is your ability to design dependent on your cognitive preferences?
- What are the characteristics of creators?
- How do individual cognitive abilities interact with the abilities of others during team activities?
- Why is a team more than a group of people?
- What can you do to help teams be successful?
- How can you measure team health?

3.1 INTRODUCTION

Since the time of the early potter's wheel, mechanical devices have become increasingly complex and sophisticated. This sophistication has evolved without much concern for how humans solve design problems. Throughout history people who were just naturally good at design were trained, through an apprentice program, to be masters in their art. The design methods they used and the knowledge of the domain in which they worked was refined through their personal experiences and passed, in turn, to their apprentices. Much of this experience was gained through experiments, through building prototypes and then going "back to the drawing board" to iterate toward the next product. The results of these experiments taught the designers what worked and what did not and pointed the way to the next refinement. With this methodology, products took many generations to be refined to the point of mature design.

However, as systems grew more complex and the world community grew more competitive, this mode of design became too time-consuming and too expensive. Designers recognized the need to find ways to deal with larger, more complex systems; to speed the design process; and to ensure that the final design

be reached with a minimum use of resources and time. In this book we discuss design techniques that meet these goals. To understand how these techniques help streamline the design process, it is important to understand how designers and design teams progress from abstract needs to final, detailed products.

To put this chapter in context, it is important to realize that design is the confluence of technical processes, cognitive processes, and social processes. We begin our discussion of how humans design mechanical objects by describing a cognitive model of how memory is structured in the individual designer. The types of information that are processed in this structure are explored, and the term *knowledge* is defined. Once we understand the information flow in human memory, we develop the different types of operations that a designer must perform in memory during the design process, and we explore creativity.

Based on this model of the individual's cognitive process, the chapter moves to the social aspect of design—working in teams. First, the structure of design teams is developed. This includes descriptions of the members of teams and how they are managed. Further, beyond the formal titles that people have, there is a more subtle, cognitive role that people play on teams. Second, an entire section is devoted to building and maintaining a design team. This includes how to start a team, inventory its health, and resolve problems as they develop. Supporting this chapter is a series of templates available at the book's website.

3.2 THE INDIVIDUAL DESIGNER: A MODEL OF HUMAN INFORMATION PROCESSING

The study of human problem-solving abilities is called *cognitive psychology*. Although this science has not yet fully explained the problem-solving process, psychologists have developed models that give us a pretty good idea of what happens inside our heads during design activities. A simplification of a generally accepted model is shown in Fig. 3.1. This model, called the *information-processing system* and developed in the late 1950s, describes the mental system used in the solution of any type of problem. In discussing that system here, we give special emphasis to the solution of mechanical design problems.

Information processing takes place through the interaction of two environments: the *internal environment* (information storage and processing inside the human brain) and the *external environment*. The external environment comprises paper and pencil, catalogs, computer output, and whatever else is used outside the human body to extend the internal environment.

In the internal environment, that is, within the human mind, there are two different types of memory: *short-term memory*, which is similar to a computer's operating memory (its random access memory or RAM), and *long-term memory*, which is like a computer's disk storage. Bringing information into this system from the external environment are *sensors*, such as the eyes, ears, and hands. Taste and smell are less often used in design. Information is output from the body with the use of the hands and the voice. There are other means of output, such

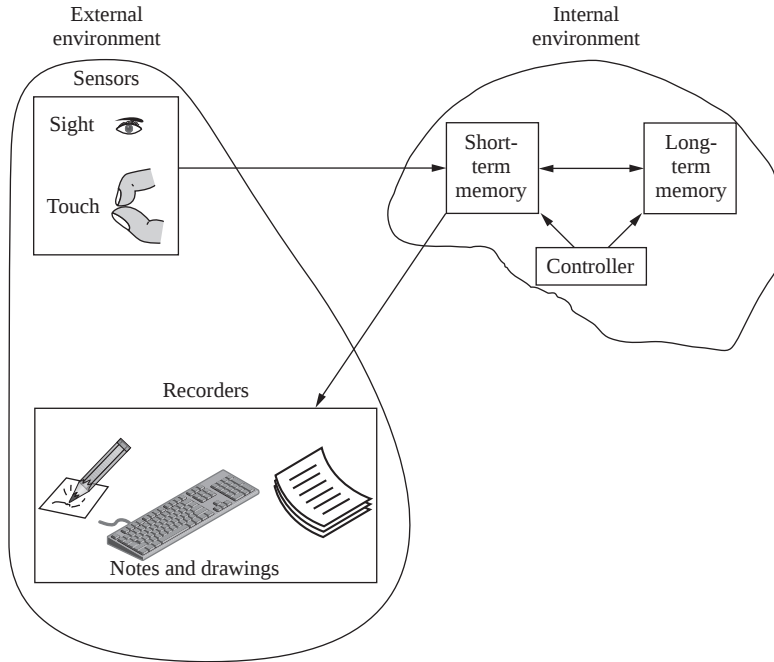


Figure 3.1 The human problem solver.

as body position, that are less often used in design. Additionally, as part of the internal processing capability, there is a *controller* that manages the information flow from the sensors to the short-term memory, between the short-term and the long-term memory, and between the short-term memory and the means of output.

Before describing short-term and long-term memory and the control of information flow, we need to describe the *information* that is processed in this system. In a computer the information is in terms of bits, or binary digits (0s and 1s), but in the human brain, information is much more complex.

In recent experiments, an orthographic drawing of a power transmission system consisting of shafts, gears, and bearings was shown to mechanical engineering students and professional engineers. The students were lower-level undergraduates who had not studied power transmission systems. The drawing was shown briefly and then removed, and the subjects were asked to sketch what they had seen. The students tended to reconstruct the drawings from the line segments and simple shapes they had seen in the original drawing. Not understanding the complexities of geared transmissions, they could not remember anything more complicated. They remembered and drew only the basic form of the components. On the other hand, the professional engineers were able to remember components grouped together by their function. In recalling a gear set, for example, the experts knew that two meshed gears and their associated shafts and bearings provide the function of changing the rpm and torque in the system. They also knew

what geometry or line segments were needed to represent the form of a gear set. Thus, the experienced engineers using functional groupings were able to include substantially more information than the students in their sketches.

The line segments remembered by the students and the functional groupings remembered by the experienced designers are called *chunks of information* by cognitive psychologists. The greater the expertise of the designer, the more content there is in the chunks of information processed. Exactly what types of information are in these chunks, however, is not always clear. Types of knowledge that might be in a chunk include

- **General knowledge**, information that most people know and apply without regard to a specific domain. For example, red is a color, the number 4 is bigger than the number 3, an applied force causes a mass to accelerate—all exemplify general knowledge. This knowledge is gained through everyday experiences and basic schooling.
- **Domain-specific knowledge**, information on the form or function of an individual object or a class of objects. For example, all bolts have a head, a threaded body, and a tip; bolts are used to carry shear or axial stresses; the proof stress of a grade 5 bolt is 85 kpsi. This knowledge comes from study and experience in the specific domain. It is estimated that it takes about ten years to gain enough specific knowledge to be considered an expert in a domain. Formal education sets the foundation for gaining this knowledge.
- **Procedural knowledge**, the knowledge of what to do next. For example, if there is no answer to problem X, then decomposing X into two independent easier-to-solve subproblems, X1 and X2, would illustrate procedural knowledge. This knowledge comes from experience, but some procedural knowledge is also based on general knowledge and some on domain-specific knowledge. We must often make use of procedural knowledge to solve mechanical design problems.

In mechanical engineering the term *feature* is synonymous with *chunks of information*. Since a design feature is some important aspect of a component, assembly, or function, the gear set discussed in the preceding example is both a chunk and a feature.

The exact language in which chunks of information are encoded in the brain is unknown. They might be dealt with as semantic information (text), graphical information (visual images), or analytical information (equations or relationships). Psychologists believe that most mechanical designers process information in terms of visual images and that these images are three-dimensional and are readily manipulated in the short-term memory.

All design and decision making is limited by human cognitive capabilities.

3.2.1 Short-Term Memory

The short-term memory is the main information processor in the human brain. It has no known specific anatomic location, yet it is known to have very specific attributes.

One important attribute of the short-term memory is its quickness. Information chunks can be processed in the short-term memory in about 0.1 second. The term *processed* implies such actions as comparing one chunk of information to another, modifying a chunk by decomposing it into smaller parts, combining two or more chunks into one new one, changing a chunk's size or distorting its shape, and making a decision about the chunk. It is unknown how much of the short-term memory is actually used to process the information. We do know that the harder it is to solve the problem, the more short-term memory is used for processing.

The capacity of the short-term memory was first described in a paper titled "The Magical Number Seven, Plus or Minus Two" (see Section 3.8), which reported that the short-term memory is effectively limited to seven chunks of information (plus or minus two). This is like having a computer RAM with only seven memory locations. These approximately seven chunks—these seven unique things—are all that a person can deal with at one time. For example, let us say we are working on a design problem and have an idea (a chunk of information, maybe just a word or maybe a visual image) that we want to compare to some constraints on the design (other chunks of information). How many constraints can we compare to the idea in our head? Only two or three at a time, since the idea itself takes one slot in the short-term memory and the constraints take two or three more. That does not leave much memory to do the processing necessary for comparison. Add any more constraints and the processing stops; the short-term memory is simply too full to make any progress on solving the problem.

A couple of quick experiments are convincing about the limits of the short term memory. Open a phone book and randomly choose a phone number in which the seven digits are unrelated to each other. (A number such as 555-2000 is not acceptable because the last four digits can be lumped together as a single chunk—two thousand.) After looking at the number briefly, close the phone book, walk across the room, and dial the number. Most people can manage to do this task if they are not interrupted or do not think about anything else. The same experiment can be tried with two unrelated phone numbers. Few people are able to remember them long enough to dial them both since they require dialing 14 pieces of information, which is beyond the capacity of the short-term memory. Granted, these 14 digits can be memorized, or stored in long-term memory, but that would take some study time.

Another example of the size limitations of short-term memory is more mechanical in nature. Consider the four-bar linkage of Fig. 3.2. It is made up of four elements: the driver A–B, the link B–C, the follower C–D, and the base D–A.

It is not difficult for most engineers to visualize the follower C–D rocking back and forth as the driver A–B is rotated. Point B makes a circle, and point C moves in an arc about point D. An expert on linkages would only use a single

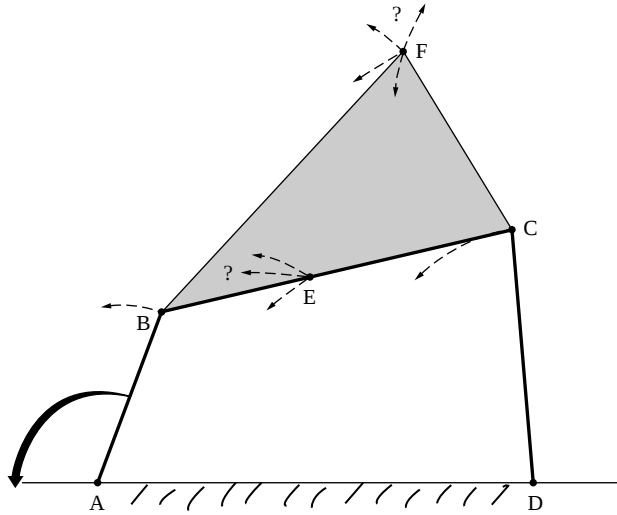


Figure 3.2 A four-bar linkage.

chunk to encode this mechanism. But a novice in the domain of four-bar linkages would need to visualize four line segments, using four chunks plus others for processing the motion. To make the task more difficult, trace the path of point E on the link. This requires more short-term memory. Harder still is tracing the path of point F. In fact, this requires so many different parameters to track that only a few linkage experts can visualize the path of point F.

Another feature of the short-term memory is the fading of information stored there. The phone number remembered earlier is probably forgotten within a few minutes. To keep from forgetting short-term information, like the phone number, many people keep repeating the information over and over. With such continuous refreshing, it is possible to retain certain objects or parts of objects within the short-term memory and to let only the unimportant information fade to make room for the processing of new chunks of information.

Last, it is impossible for us to be aware of what is happening in our short-term memory while we are solving problems. To follow our own thoughts, we need to use some of that memory to monitor and understand the problem-solving process, making that space no longer available for problem solving. Thus, you can not really observe what you are doing during problem solving without affecting what you are trying to observe.

3.2.2 Long-Term Memory

The long-term memory was earlier compared with the disk storage in a computer; like disk storage, it is for permanent retention of information. Let us look at the four major characteristics of long-term memory. First, long-term memory has seemingly unlimited capacity. Despite the cartoon in Fig. 3.3, there is no

THE FAR SIDE® BY GARY LARSON



The Far Side® by Gary Larson © 1986 FarWorks, Inc. All Rights Reserved. The Far Side® and the Larson® signature are registered trademarks of FarWorks, Inc. Used with permission.

**"Mr. Osborne, may I be excused?
My brain is full."**

Figure 3.3 Long-term memory problems.

documented case of anybody's brain becoming "full," regardless of head size. It is hypothesized that as we learn more we unconsciously find more efficient ways to organize the information by reorganizing the chunks in storage. Reconsider the difference between the student's and the expert's ways of remembering information about the power transmission system. The expert's information storage was more efficient than the student's.

The second characteristic of the long-term memory is that it is fairly slow in recording information. It takes 2 to 5 min to memorize a single chunk of information. This explains why studying new material takes so long.

The third characteristic is the speedy recovery of information from long term memory. Retrieval is much quicker than storage, the time depending on the complexity of the information and the recentness of its use. It can be as fast as 0.1 sec per chunk of information.

The fourth characteristic is that the information stored in the long-term memory can be retrieved at different levels of abstraction, in different languages, and with different features. For example, consider the knowledge an average engineer can retrieve about a car (Fig. 3.4). The sample data ranges from images of entire vehicles to semantic rules and equations for diagnosing problems. Human memory is very powerful in matching the form of the data retrieved to that which is needed for processing in the short-term memory.

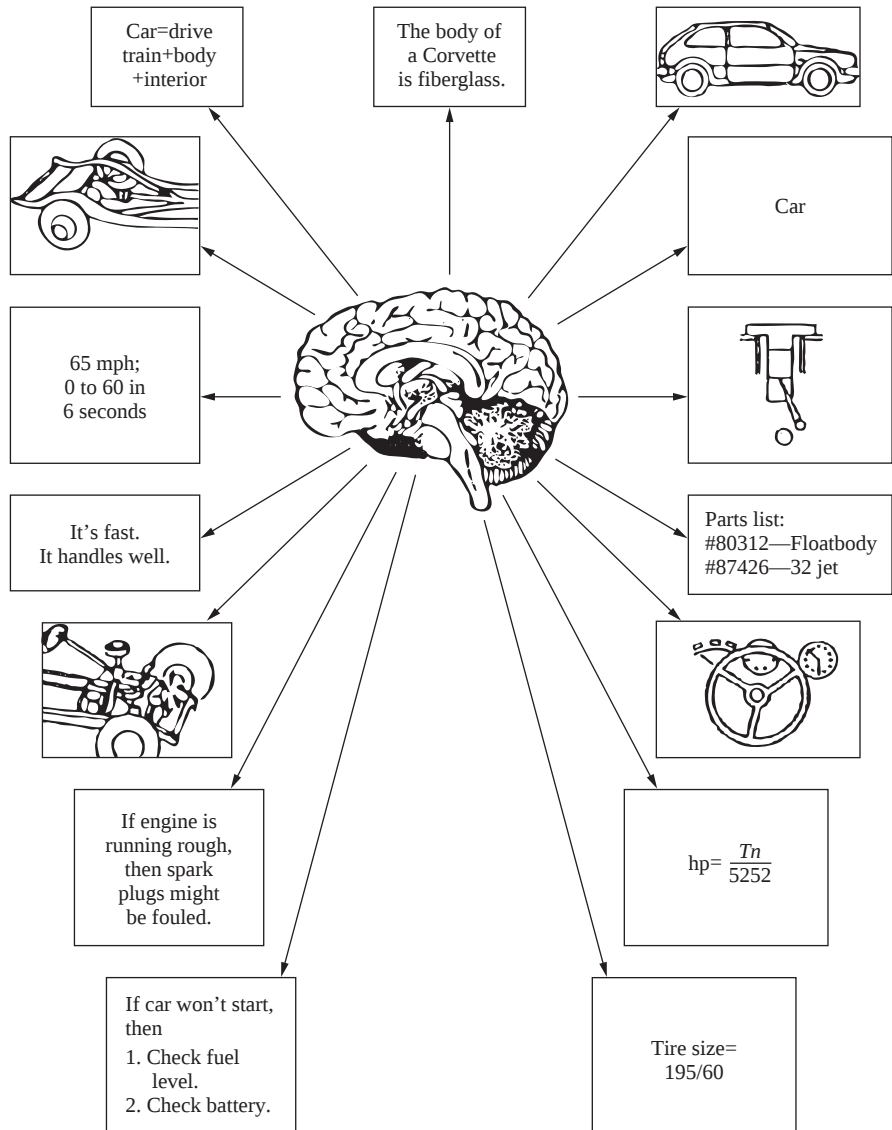


Figure 3.4 Knowledge stored in memory about cars.

3.2.3 Control of the Information-Processing System

During problem solving, the controller (Fig. 3.1) enables us to encode outside information obtained through our senses or retrieve information from long-term memory for processing in the short-term memory. Some of the information in the short-term memory is allowed to fade, and new information is input as it is needed and becomes available. Additionally, the controller can help extend the short-term memory by making notes and sketches; these need to be done quickly so that they do not bog down the problem-solving process. When we have completed manipulating the information, the controller can store the results in long-term memory, or in the external environment by describing it in text, verbally, or in graphic images.

3.2.4 External Environment

The external environment—paper and pencil, computers, books—plays a number of roles in the design process: it is a source of information; it is an analytical capability; it is a documentation/communication facility; and, most importantly for designers, it is an extension for the short-term memory. The first three of these roles seem evident; however, the last role, as an extension for the short-term memory, needs some discussion.

Because the short-term memory is a space-limited central processor, human problem solvers utilize the external environment as a short-term memory extension, much as a computer extends RAM by using cache memory. This is accomplished by making notes and sketches of ideas and other information needed in problem solving. In order to be useful to the short-term memory, any extension must share the characteristics of being very fast and having high information content. Watch any design engineer trying to solve a problem. He or she will make sketches even when not trying to communicate. These sketches serve as aids in generating and evaluating the ideas by serving as additional chunks of information to be processed. Sketches are fast to make and are information-rich.

3.2.5 Implications of the Model

One of the implications of the information-processing model of human problem solving is that the size of the short-term memory is a major limiting factor in the ability to solve problems. To accommodate this limitation we break down problems into finer and finer subproblems until we can “get our mind around it”—in other words, manage the information in our short-term memory. Typically, these fine-grained subproblems are worked on for about 1 minute before going to the next one. Thus, design of even a simple problem is the solution of many thousands of subproblems. Further, our thinking process has evolved so that, as we solve problems, our expertise about the constraints and potential solutions increases and our configuration of chunks becomes more efficient. This helps offset the “magic number” seven, but human designers are still quite limited. It would almost seem that these limitations would preclude our ability to solve

If you try to think about what you are doing while you are doing it, you stop doing it. If you don't reflect on what you just did, you are doomed to repeat it.

complex problems. As discussed in the upcoming sections, processing speed and flexibility of information storage and recovery enable designers to develop very complex products.

3.3 MENTAL PROCESSES THAT OCCUR DURING DESIGN

We can now describe what happens when a designer faces a new design problem. The problem may be the design of a large, complex system or of some small feature on a component. We will focus on how a designer understands new information such as the problem statement, how ideas are generated, and how they are evaluated.

In Section 1.6 we introduced seven basic actions of problem solving. The core actions—understand, generate, evaluate, and decide—are refined here.

3.3.1 Understanding the Problem

Consider what happens when a new problem is broached. If we think of its design state as a blackboard on which is written or drawn everything known about the device being designed, then the blackboard is initially blank, i.e., the design state is empty. Let us return to the fastening problem presented in Chap. 1 (see Fig. 1.9):

Design a joint to fasten together two pieces of 1045 sheet steel, each 4 mm thick and 6 cm wide, that are lapped over each other and loaded with 100 N.

Before any information about the problem is put on the design-state blackboard, the problem statement must be understood. If the problem is outside the realm of experience (the designer does not know what the term *lapped* means, for example), then the problem cannot be understood.

But how do we “understand” a problem? Most likely in this way: As the problem is read, it is “chunked” into significant packets of information. This happens in the short-term memory, where we naturally parse the sentence into phrases like “design a joint,” “to fasten together,” and so on. These chunks are compared with long-term memory information to see if they make sense, and then most are allowed to fade. The goal of this first pass through the problem is to try and retain only the major functions of the needed device. Usually a problem will be read or sensed a number of times until the major function(s) is identified. Unfortunately there is no guarantee that, from the usually incomplete data that

exist at the beginning of a design problem, the most important functions will be identified. In our example there is no ambiguity. The prime function is to transfer a load from one sheet of steel to another through a lapped joint.

What is important to realize is that a problem is “understood” by comparing the requirements on the desired function to information in the long-term memory. Thus, every designer’s understanding of the problem is different, because each designer has different information stored in the long-term memory. (In Chap. 6 we develop a method to ensure that the problem is fully understood with minimal bias from the designer’s own knowledge.)

3.3.2 Generating Solutions

We have seen that in trying to understand a design problem, we compare the problem to information from the long-term memory. In order to retrieve information from the long-term memory, we need a way to index the knowledge stored there. We can index that information in many ways (Fig. 3.4). As in the gearbox example at the beginning of this chapter, the most efficient indexing method is by function. What are recalled and downloaded to the short-term memory are specific (usually abstract) visual images from past experience. Thus, we search by function and recall form or graphical representations. This is not always true: we can also index our memory by shape, size, or some other form feature. However, in solving design problems, function is usually the primary index. For some problems the information recalled meets all the design requirements and the problem is solved.

If, in understanding a problem, we must recall images of previous designs, we have a predisposition to use these designs. Some designers get stuck on these initially recalled images and have difficulty evaluating them objectively and generating other, potentially better ideas. Many of the techniques discussed in Chaps. 7 and 11 are specifically designed to overcome this tendency.

On the other hand, what happens if the problem being solved is new and we find no solution to it in the long-term memory? We then use a three-step approach: decompose the problem into subproblems, try to find partial solutions to the subproblems, and finally recombine the subsolutions to fashion a total solution. The subproblems are generally functional decompositions of the total problem. The creative part of this activity is in knowing how to decompose and recombine cognitive chunks.

3.3.3 Evaluating Solutions

Often people generate ideas but have no ability to evaluate them. Evaluation requires comparison between generated ideas and the laws of nature, the capability of technology, and the requirements of the design problem itself. Comparison, then, necessitates modeling the concept to see how it performs with respect to these measures. The ability to model is usually a function of knowledge in the domain. We will address evaluation techniques in Chaps. 8, 10, and 11.

3.3.4 Deciding

At the end of each problem-solving activity, a decision is made. It may be to accept an idea that was generated and evaluated, or more likely, it will be to address another topic that is related to the problem. The rationale for how decisions are made is not well understood, but Sections 3.3.5 and 3.3.6 should help clarify what is known.

3.3.5 Controlling the Design Process

To understand how designers progress through a design problem, subjects were videotaped as they worked. In the study of these videotapes, it became evident that the path from initial problem presentation to solution was not very straightforward. It seemed like an almost random process—efforts on a subproblem made the designer aware of another subproblem, and the designer then focused attention on this second problem without having solved the first. No model for the control of focus was found. However, it was clear that the process for some designers is so chaotic that they never find solutions to their problems, while other designers rapidly proceed through the design effort. The techniques discussed in this book are intended to give structure to the design process so that the path from problem statement to solution is as controlled and direct as possible.

3.3.6 Problem-Solving Behavior

Everybody has a unique manner of problem solving. A person's problem-solving behavior affects how decisions are made individually and has a significant impact on team effectiveness. The following discussion is centered around five personal problem-solving dimensions. These five are useful for describing how an individual solves a design problem because they describe an individual's information management and decision-making preferences. Since all the team members bring their individual problem-solving processes to team activities, it is the interaction of all the individuals' solution processes that determines the team's health. For each of the five dimensions, suggestions for how to counteract extreme behavior are given. Some of these are useful to the individual working alone, and all are important in team situations and will be referenced later in the chapter when we talk about team health. A template for easily evaluating your problem-solving behavior is available.



The first personal problem-solving dimension describes an individual's **energy source** or **extraversion**. It is a measure of whether you are an *internal* or *external* problem solver. For a rough estimate of your, or a colleague's, energy source, answer five questions. If scoring a colleague, pretend you are that person. The five questions are shown in Figs. 3.5–3.9, screen shots from the template. In each of the five questions are shown with two potential responses. In Fig. 3.5 the top responses indicate an internal energy source and the bottom responses indicate an external energy source. For the example here, internal is

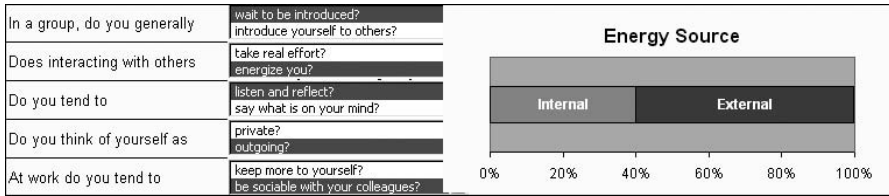


Figure 3.5 Energy source personal problem-solving dimension.

selected for the first and third questions and so the person is 2/5 or 40% internal and 60% external. In the template, the bar chart updates as you select the responses.

If a person is reflective, is a good listener, thinks and then speaks, and enjoys solving problems alone, then she is an internal problem solver. If the person's energy comes from outside through interactions with others (i.e., the person is sociable and tends to speak and then think) she is an external problem solver. About 75% of all Americans and 48% of engineering students and top executives are external problem solvers. There is no right or wrong style; this is merely the way people operate. They may show slightly different styles in different situations, but will generally not deviate very far from type.

In team settings both internals and externals have characteristics that are essential to the team but may cause difficulty—the externals tend to overwhelm the internals, who are reluctant to share their ideas. Here are some suggestions to keep the externals productive but not domineering:

- Externals need to allow others time to think. Point out to them that it is not necessary to fill in all the pauses with words.
- Externals need to practice listening to the ideas and suggestions of others and pausing before they react. Brainstorming or another creativity-support activity can help here (see Section 7.4).
- Encourage externals to recap what has been said to make sure they have heard the contributions of others.
- Externals need to realize that silence does not always mean consent. Sometimes an external will overwhelm the internals, who will become quiet rather than argue the point.

Here are some suggestions to assist internals in getting their ideas out for consideration:

- Encourage internals to share more than their final response. There is value in thinking out loud, as even the most trivial idea may be part of a good solution. The process will judge the value of the ideas.
- Try suggesting techniques that enable internals to have an equal say in selecting ideas and plans, such as the techniques in Chaps. 5–12.

- Encourage internals to develop some nonverbal, body-language signals that indicate assent or dissent. Make sure that these signals are understood by other team members.
- Encourage internals to restate their ideas. This restating signifies to the internal that his or her ideas count and forces the externals to listen.
- Get internals to push externals for more clarity and meaning.

The second dimension reflects your preference for an **information management style** or **originality**. It is a measure of whether you like working with *facts* or *possibilities*. For an estimate of your or a colleague's information management style, answer the questions in Fig. 3.6. For the example shown, the individual operates on both facts and possibilities with a slight tendency for possibilities.

People who prefer facts and details are literal, practical, and realistic; they appreciate the here and now. Those who think in terms of possibilities, patterns, concepts, and theories are looking for relationships between pieces of information and the meaning of the information. About 75% of Americans are fact-oriented, as are 66% of top executives; yet only 34% of all engineering students are fact-oriented. This is interesting in light of the heavy emphasis on math and science that is the focus of an engineering education. Other labels that could be placed on the scale are *Preserver* and *Explorer*, where the Preservers maintain the system, the Explorers are the boat rockers.

To solve most problems it is important to have a balance between the two extremes. When solving a problem alone, fact-oriented people have trouble getting started, whereas possibility-oriented people have trouble doing the details. This problem-solving dimension is the cause of most miscommunication, misunderstanding, and team problems. Design requires working with both facts and possibilities. Thus, both types of thinking are essential on a design team. However, individuals with a strong tendency toward either extreme may need help in the team setting. Some suggestions for fact-oriented team members are as follows:

- Encourage fact-oriented team members to fantasize, think wildly, and allow others to think wildly. Wild ideas can lead to good ideas. Brainstorming (Section 7.4) and thinking out loud (rambling) bring out such ideas.
- Encourage fact-oriented team members to allow the team to set goals rather than dive right into the problem and tackle the details.

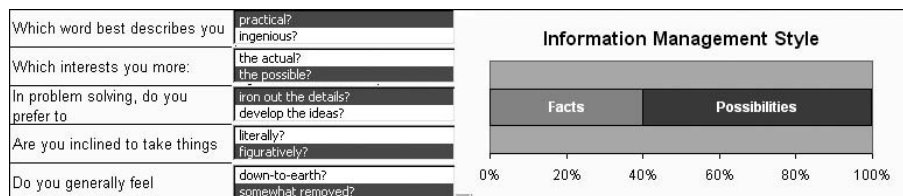


Figure 3.6 Information management style problem-solving dimension.

Here are some suggestions for team members who think in terms of possibilities:

- Encourage possibility-oriented team members to deal with details. The best idea will never reach maturity if the details are not attended to. It is frustrating to them but possibly worthwhile to have them take on the responsibility of a detail task.
- Force possibility-oriented team members to be specific and avoid generalities. They should be encouraged to try to enumerate the exact items they want to address instead of making sweeping general statements.
- Remind possibility-oriented team members to stick to the issues. Other team members can control the flow of the problem solving by clearly stating the issues being addressed. Other issues that arise during discussion should be recorded and then shelved for later consideration.

The third dimension measures which **information language** a person prefers to use, *verbal* or *visual*. For a rough idea of your or a colleague's information language style, answer the questions in Fig. 3.7. The example individual is primarily a visual problem solver, but can work verbally.

Visual information includes pictures, diagrams, graphs, and hardware. Verbal information includes written or spoken words and mathematical formulas. It is interesting to note that most people favor visual information, yet most classes in school are presented in a verbal language. This mismatch is especially striking in science and engineering classes.

When you are working alone, the language you use is not an important consideration. In teams, however, the preferred languages greatly affect the development of a shared vision of the problem and alternative solutions. Some guidelines on how to manage the two types of communication language in team situations follow.

- Help identify information that needs to be communicated, regardless of language.
- Help identify differences in team members' mental models, encouraging extra effort by both visual and verbal people to communicate clearly with other members to develop a shared understanding.

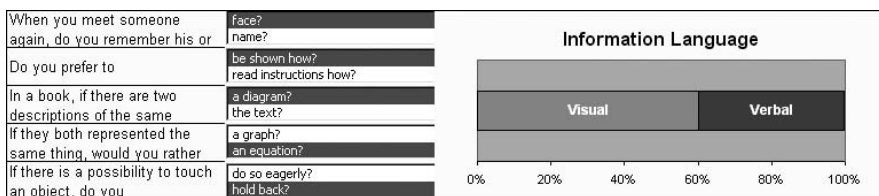


Figure 3.7 Information language personal problem-solving dimension.

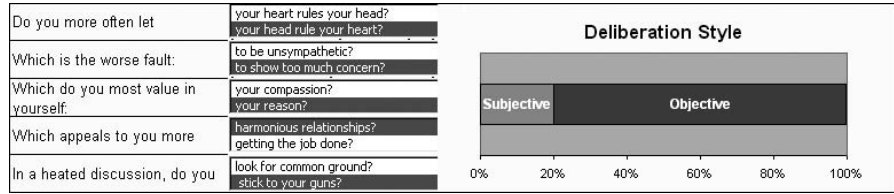


Figure 3.8 Deliberation style personal problem-solving dimension.

- If words and equations aren't working, try a diagram or picture. If the picture isn't working, try words and equations.

The fourth dimension reflects the **deliberation style** or **accommodation**, the *objectivity* or *subjectivity* with which problems are solved. To get an estimate of your or a colleague's deliberation style, answer the five questions in Fig. 3.8. In the example, the person is primarily an objective problem solver.

Some team members take a subjective approach, others an objective one. People who rely on interpersonal involvement, circumstances, and the "right thing to do" take a subjective approach to design. These team members can be referred to as "adaptors." Conversely, team members who are logical, detached, and analytical take an objective approach to problems. They challenge others when their logic tells them that they are right. About 51% of Americans are objective decision-makers, as are 68% of engineering students and 95% of top executives.

As it is important to have a variety of information-collection approaches on a design team, it is equally important to have a range of deliberation styles. Although engineers are trained to make decisions based on objective measures, the greatest number of decisions faced in every design problem have incomplete, inconsistent, qualitative information requiring subjective evaluation. For objective designers the following may help in working with the team:

- Encourage objective team members to pay attention to the feelings of others. Gut feelings are often right, and sometimes a lack of information forces one to rely on these feelings.
- Help objective team members understand that how the team functions is as important as what is accomplished. If there is acrimony, no decisions will be made.
- Remind objective team members that not everyone likes to discuss a topic merely for the sake of argument. Others may drop out from exhaustion and be taken to be conceding the point.
- Encourage objective team members to express how they feel about the outcome once in a while. Objective decision-makers may have trouble expressing feelings.

Subjective people are in a minority on most design teams. Thus, they must develop techniques to get their opinions heard and not get their sensitivities hurt.

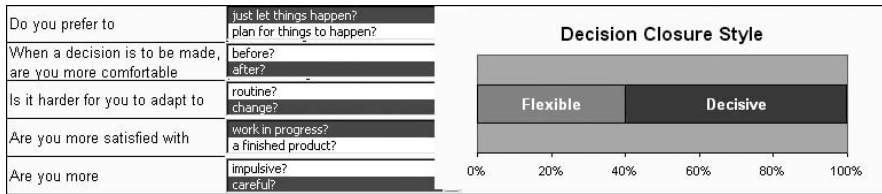


Figure 3.9 Decision closure style personal problem-solving dimension.

Here are some ideas:

- Help subjective team members to realize that it is all right to disagree and argue.
- Reassure subjective team members that while harmony is important, not every resolved issue will satisfy everyone even if consensus is reached.
- Reinforce to subjective team members that discussions about ideas are not personal attacks.

The fifth and final personality dimension relates to the need to actually come to a conclusion during decision making. **Decision closure style** ranges from **flexible** to **decisive**. For a rough estimate of your or a colleague's decision closure style, answer the questions in Fig. 3.9.

Some people are flexible and others are decisive. If a person goes with the flow; is flexible, adaptive, and spontaneous; and finds it difficult to make and stick with decisions, he is considered flexible. If, on the other hand, he makes decisions with a minimum of stress and likes an environment that is ordered, scheduled, controlled, and deliberate, then he is decisive. About half of all Americans are decisive, as are 64% of engineering students and 88% of top executives. One characteristic of flexible decision makers is that they have a tendency to procrastinate because they want to remain adaptive. This can make working with them difficult. The following are some suggestions for flexible decision-makers on the team:

- Give flexible decision-makers plans in advance so that they can think about them in their own time.
- Acknowledge the flexible decision-maker's contribution as a step toward moving to closure. Remind them that problems are solved one step at a time.
- Set clear decision deadlines in advance.
- Encourage feedback from flexible decision-makers so that they can think about the direction of their thoughts.
- Encourage flexible decision-makers to settle on something and live with it a while before redesigning. Encourage them to take a clear position and stick to it. This may be difficult for them to do.

A contrary characteristic of decisive people is that they tend to jump to conclusions. This too can adversely affect teamwork, as many ideas may be generated

and consensus may be needed to reach a decision. Here are some suggestions for slowing down decisive people:

- Ask decisive people questions about their decision process. Remind them that most problems need to be subdivided into smaller problems to be solved.
- Let decisive people organize the data collection and review process.
- Utilize techniques, such as brainstorming, that suppress judgment. Do not let them settle on the first good idea they hear.
- Remind decisive people that they are not always right.

This discussion may seem like a lot of detail for an engineering book. Research has shown, however, that paying attention to the psychological makeup of a team is critical.

3.4 CHARACTERISTICS OF CREATORS

Some people seem naturally more creative than others. Before describing the characteristics of a creative design engineer, let us clarify what we mean by “creative.” A creative solution to a problem must meet two criteria: it must solve the problem in question, and it must be original. Solving a problem involves understanding it, generating solutions for it, evaluating the solutions, deciding on the best one, and determining what to do next. Thus, creativity is more than just coming up with good ideas. The second criterion, originality, depends on the knowledge of the designer and of society as a whole. What is new and original to one person may be old hat to another. If someone who has never before experienced a wheel designs one, then it is original for that person. But it is society that assesses “originality” and labels a solution or a person “creative.”

As discussed earlier, all humans have the same cognitive, or problem-solving, structure. Why is it, then, that some engineers can generate ingenious ideas while others, who may be brilliant at complex analysis, cannot come up with new concepts no matter how hard they try? There has been a lot of research on creativity, yet this trait is still not very well understood. The best way to understand the results of the research to date is in terms the relationship of creativity to other attributes.

Creativity and intelligence. There appears to be little correlation between creativity and intelligence.

Creativity and visualization ability. Creative engineers have good ability to visualize, to generate and manipulate visual images in their heads. We have seen before that people represent information in their minds in three ways: as semantic information (words), as graphical information (visual images), and as analytical information (equations or relationships). Words and equations convey serial information. They are generally understood on the basis of word order or the order of variables and constants. Pictures, or visual images, on the other hand, contain parallel information—you can see many different

The odds are greatly against you being immensely smarter than everyone else.
— John R. Page, *Rules of Engineering*

things in a single image. Some people are very good at decomposing and manipulating visual images in their heads, whereas others are not. It appears, however, that the ability to manipulate complex images of mechanical devices can be improved with practice. This may be related to the formation of more information-rich chunks having functional information or to some other mechanism.

Creativity and knowledge. The model of the information-processing system implies that all designers start with what they know and modify this to meet the specific problem at hand. At every step of the way, the process involves small movements away from the known, and even these small movements are anchored in past experience. Since creative people form their new ideas out of bits of old designs, they must retain a storehouse of images of existing mechanical devices in their long-term memory. Thus, in order to be a creative mechanical designer, a person must have knowledge of existing mechanical products.

Additionally, part of being creative is being able to evaluate the viability of ideas. Without knowledge about the domain, the designer cannot evaluate the design. Knowledge about a domain is only gained through hard work in that domain. Thus, a firm foundation in engineering science is essential to being a creative designer of mechanical devices. For example, during World War II many people sent ideas for weapons to the Department of War. Some were very far-fetched ideas for death rays or for building 5-mile-high walls or domes over Europe to stop the bombers. These were very original but unworkable and were therefore not creative. The “inventors” had good intentions but lacked the knowledge to develop creative solutions to the war problems.

Creativity and partial solution manipulation. Since new ideas are born from the combination of parts of existing knowledge, the ability to decompose and manipulate this knowledge seems to be an important attribute of a creative designer. This attribute, more than any other so far discussed, appears to become stronger with exercise. Although there is no scientific evidence to support this contention, anecdotal evidence does support it.

Creativity and risk taking. Another attribute of creative engineers is the willingness to take an intellectual chance. Fear of making a mistake or of spending time on a design that in the end does not work is characteristic of a noncreative individual. Edison tried hundreds of different lightbulb designs before he found the carbon filament.

Creativity and conformity. Creative people also tend to be nonconformists. There are two types of nonconformists: constructive nonconformists and

obstructive nonconformists. Constructive nonconformists take a stand because they think they are right. Obstructive nonconformists take a stand just to have an opposing view. The constructive nonconformist might generate a good idea; the obstructive nonconformist will only slow down the design progress. Creative engineers are constructive nonconformists who may be hard to manage since they want to do things their own way.

Creativity and technique. Creative designers have more than one approach to problem solving. If the process they initially follow is not yielding solutions, they turn to alternative techniques. A number of books listed in Section 3.7 give methods to enhance creativity. Many of the techniques covered in these are woven into the mechanical design techniques presented in the remainder of this book. This is especially true in the chapters on concept and product generation (Chaps. 7 and 9).

Creativity and environment. If the work environment allows risk taking and nonconformity and encourages new ideas, creativity will be higher. Further, if teammates and other colleagues are creative, the environment for creativity is greatly enhanced. In the discussion of teams in Section 3.5, it is stated that, on a team, the sum is greater than the parts. This is especially true for creativity.

Creativity and practice. Creativity comes with practice. Most designers find that they have creative phases in their careers—periods when they have many good ideas. During these times the environment is supportive and one good idea builds on another. However, even with a supportive environment, practice enhances the number and quality of ideas.

To summarize, the creative designer is generally a visualizer, a hard worker, and a constructive nonconformist with knowledge about the domain and the ability to dissect things in his or her head. Even designers who do not have a strong natural ability can develop creative methods by using good problem-solving techniques to help decompose the problem in ways that maximize the potential for understanding it, for generating good solutions, for evaluating the solutions, for deciding which solution is best, and for deciding what to do next.

One final comment: There are many design tasks that require talents very different from those used to describe a creative person. Design requires much attention to detail and convention and demands strong analytic skills. Therefore, there are many good designers who are not particularly creative individuals; a design project requires people with a variety of skills and talents.

3.5 THE STRUCTURE OF DESIGN TEAMS

The material already covered describes an individual designer. However, because of the complexity of most products, design work is generally done by design teams. As shown in Fig. 3.10, the complexity of mechanical devices has grown rapidly over the last 200 years. Gone are the days when a single individual could design an entire product. Even Edison had a team of others that worked with

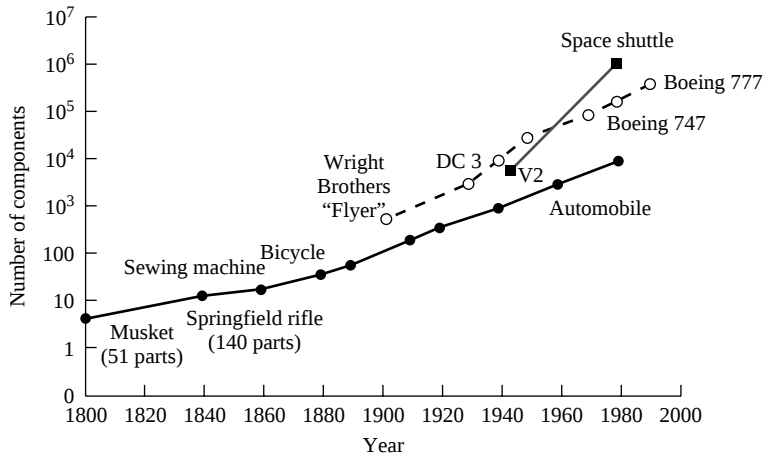


Figure 3.10 Increasing complexity in mechanical design.

him. For example, the Boeing 777 aircraft, which has over 5 million components, required over 10 thousand person-years of design time. Thousands of designers worked over a three-year period on the project. Obviously, a single designer could not approach this effort.

Modern design problems require a design *team*—a small number of people with complementary skills who are committed to a common purpose, common performance goals, and a common approach for which they hold themselves mutually accountable.

A team is a group of people with complementary skills who are committed to a common purpose, performance goals, and approach for which they hold themselves mutually accountable. A group is not necessarily a team. Groups that interact primarily to share information and to help each individual perform within his or her area of responsibility is not a team. An effective team is more than the sum of its parts. Important points about teams are the following:

1. Teamwork is central to success in engineering as most problems are made of many interdependent subparts, all of which must be solved concurrently. Teams bring together complementary skills and experiences, which are needed to solve many engineering problems.
2. Management takes risks in forming teams as a team must be empowered to make decisions, removing this responsibility from the management.
3. Teams establish communication to support real-time problem solving.
4. Teams develop decisions by consensus rather than by authority. This leads to more robust decisions.

In the most basic sense, teams solve design problems in the same way an individual does—understanding, generating, evaluating, and decision making.

A team is a group of people in search of a common understanding.

However, there are some important differences.

- Team members must learn how to *collaborate* with each other. Collaboration means more than just working together—it means getting the most out of other team members. The suggestions that follow help develop a collaborative team.
- Teams are generally empowered to make decisions. Since these are team decisions, members must *compromise* to reach them. Empowering teams to make these decisions means that management takes a risk in giving up responsibility for them. Further, developing decisions by *consensus* rather than by authority leads to more robust decisions.
- Team members must establish *communication* to support real-time problem solving. Further, members need to ensure that the others have the same understanding of design ideas and evaluations that they have. It is very difficult for people with different areas of expertise to develop a shared vision of the problem and its potential solutions. Developing this shared vision requires the development of a rich understanding of the problem.
- It is important that team members and management be *committed* to the good of the team. If they are not, it will be difficult reaching the other team goals.

To address what is special about teams, in this chapter we first itemize the different technical roles people play on teams and then, in Section 3.6, we address building teams and maintaining team health.

3.5.1 Members of Design Teams

In this section, we list the individuals who might fill a role on a product design team. The roles on a design team will vary with product development phase and from product to product, and the titles will vary from company to company. Each position on the team is described as if filled by one person. In a large design project, there may be many persons filling that role, whereas in a small project one individual may fill many roles.

Product design engineer. The major design responsibility is carried by the product design engineer (hereafter referred to as the *design engineer*). This individual must be sure that the needs for the product are clearly understood and that engineering requirements are developed and met by the product. This usually requires both creative and analytical skills. The design engineer must bring knowledge about the design process and knowledge about specific technologies to the project. The person who fills this position usually has a four-year engineering degree. In smaller companies he or she may be a nondegree designer who has extensive experience in the product area. For most product design projects, more than one design engineer will be involved.

Product manager. In many companies, this individual has the ultimate responsibility for the development of the product and represents the major link between the product and the customer. Because the product manager is accountable for the success of the product in the marketplace, he or she is also often referred to as the *marketing manager* or the *product marketing manager*. The product manager is often from the sales or customer service department.

In order to initiate a design project, management must appoint the nucleus of a design team—at a minimum, a design engineer and a product manager.

Manufacturing engineer. Design engineers generally do not have the necessary breadth or depth of knowledge about various manufacturing processes to fully support the design of most products. This knowledge is provided by the manufacturing or industrial engineer, who must have a grasp not only of in-house manufacturing capabilities but also of what the industry as a whole has to offer.

Designer. In many companies, the design engineer is responsible for specification development, planning, conceptual design, and the early stages of product design. The project is then turned over to *designers*, who finish detailing the product and developing the manufacturing and assembly documentation. Designers are often CAD experts with two-year technology degrees. At some companies designers are the same as design engineers.

Technician. The technician aids the design engineer in developing the test apparatus, performing experiments, and reducing data in the development of the product. The insights gained from the technician's hands-on experience are usually invaluable.

Materials specialist. In some products, the choice of materials is forced by availability. In others, materials may be designed to fit the needs of the product. The more a product moves away from the use of known, available materials, the more a materials specialist is needed as a member of the design team. This individual is usually a degreed materials engineer or a materials scientist. Often the materials specialist will be a vendor's representative who has extensive knowledge about the design potential and limitations of the vendor's materials. Many vendors actually provide design assistance as part of their service.

Quality control/quality assurance specialist. A quality control (QC) specialist has training in techniques for measuring a statistically significant sample to determine how well it meets specifications. This inspection is done on incoming raw materials, incoming products from vendors, and products produced in-house. A quality assurance (QA) specialist makes sure that the product meets any pertinent codes or standards. For example, for medical products, there are many FDA (Food and Drug Administration) regulations that must be met. Often QC and QA are covered by one person.

Analyst. Many engineers work as analysts. Analysts usually perform complex mathematical studies of design performance using finite-element

methods, thermal system modeling, or other advanced software. They are generally specialists who focus on one type of system or method.

Industrial designer. Industrial designers are responsible for how a product looks and how well it interacts with consumers; they are the stylists who have a background in fine arts and in human factors analysis. They often design the envelope within which the engineer has to work.

Assembly manager. Where the manufacturing engineer is concerned with making the components from raw materials, the assembly manager is responsible for putting the product together. As you will see in Chap. 11, concern for the assembly process is an important aspect of product design.

Vendor's or supplier's representatives. Very few products are made entirely in one factory. In fact, many manufacturers outsource (i.e., have suppliers provide) 70% or more of their product. Usually there will be many suppliers of both raw and finished goods. There are three types of relationships with suppliers: (1) partnership—the supplier takes part in the process beginning with requirements and concept development; (2) mature—the supplier relies on the parent company's requirements and concepts to develop needed items; and (3) parental—the supplier builds only what the parent company specifies. Often it is important to have critical suppliers on the design team, as the success of the product may be highly dependent on them.

As Fig. 3.11 illustrates, having a design team made up of people with varying views may create difficulties, but teams are essential to the success of a product.

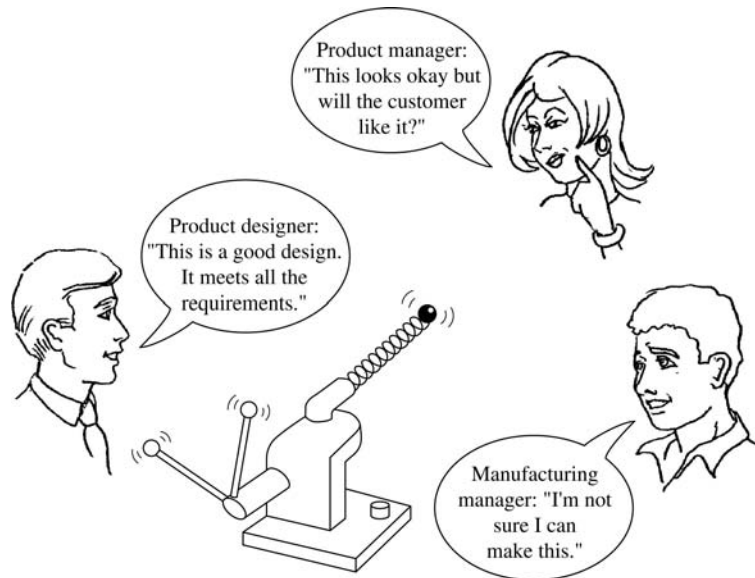


Figure 3.11 The design team at work.

The breadth of these views helps in developing a quality design. Part of the promise of PLM is to help all these different contributors communicate in a consistent and productive manner.

3.5.2 Design Team Management

Since projects require team members with different domains of expertise, it is valuable to look at the different structures of teams in an organization. This is important because product design requires coordination across the functions of the product and across the phases in the product's development process. Listed next are the five types of project structures. The number in parentheses is the percentage of development projects that use that type. These results are from a study of 540 projects in a wide variety of industries.

Functional organization (13%). Each project is assigned to a relevant functional area or group within a functional area. A functional area focuses on a single discipline. For aircraft manufacturers, Boeing, for example, the main functions are aerodynamics, structures, payload, propulsion, and the like. The project is coordinated by functional and upper levels of management.

Functional matrix (26%). A project manager with limited authority is designated to coordinate the project across different functional areas or groups. The functional managers retain responsibility and authority for their specific segments of the project.

Balanced matrix (16%). A project manager is assigned to oversee the project and shares with the functional managers the responsibility and authority for completing the project. Project and functional managers jointly direct many work-flow segments and jointly approve many decisions.

Project matrix (28%). A project manager is assigned to oversee the project and has primary responsibility and authority for completing the project. Functional managers assign personnel as needed and provide technical expertise.

Project team (16%). A project manager is put in charge of a project team composed of a core group of personnel from several functional areas or groups, assigned on a full-time basis. The functional managers have no formal involvement. Project teams are sometimes called "Tiger teams," "SWAT teams," or some other aggressive name, because this is a high-energy structure and the team is disbanded after the project is completed.

What is important about these structures is that some of them are more successful than others. Structures focused on the project are more successful than those built around the functional areas in the company (Fig. 3.12). Here the balanced matrix, project matrix, and project teams resulted in a higher percentage of success across all measures. Thus, when planning for a design project, organize the talent around the project whenever possible.

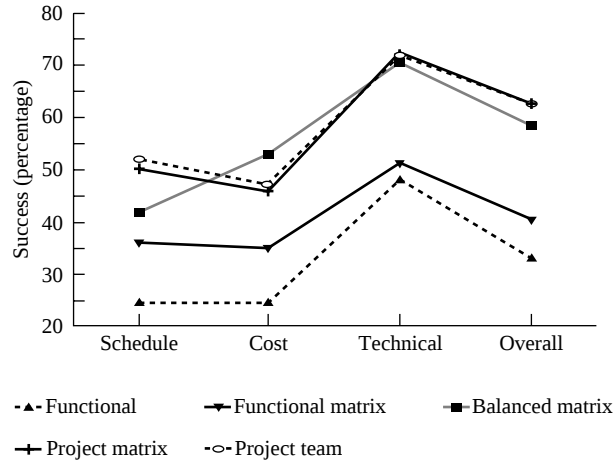


Figure 3.12 Project successes versus team structure.

3.6 BUILDING DESIGN TEAM PERFORMANCE

It can be very exciting being part of a team that is productive and is making good use of all the members. Conversely, it can be hellish working on a team that is not functioning very well. So the goal of this section is to help you build and maintain successful teams. To help ensure success, we will use *Team Contracts*, *Team Meeting Minutes*, and *Team Health Assessments*. Each of these encourages behavior that leads to a successful team experience.

According to a leading book on teams, there are ten characteristics of a successful team. Included in the description of each of these characteristics is a guide to where this text presents material to help make teams successful.

1. **Clarity in goals.** The process developed in this book focuses on goals during process planning in Chap. 4 and for the product itself in Chap. 6. Further, the Team Contract suggested later in this section encourages documenting the immediate team goals.
2. **Plan of action.** Chapter 4 is all about project planning.
3. **Clearly defined roles.** We have already discussed roles, and documenting them is part of the Team Contract.
4. **Clear communication.** Team Contracts, Team Meeting Minutes, and Team Health Assessments (all in this chapter) plus virtually all the process methods in this book are designed to help with communication.
5. **Beneficial team behaviors.** As with communication, the material in this book is designed to result in beneficial behaviors.
6. **Well-defined decision process.** The decision process is introduced in Chap. 4 and is the focus of Chap. 8.

7. **Balanced participation.** Equal division of work is very important for a successful team. This is further discussed later in this chapter.
8. **Established ground rules.** This is discussed later in this chapter.
9. **Awareness of team process.** This is what we are talking about in this entire chapter.
10. **Use of sound generation/evaluation approach.** As introduced in Chap. 1, the seven activities of the design process are: *Establish the Need, Plan, Understand, Generate, Evaluate, Decide, and Document*. Generate and Evaluate are covered in Chaps. 7–12.

To set the foundation for future work, the remainder of this chapter covers Team Contracts, Team Meeting Minutes, and Team Health Assessments.

3.6.1 Team Contract

A good starting point for a team is with a team contract. Team contracts are seldom done in industry because the basics of it are assumed in the employment contract, and it is further assumed that people know how to work to make a team successful. Here we will use a contract as both a learning tool and as a way to increase the odds of team success.

Figure 3.13 shows an example Team Contract. The first section is for the assignment of roles on the team and goals for the team. As suggested in the list of characteristics for a successful team, the goals and roles need to be known and agreed to. Roles can be developed from the list in Section 3.5 and make the goals be as specific as possible. Much in Chaps. 4 and 5 focuses on goals.

In the second section the team members sign, indicating they agree to a list of performance expectations that are shown in the example. Additionally, the form has room for other expectations to which the team may want to agree. The final section on the form is for strategies for conflict resolution. Hopefully these won't be needed, but, like any other contract, methods for problem resolution need to be addressed at the beginning to prevent difficulties later. Suggested strategies are shown in the example.

3.6.2 Team Meeting Minutes

Before a meeting begins, it is essential to have an agenda. Without an agenda, meetings wander and it is often not clear whether anything was accomplished. Thus, the purpose of the first section in the team meeting minutes (Fig. 3.14) is to itemize the agenda. Agendas should be written in terms of the goals of the meeting. Agenda items such as "Present the results of the stress analysis" are not sufficient. Why are the results being presented? What is to be accomplished by telling others the results? It is better to state this in terms of what is to be accomplished: "Decide how the stress affects the assembly's performance" or "Determine if the stress is low enough to meet the requirements of the system."

The second section itemizes the high points of the discussion. To understand why taking notes about the high points is so important, consider the results of an



Team Contract						
Design Organization: The B Team			Date: Jan. 2, 2009			
Team Member	Roles		Signature			
Jason Smathers	Lead designer		Jason Smathers			
Brittany Spars	Structural engineer		Brittany Spars			
Deon Warner	Systems engineer		Deon Warner			
Team Goals			Responsible Member			
1. Develop layout and initial input to solid model.			JS			
2. Analyze for fatigue and other failures.			BS			
3. Detail latching mechanism.			JS			
4. Develop wiring plan.			DW			
5.						
Team Performance Expectations			Initial			
• Strive to complete all assigned tasks before or by deadlines.			JS	BS	DW	
• Complete all tasks to the best of ability.			JS	BS	DW	
• Listen carefully and attentively to all comments at meetings.			JS	BS	DW	
• Accept and give criticism in a professional manner.			JS	BS	DW	
• Focus on results before the fact, rather than excuses after.			JS	BS	DW	
• Provide as much notice as possible of commitment problems.			JS	BS	DW	
• Attend and participate in all scheduled group meetings.			JS	BS	DW	
Strategies for Conflict Resolution						
• Amend contract with deadlines for agreed to tasks.						
• Reward entire team for goals met with some treat or social gathering.						
• As a team, go to a higher authority for assistance with a team problem.						
• Don't kill messengers. Seek to encourage the airing of problems.						
The Mechanical Design Process			Designed by Professor David G. Ullman			
Copyright 2008, McGraw-Hill			Form # 2.0			

Figure 3.13 Example team contract.



Team Meeting Minutes		
Design Organization: The C Team		Date: Jan. 30, 2009
Agenda 1. Finalize the plan for the exotherm system. 2. Decide on the final shape for the housing. 3. Resolve how to complete task 3. 4. Plan the postproject party. 5. 6.		
Discussion: Jason, Brittany, and Deon attended. The meeting lasted an hour. The agenda was fully covered and new issues were added to the list for the next meeting.		
Decisions Made 1. Exotherm plan finalized. See Attachment A. 2. Housing alternative 3 was chosen. 3.		
Action Items	Person Responsible	Deadline
Jason details Housing alternative 3	JS	Thursday
Brittany to plan party	BS	2/10
Deon will assist Brittany to get Task 3 completed by Thursday	BS	Thursday
Team member: Jason Smathers	Date for next meeting: Thursday	
Team member: Brittany Spars		
Team member: Deon Warner		
Team member:		
<i>The Mechanical Design Process</i> Copyright 2008, McGraw-Hill		Designed by Professor David G. Ullman Form # 3.0

Figure 3.14 Team meeting minutes.

experiment where a group was asked 2 weeks after a meeting to recall specific details of that meeting. In recounting the meeting they

- Omitted 90% of the specific points that were discussed.
- Recalled half of what they did remember incorrectly.
- Remembered comments that were not made.
- Transformed casual remarks into lengthy orations.
- Converted implicit meanings into explicit comments.

Recording the decisions made is even more important. Often decisions are clear. For example, “Choose to use 5056-T6 aluminum for the brace” or “The potential difference on anode and cathode of the X-ray tube will be 140 keV.” However, if you listen carefully to unstructured meetings, you find that they wander from topic to topic. When one topic gets difficult because some of the parties disagree or more information is needed, the conversation moves to another topic with no resolution of the initial topic. If stuck, decide what to do to get unstuck and record that call for action. For example, “A decision was made to gather more information on material x” or “We will use Belief Maps to help the team work toward agreement.” These decisions lead directly to the most important item in the meeting minutes, the action items—an itemized list of what is to happen next. State each action item as a clear deliverable, assign the responsible party, and determine by when it is to be done.

3.6.3 Team Health Assessment

One of the most important activities is assessing the team’s health. A form for assessing team health is shown in Fig. 3.15. This form includes 17 measures (with room for more) to be assessed periodically by the team to measure how it is doing. For each measure, the response ranges from strongly agree to strongly disagree, with attention needed to remedy problems in areas where at least *one person* does not agree with the measure. The team needs to devise remedies for these “problem areas.” Not doing so allows problems to fester and worsen.

This assessment should be used periodically and especially when any team members experience one of the following:

- A loss of enthusiasm
- A sense of helplessness
- A lack of purpose or identity
- Meetings in which the agenda is more important than the outcome
- Cynicism and mistrust
- Interpersonal attacks made behind peoples backs
- Floundering
- Overbearing or reluctant team members



Team Health Assessment							
Team Assessed:		Date:					
SA = Strongly Agree, A = Agree, N = Neutral, D = Disagree, SD = Strongly Disagree, NA = Not Applicable							
Measure		SA	A	N	D	SD	NA
1	Team mission and purpose are clear, consistent and attainable.	☐	☐	☐	☐	☐	☐
2	I feel that I am part of a team.	☐	☐	☐	☐	☐	☐
3	I feel good about the team's progress.	☐	☐	☐	☐	☐	☐
4	Respect has been built within the team for diverse points of view.	☐	☐	☐	☐	☐	☐
5	Team environment is characterized by honesty, trust, mutual respect, and team work.	☐	☐	☐	☐	☐	☐
6	The roles and work assignments are clear.	☐	☐	☐	☐	☐	☐
7	Team treats every member's ideas as having potential value.	☐	☐	☐	☐	☐	☐
8	Team encourages individual differences.	☐	☐	☐	☐	☐	☐
9	Conflicts within the team are aired and worked to resolution.	☐	☐	☐	☐	☐	☐
10	Team takes time to develop consensus by discussing the concerns of all members to arrive at an acceptable solution.	☐	☐	☐	☐	☐	☐
11	Decisions are made with input from all in a collaborative environment.	☐	☐	☐	☐	☐	☐
12	The environment encourages communication and does not "kill the messenger" when the news is bad.	☐	☐	☐	☐	☐	☐
13	When one team member has a problem others jump in to help.	☐	☐	☐	☐	☐	☐
14	Dysfunctional behavior is dealt with in an appropriate manner.	☐	☐	☐	☐	☐	☐
15	When someone on the team says they are going to do something, the team can count on it being done.	☐	☐	☐	☐	☐	☐
16	There is no "them and us" on the team.	☐	☐	☐	☐	☐	☐
17	Our team cultivates a "what we can learn" attitude when things do not go as expected.	☐	☐	☐	☐	☐	☐
18		☐	☐	☐	☐	☐	☐
19		☐	☐	☐	☐	☐	☐
20		☐	☐	☐	☐	☐	☐
Remedies for improving the Neutral (N), Disagree (D) and Strongly Disagree(SD) responses:							
Assessor:							
The Mechanical Design Process		Designed by Professor David G. Ullman					
Copyright 2008, McGraw-Hill		Form # 3.0					

Figure 3.15 Team health assessment.

3.7 SUMMARY

- The human mind uses the long-term memory, the short-term memory, and a controller in the internal environment when problem solving.
- Knowledge can be considered to be composed of chunks of information that are general, domain-specific, or procedural in content.
- The short-term memory is a small (seven chunks, features, or parameters) and fast (0.1-sec) processor. Its properties determine how we solve problems. We use the external environment to augment the size of the short-term memory.
- The long-term memory is the permanent storage facility in the brain. It is slow to remember, it is fast to recall (sometimes), and it never gets full.
- Creative designers are people of average intelligence; they are visualizers, hard workers, and constructive nonconformists with knowledge about the problem domain. Creativity takes hard work and can be aided by a good environment, practice, and design procedures.
- Because of the size and complexity of most products, design work is usually accomplished by teams rather than by individuals.
- Working in teams requires attention to every team member's problem-solving style (including yours)—introverted or extroverted, fact or possibility, verbal or visual, objective or subjective, or decisive or flexible.
- It is important to have team goals and roles, keep meeting minutes, and assess team health.
- Many activities can help build team health.

3.8 SOURCES

- Adams, J. L.: *Conceptual Blockbusting*, Norton, New York, 1976. A basic book for general problem solving that develops the idea of blocks that interfere with problem solving and explains methods to overcome these blocks; methods given are similar to some of the techniques in this book.
- Larson, E., and D. Gobeli: "Organizing for Product Development Projects," *Journal of Product Innovation Management*, No. 5, pp. 180–190, 1988. The study in Section 3.5.2 on design team management is from this paper.
- Koberg, D., and J. Bagnall: *The Universal Traveler: A Systems Guide to Creativity, Problem Solving and the Process of Reaching Goals*, Kaufman, Los Altos, Calif., 1976. A general book on problem solving that is easy reading.
- Miller, G. A.: "The Magical Number Seven, Plus or Minus Two: Some Limits on Our Capacity for Processing Information," *Psychological Review*, Vol. 63, pp. 81–97, 1956. The classic study of short-term memory size, and the paper with the best title ever.
- Newell, A., and H. Simon: *Human Problem Solving*, Prentice Hall, Englewood Cliffs, N.J., 1972. This is the major reference on the information processing system. A classic psychology book.
- Plous, S.: *The Psychology of Judgment and Decision Making*, McGraw-Hill, New York, 1993. The importance of meeting notes example is from this interesting book.
- Weisberg, R. W.: *Creativity: Genius and Other Myths*, Freeman, San Francisco, 1986. Demystifies creativity; the view taken is similar to the one in this book.

The next five titles are all good books on developing and maintaining teams.

Belbin, R. M.: *Management Teams*, Heinemann, New York, 1981.

Cleland, D. I., and H. Kerzner: *Engineering Team Management*, Van Nostrand Reinhold, New York, 1986.

Johansen, R., et al.: *Leading Business Teams*, Addison-Wesley, New York, 1991.

Katzenbach, J. R., and D. Smith: *The Wisdom of Teams*, Harvard Business School Press, 1993.

Scholtes, P. R., et al.: *The Team Handbook*, 3rd edition, Oriel Inc, 2003.

The problem-solving dimensions in Section 3.3.5 are based on the Myers-Briggs Type Indicator. These titles give more details on this method.

Keirsey, D., and M. Bates: *Please Understand Me*, 5th ed., Prometheus Nemesis, 1978.

Kroeger, O., and J. M. Thuesen: *Type Talk at Work*, Delta, 1992.

Kroeger, O., and J. M. Thuesen: *Type Talk*, Delta, 1989.

3.9 EXERCISES

- 3.1 Develop a simple experiment to convince a colleague that the short-term memory has a capacity of about seven chunks.
- 3.2 Think of a simple object, write about it, and sketch it in as many ways as possible. Refer to Table 2.1 and Fig. 3.4 to encourage a range of language and abstraction.
- 3.3 Describe a mechanical design problem to a colleague. Be sure to describe only its function. Have the colleague describe it back to you in different terms. Did your colleague understand the problem the same way as you? Was the response in terms of previous partial solutions?
- 3.4 During work on a team, identify the secondary roles each person is playing. Can you identify who fills each role?
- 3.5 For a new team begin with these team-building activities.
 - a. *Paired introductions.* Get to know each other by asking questions such as
 - What is your name?
 - What is your job (class)?
 - Where did you grow up (go to school)?
 - What do you like best about your job (school)?
 - What do you like least about your job (school)?
 - What are your hobbies?
 - What is your family like?
 - b. *Third-party introductions.* Have one member of the team tell another the information in (a). Then the second member introduces the first member to the rest of the team using all the information that he or she can remember. It makes no difference if the team heard the initial introduction.
 - c. *Talk about first job.* Have each member of the team tell the others about his or her first job or other professional experience. Information such as this can be included:
 - What did you do?
 - How effective was your manager?
 - What did you learn about the real world?

- d. *“What I want for myself out of this.”* Have each member of the team tell the others for 3 to 5 min what his or her goals are for participation in the project. What do they want to learn or do, and why? Consider personal goals such as getting to know other people, feeling good about oneself, learning new skills, and other nontask goals.
 - e. *Team name.* Have each person write down as many potential team names as possible (at least five). Discuss the names in the team, and choose one. Try to observe who plays which secondary role.
- 3.6 Pick an item from the team health assessment. For that item, one member of a four-person team checks “Strongly Disagrees.” Develop a list of actions you would take as a team leader or team member.



3.10 ON THE WEB

Templates for the following documents are available on the book’s website: www.mhhe.com/Ullman4e

- Personal Problem Solving Dimensions
- Team Contract
- Team Meeting Minutes
- Team Health Assessment

4

C H A P T E R

The Design Process and Product Discovery

KEY QUESTIONS

- What are the six phases of the mechanical design process?
- What are the three prime sources for new products?
- What does it mean for a product to be “mature”?
- How can a SWOT analysis help choose which products to develop?
- How did Benjamin Franklin contribute to decision making?
- What are the six basic decision-making activities?

4.1 INTRODUCTION

In this chapter, we introduce the major phases in the design process and tackle the first of them, discovering the need. The six-phase design process established here sets the structure for the rest of this book. Since design is fundamentally the effort to fulfill a need, discovering the need is always the first phase in the process. Because there are always more needs than there are resources to meet them, key here is deciding which product ideas to develop. Thus, in this chapter we also introduce the basics of decision making. Making good decisions is probably the most important and least studied engineering skill. We will refine decision making when choosing a concept and again when making Product Development decisions.

4.2 OVERVIEW OF THE DESIGN PROCESS

Regardless of the product being developed or changed, or the industry, there is a generic set of phases that must be accomplished for all projects. These are listed

Design is a process—not just building hardware.
—Tim Carver, OSU student, 2000

in Fig. 4.1. They are a refinement of the phases in a product's life cycle (Fig. 1.8) that are of concern to the designer. For each phase, there are a series of activities that need to be accomplished. The phases and activities are briefly introduced in this chapter and refined throughout the rest of the book. After this introduction, the first phase, Product Discovery, is explained in detail.

This design process, as shown, applies to design of systems, subsystems, assemblies, and components. It applies to new, innovative products and to changes in existing products. Of course, the detail and emphasis will change with the level of decomposition and with the amount of change needed. To help introduce the phases and how they are used at all levels in a product's decomposition consider the design of a General Electric CT Scanner.

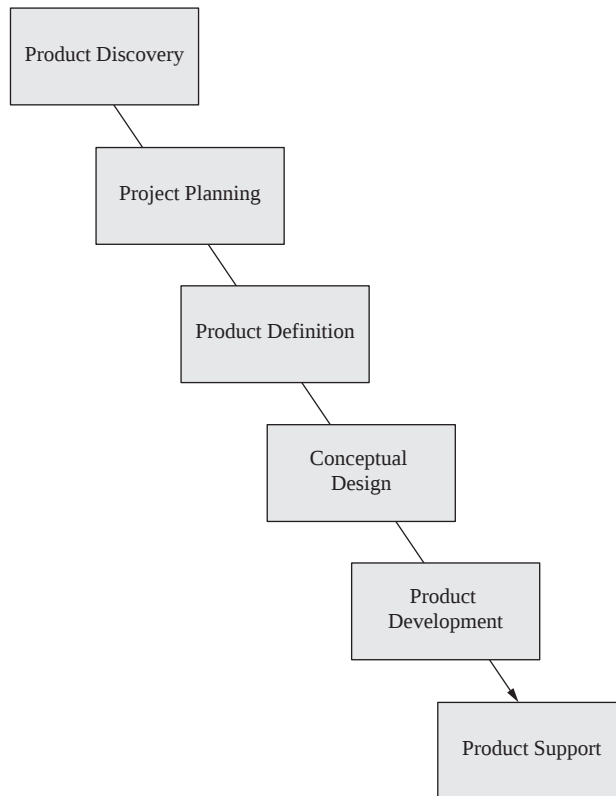


Figure 4.1 The mechanical design process.

General Electric designs and manufactures many different types of products including home appliances, lightbulbs, jet engines, and a host of medical products. One of the products developed by GE's healthcare business is the CT Scanner shown in Fig. 4.2. The full name of the technology used in this scanner is X-ray Computed Tomography (CT). CT is a diagnostic imaging technique that can produce solid images of the organs inside patients. A CT system consists of a patient table that can be positioned and moved through the bore of the gantry. Beneath the sleek outer casing, the gantry houses a frame that holds an X-ray tube and a detector. The X-ray tube is on the top at the 1 o'clock position in Fig. 4.3 and the arc-shaped detector is on the bottom at the 7 o'clock position. The frame, X-ray tube, and detector rotate around the patient at 120 rpm. This means that there is a centrifugal acceleration on the components of more than 10gs. Thus, the X-ray tube components experience very large radial body loads and convey centrifugal loading to the gantry support of approximately 2000 N of radial force.

In order to generate images of organs the tube emits rays that pass through the patient, are sensed by the detector, and are processed by a computer, as shown in Fig. 4.4. To accomplish this, the X-ray tube emits bursts of X-rays. During emission, the tube requires 60–100 kW of power. This power must be transmitted to the rotating tube, where the majority of the power is converted into waste heat that must be transferred out of the gantry. Making the design task even more



Figure 4.2 GE CT Scanner. (Source: Reprinted with permission of GE Medical.)

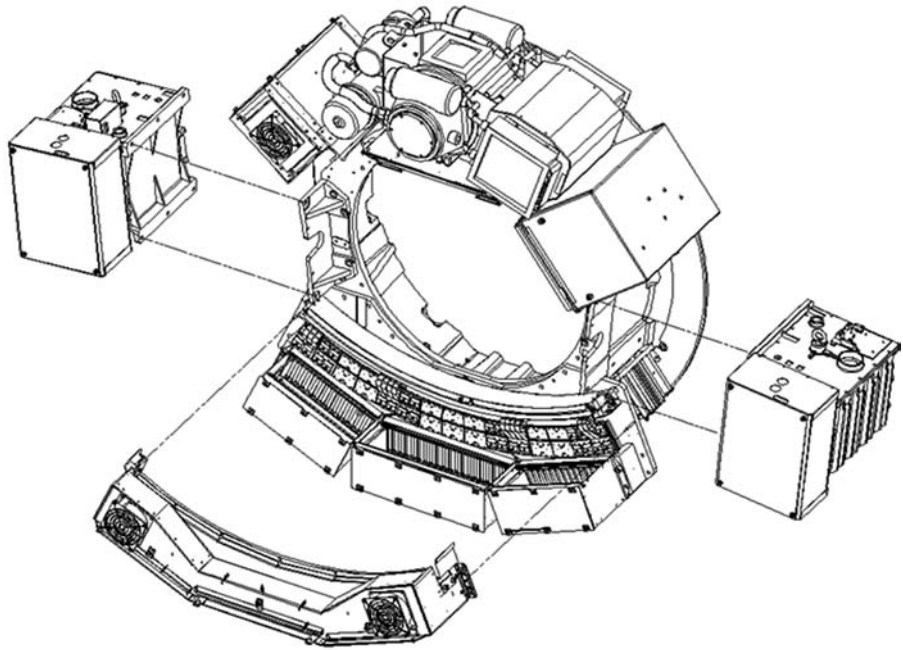


Figure 4.3 The insides of a CT gantry. (Source: Reprinted with permission of GE Medical.)

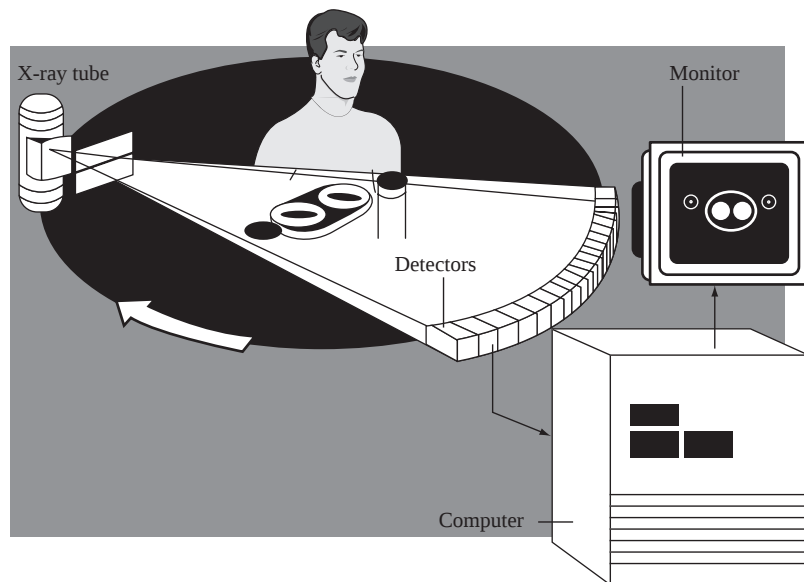


Figure 4.4 How a CT works.

difficult, the anode in the X-ray tube is rotating on an axis perpendicular to the plane of the gantry at 7000–10,000 rpm. The bearings for the anode are operating in a vacuum at a temperature of 450°C.

The design of the X-ray tube is a tremendous undertaking requiring hundreds of design and manufacturing engineers, materials scientists, technicians, purchasing agents, drafters, and quality-control specialists, all working over several years. To recap:

- The *system* is the CT Scanner.
- Major *subsystems* are the patient table and the gantry.
- A major *assembly* in the gantry is the frame with the X-ray tube and detector.
- The X-ray tube itself is a *subsystem* in the frame assembly.
- Two *components* in the X-ray tube are the anode and its bearings.

Regardless of which of these are being designed or changed, there will be a Plan, Product Definition, and Conceptual Design before there are products. These phases, itemized in Fig. 4.1, are common to the design of every system, subsystem, assembly, and component. Let's expand each of the phases.

4.2.1 Product Discovery

Before the original design or redesign of a product can begin, the need for it must be established. As shown in Fig. 4.5, there are three primary sources for design projects: technology, market, and change. We will delve into these sources later in this chapter. Regardless of the source, a common activity at most companies is maintaining a list of potential projects. Since companies have limited people and money, the second activity, after identifying the products, is choosing which

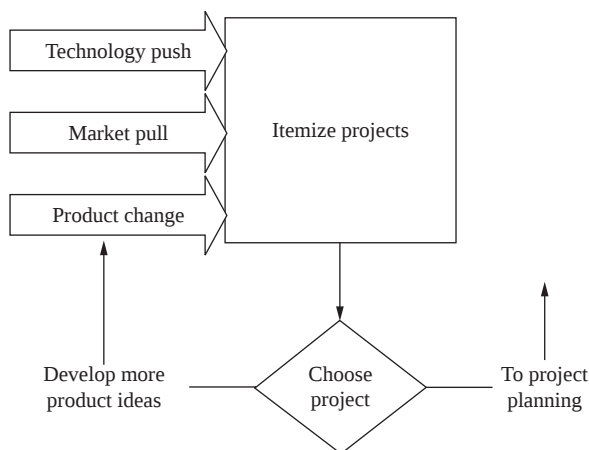


Figure 4.5 The Discovery phase of the mechanical design process.

of them to work on. Sometimes this decision comes before Project Planning as structured in this book, and sometimes it is postponed until later, after Planning, Product Definition, and Conceptual Design has been done and more is known about each of the options. The ordering of work on these phases will be further discussed in Chap. 5.

The GE CT Scanner is a mature product, but new products are advancing the state of the art at a rapid rate. For mature products, design changes focus on improved reliability, cost, and supply chain management. New products are pulled by new imaging applications and performance capability. For the X-ray tube itself, changes are usually in response to market pull for more detailed X-rays and faster times. These system-level needs are projected down as projects to redesign the X-ray tube. Within these projects, the needs are communicated as specifications for higher power, more rotational speed, better heat removal, and other technical changes.

4.2.2 Project Planning

The second phase is to plan so that the company's resources of money, people, and equipment can be allocated and accounted for (Fig 4.6). Planning needs to precede any commitment of resources; however, as with much design activity, this requires speculating about the unknown—and that makes the planning for a product that is similar to an earlier product easier than planning for a totally new one. Since planning requires a commitment of people and resources from all parts of the company, part of the planning is forming the *design team*. As discussed in Chap. 3, few products or even subsystems of products are designed by one person. Additionally, much planning work goes into developing a schedule and estimating the costs. The final goal of the activities in this phase is generating a set of tasks that need to be performed and a sequence for them. Planning is covered in detail in Chap. 5.

The plan for redesigning the X-ray tube is very complex as it is usually only a small part of the plan to redesign the entire CT Scanner to create the next model. Thus, the tasks, schedule, and budget must integrate with many other similar plans.

4.2.3 Product Definition

During the product definition phase (Fig. 4.7), the goal is to understand the problem and lay the foundation for the remainder of the design project. Understanding the problem may appear to be a simple task, but since most design problems are poorly defined, finding the definition can be a major undertaking. In Chap. 6, we will look at a technique to accomplish this. Using this technique, the first activity will be to *identify the customers* for the product. This activity serves as the basis to *generate the customers' requirements*. These requirements are then used to *evaluate the competition* and to *generate engineering specifications*, measurable behaviors of the product-to-be that, later in the design process, will help in

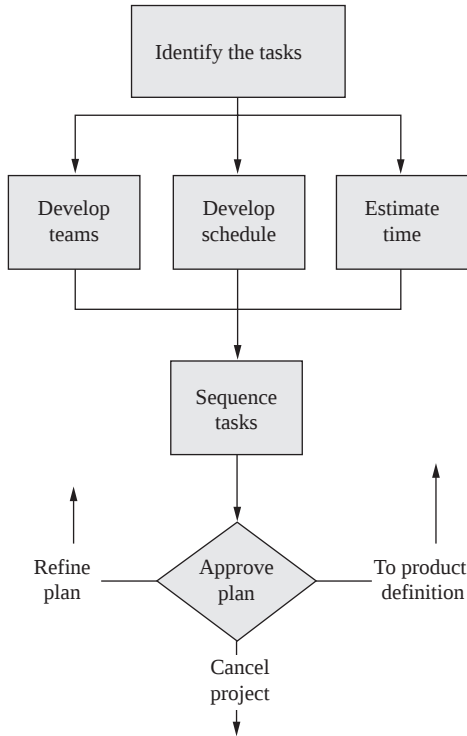


Figure 4.6 The Project Planning phase of the mechanical design.

determining product quality. Finally, in order to measure the “quality” of the product, we *set targets for its performance*.

Often, the results of the activities in this phase determine how the design problem is decomposed into smaller, more manageable design subproblems. Sometimes not enough information is yet known about the product, and decomposition occurs later in the design process.

In redesigning the X-ray tube, the needs are translated into realizable targets for power, rotational speed, heat removal, and other technical specifications. These specifications are developed in concert with other design teams that need to supply the power, structurally support and power the rotating X-ray tube, and dispose of the waste heat.

4.2.4 Conceptual Design

Designers use the results of the Planning and Product Definition phases to generate and evaluate concepts for the product or product changes (Fig 4.8). When we *generate concepts*, the customer’s requirements serve as a basis for developing a

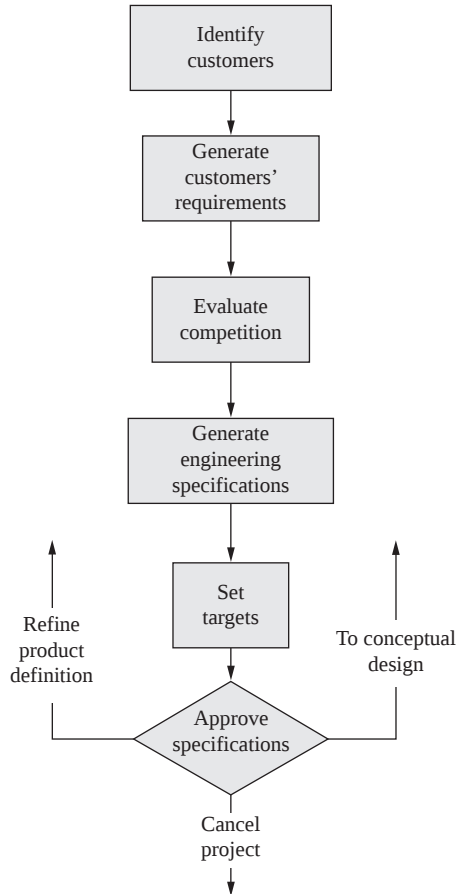


Figure 4.7 The Product Definition phase of the mechanical design.

Developing a concept into a product without prior effort on the earlier phases of the design process is like building a house with no foundation.

functional model of the product. The understanding gained through this functional approach is essential for developing concepts that will eventually lead to a quality product. Techniques for concept generation are given in Chap. 7.

After we *evaluate concepts*, the goal is to compare the concepts generated to the requirements developed during Product Definition and make decisions.

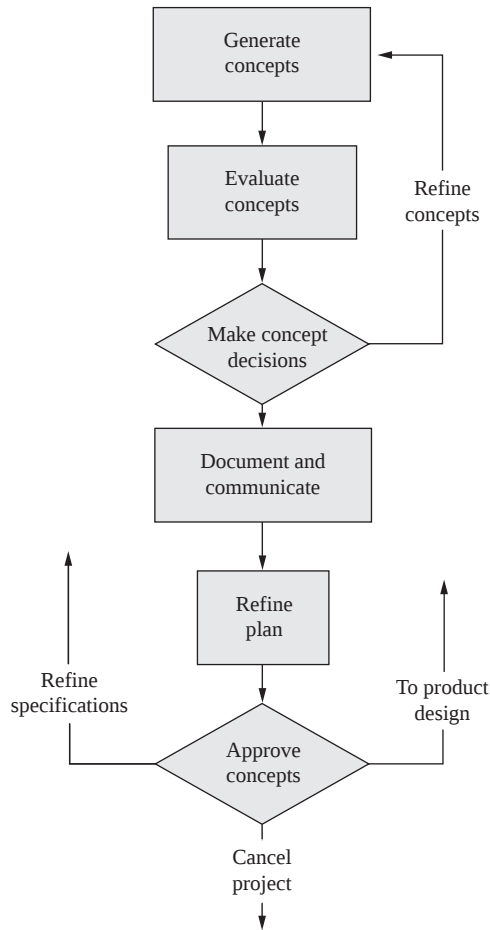


Figure 4.8 The Conceptual Design phase of the mechanical design process.

Concept decisions are made with limited knowledge. As shown in Fig. 1.11 knowledge increases with time and effort. One goal in Conceptual Design is choosing the best alternatives with the least expenditure of time and other resources needed to gain knowledge. Techniques helpful in concept evaluation and decision making are in Chap. 8.

During projects to redesign the X-ray tube, concepts are small changes to existing products, and the X-ray design team at GE uses detailed analytical models to evaluate them. However, for the Mars Rover, introduced in Chap. 2, new wheel concepts were dramatically different from those previously used and concept evaluation was much less analytical than at GE.

4.2.5 Product Development

After concepts have been generated and evaluated, it is time to refine the best of them into actual products (see Fig. 4.9). The Product Development phase is discussed in detail in Chaps. 9–11. Unfortunately, many design projects are begun here, without benefit of prior specification or concept development. This design approach often leads to poor-quality products and in many cases causes costly changes late in the design process. It cannot be overemphasized: *Starting a project by developing product, without concern for the earlier phases, is poor design practice.*

At the end of the Product Development phase, the product is released for production. At this time, the technical documentation defining manufacturing,

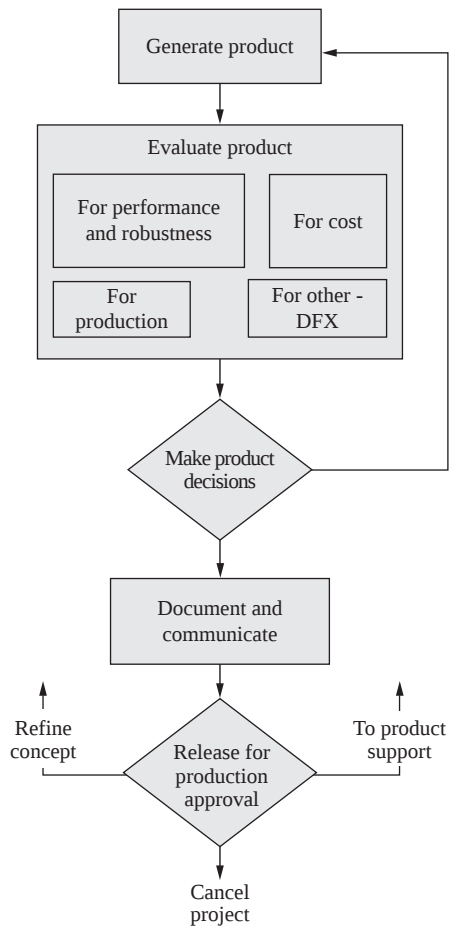


Figure 4.9 The Product Development phase of the mechanical design process.

assembly, and quality control instructions must be complete and ready for the purchase, manufacture, and assembly of components.

The GE design team used refined analytical models and component system testing during Product Development. Their final prototypes generally use actual production processes and production lines to fabricate final prototypes. This helps to ensure they capture the expected product quality and not be misled by “laboratory” produced prototypes.

4.2.6 Product Support

The design engineer’s responsibility may not end with release to production. Often there is continued need for manufacturing and assembly support, support for vendors, and help in introducing the product to the customer (see Fig. 4.10). Additionally, design engineers are usually involved in the engineering change process. This is the process where changes made to the product, for whatever reason, are managed and documented. This is one of the Product Support topics discussed in Chap. 12.

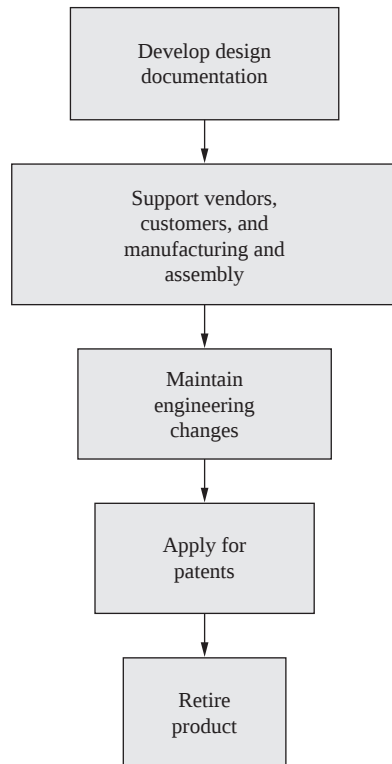


Figure 4.10 The Product Support phase of the mechanical design process.

Finally, the designers may be involved in the retirement of the product. This is especially true for products that are designed for specialized short-term use and then decommissioning. But, as pointed out in the Hannover Principles, this should be a concern regardless of the product throughout the design process.

Whereas the GE team must continue to support the X-ray tubes that are in use, a better example of postproduction support is the Mars Rover. The two Rovers were designed for 90 Mars days of operation. As of this writing, they have both lasted over 3.5 years. One of the Rovers is operating with only five of its original six legs providing power as one drive motor has ceased to function. The other has lost one of its four steering motors. In both cases, the engineers had to figure out how to change the Rovers from Earth to compensate for the failures. This is an extreme example of postproject Product Support.

Before refining the first phase, Product Discovery, later in the chapter, some justification is in order for why a product needs to be developed carefully through these six phases.

4.3 DESIGNING QUALITY INTO PRODUCTS

A good design process will support designing quality into the product. Traditionally, quality has been the concern of Quality Control (QC) or Quality Assurance (QA). QC/QA specialists inspect products as they are being manufactured and assembled. They check for conformance with the technical documentation (i.e., drawings, material properties, and other specifications) developed during design. They check dimensions, material properties, surface finishes, and other factors that are critical for form and function. This is often referred to as “inspecting quality into a product.”

It is less expensive and much more effective to *design quality into a product*. This implies not only designing a product that works as it should, lasts a long time, and meets the other customer desires listed in Table 1.1, but it also means designing the components and assemblies so they are easy to make, they have few or no tightly toleranced dimensions, and they have few critical (i.e., prone to failure) features. Finally, designing quality into a product also implies designing the product so that it is easy and foolproof to assemble.

Many engineering best practices help design quality into a product. Table 4.1 itemizes techniques generally considered as best practice and discussed in this text. They appear in the order in which they are generally applied to a typical design problem. However, each design problem is different, and some techniques may not be applicable to some problems. Additionally, even though the techniques are described in an order that reflects sequential and specific design phases, they

Quality cannot be manufactured or inspected into a product,
it must be designed into it.

Table 4.1 Best practices presented in this text

Project Planning (Chap. 5)	Product Development
Generating a product development plan	Product generation (Chap. 9)
Managing the project	Form generation from function
	Form representation
Specification Development (Chap. 6)	Materials and process selection
Understanding the design problem	Vendor development
Developing customer's requirements	Product evaluation (Chaps. 10 and 11)
Assessing the competition	Functional evaluation
Generating engineering specifications	Evaluating performance
Establishing engineering targets	Tolerance analysis
	Sensitivity analysis
Conceptual Design	Robust design
Generating concepts (Chap. 7)	Design for cost
Functional decomposition	Design for value
Generating concepts from functions	Design for manufacture
Evaluating concepts (Chap. 8)	Design for assembly
Judging feasibility	Design for reliability
Assessing technology readiness	Design for test and maintenance
Using the decision matrix	Design for the environment
Robust decision making	
	Product Support (Chap. 12)
	Developing design documentation
	Maintaining engineering changes
	Applying for a patent
	Design for end of product life

are often used in different order and in different phases. Understanding the techniques and how they add quality to the product aids in selecting the best technique for each situation.

The techniques described in this text comprise a design strategy that will help in the development of a quality product that meets the needs of the customer. Although these techniques will consume time early in the design process, they may eliminate expensive changes later. The importance of this design strategy is clearly shown in Fig. 4.11, a reprint of Fig. 1.5.

Figure 4.11 shows that Company A structures its design process so that changes are made early, while Company B is still refining the product after it has been released to production. At this point, changes are expensive, and early users are subjected to a low-quality product. The goal of the design process is not to eliminate changes but to manage the evolution of the design so that most changes come through iterations early in the process. The techniques listed in Table 4.1 also help in developing creative solutions to design problems. This may sound paradoxical, as lists imply rigidity and creativity implies freedom, however, creativity does not spring from randomness. Thomas Edison, certainly one of the most creative designers in history, expressed it well: "Genius," he said, "is 1% inspiration and 99% perspiration." The inspiration for creativity can only occur if the perspiration is properly directed and focused. The techniques presented here

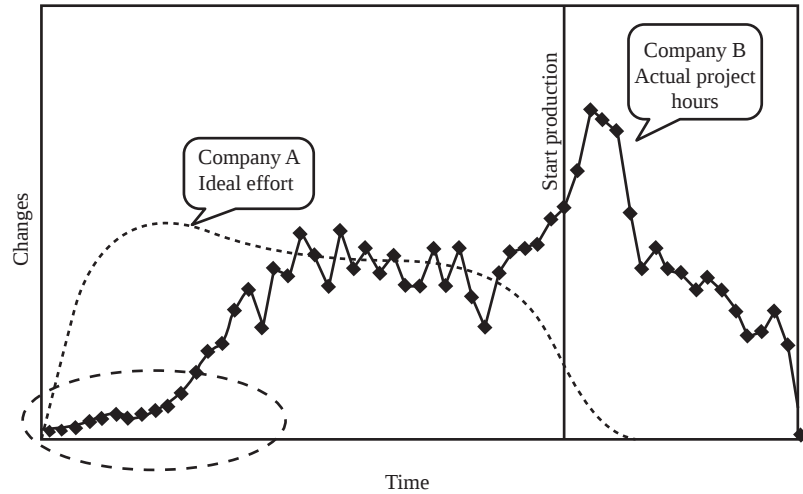


Figure 4.11 Engineering changes during automobile development.

help the perspiration occur early in the design process so that the inspiration does not occur when it is too late to have any influence on the product. Inspiration is still vital to good design. The techniques that make up the design process are only an attempt to organize the perspiration.

These techniques also force documentation of the progress of the design, requiring the development of notes, sketches, informational tables and matrices, prototypes, and analyses—records of the design’s evolution that will be useful later in the design process.

In the 1980s, it was realized that the process was as important as the product. One result of this realization is that Product Development is now often referred to as integrated product and process development or IPPD. Note that the term *process* is on equal footing with the *product*. Note also that IPPD implies that the product and process are under development. They are evolving.

Another result of this awareness is the increasing use of the International Standard Organization’s ISO 9000, the quality management system. ISO 9000 was first issued in 1987 and now has been adopted by most countries. There are millions of companies with ISO-9000 certification worldwide. All major manufacturing companies are ISO-9000 certified and, regardless of size, any company involved in international Product Development or manufacturing is also.

Prior to 2000 there were five standards numbered 9000 through 9005. In 2000, these were reduced to: ISO 9000, fundamentals and vocabulary; ISO 9001, requirements; and ISO 9004, guidance for performance improvement. ISO-9000 registration means that the company has a quality system that

1. Standardizes, organizes, and controls operations.
2. Provides for consistent dissemination of information.

3. Improves various aspects of the business-based use of statistical data and analysis.
4. Enhances customer responsiveness to products and service.
5. Encourages improvement.

Companies decide to seek ISO-9000 certification because they feel the need to control the quality of their products and services, to reduce the costs associated with poor quality, or to become more competitive. Also, they may choose this path simply because their customers expect them to be certified or because a regulatory body has made it mandatory.

In order to receive the certification, they must first develop a process that describes how they develop products, handle product problems, and interact with customers and vendors. Among the materials that must be prepared are written procedures that

- Describe how most work in the organization gets carried out (i.e., the design of new products, the manufacture of products, and the retirement of products).
- Control distribution and reissue of documents.
- Design and implement a corrective and preventive action system to prevent problems from recurring.

Once this material is developed the company invites an accredited external auditor (registrar) to evaluate the effectiveness of the process. If the auditors like what they see, they will certify that the quality system has met all of the ISO's requirements. They will then issue an official certificate. The company can then announce to the world that the quality of their products and services is managed, controlled, and assured by a registered ISO-9000 quality system. The certification typically expires after three years. Also, the registration agency typically requires surveillance audits at six-month intervals to maintain the currency of the certificate.

It must be made clear that ISO 9000 does not give a plan or process for developing products. It only requires a company to have a documented Product Development process on which the plan for a particular product can be based. The certification is not on the quality of the process itself, but that it exists, is maintained, and is used. Thus, a company can have a very poor methodology for developing products and still be certified. However, it is assumed that if a company is going to go to the trouble to get certified and wants to remain competitive in its markets, it will work to make this process and its Product Development plans as good as it can.

4.4 PRODUCT DISCOVERY

The goal of Product Discovery (Fig. 4.12), the first phase in the design process, is to develop a list of design projects that includes new products and product changes, and to choose which projects to work on. The term “discovery” may sound odd,

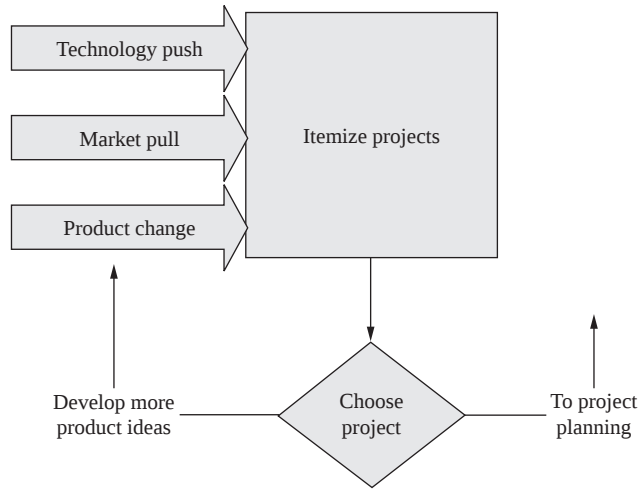


Figure 4.12 The Product Discovery phase of the mechanical design process.

but every design effort begins with the discovery of a need for product. There are three prime sources for new products: *market pull*, *technology push*, and *product change*.

Market pull occurs when there is customer demand for new products or product features. About 80% of new product development is market-driven. Without a customer for the product, there is no way to recover the costs of design and manufacture. Conversely, technology push is when a new technology is developed before there is customer demand. Let us refine these two product sources.

To manage market pull, the sales and marketing departments of most companies have a long list of new products or product improvements that they would like. When they see customers purchase a competitor's product, they wish their products had the unique features found on that product. Further, if they are doing a good job, they project the customer demand into the future. If sales and marketing had their way, there would be a continuous flow of product improvements and new products so that all potential customers could be satisfied. In fact, this is the direction that product development has been taking for the last few years—near-custom products with short development time.

At the same time, engineers and scientists have ideas for new products and product improvements based on technology. Rather than being driven by the customer, these ideas are driven by new technologies and what is learned during the design process. In fact, most product-producing companies spend from 2% to 10% of their revenue on research and development. And, since design is learning, by the time a designer finishes with a project, she or he knows enough to improve it. Most engineers would like to have a second chance at each project so they can, based on their new understanding, do it better the second time.

When a company wants to develop a product without market demand, utilizing a new technology, they are forced to commit capital investment and possibly years of scientific and engineering time. Even though the resulting ideas may be innovative and clever, they are useless unless they can be matched to a market need or a new market can be developed for them. Of course, devices such as sticky notes and many other products serve as examples of products that have been successfully introduced without an obvious market need. While these types of products have high financial risk, they can reap a large profit because of their uniqueness.

4.4.1 Product Maturity

Let's explore the need for new products further by examining the technology maturity "S" curve shown in Fig. 4.13. This shows the stages a technology matures through as it goes from a new product to a mature product. Products are often introduced to the market while some of the technologies it uses are still in the "make it work properly" stage, some even sooner. Product changes and improvements occur as technologies mature over time. Think of each of these improvements as redesign projects—they are. By the time a technology begins to reach maturity, the market is saturated with competition and companies need to decide if they are going to continue to develop using the existing technologies or innovate, develop new technologies, and begin the "S" curve again, as shown in Fig. 4.14.

If companies stay with the current technologies and further refine them, they probably have much competition and little room for improvement. If they innovate, they are taking a risk as the product matures.

4.4.2 Kano's Model of Customer Satisfaction

Another way to look at the need for product development is to examine Kano's Model of Customer Satisfaction. The Kano model was developed by Dr. Noriaki Kano in the early 1980s to describe customer satisfaction. This model will help us understand how and why features mature. Kano's model plots customer

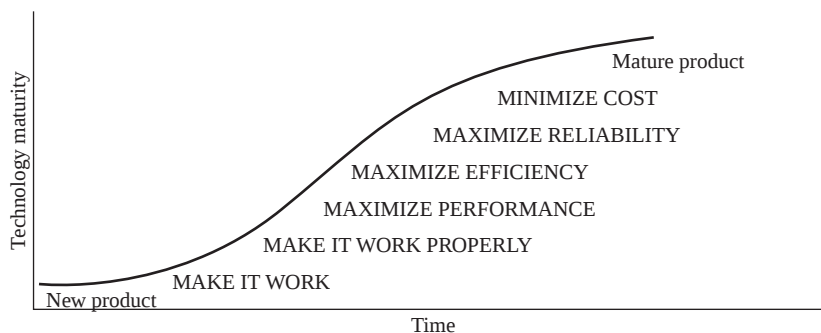


Figure 4.13 Product maturity "S" curve.

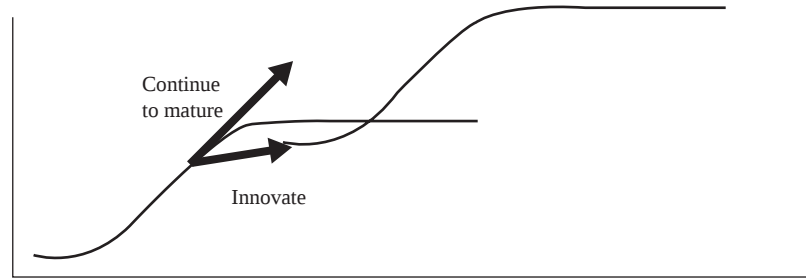


Figure 4.14 A decision point on the “S” curve.

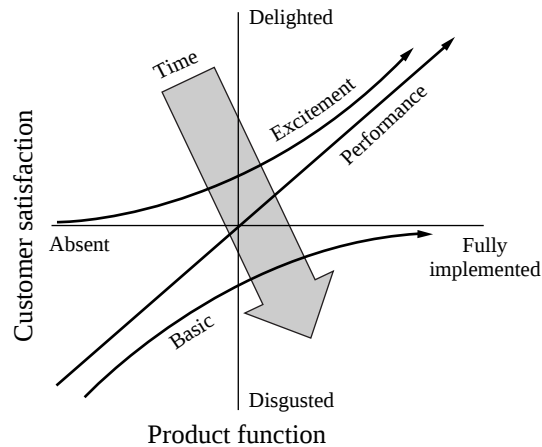


Figure 4.15 The Kano diagram for customer satisfaction.

satisfaction, from disgusted to delighted, versus product function, from absent to fully implemented, as shown in Fig. 4.15. This plot shows three lines representing basic features, performance features, and excitement features.

Basic features refer to customers’ requirements that are not verbalized as they specify assumed functions of the device. The only time a customer will mention them is if they are missing. If they are absent in the final product, the customer will be disgusted with it. If they are included, the customer will be neutral. An example is the requirement that a car should have brakes. If there are no brakes, then the customer is going to be disgusted with the product (and may be injured also). Brakes are expected on cars and so, just being there is not a cause for delight, just a neutral reaction from the customer. However, how well the brakes perform is a concern.

Requirements for *performance* features are verbalized in the form that the better the performance, the better the product. For example, a requirement on brake stopping distance is clearly a performance requirement. Generally, the shorter the distance, the more delighted the customer.

But what if your car would apply the brakes when you said so? What if you said, “Slow down” and the car gently decelerated, and when you yelled, “STOP!” it braked hard? This capability is unexpected. If it is absent, the customers are neutral because they don’t expect it anyway. However, if customers’ reactions to the final product are surprise and delight at the additional functions, then the product’s chance of success in the market is high. Requirements for excitement-quality features are often called “wow requirements.” If you went to a car showroom and test-drove a car with voice-activated brakes, this would be unexpected. Your reaction to the system would be “wow.” If the system worked well, you would be delighted, if it were not there at all, you wouldn’t know the difference and so would be neutral. Excitement-level features on a product generally require new technologies.

Over time, excitement-level features become performance-level features and, ultimately, basic features. This is true for most features of home entertainment systems, cars, and other consumer products. When first introduced, a new feature is special in one brand and consumers are surprised and delighted. The next year, as the technology matures, every brand has the feature and some perform better than others. Companies then work the “S” curve to improve performance, efficiency, reliability, and cost. After a few years, the feature is not even mentioned in advertising because it is an expected feature of the product.

The Kano model is just another view of technology maturity. Companies need to make decisions about whether to invest in innovation to “wow” customers or improve performance, efficiency, reliability, and cost and work their way farther up the “S” curve.

In addition to market pull and technology push, the third source of design projects is in response to the need for a change. There are three major sources for product changes:

- A vendor can no longer supply materials or components used in the product or has recommended improved ones. This may require the development of new plans, specifications, and concepts.
- Manufacturing, assembly, or another downstream phase in the product’s life cycle has identified a quality, time, or cost improvement that results in a cost-effective change in the product.
- The product fails in some way and the design needs to be changed. This type of change can be very costly. Reflect back to Fig. 4.11, where the automobile manufacturer was still making design changes after release for production. As discussed there, these changes are very expensive.

Change-driven projects are so important that an entire section will be devoted to them in Chap. 12.

4.4.3 Product Proposal

Regardless of the source, one deliverable from this phase of the design process is the product proposal. A template for developing such a proposal is available and is shown with a simplistic example in Fig. 4.16.

Note in this example that there is sufficient information to at least initiate discussions about how much resources should be allocated to following up on this proposal. In a real situation, much more documentation would be needed on each of these items.



Product Proposal	
Design Organization: xxxxxxxx	Date: June 23, 2010
Proposed Product Name: The Toastalator	
Summary: Customers who live in small spaces and have the need in the morning to both make coffee and cook toast. The concept here is for a device that combines these two products in a small space.	
Background of the Product: Observations of people living in small apartments have revealed an opportunity to minimize the space used when preparing breakfast. Since we manufacture both coffee makers and toasters this seems like a reasonable opportunity to pursue.	
Market for the Product: Although there is no firm evidence, there is anecdotal demand for this product. Studies of space availability and market size are needed. An initial survey shows the potential for up to 10 million customers.	
Competition: There is no known product such as this on the market today. And an initial patent survey has shown no recent activity with similar products.	
Manufacturing Capability: XXXXX currently manufactures similar products independently.	
Distribution Details: XXXX as distribution channels for similar products.	
Proposal Details:	
Task 1: Develop better market numbers.	
Task 2: Develop project plans through the Conceptual Design phase.	
Task 3: Develop product definition.	
Task 4: Develop and evaluate a proof-of-concept prototype.	
Team member:	Prepared by:
Team member:	Checked by:
Team member:	Approved by:
Team member:	
<i>The Mechanical Design Process</i> Copyright 2008, McGraw-Hill	
Designed by Professor David G. Ullman Form # 8.0	

Figure 4.16 The product proposal template.

4.5 CHOOSING A PROJECT

The hard part of this phase of the design process is deciding which projects to undertake and which to leave for later. We all think we make good decisions. It has been estimated, however, that over half of all decisions fail! A failed decision is later remade or the results ignored altogether. Failed decisions result in lost time and cost money. Any time you revisit a decision, all the work, tooling, prototypes, and CAD models made in the interim have little value.

During the Product Definition phase, there are usually more product ideas than there are time, people and money to do them all. The goal here is to choose which projects to undertake and which to leave for later or not attempt at all. This effort is commonly called *project portfolio management*, where a portfolio is a list of potential projects and the goal is to decide which of them to undertake.

To choose the best we need to know how to make decisions. We will introduce good decision-making practice and then specifically address portfolio decisions, the key decision needed during discovery. We will revisit decision making in Conceptual Design and then again during Product Development, adding to what we learn here.

In the remainder of this section, three methods will be presented that can help in choosing a project from the portfolio. The first two are simple, but somewhat limited. The third sets the foundation for decision-making processes that will be developed later in the book.

4.5.1 SWOT Analysis

The first decision support method we will use to help us choose a project is called a SWOT analysis. SWOT stands for Strengths, Weaknesses, Opportunities, and Threats. This method is commonly used in business, can be applied to the evaluation of single projects, and is easy to do. The basics of the method are to list the four SWOT items on a quadchart (each of four quadrants filled in with SWOT entries), as shown in Fig. 4.17 and then informally weigh the strengths versus the weaknesses and the opportunities versus the threats. As an example in the figure, a bicycle manufacturing company is considering adding a tandem bicycle to its product line.

Filling out a SWOT analysis makes it easier to judge whether or not a single potential project should be undertaken. Although this method does lay out the major points to consider when for decision making, it does not actually help in making the decision. It is still not clear whether or not BURL should undertake building a tandem bicycle.

Design is the technical and social evolution of information
punctuated by decision making.



SWOT Analysis	
Design Organization: BURL Bicycles	Date: Nov. 11, 2007
Topic of SWOT Analysis: Explore the potential for adding a tandem bicycle to the product line in 2008.	
Strengths: • BURL has the technology to design a top quality tandem bicycle. • BURL's engineers want to do this project. • It will expand the product line. • Market for tandems is growing, although no exact market numbers have been collected. • For the most part, they can be made with current equipment and processes. • We can use our patented suspension to differentiate BURL's tandem from the rest.	Weaknesses: • Market for tandems is small, <1% of all bicycle sales. • The profit margin may be smaller than on traditional bikes. • Cost to develop may exceed \$40,000. • Pay back time is estimated at 3 years. • It will take 6 months to get to market, missing the current sales season. • A tandem is just different enough to need unique marketing and shipping.
Opportunities: • A tandem will open BURL into new markets. • A tandem might allow bike shops that carry BURL to expand business and order more bikes.	Threats: • The product is not unique enough to attract customers. • We can't get bike shops to carry them. • It will cost more than \$40,000 to develop. • Engineering can't get it to ride like a CLIEN.
Team member: Fred Flemer	Prepared by: Fred Flemer
Team member: Bob Ksaskins	Checked by: Bob Ksaskins
Team member:	Approved by: Betty Booper
<i>The Mechanical Design Process</i> Copyright 2008, McGraw-Hill	
Designed by Professor David G. Ullman Form # 11.0	

Figure 4.17 SWOT diagram example.

4.5.2 Pro-Con Analysis

To take the SWOT analysis one step further, consider a pro-con analysis. An early, recorded use of this type of analysis is by Ben Franklin. Besides being a statesman, he was a designer of stoves, bifocals and many other inventions. In a 1772 letter to Joseph Priestly (the discoverer of oxygen), Franklin explained how he analyzed his problems when intuition failed him.

Dear Sir:

In the affair of so much importance to you, where in you ask my advice, I cannot, for want of sufficient premises, advise you what to determine, but if you please I will tell you how. When those difficult cases occur, they are difficult, chiefly because while we have them under consideration, all the reasons pro and con are not present to the mind at the same time; but sometimes one set present themselves, and at other times another, the first being out of sight. Hence the various purposes or information that alternatively prevail, and the uncertainty that perplexes us.

To get over this, my way is to divide a sheet of paper by a line into two columns; writing over the one Pro, and over the other Con. Then, during three or four days consideration, I put down under the different heads short hints of the different motives, that at different times occur to me, for or against the measure.

When I have thus got them all together in one view, I endeavor to estimate their respective weights; and when I find two, one on each side, that seem equal, I strike them both out. If I find a reason pro equal to some two reasons con, I strike out the three. If I judge some two reasons con, equal to three reasons pro, I strike out the five; and thus proceeding I find at length where the balance lies; and if, after a day or two of further consideration, nothing new that is of importance occurs on either side, I come to a determination accordingly.

And, though the weight of the reasons cannot be taken with the precision of algebraic quantities, yet when each is thus considered, separately and comparatively, and the whole lies before me, I think I can judge better, and am less liable to make a rash step, and in fact I have found great advantage from this kind of equation . . .

Franklin considers whether to accept or reject a single alternative. This is really a choice between two alternatives: do this or do something else (including nothing). Franklin advises five steps for making a decision:

- Step 1:** Make two columns on a sheet of paper and label one “Pros” and the other “Cons.”
- Step 2:** Fill in the columns with all the pros and cons of an alternative.
- Step 3:** Estimate the importance of each pro and each con.
- Step 4:** Eliminate pros and cons this way:
 - a. When two are of about equal importance, cross them both out and
 - b. Find other importance equalities of pros and cons—for example, the importance of two pros equals three cons—and then strike them out.
- Step 5:** When one or the other column becomes dominant, then “come to the determination accordingly.”

You can extend the idea of using pro-con lists to include more than one alternative, but the balancing step quickly becomes complex. Still, NASA frequently uses this approach to help organize experts when evaluating multiple project proposals. For each proposal, the experts list the pros and cons. They then informally balance the pros and cons to differentiate among the alternatives. This helps to tease out the good and bad points.



Pro-Con Analysis	
Design Organization: BURL Bicycles	Date:
Topic of Pro-Con Analysis: Should BURL market a tandem bicycle?	
Pro: • BURL has the technology to design a top-quality tandem bicycle. • BURL's engineers want to do this project. • It will expand the product line. • Market for tandems is growing, although no exact market numbers have been collected. • For the most part they can be made with current equipment and processes. • We can use our patented suspension to differentiate BURL's tandem from the rest. A tandem will open BURL into new markets. • A tandem might allow bike shops that carry BURL to expand business and order more bikes.	Con: • Market for tandems is small, <1% of all bicycle sales. • The profit margin may be smaller than for traditional bikes. • Cost to develop may exceed \$40,000. • Pay-back time is estimated at 3 years. • It will take 6 months to get to market, missing the current sales season. • A tandem is just different enough to need unique marketing and shipping. • The product is not unique enough to attract customers. • We can't get bike shops to carry them. • It will cost more than \$40,000 to develop. • Engineering can't get it to ride like a BURL.
Team member: Fred Flemer	Prepared by: Fred Flemer
Team member: Bob Ksaskins	Checked by: Bob Ksaskins
Team member:	Approved by: Betty Booper
<i>The Mechanical Design Process</i> Copyright 2008, McGraw-Hill	
Designed by Professor David G. Ullman Form # 9.0	

Figure 4.18 Pro-con analysis example.

If you look back at the SWOT analysis, the statements there are all an argument either for or against designing and marketing a tandem bicycle. In Fig. 4.18, these are reordered on Pro-Con Analysis Template. Thus, we have already completed steps 1 and 2 of Franklin's method.

Step 3 forces you to put a value on how important each of the pro and con statements is to the success of the project in preparation for step 4. For example, in looking down the list, it appears that

Market for tandems is growing although no exact market numbers have been collected.

is about as important as

A tandem is just different enough to need unique marketing and shipping.

So, according to step 4, they need to be crossed out. Then,

BURL has the technology to design a top-quality tandem bicycle.
and
BURL's engineers want to do this project.

is about as important as

Engineering can't get it to ride like a BURL.

So they too need to be crossed out. Continuing this way BURL ultimately sees that the cons outweigh the pros and decides not to undertake a tandem project.

4.5.3 Basics of Decision Making

Although the two methods just presented begin to get the information organized for good decision making, they are both limited to one alternative. In this section, we will formalize the entire decision-making process and make a protocol decision.

The basic structure of decision making is the same, whether addressing discovery issues or concept selection or choosing product details. In each case, there are six basic activities. Let's look at these activities in more detail:

1. **Clarify the issue** that needs a satisfactory solution.
2. **Generate alternatives**—itemize the potential solutions for the issue.
3. **Develop criteria** as they measure a satisfactory solution for the issue.
4. **Identify criteria importance** of each criterion relative to the others.
5. **Evaluate** the value of the alternatives by comparing them to the criteria.
6. Based on the evaluation results, **decide what to do next**. This decision will direct the process to
 - a. Add, eliminate, or refine alternatives.
 - b. Refine criteria.
 - c. Refine evaluation—work to gain consensus and reduce uncertainty.
 - d. Choose an alternative—you've made a decision, document it and address other issues.

These are shown in a flow diagram in Fig. 4.19.

We will reuse this list of activities and this diagram numerous times throughout the book.

Comparing the SWOT analysis to this ideal flow, SWOT is limited to activities 1, 2, and 5. It addresses only a single alternative and never actually itemizes the criteria for evaluation, even though they are inherent in the SWOT statements. (As we shall see in a moment.) SWOT focuses informally on the evaluation and never really gets to “what to do next.” Thus, it is not really a decision-making method by our definition, even though it supports some of the activities.

The pro-con method adds concern for the importance (activity 4) of the statements and gives a limited idea of “what to do” to the process (activity 6).

4.5.4 Making a Portfolio Decision

Here we will apply the activities listed in the previous section to the bicycle example.

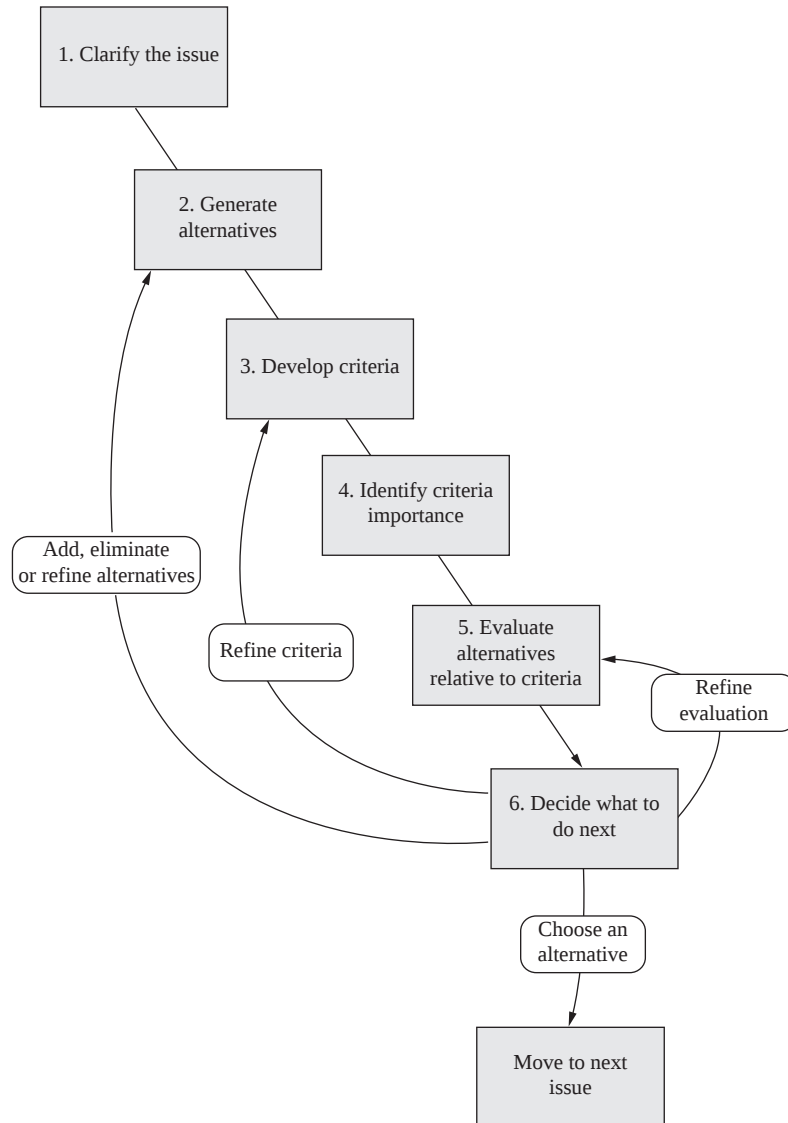


Figure 4.19 The decision-making flow.

Activity 1. BURL clarifies the issue. This was already done earlier, but we will make it broader here: “Choose, from a list of alternative product development projects, which one should be undertaken first?” In general, an issue is a question that needs to be addressed with some object or course of action chosen to answer the question and resolve the issue.

Activity 2. BURL itemizes the alternatives to be considered. This list can be as few as two items or in the hundreds, and spans all the way from minor product

changes to new models of existing products to innovative new products. For our bicycle company example, the options are

- Upgrade the current road bike.
- Introduce a tandem (as already considered).
- Add a front suspension to a soft-tail mountain bike product.

Activity 3. BURL develops criteria that are the basis for evaluating the alternatives. This is such an important activity that all of Chap. 6 is devoted to developing engineering specifications, the criteria for evaluating concepts and products. For many types of issues, those that are commonly repeated, a generic set of criteria can be used, at least as a starting place. For portfolio issues, the following list of criteria have evolved over time and can be used here:

- Acceptable program complexity: The complexity of the effort is within the experience of the organization or vendors. People are available with the skill sets needed to do the work.
- Clear market need: There is an established need in a market. (If evaluating innovative products, this may not be important.)
- Acceptable competitive intensity: The competitive intensity is reasonable and the alternative is not so new to the organization to impede commercialization.
- Acceptable five-year cash flow: The cash needed or generated over a five-year period is within reason.
- Reasonable payback time: The payback period for the needed investment and costs is acceptable.
- Acceptable start-up time: The time to realize cash flow is within the means of the organization.
- Good company fit: The newness or impact on the organization is acceptable—the new product or improvement fits the organization’s image.
- Strong proprietary position: The ability to withstand the competition’s efforts to erode the unique features that discriminate is good.
- Good platform for growth: The effort leads to future products or services.

In the SWOT analysis, we can see that the strengths, weaknesses, opportunities, and threats listed are evaluations of the criteria just listed. For example, the SWOT statements: “Market for tandems is growing although no exact market numbers have been collected” and “Market for tandems is small, < 1% of all bicycle sales” are qualitative evaluation statements for “Clear market need.”

One way to develop a list of criteria is to begin with a SWOT or pro-con analysis and then group the statements in categories. In fact, the list of protocol criteria was developed by examining many different protocol decisions, SWOT analyses and pro-con lists to find the common measures.

Activity 4. BURL decides what is most important. Not all of the nine criteria listed above are of equal importance. Complicating this is that the importance is in the eye of the beholder. For example, BURL’s financial people think that

Table 4.2 The portfolio scoring by BURL

Criteria	Alternatives					
	Upgrade road bike		Tandem		Front suspension for mountain bike	
	Agreement	Certainty	Agreement	Certainty	Agreement	Certainty
Acceptable program complexity	SA	C	N	C	D	VU
Clear market need	N	VC	D	U	SA	VC
Acceptable competitive intensity	A	C	A	N	N	N
Acceptable five-year cash flow	D	C	D	C	A	C
Reasonable payback time	N	C	D	U	A	U
Acceptable start-up time	A	VC	A	VC	N	C
Good company fit	A	C	D	C	N	C
Strong proprietary position	SD	C	SA	C	A	C
Good platform for growth	D	C	A	U	A	C

“acceptable five-year cash flow” and “reasonable payback time” are most important, and marketing wants to see “Good company fit.” Engineering wants a “Strong proprietary position” and a “Good platform for growth.”

For now, we will assume they are all equally important and address this activity further in Chap. 8.

Activity 5. BURL evaluates the alternatives relative to the criterion. These evaluations can range from qualitative assessments to the results of analytical simulations. For now, we will work with the qualitative statements made in the SWOT analysis and use a very simplified method to evaluate and decide what to do next. This will be refined as the product matures and more numerical analyses and simulation become possible.

To support this evaluation BURL used a *Decision Matrix*, a table with the alternatives in columns and the criteria in rows (Table 4.2). The cells of the matrix contain the evaluation results. For this qualitative assessment, BURL evaluated each alternative relative to each criterion using two measures. The first is how well the alternative meets the criterion in terms of level of agreement with the statement “I <X> that the <alternative> has <criterion>” where <X> equals

- Strongly agree (SA)
- Agree (A)
- Neutral (N)
- Disagree (D)
- Strongly disagree (SD)

For example “I <disagree> that the <Introduce tandem> has <clear market need>.”

Further, a second score will also be used, the level of certainty with which the evaluation is made. This is in terms of

- Very certain (VC)
- Certain (C)

- Neutral (N)
- Uncertain (U)
- Very uncertain (VU)

Certainty is a measure of how much you know or how much variation you expect. From an engineering standpoint the program complexity range is from disagree to strongly agree. However, for the development of the front suspension, the disagree assessment is very uncertain. Thus, in considering this alternative it will be hard to make use of the program complexity to judge whether or not to undertake a project to develop the front suspension. Further, note that only the front suspension option has an acceptable five-year cash flow. This implies that from a financial viewpoint none of these projects may be acceptable. We will do much more with tables of this type as the book evolves.

Activity 6. Based on the evaluation results, BURL must decide what to do next. It seems clear that none of these alternatives are outstanding. The financial picture of the first two alternatives looks weak. The complexity of the third alternative is questionable but knowledge about it is uncertain. So, one activity should be to develop other alternatives that overcome the drawbacks of the current portfolio. Additionally, it may be worthwhile to better understand the program complexity for the front suspension system.

Although the decision matrix has not given BURL a definitive decision, it has provided them a window on which to base a decision and has directed them about what to do next. This methodology will be refined as the book progresses.

4.6 SUMMARY

- There are six phases of the mechanical design process: Product Discovery, Project Planning, Product Definition, Conceptual Design, Product Development, and Product Support.
- The design process focuses effort on early phases, when the major decisions are made and quality is initiated. Additionally, a good process encourages communication, forces documentation, and encourages data gathering to support creativity.
- There are specific design process best practices that have been proven to improve product quality.
- New products originate from technology push, market pull, and product change.
- Products mature over time and new products emerge during maturation.
- A SWOT analysis can help choose which products to develop.
- Benjamin Franklin developed one of the earliest examples of using a pro-con analysis to make simple decisions.
- There are six basic decision-making activities: clarify the issue, generate alternatives, develop criteria, identify criteria importance, evaluate the value of the alternatives, and decide what to do next.
- The decision matrix can help in deciding what to do next.

4.7 SOURCES

“Letter to Joseph Priestley,” *Benjamin Franklin Sampler*, New York, Fawcett, 1956.

Ullman, David: *Making Robust Decisions*, Trafford, 2006. A complete book on design decision making.

4.8 EXERCISES

- 4.1 Develop a list of original design problems that you would like to do (at least 3). Choose one to work on that is within the time and knowledge available.
- 4.2 Make a list of features you don’t like about products you use. One way to develop this list is to note every time a device you use does not have a feature that is easy to use, doesn’t work like you think it should, or is missing as you go through your day. If you pay attention, a list like this will be easy to develop. Once the list has at least five items on it, choose one to improve through a redesign project.
- 4.3 Do a SWOT analysis on
 - The idea of taking Philosophy 101.
 - Buying an electric car.
 - Adding solar hot water heater to your parent’s house.
 - Adding a new feature to your backpack or briefcase.
- 4.4 Use Ben Franklin’s pro-con method to decide
 - Whether or not to go to coffee with the person next to you.
 - Whether or not to buy a new cell phone (pick the latest and greatest).
 - If the fix on your latest idea (e.g., bookcase, car repair, code, etc.) is worth pursuing.
- 4.5 Use a decision matrix to decide what to do next for
 - Purchasing one of three specific bicycles (or cars, electronic equipment) that you are interested in.
 - Choosing a ball bearing, a bronze bushing, or a nylon bearing for a pivot on the rear suspension of a bicycle.
 - Specifying a heating system for a house you are designing. The options are an air-to-air heat pump, air-to-water heat pump, or water-to-water heat pump.



4.9 ON THE WEB

Templates for the following documents are available on the book’s website: www.mhhe.com/Ullman4e

- Product Proposal
- Pro-Con Analysis
- SWOT Analysis

CHAPTER

5

Planning for Design

KEY QUESTIONS

- How does planning help in completing the five phases of the mechanical design process in a timely, cost-effective manner?
- Does one type of plan fit all design projects?
- What is the difference between a waterfall and a spiral plan?
- Why are deliverables so important?
- How can a plan be developed when the future is so uncertain?

5.1 INTRODUCTION

The goal of project planning is to formalize the process so that a product is developed in a timely and cost-effective manner. Planning is the process used to develop a scheme for scheduling and committing the resources of time, money, and people, as shown in Fig. 5.1. Planning results in a map showing how product design process activities are scheduled. The phases shown in Fig. 4.1—specification definition, conceptual design, and product development—must be scheduled and have resources committed to them. The flow shown in the figure is only schematic; it is not sufficient for allocating resources or for developing a schedule.

Planning generates a procedure for developing needed information and distributing it to the correct people at the correct time. Important information includes product requirements, concept sketches, system functional diagrams, solid models, drawings, material selections, and any other representation of decisions made during the development of the product.

The activity of planning results in a blueprint for the process. The terms *plan* and *process* are often used interchangeably in industry. Most companies have a generic process (i.e., a master plan) that they customize for specific products. This

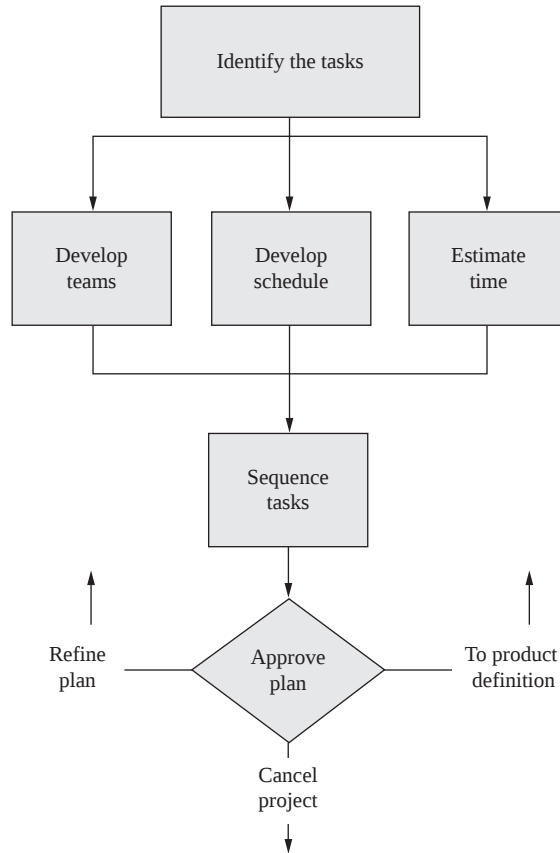


Figure 5.1 Project planning activities.

master plan is called the product development process, product delivery process, new product development plan, or product realization plan. In this book, we will refer to this generic process as the Product Development Process and use the acronym PDP.

Changing the design process in a company requires breaking down the way things have always been done. Although it can be quite difficult, many companies have accomplished it during recent decades. Generally, companies that have enjoyed good markets for their products and have begun to see these markets erode begin to look at their product development process as part of their effort to reengineer themselves to meet the competition. Most successful companies put emphasis on the continuous improvement of both the product and the process for developing the product.

If you do not know where you are going, you can not know when you get there. (Modernized from “Our plans miscarry because they have no aim. When a man does not know what harbor he is making for, no wind is the right wind” Lucius Annaeus Seneca [4 BC–AD 65].)

5.2 TYPES OF PROJECT PLANS

There are many different types of project plans. The simplest is the Stage-Gate or Waterfall plan. As shown in Fig. 5.2, work done in each stage is approved at a decision gate before progressing to the next stage. In its simplest form, the stage-gate methodology is very simple: Stage 1 = Product discovery, Stage 2 = Develop concepts, Stage 3 = Evaluate concepts, and so on. More likely, the stages are focused on specific systems or subsystems. Further, each stage may contain a set of concurrent activities executed in parallel, not in sequence.

The Stage-Gate Process can also be represented as a waterfall (Fig. 5.3) with each stage represented like a flat area where the water pools before falling to the next pool. The Stage-Gate method was formalized by NASA in the 1980s for managing massive aerospace projects.

The gates are often referred to as *design reviews*, formal meetings during which the members of the design team report their progress to management. Depending on the results of the design review, management then decides to either continue the development of the product, perform more work in the previous stage, or to terminate the project before any more resources are expended.

A major assumption in stage-gate or waterfall plans is that work can be done sequentially. This means that the product definition can be determined early in the process and that it will flow through concept to product. This is true for most mature types of products. A good example is the process used by Irwin in the design of new tools such as the Quick-Grip Clamp introduced in Chap. 2. Figure 5.4 shows the process used for the development of the clamp. At each stage, Irwin refines the definition into the objective and the deliverables. For example, the objective of “MS2-Design” is “Concept feasibility and robust business case.” In order to know that the objective has been achieved, there must be a set of deliverables. These include

- Concept development
- Technical feasibility
- Cost targets and financials
- Concept validation by consumers
- Legal assessment of intellectual property

The gate that follows Design is refined with the decisions made, who makes the decisions, and the criteria for the decisions. At Irwin, for example, the decisions made at the gate following MS 2 are select concept, approve business cases, accept

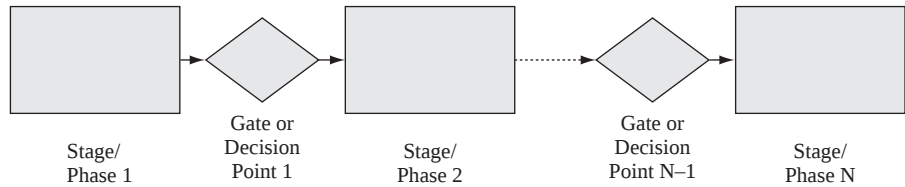


Figure 5.2 The Stage-Gate process.

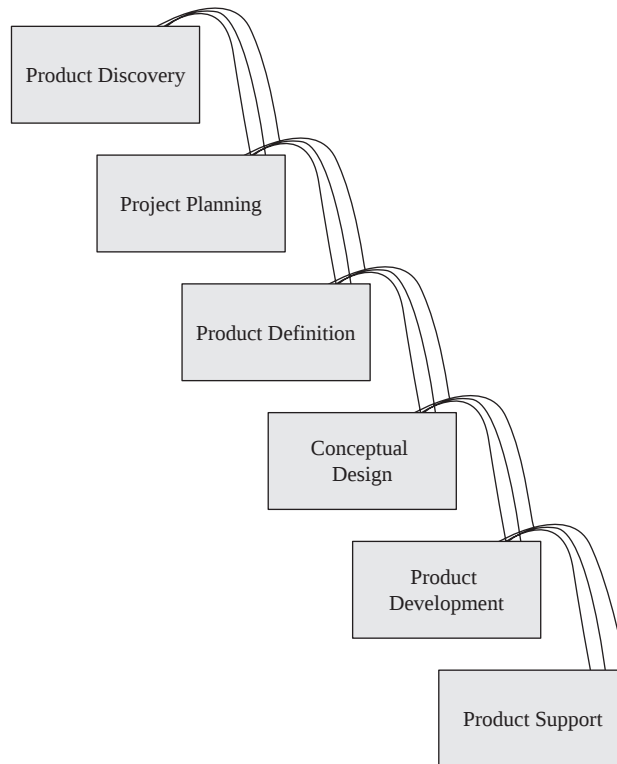


Figure 5.3 The Waterfall model.

prototype development results. These decisions are made by the leadership team including the President, the Vice President of Manufacturing, the Vice President for Research and Development, the Chief Financial Officer, and others. Decisions at this level may seem extreme for something as simple as a clamp, but this is a major product for a company such as Irwin, and thus concern goes all the way to the top of the organization.

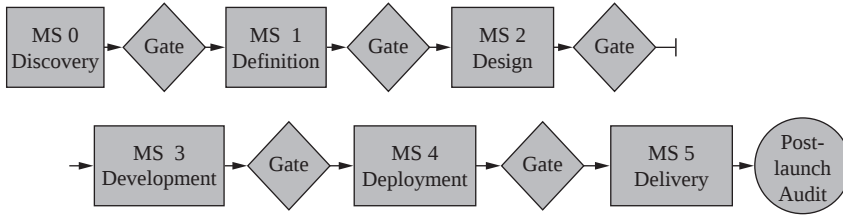


Figure 5.4 Irwin Tools product development process. (Reprinted with permission of Irwin Industrial Tools.)

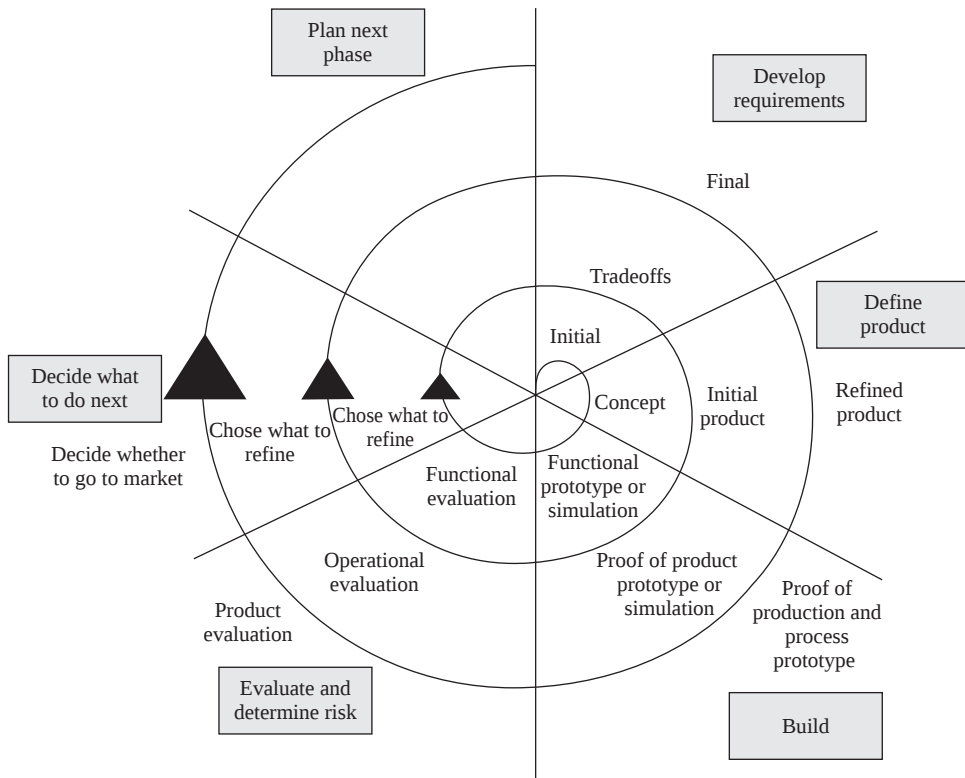


Figure 5.5 Spiral development of mechanical systems.

More recently, a spiral process has become very popular in software design. The spiral (Fig. 5.5) begins at the center with the basic concept and the rapid development of the first prototype; this is then evaluated by customers. The requirements for the product are revisited, and, in the second spiral, a new design prototype is tested. This methodology works in software because code is easier to prototype than most mechanical products. However, the continuing development

of rapid prototyping (see Section 5.3) is making this more realistic for hardware development. Primary characteristics of the spiral process are

- The iterative approach enables each task to be revisited during each cycle.
- Requirements can be reassessed.
- Prototypes and simulations can be elaborated and improved.
- The process enables “good enough for the moment” implementations.
- There is a clear decision point in each cycle.
- Each cycle provides objectives, constraints, alternatives, risks, review, and commitment to proceed.
- The level of effort is driven by risk considerations.

The spiral in Fig. 5.5 has been modified to show important activities in the mechanical design process. The spiral begins with initial requirements and progresses through concepts to functional prototypes and simulations, evaluation of these for how well they meet the initial requirements and for the risks incurred with future development, and helps to determine what to do next. Once this is understood, planning for the next cycle can occur. The second level of the spiral shows requirements traded off against each other as an initial product and its evaluation occur. Again, what to do next is determined and plans are made for the next phase, continuing the outward spiral toward the product. There may be more spirals than are shown here. Much of the terminology used in Fig. 5.5 will be defined later in this chapter.

Even more recent than the spiral process in software development is Extreme Programming. Extreme Programming is built around many small releases and integrated testing. One goal is a daily building of new code on the customer’s site for easy testing. This methodology harks back to the early days of mechanical engineering when something would be tried, broken, fixed, and tried again. In the early days of aircraft development, a test pilot would crash, the crew would fix the airplane, and assuming the pilot could still fly, he’d take it up again. As systems became more complex, the ability to make rapid changes in mechanical systems became more difficult. With rapid prototyping, this ability to make rapid changes is beginning to reappear. The down side of Extreme Programming is that there is no set target and you never know when you’re done. This problem is a major topic in Chap. 6.

In this book, we follow the waterfall process for a number of reasons. First spiral or extreme methods are better suited for software, where the development of prototypes usually takes far less time. Second, it is best to know where you’re going before you start or you don’t know when you get there. The flexibility of changing requirements needs to be weighed against not knowing when you’re finished. This is not to say that there is not iteration in the waterfall, just that it is built-in and planned. Third, the spiral process is best for new technologies when there is only a weak market pull and requirements are not clear. Finally, books are by nature serial, one chapter following another. There is no choice but to present the material in this manner. However, any particular project may be a

Design is an iterative process. The necessary number of iterations is one more than the number you currently have done.

This is true at any point in time.

—John R. Page, *Rules of Engineering*

combination of the linear, stage-gate or waterfall, and the more recent spiral and extreme processes.

5.3 PLANNING FOR DELIVERABLES— THE DEVELOPMENT OF INFORMATION

Progress in a design project is measured by deliverables such as drawings, prototypes, bills of materials (e.g., parts lists), results of analysis, test results, and other representations of the information generated in the project. These deliverables are all models of the final product. During product development, many models (i.e., design information representations) are made of the evolving product. Some of these models are analytical models—quick calculations on a bit of paper or complex computer simulations; some will be graphical representations—simple sketches or orthographic mechanical drawings; some will be CAD solid models and some will be physical models—prototypes.

Each of these models or prototypes is a representation of information that describes the product. In fact, *design is the evolution of information punctuated by decisions*. Each model or prototype is not only the embodiment of what is known about the product, but knowledge is gained in building or developing it. So the deliverables serve two purposes—they are the embodiment of the information that describes the product and they are a means to communicate that information to others. Thus, it is important to understand the information developed during the design process.

5.3.1 Physical Models—Prototypes

Physical models of products are often called *prototypes*. The characteristics of prototypes that must be taken into account when planning when to use them and what types to use are their *purpose*, the *phase* in the design process when they are used, and the *media* used to build them.

The four *purposes* for prototypes are proof-of-concept, proof-of-product, proof-of-process, and proof-of-production. These terms are traditionally applied only to physical models; however, solid models in CAD systems can often replace these prototypes with less cost and time.

- A **proof-of-concept or proof-of-function prototype** focuses on developing the function of the product for comparison with the customers' requirements or engineering specifications. This kind of prototype is intended as

a learning tool, and exact geometry, materials, and manufacturing process are usually not important. Thus, proof-of-concept prototypes can be built of paper, wood, parts from children's toys, parts from a junkyard, or whatever is handy.

- A **proof-of-product prototype** is developed to help refine the components and assemblies. Geometry, materials, and manufacturing process are as important as function for these prototypes. The recent development of *rapid prototyping* or *desktop prototyping*, using stereo lithography or other methods to form a part rapidly from a CAD representation, has greatly improved the time and cost efficiency of building proof-of-product prototypes.
- A **proof-of-process prototype** is used to verify both the geometry and the manufacturing process. For these prototypes, the exact materials and manufacturing processes are used to manufacture samples of the product for functional testing.
- A **proof-of-production prototype** is used to verify the entire production process. This prototype is the result of a *preproduction run*, the products manufactured just prior to production for sale.

In *Star Trek*, the science fiction series and movies, physical objects were produced in a “replicator.” Using just voice commands, this device could produce food, weapons, and just about anything else that could be imagined. Mechanical design is moving toward having replicators. Designers can conceive of a part, represent it in a solid-modeling CAD system, and “print” it out as a solid object using a *rapid prototyping system*. Rapid prototyping or solid printing produces solid parts useful for physical part evaluation, as patterns for molding or casting parts, or as visual models to gain customer feedback. In the 1980s and early in the 1990s, rapid prototype parts were usually made of wax, plastic, or cellulose. By 2000, some methods could make metal parts directly usable for small production runs and as molds for plastic parts capable of making tens of thousands of parts. Some rapid prototyping systems make parts using a laser to cut and glue thin layers of material together. Others use a laser to solidify liquid resins in places where solid material is desired. Still other systems deposit small amounts of materials much like building a part from small bits of clay. In the future, these systems may be able to make parts by building at the atomic level, enabling variations in material properties throughout a single component. These systems will approach science fiction by enabling a component or an entire product to be made in any place, on demand.

5.3.2 Graphical Models and CAD

Some companies rely solely on computer-generated solid models, others still rely on traditional drawings made either with a 2-D CAD package, or output from a solid model. Regardless of how they are produced, the graphical models are not

only the preferred form of data communication for the designer, they are also a necessary part of the design process. Specifically, drawings and solid models are used to

1. Archive the geometric form of the design.
2. Communicate ideas between designers and between designers and manufacturing personnel.
3. Support analysis. Missing dimensions and tolerances are determined as the drawing or model is developed.
4. Simulate the operation of the product.
5. Check completeness. As sketches or other drawings are being made, the details left to be designed become apparent to the designer. This, in effect, helps establish an agenda of design tasks left to accomplish.
6. Act as an extension of the designer's short-term memory. Designers unconsciously use drawings as part of their problem-solving process and often consciously use drawings to store information they might otherwise forget.
7. Act as a synthesis tool. Sketches and formal drawings enable the piecing together of unconnected ideas to form new concepts.

During the design process, many types of drawings are generated. Sketches used during conceptualization must evolve to final drawings that give enough detail to support production. This evolution usually begins with a layout drawing of the entire product to help define the geometry of the developing assemblies and components. The details of the components and assemblies are partially specified by the information developed on the layouts. As the product is refined, this information is transferred to detail and assembly drawings.

The development of modern solid-modeling CAD systems has blurred the differentiations between the types of drawings. These systems enable the co-evolution of details and assemblies in a layout environment. Further, they have automated many of the drawing standards. That being said, the traditional types of drawings will be introduced because they have specific characteristics important to even the most modern CAD systems.

The development of the drawings is synergistic with the evolution of the product geometry and further refinement of its function. As drawings are produced, more knowledge about the product is developed. Some of the major characteristics of the different types of drawings produced during product design and their role in the design process are itemized next.

Sketches. Sketching as a form of drawing is an extension of the short-term memory needed for idea generation (see Chap. 3). As the shape of components and assemblies evolve, drawings that are more formal are used to keep the information organized and easily communicated to others. Thus, a well-trained engineer has CAD skills and the ability to represent concepts that are more abstract and best represented as sketches.

Layout Drawings. A layout drawing is a working document that supports the development of the major components and their relationships. A typical layout drawing is shown in Fig. 5.6. Consider the characteristics of a layout drawing:

- A layout drawing is a working drawing and as such is frequently changed during the design process. Because these changes are seldom documented, information can be lost. Good records in the design notebook can compensate for this loss.
- A layout drawing is made to scale.
- Only the important dimensions are shown on a layout drawing. In Chap. 10, we see that starting with the spatial constraints sets the stage for developing the architecture and individual components in the product generation process. These constraints are best shown on a layout drawing.
- Tolerances are usually not shown, unless they are critical.
- Notes on the layout drawing are used to explain a design feature or the function of the product.
- A layout drawing often becomes obsolete. As detail drawings and assembly drawings are developed, the layout drawing becomes less useful. If the

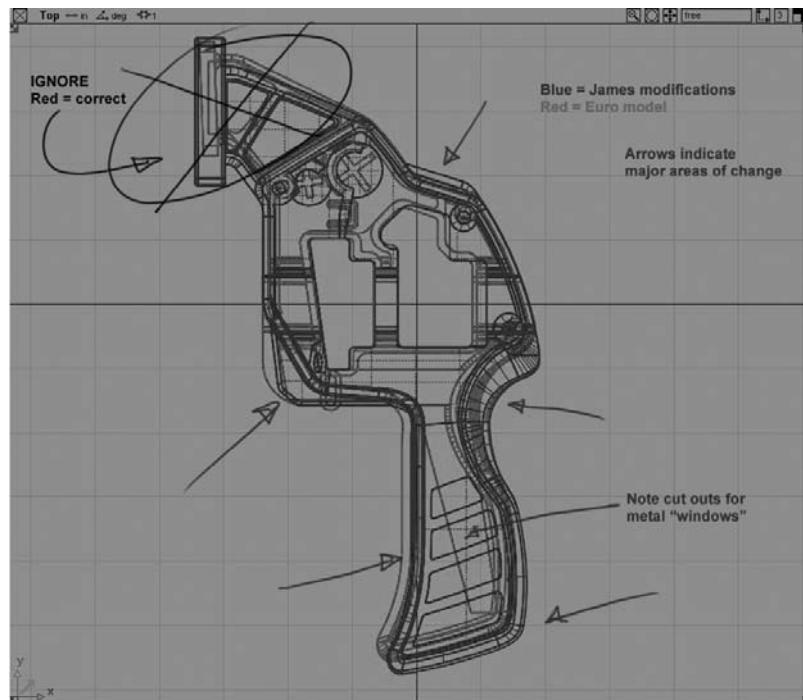


Figure 5.6 Typical layout drawing. (Reprinted with permission of Irwin Industrial Tools.)

product is being developed on a CAD system, however, the layout drawing's data file becomes the basis for the detail and assembly drawings.

The layout drawing shown in Fig. 5.6 was done on a solid-modeling system. This system enables the exploration of changes. The good news is that the solid model enables accurate visualization of the important geometry being studied, and the model provides much of what is needed for detail and assembly drawings. The bad news is that there is much time involved in this model, so changes in the configuration are expensive and discouraged.

Detail Drawings. As the product evolves on the layout drawing, the detail of individual components develops. These are documented on detail drawings. A typical detail drawing is shown in Fig. 5.7. Important characteristics of a detail include the following:

- All dimensions must be toleranced. In Fig. 5.7, many of the dimensions are made with unstated company-standard tolerances. Most companies have standard tolerances for all but the most critical dimensions. The upper and lower limits of the critical dimensions in Fig. 5.7 are given.
- Materials and manufacturing detail must be in clear and specific language. Special processing must be spelled out clearly.

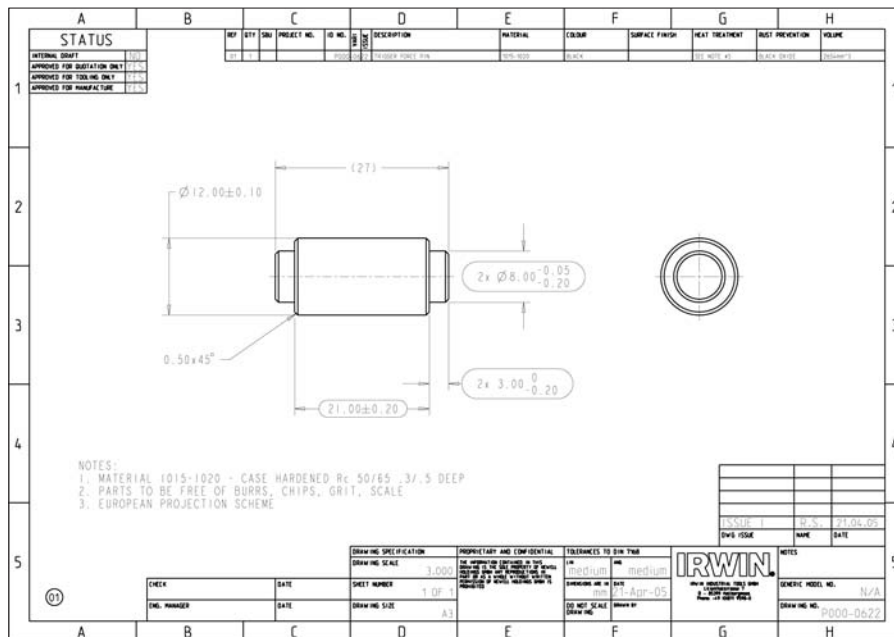


Figure 5.7 Typical detail drawing. (Reprinted with permission of Irwin Industrial Tools.)

- Drawing standards such as those given in ANSI Y14.5M-1994, *Dimensions and Tolerancing*, and in DOD-STD-100, *Engineering Drawing Practices*, or company standards should be followed.
- Since the detail drawings are a final representation of the design effort and will be used to communicate the product to manufacturing, each drawing must be approved by management. A signature block is therefore a standard part of a detail drawing.

Layout and assembly drawing focus on systems or subsystems, detail drawings address single components.

Assembly Drawings. The goal in an assembly drawing is to show how the components fit together. There are many types of drawing styles that can be used to show this. Assembly drawings are similar to layout drawings except that their purpose, and thus the information highlighted on them, is different. An assembly drawing has these specific characteristics:

- Each component is identified with a number or letter keyed to the Bill of Materials (BOM). Some companies put their Bill of Materials on the assembly drawings; others use a separate document. (The contents of the Bill of Materials are discussed in Section 9.2.)
- References can be made to other drawings and specific assembly instructions for additional needed information.

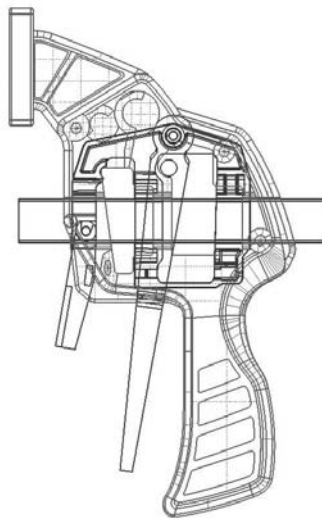


Figure 5.8 Typical assembly drawing. (Reprinted with permission of Irwin Industrial Tools.)

- Necessary detailed views are included to convey information not clear in the major views.
- As with detail drawings, assembly drawings require a signature block.

Graphical Models Produced in Modern CAD Systems. As mentioned in the introduction to this section, in modern solid-modeling CAD systems, layout, detail, and assembly drawings are not distinct. These systems enable the designer to make a solid model of the components and assemblies and, from these, semi-automatically make detail and assembly drawings. In these systems, the layout of components and assemblies and the details of the components and how they fit together into assemblies, all coevolve. This is both a blessing and a curse. On the positive side:

- Solid models enable rapid representation of concepts and the ability to see how they assemble and operate without the need for hardware.
- The use of solid-modeling systems improves the design process because features, dimensions, and tolerances are developed and recorded only once. This reduces the potential for error.
- Interfaces between components are developed so that components share the same features, dimensions, and tolerances, ensuring that mating components fit together.
- Detail and assembly drawings are produced semiautomatically, reducing the need to have expert knowledge of drafting methods and drawing standards.
- Files created are usable for making prototypes using rapid prototyping methods; developing figures for manufacturing and assembly; and providing diagrams for sales, service, and other phases of the product life cycle.

However, these tools also have a negative side:

- There is a tendency to abandon sketching. Sketches are a rapid way to develop a high number of ideas. The time required to develop a solid model is much longer than the time to make a sketch. This means the number of alternatives developed may be lower than it should be.
- Too much time is often spent on details too soon. Solid-modeling systems usually require details in order to even make a “rough drawing.” Thinking through these details in conceptual design may not be a good use of time, and once drawn there is a reluctance to abandon poor designs because of the time invested.
- Often valuable design time is spent just using the tool. Learning a solid-modeling system takes time and using it often requires time-consuming control of the program. This design time is lost.
- Many solid-modeling systems require the components and assemblies to be planned out ahead of time. These systems are more like an automated drafting system than a design aid.

In Table 5.1 (in Section 5.3.4), the different types of models used in mechanical design are itemized. Solid modeling and rapid prototyping are making it so that not only are layout, detailed, and assembly drawings merging, but so is the production of proof-of-concept, proof-of-product, proof-of-process, and proof-of-production prototypes. This merging is making it easier to produce more products in a shorter time.

5.3.3 Analytical Models

Often the level of approximation of an analytical model is referred to as its fidelity. *Fidelity* is a measure of how well a model or simulation analysis represents the state and behavior of a real-world object. For example, up until the late seventeenth century, all military calculations of cannonball trajectories were computed as if the projectile went up in a straight line, then followed a circular arc and another straight line straight down to the target (Fig. 5.9). These were low-fidelity simulations. However, in the late fifteenth century Leonardo da Vinci knew this model was wrong—that the trajectory was actually parabolic—and developed more accurate methods to compute the impact point. Even though he didn't have the mathematics to write the equations to describe his conclusions, his simulations were of better fidelity than preceding ones. It wasn't until Galileo that the parabolic model was developed and higher fidelity estimates could be made. These were later refined by Newton, and even later by the addition of the effects of aerodynamic drag and higher order dynamics.

Back-of-the-envelope calculations are low fidelity, whereas detailed simulations—hopefully—have high fidelity (it depends on the accuracy of the information input into them). Experts often run simulations to predict performance and cost. At the early stages of their projects, these simulations are usually

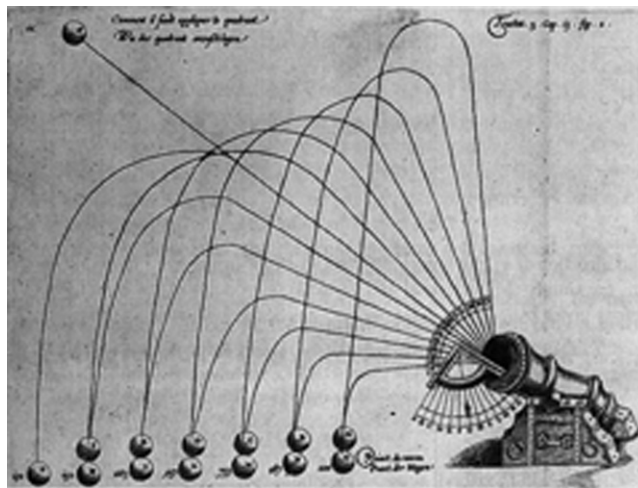


Figure 5.9 Pre-da Vinci trajectory estimations.

at low levels of fidelity, and some may be qualitative. Increasing fidelity requires increased refinement and increased project costs. Increased knowledge generally comes with increased fidelity, but not necessarily; it is possible to use a high-fidelity simulation to model “garbage” and thus do nothing to reduce uncertainty. In Chap. 10, we will talk more about analytical modeling.

5.3.4 Choosing the Best Models and Prototypes

Table 5.1 lists many types of models and prototypes that can be used in developing a product. These are listed by the medium used to build the model and the phase in the design process. There are two columns for drawings as many companies still use traditional layout and detail drawings, whereas others rely totally on CAD solid models.

There are trade-offs to be considered in developing models and prototypes: On one hand, they help verify the product, while on the other, they cost time and money. Further, there is a tension between the specifications for the product (what is supposed to happen) and the prototype (the current reality). In general, small companies are physical model-driven; they develop many prototypes and work from one to the next, refining the product. Large companies, ones that coordinate large volumes of information, tend to try to meet the specification through CAD and analytical modeling, building only a few physical prototypes.

An important decision made during planning is how many models and prototypes to schedule in the design process. There is currently a strong move toward replacing physical prototypes with computer models because simulation is cheaper and faster. This move will become stronger as virtual reality and rapid prototyping are further developed. Toyota has resisted these technologies in favor of developing physical prototypes, especially in the design of components that are primarily visual (e.g., car bodies). In fact, Toyota claims that using many simple prototypes, it can develop cars with fewer people and less time than companies that rely heavily on computers. GE, in its development of X-ray tubes for CT Scanners does much analysis, but moves to physical prototypes for

Table 5.1 Types of models

Phase	Medium			
	Physical (form and function)	Analytical (mainly function)	Graphical (Traditional) (mainly form)	Graphical (CAD) (form and function)
Concept	Proof-of-concept prototype	Back-of-the-envelope analysis	Sketches	Hand sketches and solid models
↓	Proof-of-product prototype	Engineering science analysis	Layout drawings	↓
	Final product	Proof-of-process and proof-of-production prototypes	Finite element analysis; detailed simulation	Detail and assembly drawings
				Solid models

proof-of-concept. The number of models and prototypes to schedule is dependent on the company culture and the ability to produce usable prototypes rapidly.

Finally, when planning for models and prototypes, be sure to set realistic goals for the time required and the information learned. One company had a series of four physical prototypes in its product development plan. But it turned out that the engineers were designing the second prototype (P2) while P1 was still being tested. Further, they developed P3 while P2 was being tested, and they developed P4 while P3 was being tested. Thus, what was learned from P1 influenced P3 and not P2, and what was learned from P2 influenced only P4. This waste of time and money was caused by a tight time schedule developed in the planning stage. The engineers were developing the prototypes on schedule, but since the tasks were not planned around the information to be developed, they were not learning from them as much as they should have. They were meeting the schedule for deliverable prototypes, not for the information that should have been gained.

5.4 BUILDING A PLAN

A project plan is a document that defines the tasks that need to be completed during the design process. For each task, the plan states the objectives; the personnel requirements; the time requirements; the schedule relative to other tasks, projects, and programs; and, sometimes, cost estimates. In essence, a project plan is a document used to keep that project under control. It helps the design team and management to know how the project is actually progressing relative to the progress anticipated when the plan was first established or last updated. There are five steps to establishing a plan. A template such as that in Fig. 5.10 can be used to support these steps. In this example, one task is detailed for a plan to develop a Baja car for an SAE (Society of Automotive Engineers) student contest. The plan is detailed in Fig. 5.16.

5.4.1 Step 1: Identify the Tasks

As the design team gains an understanding of the design problem, the tasks needed to bring the problem from its current state to a final product become clearer. Tasks are often initially thought of in terms of the activities that need to be performed (e.g., “generate concepts” or other terms used in Figs. 4.5–4.10). The tasks should be made as specific as possible, and as detailed in the next step, they should focus on what needs to be achieved rather than the activities. In some industries, the exact tasks to be accomplished are clearly known from the beginning of the project. For example, the tasks needed to design a new car are similar to those that were required to design the last model; the auto industry has the advantage

A task that only describes an activity, is done when you run out of time.



Project Planning	
Design Organization: Oregon State University Baja Team	
Date: Oct. 2, 2007	
Proposed Product Name: Killer Beaver	
Task 6	Name of Task: Preliminary Engine Compartment Design
	Objective: Develop solid model of the engine compartment Run initial FEM Analyze human factors for assembly and maintenance
	Deliverables: CAD solid model FEM results showing weak points based on static and fatigue analysis Simulation of assembly of engine and components Simulation of routine maintenance
	Decisions needed: Decision 1: Choose configuration for compartment Decision 2: Identify work needed to finalize the design
	Personnel needed: Title: student Hours: 75 Percent full time: 20% Title: Hours: Percent full time:
	Time estimate: Total hours: 75 Elapsed time (include units): 3 weeks
	Sequence: Predecessors: Task 4, Preliminary roll cage design Successors: Task 7, Final Engine Compartment Design Start Date: Oct. 12 Finish Date: Nov. 2
	Costs: Capital Equipment Disposables:
	Team member: James
Team member: Tim	Checked by: Pat
Team member: Pat	Approved by:
Team member:	
<i>The Mechanical Design Process</i>	
Copyright 2008, McGraw-Hill	
Designed by Professor David G. Ullman	
Form # 10.0	

Figure 5.10 Example plan template.

of beginning with a clear picture of the tasks needed to complete a new design. However, for a totally new product, the tasks may not be so clear.

5.4.2 Step 2: State the Objective for Each Task

Each task must be characterized by a clearly stated objective. This objective takes some existing information about the product—the input—and, through some activity, refines it for output to other tasks. Even though tasks are often initially conceived as activities to be performed, they need to be refined so that the results of the activities are the stated objectives. Although the output information can be only as detailed and refined as the present understanding of the design problem, each task objective must be

- Defined as information to be refined or developed and communicated to others, not as activities to be performed. This information is contained in *deliverables*, such as completed drawings, prototypes built, results of calculations, information gathered, or tests performed. If the deliverables cannot be itemized, the objective is not clear—then you know you are done only when you run out of time.
- Presented in terms of the decisions that need to be made and who will be involved in making them.
- Easily understood by all on the design team.
- Specific in terms of exactly what information is to be developed. If concepts are needed, then tell how many are sufficient.
- Feasible, given the personnel, equipment, and time available. See step 3.

5.4.3 Step 3: Estimate the Personnel, Time, and Other Resources Needed to Meet the Objectives

For each task, it is necessary to identify who on the design team will be responsible for meeting the objectives, what percentage of their time will be required, and over what period they will be needed. In large companies, it may only be necessary to specify the job title of the workers on a project, as there will be a pool of workers, any of whom could perform the given task. In smaller companies or groups within companies, specific individuals might be identified.

Many of the tasks require virtually a full-time commitment; others require only a few hours per week over an extended period. For each person on each task, it is necessary to estimate not only the total time requirement but also the distribution of this time. Finally, the total time to complete the task must be estimated. Some guidance on how much effort and how long a design task might take is given in Table 5.2. (The values given are only for guidance and can vary greatly.)

Similar comments apply to other resources needed to complete the task, especially those used for simulation, testing, and prototype manufacture. These resources and personnel are the means to complete the task.

Notice in Table 5.2 that no entry estimates the required time to be less than a week. *Design takes time*. Often it takes twice as long as the original estimate,

Table 5.2 The time it takes to design

Task	Personnel/time
Design of elemental components and assemblies. All design work is routine or requires only simple modifications of an existing product.	One designer for one week
Design of elemental devices such as mechanical toys, locks, and scales, or complex single components. Most design work is routine or calls for limited original design.	One designer for one month
Design of complete machines and machine tools. Work involved is mainly routine, with some original design.	Two designers for four months
Design of high-performance products that may utilize new (proven) technologies. Work involves some original design and may require extensive analysis and testing.	Five designers for eight months

especially if the design project is not routine or new technologies are used. Some pessimists claim that after making the best estimate of time required, the number should be doubled and the units increased one step. For example, an estimate of one day should really be two weeks.

A more accurate method for estimating the total time required for a project is based on the complexity of the product's function. The theory is that the more complex the function, the more complex the product and the longer the time needed to design the product. Product function development is a key part of concept generation and is covered in detail in Chap. 7. Thus, in order to use this method for time estimation, there has to be some understanding of the functions of the product. During the product development process, often a task in the conceptual design phase is titled "refine plans" to reflect the dependence of the plan on the concept being developed.

The total time required for a project can be estimated by

$$\text{Time (in hours)} = A * PC * D^{0.85}$$

where

A = a constant based on past projects in the company. This constant is dependent on the size of the company and how well information is communicated among the various functions. Typically, $A = 30$ for a small company with good communication and $A = 150$ for a large company with average communication. Note that communication and thus time is estimated at five times greater in a large organization.

PC = product complexity based on function (discussed shortly).

D = project difficulty: $D = 1$, not too difficult (i.e., using well-known technologies); $D = 2$, difficult (i.e., some new technologies); $D = 3$, extremely difficult (i.e., many new technologies).

Everything takes twice as long.

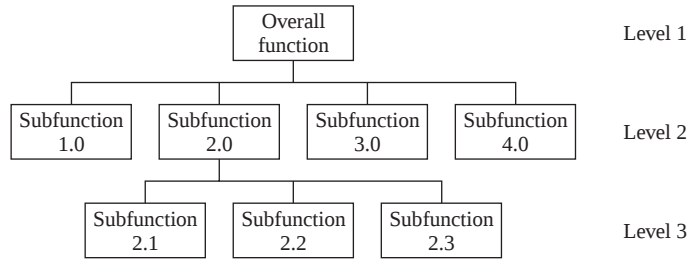


Figure 5.11 Example of a function diagram.

Product complexity is based on the functions of the product. A function diagram will typically look as shown in Fig. 5.11. Details on how to develop such a diagram will be covered in Chap. 7.

The product complexity is estimated by

$$PC = \sum j * F_j$$

where

j = the level in the function diagram

F_j = the number of functions at that level

For the example in Fig. 5.11, there is 1 function on the top layer (always there), 4 on the second level, and 3 on the third:

$$PC = 1 * 1 + 2 * 4 + 3 * 3 = 18$$

For example, a small company with good communication ($A = 30$) is designing a difficult product ($D = 2$) that has $PC = 18$, then an estimate of the total time is 973 hours, or two designers working for 3 months. This method has been shown to be fairly accurate within a single company that has calibrated the value for A , and models function in a consistent manner.

Time estimation is very difficult and subject to error. Thus, it is recommended that task time be based on three estimates: an optimistic estimate o , a most-likely estimate m , and a pessimistic estimate p . From these three, the statistical best estimate of task time is

$$\text{Time estimate} = \frac{o + 4m + p}{6}$$

This formula is used as part of the PERT (Program Evaluation and Review Technique) method. See the sources in Section 5.8 for more details on PERT.

Finally, note that the distribution of time across the phases of the design process is generally in the following ranges:

Project planning: 3 to 5%

Specification definition: 10 to 15%

Conceptual design: 15 to 35%

Product development: 50 to 70%

Product support: 5 to 10%

These percentages are based on studies of actual projects. The exact proportion in each phase greatly depends on the type of product, the amount of original design work, and the structure of the design process within the company.

5.4.4 Step 4: Develop a Sequence for the Tasks

The next step in working out the plan is to develop a task sequence or schedule. Scheduling tasks can be complex. The goal is to have each task accomplished before its result is needed and, at the same time, to make use of all of the personnel, all of the time. Additionally, it is necessary to schedule design reviews or other forms of approval to continue the project. The tasks and their sequence is often referred to as a *work breakdown structure*.

For each task, it is essential to identify its *predecessors*, which are the tasks that must be done before it, and the *successors*, the tasks that can only be done after it. By clearly identifying this information, the sequence of the tasks can be determined. A method called the *CPM (Critical Path Method)* helps determine the most efficient sequence of tasks. The CPM is not covered in this book.

Often tasks are interdependent—two tasks need decisions from each other in order to be completed. Thus, it is important to explore how tasks can be started with incomplete information from predecessors and how they can supply incomplete information to successors.

The best way to develop a schedule is to use a bar chart, shown in Fig. 5.12. (This type of chart is often called a *milestone* or *Gantt chart*.) On the chart, (1) each task is plotted against a time scale (time units are usually weeks, months, or quarters of a year); (2) the total personnel requirement for each time unit is plotted; and (3) the schedule of design reviews is shown. The Gantt chart in Fig. 5.12 was developed on a spread sheet (there are templates available for this). Many Gantt charts are developed using Microsoft Project™, as shown in Fig. 5.16.

In developing the task sequence pay attention to task dependencies. Step 1 emphasized concern for the information needed by the task and the information generated by the task. If a series of tasks simply build on each other, the information developed by one is the information needed by the next and the tasks are *sequential*. If two or more tasks must be accomplished at the same time to

A plan is a “work breakdown structure” because without one the Work remaining will grow until you have a Breakdown unless you enforce some Structure on it.

—Taken from John R. Page, *Rules of Engineering*

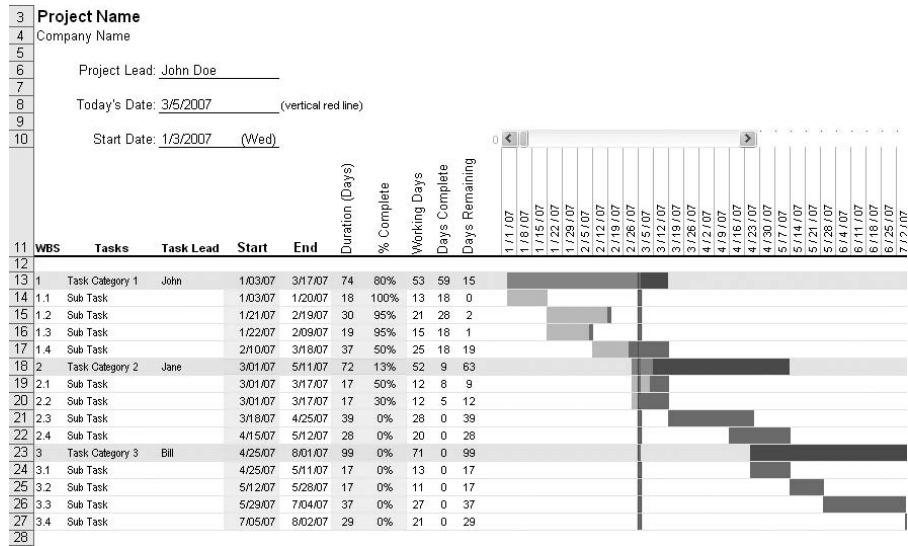


Figure 5.12 Gantt chart built on a spreadsheet.

produce information for a future task, then they are *parallel*. There are two types of parallel tasks: *uncoupled* and *coupled*.

For example, in designing the MER, a decision was made early on to use the same type of motor and reduction gears both to power the MER and to steer it. Thus, tasks to develop the steering and drive train were closely coupled. Many other tasks occurred at the same time as the development of the drive train and the steering that were not coupled, for example, the Inertial Measurement Unit (IMU), Warm Electronics Box (WEB), and many other systems (see Fig. 2.7).

The three types of task sequences (serial, parallel coupled, and parallel uncoupled) can be discovered by using a *Design Structure Matrix (DSM)*. A DSM is a simple diagram that helps sequence tasks, as shown in Fig. 5.13. Consider in the DSM shown here a subset of the tasks that may be required to develop a new bicycle seat. Each task is assigned a row and given a letter name. These letter names also appear as the names of the columns, in the same order. To develop a DSM, consider the tasks, one at a time. In the task's row put an X for every other task on which it is dependent. In the diagonal put the letter name to make reading easier.

Task A is not dependent on another task and so there are no Xs in the first row. The generation of concepts, Task B, needs the specifications developed in Task A and sequentially follows it. Similarly, Task C follows Task A but is not dependent on Task B. Thus, it can be done in parallel with Task B, but is uncoupled from it. Task D is dependent on Tasks B and C. Tasks E, F, and G are coupled as is evidenced by the Xs in the lower right corner of the matrix. Task E is dependent on Tasks F and G; Task F is dependent on Tasks E and G; and Task G

	A	B	C	D	E	F	G
Generate specifications	A	A					
Generate two concepts	B	X	B				
Develop test plan	C	X		C			
Test the concepts	D		X	X	D		
Design production parts	E	X	X		X	E	X
Design plastic injection mold	F	X				X	F
Design assembly tooling	G				X	X	G

Figure 5.13 Design Structure Matrix.

is dependent on Tasks E and F. Further, Tasks E and F are dependent on other tasks as well.

Reading down a column it is easy to see which tasks are dependent on the information developed. For example, reading down column “B” it is easy to see that Tasks D and E are dependent on the concepts being developed in Task B.

The DSM is very useful when the order of the tasks is not evident. The initial task order can be rearranged so the sequence flows in a manageable fashion.

5.4.5 Step 5: Estimate the Product Development Costs

The planning document generated here can also serve as a basis for estimating the cost of designing the new product. Even though design costs are only about 5% of the manufacturing costs of the product (Fig. 1.2), they are not trivial.

The cost estimate needed here is for the project, not the product. Product cost estimates are covered in Chap. 11. A majority of project costs are in salaries. Some basic guidelines for making a project cost estimate are

- Engineer salaries range from \$50k to \$100k per year, or assuming 2000 work hours year, \$25–\$50/hour. However, the cost to the project is more than just salaries, as all companies add on a “burden” that covers the costs of buildings, utilities, support personnel, and general equipment. Burden rates range from 100% in industry up to 300% in a government lab. Thus, the least expensive engineering in an industrial organization will cost \$50 an hour, and a senior engineer in a government lab will cost \$200 an hour.
- Most mechanical design projects require physical prototypes and test facilities. Each organization has a method to account for these costs. They may be lumped into the burden rate, or may be a separate item paid for by the hour. The same consideration must be given to computer costs to support CAD, simulation, meeting support, PLM, and other needs.
- For many projects, there is the need to travel to meet with other members of the design team, vendors, and suppliers. Travel costs can add up fast and must be included in planning.

5.5 DESIGN PLAN EXAMPLES

5.5.1 A Very Simple Plan

We will now look at two simple problems to see how different problems require different design processes. Recall the problem statements from Chap. 1 (see Fig. 1.9):

What size SAE grade 5 bolt should be used to fasten together two pieces of 1045 sheet steel, each 4 mm thick and 6 cm wide, which are lapped over each other and loaded with 100 N?

and

Design a joint to fasten together two pieces of 1045 sheet steel, each 4 mm thick and 6 cm wide, which are lapped over each other and loaded with 100 N.

The solution of the first joint design problem is fairly straightforward (Fig. 5.14). It is fully defined, and understanding the problem is not hard. Since the problem statement actually defines the product, there is no need to generate and evaluate concepts or to generate a product design since it already exists. The only real effort involved in this design problem is to evaluate the product. This is done using standard equations from a text on machine component design or using company or industrial standards. In a component-design text, we find analysis methods for several different failure modes: the bolt can shear, the sheet steel can crush, and so on. After completing the analysis, you will make a decision as to which of the failure modes is most critical and then specify the smallest size of bolt that will not permit failure. This decision, part of the evaluation, is documented as the answer to the problem. In a classroom situation, you will undergo a “design review” when your answer is graded against a “correct” answer.

Very few real design problems have a single correct answer. In fact, reality can cause quite a shift from the design process illustrated in Fig. 5.14. Consider one example: An experienced design engineer began a new job with a company that manufactured machines in an industry new to him. One of his first projects included the subproblem of designing a joint similar to bolt analysis problem. He followed the process in Fig. 5.14 and documented his results on an assembly drawing of the entire product. His analysis told him that a 1/4-in.-diameter bolt would carry the load with a generous factor of safety. However, his manager, an experienced designer in the industry, on reviewing the drawing, crossed off the 1/4-in. bolt and replaced it with a 1/2-in bolt, explaining to the new designer that it was an unwritten company standard based on years of experience never to use bolts of less than 1/2-in. diameter. The standard was dictated by the fact that

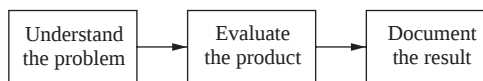


Figure 5.14 Simple plan for a lap joint.

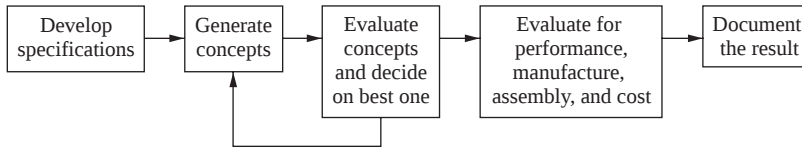


Figure 5.15 Design process for a more complex lap joint.

service personnel could not see anything smaller than a 1/2-in. bolt head in the dirty environment in which the company's equipment operated. On all subsequent products, the designer specified 1/2-in. bolts without performing any analysis.

For the second joint design problem, the process is more complex (Fig. 5.15). There are a number of concepts that might fasten the sheets. Typical options include using a bolt, welding the pieces together, using an adhesive, or folding the metal to make a seam. You might perform an analysis on each of these options, but that would be a waste of time because the results would still provide no clear way of knowing which joint design might be best. What is immediately evident is that the requirements on this joint are not well articulated. In fact, if they were, perhaps none of the earlier concepts would be acceptable.

So the first step in solving this problem should be specification development for the joint. Various questions should be addressed: Does the joint need to be easily disassembled or leak-resistant? Does it need to be less than a certain thickness? Can it be heated? After all the specifications are understood, it will be possible to generate concepts (maybe ones previously thought of, maybe not), evaluate these concepts, and limit the potential designs for the joint to one or two concepts. Thus, before performing analysis on all of the joint designs (evaluating the product), it may be possible to limit the number of potential concepts to one or two. With this logic, the design process would follow the flow of Fig. 5.15, a process similar to that in Fig. 4.1, except there may be no need to generate product. The problem solved here is so mature that the concepts developed are fully embodied products. The concept, a "welded lap joint," is fairly refined. The only missing details are the materials, the weld depth, the length of the weld leg, and other details requiring expertise in welding design. However, if the requirements on the joint were out of the ordinary, then the concepts generated might be more abstract and have many possible product embodiments.

5.5.2 Development of a New Product for a Single or Small Run

Many products are made only once or at most a few times. Planning for manufacture and assembly is different for these products than for those that are mass-produced. Specifically, for a small run there is less latitude for choosing manufacturing methods. Methods such as forging metal, injection-molding plastics, and manufacturing custom control circuits require mass production to amortize the tooling costs required. This restricts the types of components that can

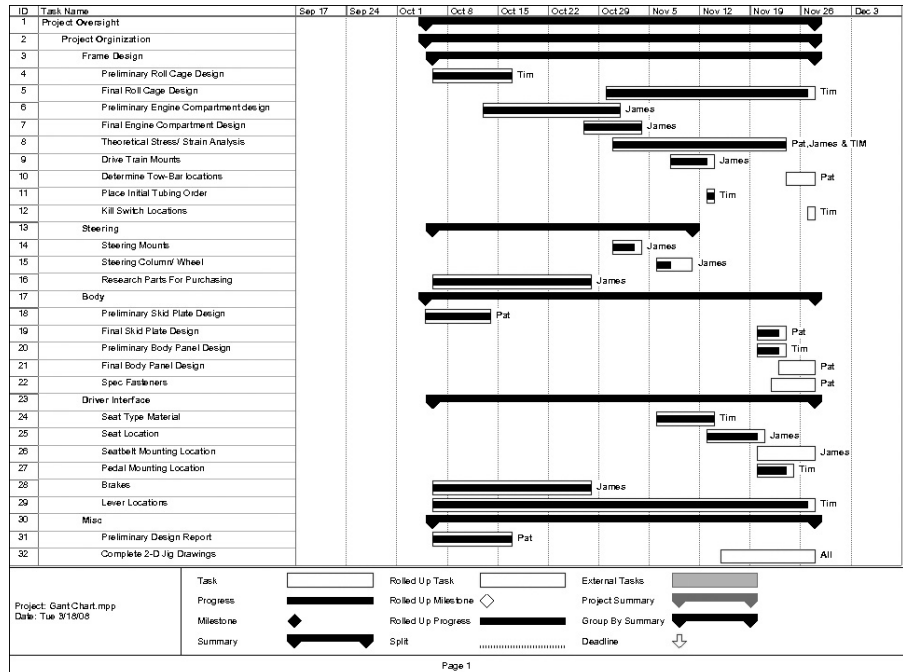


Figure 5.16 Project plan for Baja car.

be designed. Often the first item built is both a prototype and the final product delivered to the customer. There is more of a tendency to buy off-the-shelf components for short-run products. There is also less concern for assembly time than for mass-produced products.

Figure 5.16 is the project plan of Oregon State University's 2007 SAE Baja car. This team consistently places in the top 10 in races they enter. This plan was done on Microsoft ProjectTM and it shows the major tasks during a six-week period in the fall term. Note that some of the tasks did not take as long as planned, and others were not even done at all. Keep in mind that a plan is just that, a plan for doing work and developing deliverables. Reality seldom fits the plan precisely, even for this team that based their estimates on those developed in prior years.

5.5.3 Development of a New Product for Mass Production

Planning for new products can range from very simple to nearly impossible. Consider these two examples: A toy manufacturer is to develop a new toy that is similar to other toys they currently make (e.g., new action figures and toy

A plan is only valid until you start working.

cars with cosmetic or minor functional changes). Thus, the product development plan is similar to that for the previously designed toys. At the other end of the planning spectrum is a company that has just developed a new technology and has never made a similar product before. For example, when first producing the iPod, Apple's planning required many tasks that were highly uncertain. The ability to plan for a new product in this situation is much more challenging than it is for the toy company.

Designing products for mass production requires careful planning for manufacture and assembly. These projects give the design engineer more flexibility in selecting materials and manufacturing processes and increase the project's dependence on manufacturing engineers.

5.6 COMMUNICATION DURING THE DESIGN PROCESS

Communication of the right information to the right people at the right time is one of the key features of a successful design project and a key reason for the existence of PLM. All communication begins with informal, face-to-face discussions and notes on scraps of paper. An engineering design paradox arises with these informal forms of communication. First, they are essential and must be informal if information is to be shared and progress to be made. Second, for the most part the information is not in a form that is documented for future use. In other words, the information and arguments used to reach many decisions are not recorded as part of any permanent design record and can be lost or easily misinterpreted. Thus, it is important to make the effort to record important discussions and decisions.

Formal communication generally is in the form of *design notebooks*, *design records*, *communications to management*, and *communication of the final design to downstream phases*.

5.6.1 Design Notebooks and Records

Each technique discussed in this book produces documents that will become part of a design file for the product. The company keeps this file as a record of the product's development for future reference, perhaps to prove originality in case of patent application or to demonstrate professional design procedures in case of a lawsuit. However, a complete record of the design must go beyond these formal documents.

In solving any design problem, it is essential to keep track of the ideas developed and the decisions made in a *design notebook*. Some companies require these, with every entry signed and dated for legal purposes. In cases where a patent may be applied for or defended against infringement, it is necessary to have complete documentation of the birth and development of an idea. A design notebook with sequentially numbered, signed, and dated pages is considered good documentation; random bits of information scrawled on bits of paper are not. Additionally, a

lawsuit against a designer or a company for injury caused by a product can be won or lost on the basis of records that show that state-of-the-art design practices were used in the development of the product. Design notebooks also serve as reference to the history of the designer's own work. Even in the case of a simple design, it is common for designers to be unable to recall later why they made a specific decision. Also, it is not uncommon for an engineer to come up with a great idea only to discover it in earlier notes.

The design notebook is a diary of the design. It does not have to be neat, but it should contain all sketches, notes, and calculations that concern the design. Before starting a design problem, be sure you have a bound notebook—one with lined paper on one page and graph paper on the other is preferable. The first entry in this notebook should be your name, the company's name, and the title of the problem. Follow this with the problem statement, as well as it is known. Number, date, and sign each page. If test records, computer readouts, and other information are too bulky to be cut and pasted into the design notebook, enter a note stating what the document is and where it is filed.

There have been efforts to keep design notebooks on computers. It is still difficult for computer-based systems to manage the sketches and notes, and they lack the permanence to hold up in court.

More formal design records are created with each step of the design process. In this book, there are over 20 templates used that give an outline for the needed records. The information contained in these is what is managed and integrated in a PLM system.

5.6.2 Documents Communicating with Management

During the design process, periodic presentations to managers, customers, and other team members will be made. These presentations are usually called *design reviews* and are shown as an “approve plan” decision point in Fig. 5.1. Although there is no set form for design reviews, they usually require both written and oral communication. Whatever the form, these guidelines are useful in preparing material for a design review.

Make it understandable to the recipient. Clear communication is the responsibility of the sender of the information. It is essential in explaining a concept to others that you have a clear grasp of what they already know and do not know about the concept and the technologies being used.

Carefully consider the order of presentation. How should a bicycle be described to someone who has never seen one? Would you describe the wheels first, then the frame, the handlebars, the gears, and finally the whole assembly? Probably not, as the audience would understand very little about how all these bits fit together. A three-step approach is best: (1) Present the whole concept or assembly and explain its overall function, (2) describe the major parts and how they relate to the whole and its function, and (3) tie

the parts together into the whole. This same approach works in trying to describe the progress in a project: *Give the whole picture; detail the important tasks accomplished; then give the whole picture again.* There is a corollary to this guideline: *New ideas must be phased in gradually.* Always start with what the audience knows and work toward the unknown. Above all, do not use jargon or terms with which the audience is not familiar. If in doubt about a concept or TLA (Three Letter Acronym), define it.

Be prepared with quality material. The best way to make a point, and to have any meeting end well, is to be prepared. This implies (1) having good visual aids and written documentation, (2) following an agenda, and (3) being ready for questions beyond the material presented.

Good visual aids include diagrams and sketches specifically prepared to communicate a well-defined point. In cases in which the audience in the design review is familiar with the design, mechanical drawings might do, but if the audience is composed of nonengineers who are unfamiliar with the product, such drawings communicate very little. It is always best to have a written agenda for a meeting. Without an agenda, a meeting tends to lose focus. If there are specific points to be made or questions to be answered, an agenda ensures that these items are addressed.

5.6.3 Documents Communicating the Final Design

The most obvious form of documentation to result from a design effort is the material that describes the final design. Such materials include computer solid models, drawings (or computer data files) of individual components (detail drawings) and of assemblies to convey the product to manufacturing. They also include written documentation to guide manufacture, assembly, inspection, installation, maintenance, retirement, and quality control. These topics will be covered in Chaps. 9 and 12.

Often it is necessary to produce a design report. The following format is a good outline to follow.

1. **Title page:** The title of the design project is to be in the center of the page. Below it, list the following items:
 - a. Date:
 - b. Course/Section:
 - c. Instructor:
 - d. Team Members:
2. **Executive summary:**
 - a. The purpose of the Executive Summary is to provide key information up front, such that while reading the report, a reader has expectations that are fulfilled on a continuous basis. Key to a good summary is the *first* sentence, which *must* contain the most essential information that you wish to convey.



- b. The summary is to be written as if the reader is totally uninformed about your project and is not necessarily going to read the report itself.
 - c. It must include a short description of the project, the process and the results.
 - d. The Executive Summary is to be one page or less with one figure maximum.
- 3. **Table of contents:** Include section titles and page numbers.
- 4. **Design problem and objectives:** Give a clear and concise definition of the problem and the intended objectives. Outline the design constraints and cost implications.
 - a. Include appropriate background on the project for the reader to be able to put the information provided in context.
 - b. The final project objectives *must* also be presented in the form of a set of engineering specifications.
- 5. **Detailed design documentation:** Show all elements of your design including an explanation of
 - a. Assumptions made, making sure to justify your design decisions.
 - b. Function of the system.
 - c. Ability to meet engineering specifications.
 - d. Prototypes developed, their testing and results relative to engineering specifications.
 - e. Cost analysis.
 - f. Manufacturing processes used.
 - g. DFX results.
 - h. Human factors considered.
 - i. All diagrams, figures, and tables should be accurately and clearly labeled with meaningful names and/or titles. When there are numerous pages of computer-generated data, it is preferable to put this information in an appendix with an explanation in the report narrative. For each figure in the report, ensure that every feature of it is explained in the text.
- 6. **Laboratory test plans and results** for all portions of the system that you built and tested. Write a narrative description of test plan(s). Use tables, graphs, and whatever possible to show your results. Also, include a description of how you plan to test the final system, and any features you will include in the design to facilitate this testing. This section forms the written record of the performance of your design against specifications.
- 7. **Bills of materials:** Parts costs include only those items included in the final design. A detailed bill of materials includes (if possible) manufacturer, part number, part description, supplier, quantity, and cost.
- 8. **Gantt chart:** Show a complete listing of the major tasks to be performed, a time schedule for completing them, and which team member has the primary responsibility (and who will be held accountable) for each task.

9. **Ethical consideration:** Provide information on any ethical considerations that govern the product specifications you have developed or that need to be taken into account in potentially marketing the product.
10. **Safety:** Provide a statement of the safety consideration in your proposed design to the extent that is relevant.
11. **Conclusions:** Provide a reasoned listing of only the most significant results.
12. **Acknowledgments:** List individuals and/or companies that provided support in the way of equipment, advice, money, samples, and the like.
13. **References:** Including books, technical journals, and patents.
14. **Appendices:** As needed for the following types of information:
 - a. Detailed computations and computer-generated data.
 - b. Manufacturers' specifications.
 - c. Original laboratory data.

5.7 SUMMARY

- Planning is an important engineering activity.
- The use of prototypes and models is important to consider during planning.
- Every product is developed through five phases: discovery, specification development, conceptual design, product development, and product support. Planning is needed to get through these phases in a timely, cost-effective manner.
- There are five planning steps: identify the tasks, state their objectives, estimate the resources needed, develop a sequence, and estimate the cost.
- There are many types of project plans. A goal is to design a plan to meet the needs of the project.
- Communication through reports and drawings are key to the success of any project.

5.8 SOURCES

- Bashir, H., and V. Thompson: "Estimating Design Complexity," *Journal of Engineering Design*, Vol. 16, No. 3, 1999, pp. 247–256. Estimates on project time are based on this paper.
- Boehm, B.: "The Spiral Model as a Tool for Evolutionary Acquisition," Software Engineering Institute, Pittsburgh, Pa. www.sei.cmu.edu/pub/documents/00.reports/pdf/00sr008.pdf
- Boehm, B.: "The Spiral Model as a Tool for Evolutionary Acquisition," *Crosstalk*, May 2001. <http://www.stsc.hill.af.mil/crosstalk/2001/may/boehm.asp>
- Cooper, Robert G.: *Winning at New Products: Accelerating the Process from Idea to Launch*, Third Edition, Perseus Books Group, 2001. The basic book on Stage-Gate methods.
- Meredith, D. D., K. W. Wong, R. W. Woodhead, and R. H. Wortman: *Design Planning of Engineering Systems*, Prentice-Hall, Englewood Cliffs, N.J., 1985. Good basic coverage of mathematical modeling, optimization, and project planning, including CPM and PERT.

MicroSoft Project™. Software that supports the planning activity. There are many share-ware versions available.

For details on the Design Structure Matrix see *The DSM Website* at MIT, <http://www.dsmweb.org/>. A tutorial there is instructive.

The Design Report format is used, with permission, from the Electrical Engineering Program at The Milwaukee School of Engineering.

5.9 EXERCISES

- 5.1 Develop a plan for the original or redesign problem identified in Exercise 4.1 or 4.2.
 - a. Identify the participants on the design team.
 - b. Identify and state the objective for each needed task.
 - c. Identify the deliverables.
 - d. Justify the use of prototypes.
 - e. Estimate the resources needed for each task.
 - f. Develop a schedule and a cost estimate for the design project.
- 5.2 For the features of the redesign problem (Exercise 4.2) develop a plan as in Exercise 5.1.
- 5.3 Develop a plan for making a breakfast consisting of toast, coffee, a fried egg, and juice. Be sure to state the objective of each task in terms of the results of the activities performed, not in terms of the activities themselves.
- 5.4 Develop a plan to design an orange ripeness tester. In a market, people test the freshness of oranges by squeezing them, and based on their experience, how much they compress when squeezed gives an indication of ripeness. There are some sophisticated methods used in industry, but the goal here is to develop something simple, that could be built for low cost.



5.10 ON THE WEB

Templates for the following documents are available on the book's website: www.mhhe.com/Ullman4e

- Project Plan
- Design Report

Understanding the Problem and the Development of Engineering Specifications

KEY QUESTIONS

- Why emphasize developing engineering specifications?
- How can you identify the “customers” for a product?
- Why is it so important to understand the voice of the customer and work to translate this into engineering specifications?
- How can you best benchmark the competition to understand design and business opportunities?
- How can you justify taking time at the beginning of a project to do specification development instead of developing concepts immediately?

6.1 INTRODUCTION

Understanding the design problem is an essential foundation for designing a quality product. “Understanding the design problem” means to *translate customers’ requirements into a technical description of what needs to be designed*. Or, as the Japanese say, “Listen to the voice of the customer.” This importance is made graphically clear in the cartoon shown in Fig. 6.1. Everyone has a different view of what is needed by the customer and it takes work to find out what this really is.

Surveys show that poor product definition is a factor in 80% of all time-to-market delays. Further, getting a product to market late is more costly to a company than being over cost or having less than optimal performance. Finding the “right” problem to be solved may seem a simple task; unfortunately, often it is not.

Besides finding the right problem to solve, an even more difficult and expensive problem for most companies is what is often called “creeping specifications.”

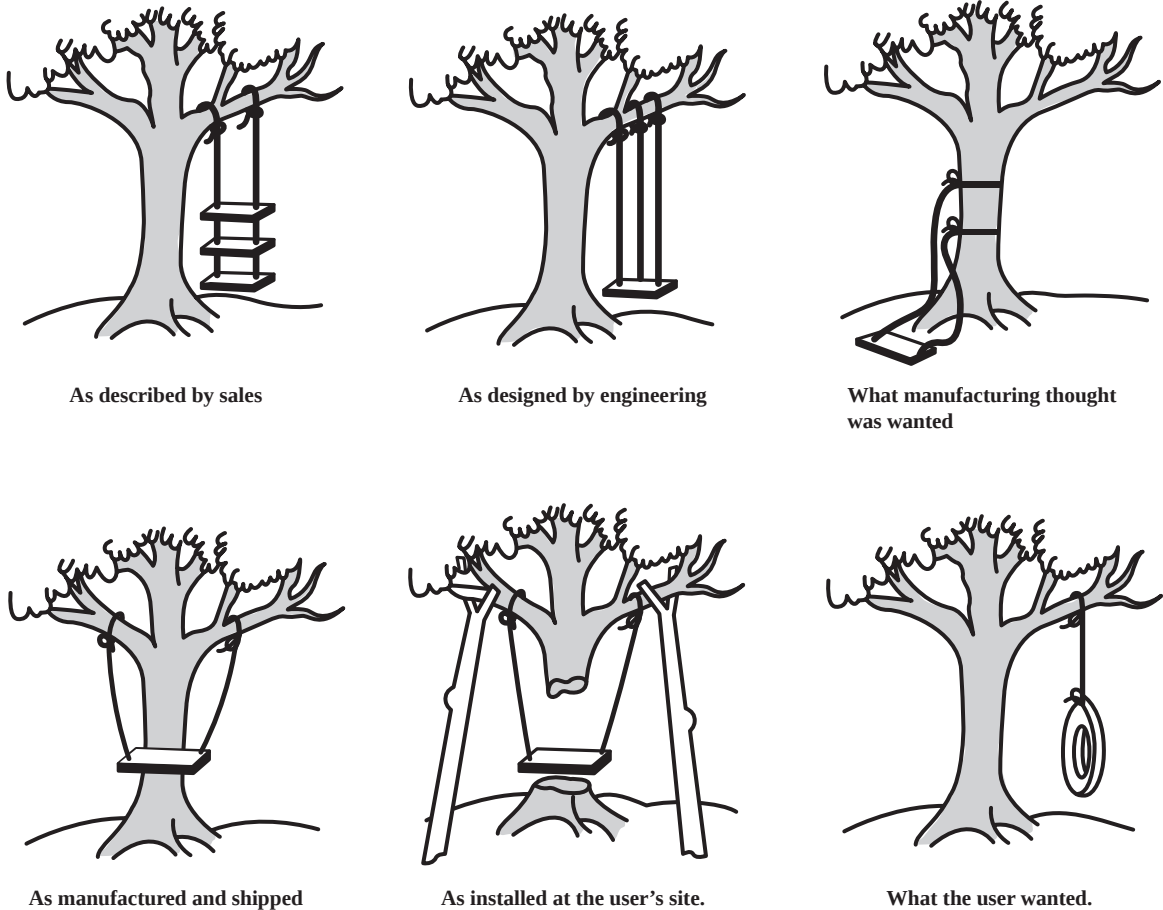


Figure 6.1 Understanding the product need.

Creeping specifications change during the design process. It is estimated that fully 35% of all product development delays are directly caused by such changes. There are three factors that cause creeping specifications. First, as the design process progresses, more is learned about the product and so more features can be added. Second, since design takes time, new technologies and competitive products become available during the design process. It is a difficult decision whether to ignore these, incorporate them (i.e., change the specifications), or start all over (i.e., decide that the new developments have eliminated the market for what you are designing). Third, since design requires decision making, any specification change causes a readdressing of all the decisions dependent on that specification. Even a seemingly simple specification change can cause redesign of virtually the whole product. The point is that when specification changes become necessary, they should be done in a controlled and informed manner.

All design problems are poorly defined.

The importance of the early phases of the design process has been repeatedly emphasized. As pointed out in Chap. 1, careful requirements development is a key feature of an effective design process. In this chapter, the focus is on understanding the problem that is to be solved. The ability to write a good set of engineering specifications is proof that the design team understands the problem.

There are many techniques used to generate engineering specifications. One of the best and currently most popular is called *Quality Function Deployment (QFD)*. What is good about the QFD method is that it is organized to develop the major pieces of information necessary to understanding the problem:

1. Hearing the voice of the customers
2. Developing the specifications or goals for the product
3. Finding out how the specifications measure the customers' desires
4. Determining how well the competition meets the goals
5. Developing numerical targets to work toward

The QFD method was developed in Japan in the mid-1970s and introduced in the United States in the late 1980s. Using this method, Toyota was able to reduce the costs of bringing a new car model to market by over 60% and to decrease the time required for its development by one-third. It achieved these results while improving the quality of the product. A recent survey of 150 U.S. companies shows that 69% use the QFD method and that 71% of these have begun using the method since 1990. A majority of companies use the method with cross-functional teams of ten or fewer members. Of the companies surveyed, 83% felt that the method had increased customer satisfaction and 76% indicated that it facilitated rational decisions.

Before itemizing the steps that comprise this technique for understanding a design problem, consider some important points:

1. No matter how well the design team thinks it understands a problem, it should employ the QFD method for all original design or redesign projects. In the process, the team will learn what it does not know about the problem.
2. The customers' requirements must be translated into *measurable design targets for identified critical parameters*. You cannot design a car door that is "easy to open" when you do not know the meaning of "easy." Is easiness measured by force, time, or what? If force is a critical parameter, then is "easy" 20 N or 40 N? The answer must be known before much time and resources are invested in the design effort.
3. The QFD method can be applied to the entire problem and any subproblem. (Note that the design of a door mechanism in the previous point is a subproblem in automobile design.)

4. It is important to first worry about *what* needs to be designed and, only after that is understood, to worry about *how* the design will look and work. Our cognitive capabilities generally lead us to try to assimilate the customers' functional requirements (what is to be designed) in terms of form (how it will look); these images then become our favored designs and we get locked onto them. The QFD procedure helps overcome this cognitive limitation.
5. This method takes time to complete. In some design projects, about one-third of the total project time is spent on this activity. Ford spends 3–12 months developing the QFD for a new feature. Experimental evidence has shown that designers who spend time here end up with better products and do not use any more total time when compared to others who do a superficial job here. Time spent here saves time later. Not only does the technique help in understanding the problem, it also helps set the foundation for concept generation.

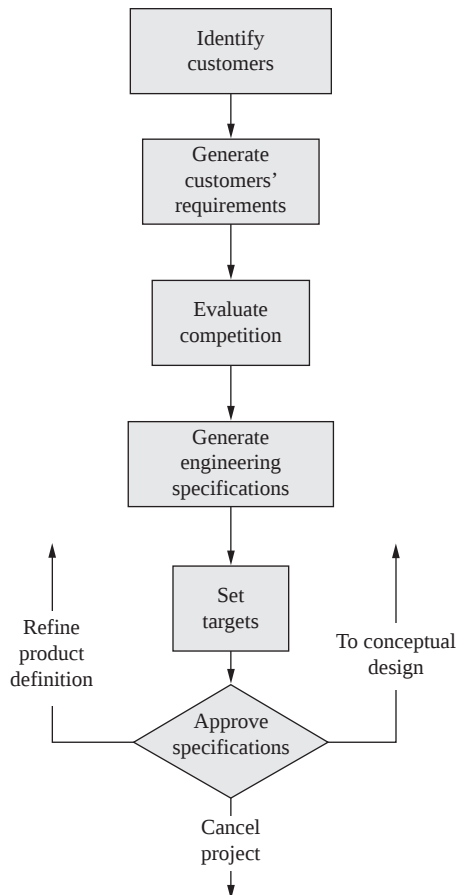


Figure 6.2 The Product Definition phase of the mechanical design process.

The QFD method helps generate the information needed in the engineering Product Definition phase of the design process (Fig. 4.1). That phase is reproduced in Fig. 6.2. Each block in the diagram is a major section in this chapter and a step in the QFD method.

Applying the QFD steps builds the *house of quality* shown in Fig. 6.3. This house-shaped diagram is built of many rooms, each containing valuable information. Before we describe each step for filling in Fig. 6.3, a brief description of the figure is helpful. The numbers in the figure refer to the steps that are detailed in the sections below. Developing information begins with identifying *who* (step 1) the customers are and *what* (step 2) it is they want the product to do. In developing this information, we also determine to whom the “what” is important—*who versus what* (step 3). Then it is important to identify how the problem is solved *now* (step 4), in other words, what the competition is for the product being designed. This information is compared to what the customers desire—*now versus what* (step 4 continued)—to find out where there are opportunities for an improved product. Next comes one of the more difficult steps in developing the house, determining *how* (step 5) you are going to measure the product’s

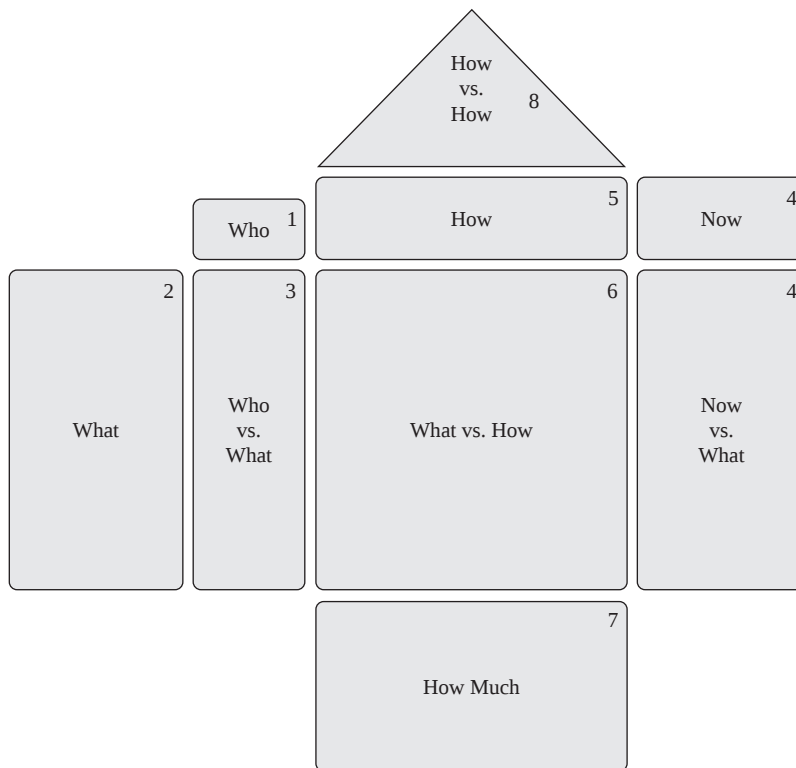


Figure 6.3 The house of quality, also known as the QFD diagram.

ability to satisfy the customers' requirements. The *hows* consists of the engineering specifications, and their correlation to the customers' requirements is given by *whats versus hows* (step 6). Target information—*how much* (step 7)—is developed in the basement of the house. Finally, the interrelationship between the engineering specifications are noted in the attic of the house—*how versus how* (step 8). Details of all these steps and why they are important are developed in Sections 6.2 through 6.9. Postage stamp-size versions of Fig. 6.3 tie the steps together.

The QFD method is best for collecting and refining functional requirements, hence the “F” in its name. However, in the material presented here, it will be used to help ensure that all requirements are collected and refined. In each step, the design of an “aisle chair” will be used as an example. This example is taken from a project to design a wheelchair to rapidly help passengers board and deplane from a Boeing 787 Dreamliner. This type of wheelchair is brought into the waiting area, the passenger transfers from their regular wheelchair to the aisle chair, which is then wheeled to the plane and down the aisle to the assigned seat where the passenger transfers out of the aisle chair into their seat. The process is reversed at the end of the flight. Aisle chairs are narrower than regular chairs so they can fit between the rows on an aircraft. A typical aisle chair is shown in Fig. 6.4.

The design effort for the Dreamliner chair resulted in the QFD shown in Fig. 6.5. This House of Quality developed during this project contained over 60 customer requirements and over 50 engineering specifications. This effort, although time consuming, resulted in the increased project understanding that was essential to develop a product that was superior to those already on the market.

The entire House is too large to read or make for a good example, so a reduced version of it will be used (Fig. 6.6). This example contains all the important points used in the larger, complete QFD. The contents of this house are developed in the following sections.



Figure 6.4 A typical aisle chair. (Reprinted with permission of Columbia Medical.)

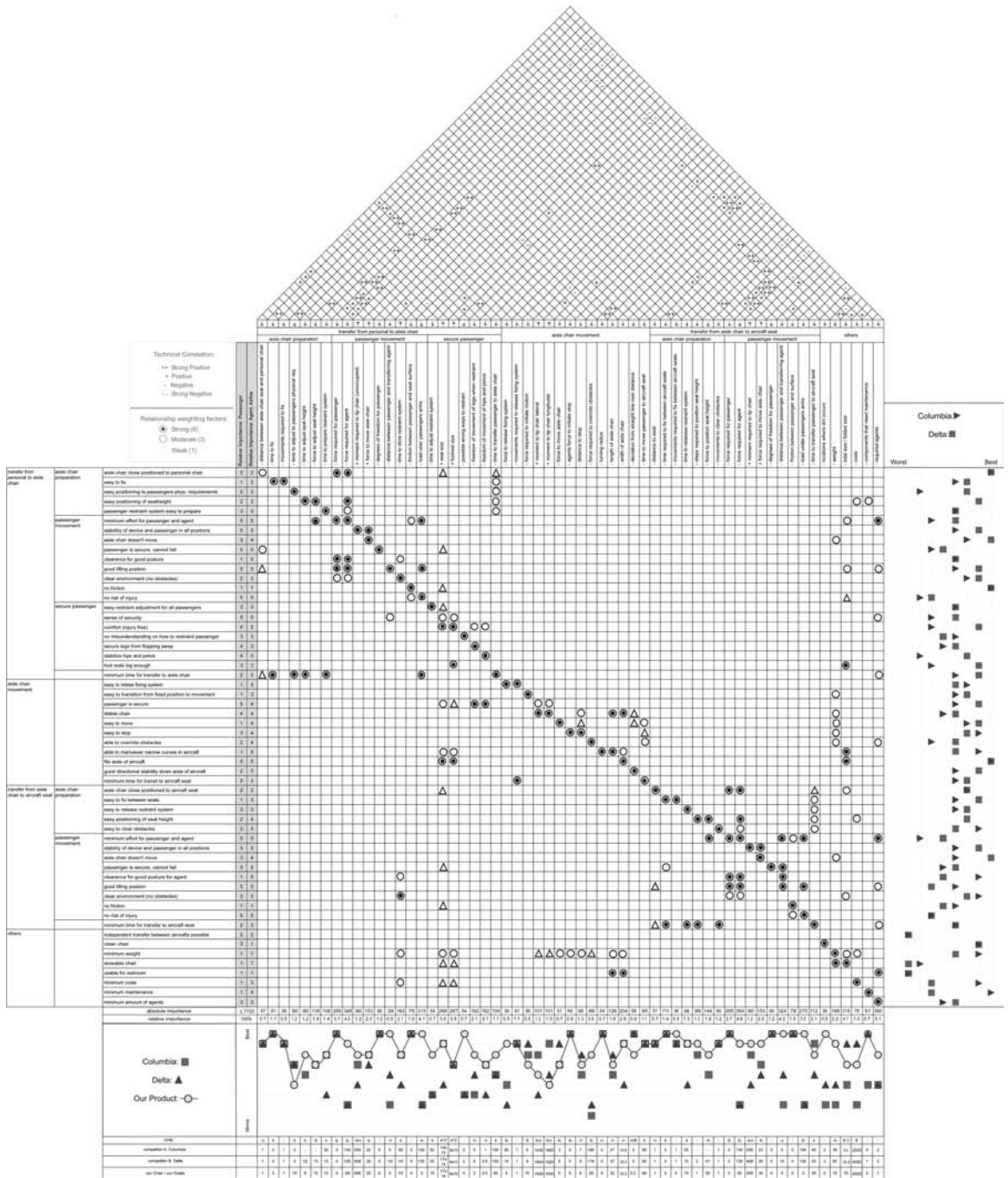


Figure 6.5 Aisle chair QFD (original available on book web site).

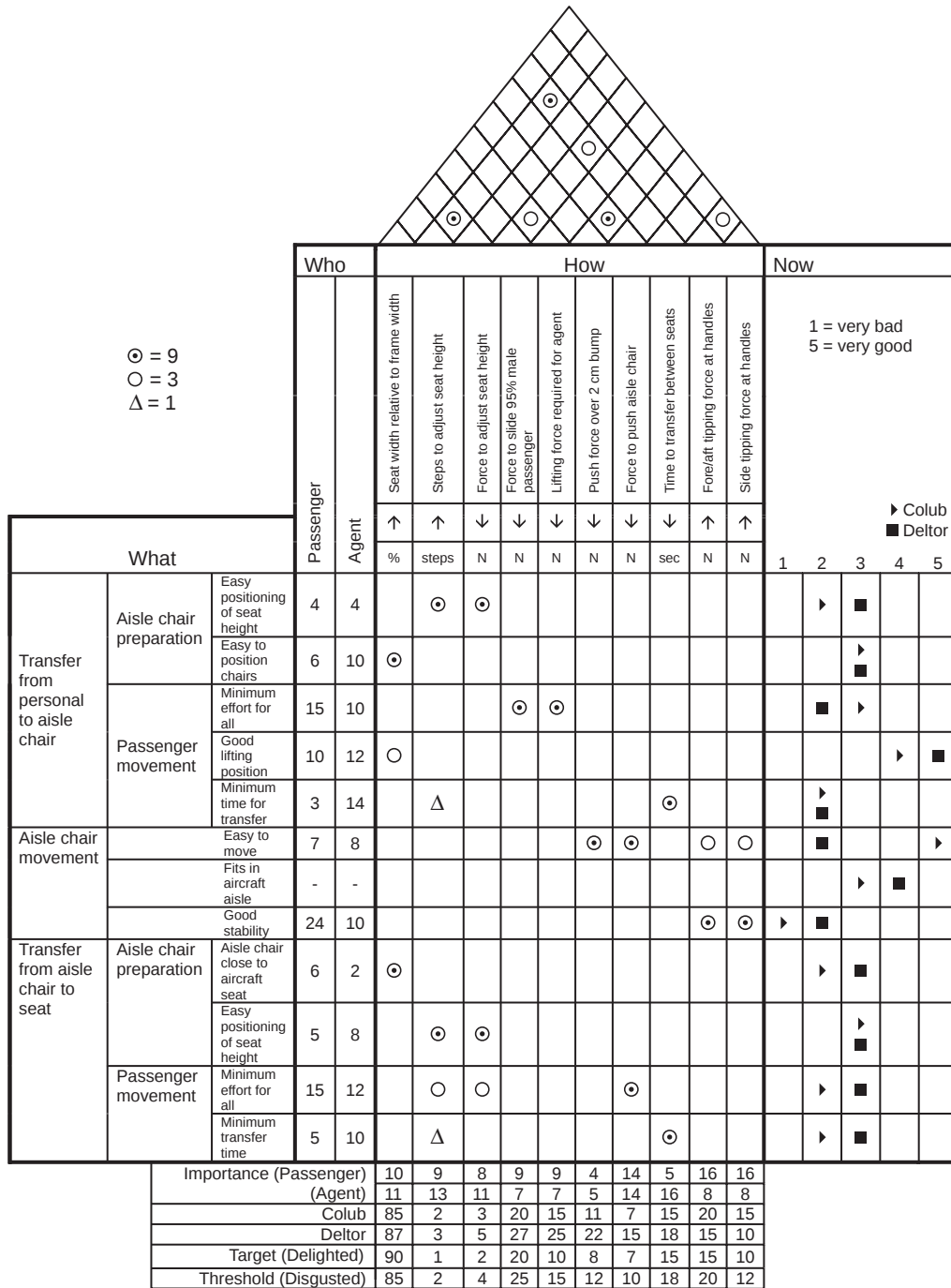


Figure 6.6 Example aisle chair QFD.

Your decisions, good or bad, affect everyone downstream.

The House of Quality can be easily built on a spreadsheet with the exception of the roof portion at the top. A simple method to construct this also on a spreadsheet is given in step 8.

6.2 STEP 1: IDENTIFY THE CUSTOMERS: WHO ARE THEY?

For most design situations, there is more than one customer; for many products, the most important customers are the consumers, the people who will buy the product and who will tell other consumers about its quality (or lack thereof). Sometimes the purchaser of the product is not the same as its user (e.g., gym equipment, school desks, and office desks). Some products—a space shuttle or an oil drill head—are not consumer products but still have a broad customer base.

For all products it is important to consider customers both outside the organizations that design, manufacture, and distribute the product—external customers—and those inside of them—internal customers. For example, beyond the consumer, the designer's management, manufacturing personnel, sales staff, and service personnel must also be considered as customers. Additionally, standards organizations should be viewed as customers, as they too may set requirements for the product. For many products, there are five or more classes of customers whose voices need to be heard.

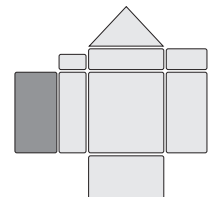
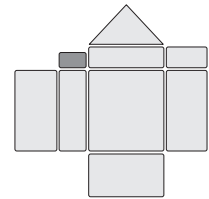
One method to make sure you have identified all the customers is to consider the entire life of the product (see Fig. 1.7). Pretend you are the product; visualize all the people that encounter you as you go through the internal and external phases itemized in life cycle diagram.

For the aisle chair, the main customers are the passengers being transported and the airline agents who assist in transporting the passengers on and off the airplane. Note that neither of these two customers purchases the aisle chair. Nor do they maintain it, clean it, or disassemble it. In Fig. 6.6 the only customers shown are the passenger and agent as “who” examples. The area below the “passenger” and “agent” will be filled in during Step 3.

6.3 STEP 2: DETERMINE THE CUSTOMERS' REQUIREMENTS: WHAT DO THE CUSTOMERS WANT?

Once the customers have been identified, the next goal of the QFD method is to determine *what* is to be designed. That is, what is it that the customers want?

- Typically, as shown by the customer survey in Table 1.1, the *consumers* want a product that works as it should, lasts a long time, is easy to maintain, looks attractive, incorporates the latest technology, and has many features.



You only think you know what your customers want.

- Typically, the *production customer* wants a product that is easy to produce (both manufacture and assemble), uses available resources (human skills, equipment, and raw materials), uses standard parts and methods, uses existing facilities, and produces a minimum of scraps and rejected parts.
- Typically, the *marketing/sales customer* wants a product that meets consumers' requirements; is easy to package, store, and transport; is attractive; and is suitable for display.

The key to this QFD step is collecting information from customers. There are essentially three methods commonly used: observations, surveys, and focus groups.

Fortunately, most new products are refinements of existing products, so many requirements can be found by *observing* customers using the existing product. For example, automobile manufacturers send engineers into shopping center parking lots to observe customers putting purchases into cars to better understand one aspect of car door requirements.

Surveys are generally used to gather specific information or ask people's opinions about a well-defined subject. Surveys use questionnaires that are carefully crafted and applied either through the mail, over the telephone, or in face-to-face interviews. Surveys are well suited for collecting requirements on products to be redesigned or on new, well-understood product domains. For original products or to gather the customers' ideas for product improvement, focus groups are best.

The *focus-group* technique was developed in the 1980s to help capture customers' requirements from a carefully chosen group of potential customers. The method begins by identifying seven to ten potential customers and asking if they will attend a meeting to discuss a new product. One member of the design team acts as moderator and another as note taker. It is also best to electronically record the session. The goal in the meeting is to find out what is wanted in a product that does not yet exist, and so it relies on the customers' imaginations. Initial questions about the participants' use of similar products are followed with questions designed to find performance and excitement requirements. The goal of the moderator is to use questions to guide the discussion, not control it. The group should need little intervention from the moderator, because the participants build on each other's comments. One technique that helps elicit useful requirements during interviews is for the moderator to repeatedly ask "Why?" until the customers respond with information in terms of time, cost, or quality. Eliciting good information takes experience, training, and multiple sessions with different participants. Usually the first focus group leads to questions needed for the second group. It often takes as many as six sessions to obtain stable information.

Later in the design process, surveys can be used to gather opinions about the relative merit of different alternatives. Observation and focus groups can be used both to generate ideas that may become alternatives and to evaluate

alternatives. All these types of information gathering rely on questions formulated ahead of time. With a survey, the questions and the answers must be formalized. Both surveys and observations usually use closed questions (i.e., questions with predetermined answers); focus groups use open-ended questions.

Regardless of the method used, these steps will help the design team develop useful data:

Step 2.1: Specify the Information Needed Reduce the problem to a single statement describing the information needed. If no single statement represents what is needed, more than one data-collecting effort may be warranted.

Step 2.2: Determine the Type of Data-Collection Method to Be Used Base the use of focus groups, observations, or surveys on the type of information being collected.

Step 2.3: Determine the Content of Individual Questions A clear goal for the results expected from *each question* should be written. Each question should have a single goal. For a focus group or observation, this may not be possible for all questions, but it should be for the initial questions and other key questions.

Step 2.4: Design the Questions Each question should seek information in an unbiased, unambiguous, clear, and brief manner. Key guidelines are

- Do not assume the customers have more than common knowledge.
- Do not use jargon.
- Do not lead the customer toward the answer you want.
- Do not tangle two questions together.
- Do use complete sentences.

Questions can be in one of four forms:

- Yes–no–don't know. (Poor for focus groups.)
- Ordered choices (1, 2, 3, 4, 5; strongly agree, mildly agree, neither agree nor disagree, mildly disagree, strongly disagree; or A = absolutely important, E = extremely important, I = important, O = ordinary, or U = unimportant [AEIOU]). Be sure that any ordered list is complete (i.e., that it covers the full range possible and that the choices are unambiguously worded). Scales with five gradations, as in the examples here, have proven best.
- Unordered choices (a, b, and/or c).
- Ranking (a is better than b is better than c).

The best questions ask about attributes, not influences. Attributes express what, where, how, or when. *Why* questions should lead to what, where, how, or when as they describe time, quality, and cost.

Step 2.5: Order the Questions Order the questions to give context. This will help participants in focus groups or surveys follow the logic.

Step 2.6: Take Data It usually takes repeated application to generate usable information. The first application of any set of questions should be considered a test or verification experiment.

Step 2.7: Reduce the Data A list of customers' requirements should be made in the customers' own words, such as "easy," "fast," "natural," and other abstract terms. A later step of the design process will be to translate these terms into engineering parameters. The list should be in positive terms—what the customers want, not what they don't want. We are not trying to patch a poor design; we are trying to develop a good one.

To gather information for the aisle chairs, focus groups of passengers were used. These began with a free discussion of people's experiences traveling by air. There is no way that an able-bodied person can understand the challenges of traveling when a wheelchair is involved, and once a group of wheelchair-bound travelers start trading stories, much is learned about what will be needed to make their experience tolerable. It is better that travel should be a "Wow" experience, as discussed in Kano's model (Section 4.4.2) than a "tolerated" experience. A similar focus group was held with agents. Finally, a researcher went to the airport and observed over 20 people boarding and off-loading using wheelchairs.

A sampling of the results of the focus groups and observations are (in no particular order)

- Easy positioning of seat height of the aisle chair so that it matches the wheelchair and the plane's seat so that the passenger can easily slide from on to the other.
- Once in the aisle chair it should be easy to move and stable.
- The aisle chair should fit in all aircraft aisles
- When transferring between chairs, the passenger with possibly some help from the agent must lift their weight enough to slide from chair to chair, so there needs to be a good lifting position for both of them so they can exert minimal effort.
- All want the transfer from seat to seat to be as fast as possible.
- It should be easy to position chairs next to each other and have them not slide apart.

To make sense of these results it is best to organize them into a hierarchical structure. In reviewing the observations it is evident that there are three main phases to the use of the aisle chair: (1) transfer the passenger from their personal wheelchair to the aisle chair, (2) move the aisle chair from the waiting area to the assigned seat, and (3) transfer the passenger from the aisle chair to the assigned seat. The same basic functions have to occur when deplaning. This is a simple form of functional modeling, which will be covered in detail in Chap. 7. Further, the action of transferring to the aisle chair requires two steps, prepare the chair and move the passenger. This decomposition of the function leads to a structure for organizing the results of voice of the customer. This can be organized like an outline (below) and also entered into the QFD as shown in Fig. 6.6.



Transfer from personal to aisle chair

1. Aisle chair preparation
 - a. Easy positioning of seat height
 - b. Easy to position chairs
2. Passenger movement
 - a. Minimum effort for all
 - b. Good lifting position
 - c. Minimum time for transfer

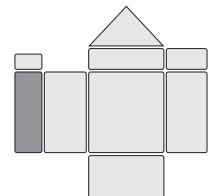
Building a hierarchy like this can help you look for completeness. If the structure has discontinuities, these may be indicators of needed information. The aisle chair has only one major function and thus the hierarchy is fairly simple. Products that have multiple uses may have multiple hierarchies.

Suggestions to get the best possible customers' requirements

- #1—Do not assume you know what the customer wants.
- If customer requirements are too vague (e.g., product must be durable), then go back to the customer and flesh these out a little more in the customer's words. What is "durability"? Does that mean you can jump up and down on it? Does it mean that it lasts more than a minute?
- Frequently, customers will try to express their needs in terms of *how* the need can be satisfied and not in terms of *what* the need is. This limits consideration of development alternatives. You should ask *why* until you truly understand what the root need is. Do keep in mind that the only way they may have of expressing what they want is in terms of analogies and comparisons to other products.
- Use Kano's model to help you steer away from basic requirements to performance and excitement requirements.
- Break down general requirements into more specific requirements by probing what is needed.
- Challenge, question, and clarify requirements until they make sense and you can put them in an outline format—a hierarchy. This helps understand function and look for completeness.
- Document situations and circumstances to illustrate a customer need.

6.4 STEP 3: DETERMINE RELATIVE IMPORTANCE OF THE REQUIREMENTS: WHO VERSUS WHAT

The next step in the QFD technique is evaluating the importance of each of the customers' requirements. This is accomplished by generating a weighting factor



for each requirement and entering it in Fig. 6.6. The weighting will give an idea of how much effort, time, and money to invest in achieving each requirement. Two questions are addressed here: (1) to whom is the requirement important? and (2) how is a measure of importance developed for this diverse group of requirements?

Since a design is “good” only if the customers think it is good, the obvious answer to the first question is, the customer. However, we know that there may be more than one customer. In the case of a piece of production machinery, the desires of the workers who will use the machine and those of management may not be the same. This discrepancy must be resolved at the beginning of the design process or the requirements may change partway through the job. Sometimes a designer’s hardest job is determining whom to please.

The region of the house of quality labeled “who vs. what” in Fig. 6.3 is for the input of the importance of each requirement. It is essential to understand which requirements each type of customer thinks is important. Note that, in most cases, less than half of the requirements have most of the importance. The best way to represent importance is with a number showing its *weight* relative to the other requirements.

Traditionally, weighting has been done by instructing the customers to rate the requirements on a scale of 1 to 10 with 10 being important and 1 being unimportant. Unfortunately, often these methods result in everything being scored 8, 9, or 10—everything is important.

A better method, the fixed sum method, is to tell each customer that they have 100 points to distribute among the requirements. Using the fixed sum of 100 forces the customer to rate some of the requirements low if they want others to be high. This method works much better than just telling them to rate requirements on a scale of 1 to 10.

To aid in weighting, write each requirement on a piece of self-stick note paper, put the notes on a wall, and ask each customer to arrange them in order of importance. If two or more requirements seem to be equally important, be sure that they don’t measure the same thing, that they are independent. Once the notes are in order, allocating the 100 points should be easier.

If there are more than 30 requirements, allocating weights can be very difficult. It is suggested that the large group of requirements be broken into smaller groups using the hierarchy, weighting each, and then renormalizing across all the requirements.

If you collect weightings from more than one representative of a customer group and they are in fairly good agreement with each other, then just average them. If weightings are significantly different from each other, then this is a signal that you have two different types of customers and you need to revisit the step 1.

The results of weighting the requirements for the aisle chair are shown in Fig. 6.6 for the passenger and the agent. The fixed sum method was used to set the weights. Note that the requirement “Fits in aircraft aisle” was not weighted. It was realized that this was a basic requirement (in Kano’s terminology) as an

One man's treasure is another's trash.
Both will judge your work.

aisle chair that does not fit in the aisle is not a viable product. Requirements that measure basic needs are not helpful. Before you eliminate them, however, go back and ask if the requirement can be reworded so that it addresses performance or excitement. Also note that the passenger is more concerned about ease of use and the agent more focused on time. This is as expected.

6.5 STEP 4: IDENTIFY AND EVALUATE THE COMPETITION: HOW SATISFIED ARE THE CUSTOMERS NOW?

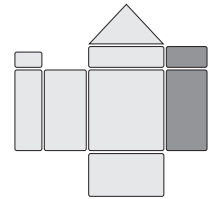
The goal here is to determine how the customer perceives the competition's ability to meet each of the requirements. Even though you may be working with a totally new design, there is competition, or at least products that come close to filling the same need that your product does. The purpose for studying existing products is twofold: first, it creates an awareness of what already exists (the "now"), and second, it reveals opportunities to improve on what already exists. In some companies, this process is called *competition benchmarking* and is a major aspect of understanding a design problem. In benchmarking, each competing product must be compared with customers' requirements (now versus what). Here we are concerned only with a subjective comparison that is based on customer opinion. Later, in step 8, we will do a more objective comparison. For each customer's requirement, we rate the existing design on a scale of 1 to 5:

1. The product does not meet the requirement at all.
2. The product meets the requirement slightly.
3. The product meets the requirement somewhat.
4. The product meets the requirement mostly.
5. The product fulfills the requirement completely.

Though these are not very refined ratings, they do give an indication of how the competition is perceived by the customer.

This step is very important as it shows opportunities for product improvement. If all the competition rank low on one requirement, this is clearly an opportunity. This is especially so if the customers ranked that specific requirement highly important in step 3. If one of the competitors meets the requirement completely, this product should be studied and good ideas used from it (note patent implications as discussed in Section 7.5).

If your organization already makes a product and you are redesigning this product, then the current product is one benchmark. If it ranks high on an important requirement, don't change the features that helped it meet that requirement. In



To steal from one person is plagiarism, to be influenced by many is good design.

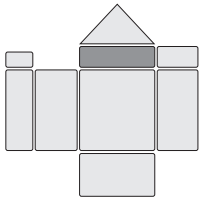
other words—Don't fix what aint broken! This step in the QFD method can help avoid needless work and product weakening.

The results of this step for the aisle chair are shown in Fig. 6.6. Here two competitor's chairs were evaluated (note that names have been changed). To determine how well the competitors met the requirements, the design team used questionnaires to evaluate them. The average results from passengers are shown in the "now vs. what" section of Fig. 6.6.

Important points to note are that

1. Both competitors have good lifting position when transferring the passenger from the personal chair to the aisle chair—study what makes this work well.
2. Both products have poor stability. Clearly, this is a market opportunity.
3. The Colub is easy to move and Delton is not, need to determine why and do what Colub does or better.
4. For most of adjustment requirements, neither of the competitors score above 3, leaving room for the development of a superior product in these areas.

There used to be a commercial on television for a family van in which the manufacturer bragged that its product was so good that one of its competitors bought and studied it. The commercial showed the competitor's technicians in white coats disassembling the van. What the commercial did not say was that the advertiser also bought and studied its competitor's product and that this is just good design practice.



6.6 STEP 5: GENERATE ENGINEERING SPECIFICATIONS: HOW WILL THE CUSTOMERS' REQUIREMENTS BE MET?

The goal here is to develop a set of *engineering specifications* from the customers' requirements. These specifications are the restatement of the design problem in terms of parameters that can be measured and have target values. Without such information the engineers cannot know if the system being developed will satisfy the customers. Engineering specifications consist of parameters of interest and targets for parameters. The parameters are developed in this step, and the target values for them are developed in step 8. In reality this step and the following one happen concurrently as will be made clear.

These specifications are a translation of the voice of the customer into the voice of the engineer. They serve as a vision of the ideal product and are used as criteria for design decisions. Conversely, this part of the QFD also builds a picture

Find the target before you empty your quiver.

of how design decisions affect the customer's perception of the quality of their product. We will make use of this in Chap. 10, where the effect of trading off the ability to meet one specification for the inability to meet another is addressed.

In this step, we develop parameters that tell *how* we know if customers' requirements have been met. We begin by finding as many engineering parameters as possible that indicate a level of achievement for customers' requirements. For example, a requirement for "easy to attach" can be measured by (1) the number of steps needed to attach it, (2) the time to attach it, (3) the number of parts, and (4) the number of standard tools used. Note that a set of units is associated with each of these measures—step count, time, part count, and tool count. *If units for an engineering parameter cannot be found, the parameter is not measurable and must be readdressed.* Each engineering parameter must be measurable and thus must have units of measure. However, "time to attach" may not be a reliable measure as it will be dependent on the skill and training of the customer. Either the customer's skill level needs to be defined or this parameter eliminated.

An important point here is that every effort must be made to find as many ways as possible to measure customers' requirements. If there are no measurable engineering parameters for customers' requirements, then the customer's requirement is not well understood. Possible solutions are to break the requirement into finer independent parts or to redo step 2 with specific attention to that specific requirement.

When developing the engineering specifications, carefully check each entry to see what nouns or noun phrases have been used. Each noun refers to an object that is part of the product or its environment and should be considered to see if new objects are being assumed. For example, if one specification in the aisle chair problem was for "easy to adjust seat height" then an adjustable seat height (a noun phrase) has been assumed as part of the solution. If the design team has made a decision that there is to be an adjustable seat height, this is acceptable. However, if no such assumption has been made, the product solution has been unknowingly limited. Paying attention to the objects that are part of the product is a major topic in concept generation.

Also shown on Fig. 6.6 are the units for each specification and the direction of improvement—the "sense" where either more is better (↑) or less is better (↓). These arrows tell whether more of the feature or parameter measured good, or bad. For example less "force required for agent" is good (↓). More "side tipping force" is desired (↑). A third option, not shown in the example is whether a specific target is best. Targets will be further discussed in step 7.

To help find specifications a checklist of the major types is given in Table 6.1. Comparing this list with the list of specifications developed for a product can reveal missing information. The major types of specifications in this list are detailed next.

Table 6.1 Types of engineering specifications

Functional performance	Life-cycle concerns (continued)
Flow of energy	Diagnosability
Flow of information	Testability
Flow of materials	Reparability
Operational steps	Cleanability
Operation sequence	Installability
Human factors	Retirement
Appearance	Resource concerns
Force and motion control	Time
Ease of controlling and sensing state	Cost
Physical requirements	Capital
Physical properties	Unit
Available spatial envelope	Equipment
Reliability	Standards
Mean time between failures	Environment
Safety (hazard assessment)	Manufacturing/assembly requirements
Life-cycle concerns	Materials
Distribution (shipping)	Quantity
Maintainability	Company capabilities

Functional performance requirements are those elements of the performance that describe the product's desired behavior. Although the customers may not use technical terms, function is usually described as the flow of energy, information, and materials or as information about the operational steps and their sequence. In Chap. 7 we develop concepts by building a functional model, based on the flow of energy, information, and materials. We will see that *developing functional requirements with the QFD and building a functional model of the product are often iterative*. The more the function is understood, the more complete are the requirements that can be developed.

Any product that is seen, touched, heard, tasted, smelled, or controlled by a human will have *human factors requirements* (see App. D for details on human factors). This includes nearly every product. One frequent customers' requirement is that the product "looks good" or looks as if it has a certain function. These are areas in which a team member with knowledge about industrial design is essential. Other requirements focus on the flow of energy and information between the product and the human. Energy flow is usually in terms of force and motion, but can take other forms as well. Information flow requirements apply to the ease of controlling and sensing the state of the product. Thus, human factors requirements are often functional performance requirements.

Physical requirements include needed physical properties and spatial restrictions. Some physical properties often used as requirements are weight; density; and conductivity of light, heat, or electricity (i.e., flow of energy). Spatial constraints relate how the product fits with other, existing objects. Almost all new design efforts are greatly affected by the physical interface with other objects that cannot be changed.

In the *Time* magazine survey on quality quoted in Chap. 1, the second most important consumer concern was “Lasts a long time,” or the product’s *reliability*. It is important to understand what acceptable reliability means to the customer. The product may only have to work once with near-absolute certainty (e.g., a rocket), or it may be a disposable product that does not need much reliability. As discussed in Chap. 11, one measure of reliability is the *mean time between failures*.

A part of reliability involves the questions, what happens when the product does fail? and, what are the *safety* implications? Product safety and hazard assessment are very important to the understanding of the product, and they are covered in Chap. 8.

An often overlooked class of requirements is the class of those relating the product life cycle other than product use. All specification types listed in Table 6.1 were taken from life cycle phases in Fig. 1.7. In designing the first BikeE, one of the design requirements set by sales/marketing was that the bicycle had to be shipped by a commercial parcel service. Such services have weight and size limits, which greatly affected the design of the product. If the advantages of distributing the product by commercial parcel service had not been realized early, extensive redesign might have been necessary. The same applies to the other life-cycle phases listed in Table 6.1 and Fig. 1.7.

A limited resource on every design project is time. *Time requirements* may come from the consumer; more often they originate in the market or in manufacturing needs. In some markets there are built-in time constraints. For example, toys must be ready for the summer buyer shows so that Christmas orders can be taken; new automobile models traditionally appear in the fall. Contracts with other companies might also determine time constraints. Even for a company without an annual or contractual commitment, time requirements are important. As discussed earlier, in the 1960s and 1970s Xerox dominated the copier market, but by 1980 its position had been eroded by domestic and Japanese competition. Xerox discovered that one of the problems was that it took it twice as long as some of its competitors to get a product to market, and Xerox put new time requirements on its engineers. Fortunately, Xerox helped its engineers work smarter, not just faster, by introducing techniques similar to those we talk about here.

Cost requirements concern both the capital costs and the costs per unit of production. Included in capital costs are expenditures for the design of the product. For a Ford automobile, design costs make up 5% of the manufacturing cost (Fig. 1.2). Many product ideas never get very far in development because the initial requirements for capital are more than the funds available. (Cost estimating will be covered in detail in Section 11.2.)

Standards spell out current engineering practice in common design situations. The term *code* is often used interchangeably with *standard*. Some standards serve as good sources of information. Other standards are legally binding and must be adhered to—for example, the ASME pressure vessel codes. Although the actual information contained in standards does not enter into the design process in this

early phase, knowledge of which standards apply to the current situation are important to requirements and must be noted from the beginning of the project.

Standards that are important to design projects generally fall into three categories: performance, test methods, and codes of practice. There are *performance standards* for many products, such as seat-belt strength, crash-helmet durability, and tape-recorder speeds. The *Product Standards Index* lists U.S. standards that apply to various products; most of those referenced are also covered by ANSI (American National Standards Institute), which does not write standards but is a clearinghouse for standards written by other organizations.

Test method standards for measuring properties such as hardness, strength, and impact toughness are common in mechanical engineering. Many of these are developed and maintained by the American Society for Testing and Materials (ASTM), an organization that publishes over 4000 individual standards covering the properties of materials, specifying equipment to test the properties, and outlining the procedures for testing. Another set of testing standards that are important to product design are those developed by the Underwriters Laboratories (UL). This organization's standards are intended to prevent loss of life and property from fire, crime, and casualty. There are over 350 UL standards. Products that have been tested by UL and have met their standards can display the words "Listed UL" and the standard number. The company developing the product must pay for this testing. Consumer products are usually not marketed without UL listing because the liability risk is too high without this proof of safe design.

Codes of practice give parameterized design methods for standard mechanical components, such as pressure vessels, welds, elevators, piping, and heat exchangers.

It is important for the design team to ensure that requirements imposed by *environmental concerns* have been identified. Since the design process must consider the entire life cycle of the product, it is the design engineer's responsibility to establish the impact of the product on the environment during production, operation, and retirement. Thus, requirements for the disposal of wastes produced during manufacture (whether hazardous or not), as well as for the final disposition of the product, are the concern of the design engineer. This topic is further discussed in Chap. 11.

Some of the *manufacturing/assembly requirements* are dictated by the quantity of the design to be produced and the characteristics of the company producing it. The quantity to be produced often affects the kind of manufacturing processes to be used. If only one unit is to be produced, then custom tooling cannot be amortized across a number of items and off-the-shelf components should be selected when possible (see Chap. 9). Additionally, every company has internal manufacturing resources whose use is preferable to contracting work outside the company. Such factors must be considered from the very beginning.

Guidelines for good specifications are

1. Each specification should measure at least one customers' requirement at the strong relationship level (see step 7). Ideally, each specification should

measure multiple requirements. If you have a diagonal of scores in step 7, you need to revisit the specifications.

2. Each specification should be measurable. Every specification should be written as if you were going to give instructions to someone to go down to the lab and measure something. It should be clear what they are going to measure. For example, the specification "Fore/aft tipping force" is a good title for a specification, but to be measurable it needs many more words. Thus, it is suggested that for each specification list, a full description of how to measure it also be developed. For example:

Fore aft tipping force = The force needed at the push handles to tip over the aisle chair when moving forward at 1 km/hr with 78.5-kg passenger (a 50% male, see App. D).

If a good statement like this cannot be developed, then the specification is not clear and needs to be reworked.

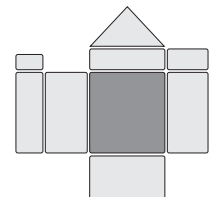
3. If the units are not clear, the specification is not clear.
4. If the sense ($\uparrow$ or $\downarrow$) is not obvious, then the specification is not clear.
5. If you need to measure something like "looks good" try transforming it into a testable measure such as "High score on 5-point attractiveness scale by $>65\%$ of passengers." This means that you set up a 5-point attractiveness scale (units = "points") such as 1 = ugly, 2 = tolerable, 3 = acceptable, 4 = attractive, 5 = captivating. Obviously the sense is ($\uparrow$). And the target (to be set in Step 7 will be ≥ 4).

Specifications for the aisle chair are shown in Fig. 6.6. Some comments about them in light of the guidelines are

1. The first specification "seat width relative to frame width" is not clear. What is to be measured here?
2. Two points about specifications that are in terms of "number of steps": (1) steps are better than time as time varies from individual to individual, and (2) you need to clearly define what a step is. A good guide for determining steps is in Section 11.5.
3. "Seat size" is not clear. What exactly needs to be measured?

6.7 STEP 6: RELATE CUSTOMERS' REQUIREMENTS TO ENGINEERING SPECIFICATIONS: HOW TO MEASURE WHAT?

To complete this step, we fill in the center portion of the house of quality. This relationship matrix is completed in parallel to Step 5, and it yields additional knowledge. Each cell of the form represents how an engineering specification



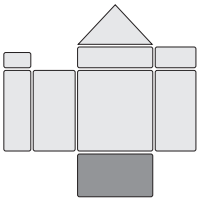
relates to a customer's requirement. Many specifications will measure more than one customer's requirement. The strength of this relationship can vary, with some engineering specifications, providing strong measures for a customer's requirement and others providing no measure at all. The relation is conveyed through specific symbols or numbers:

- = 9 = strong relationship
- = 3 = medium relationship
- △ = 1 = weak relationship
- Blank = 0 = no relationship at all

The 0-1-3-9 values are used to reflect the dominance of strong relationships. The symbols are used in the example (Fig. 6.6) and the number is used in the math that follows for the aisle chair.

Some guidelines for this step are as follows:

- Each customer's requirement should have at least one specification with a strong relationship.
- There is the temptation to make this a diagonal matrix of ●s or 9s—one engineering specification for each customer requirement. This is a weak use of the method. Ideally, each specification should measure more than one customer requirement.
- If a customer's requirement has only weak or medium relationships (see “Fits in aircraft aisle” or “good lifting position”), then it is not well understood or the specification has not been well thought through. It is evident what is meant by “fits in aircraft aisle.” The specification needs work. It is not so evident what “good lifting position” means and thus the customer's requirement needs more effort.



6.8 STEP 7: SET ENGINEERING SPECIFICATION TARGETS AND IMPORTANCE: HOW MUCH IS GOOD ENOUGH?

In this step we fill in the basement of the house of quality. Here we set the targets and establish how important it is to meet each of them. There are three parts to this effort, as shown in Fig. 6.6, calculate the specification importance, measure how well the competition meets the specification, and develop targets for your effort.

6.8.1 Specification Importance

The first goal in this step is determining the importance for each specification. If a target is important, then effort needs to be expended to meet the target. If it is not important, then meeting the goal can be more easily relaxed. In the development

of products, it is seldom that all targets can be met in the time available and so this effort helps guide what to work on. The method to find importance is as follows:

Step 2.1: For each customer multiply the importance weighting from step 3 with the 0-1-3-9 relationship values from step 6 to get the weighted values.

Step 2.2: Sum the weighted values for each specification. For specification “steps to adjust seat height” in Fig. 6.6, the passenger score is:

$$4*9+6*0+15*0+10*0+3*1+7*0+24*0+6*0+5*9+15*3+5*1 = 134.$$

Step 2.3: Normalize these sums across all specifications. The sum across all the specifications is 1475 so this specification has importance of $134/1475 = 9\%$.

Figure 6.6 shows the importance from both the passengers’ and agents’ viewpoints. Note that for the passenger specifications revolving around moving from their chair to the aisle chair are most important. From the agents’ viewpoint both these specifications and time measures are important.

6.8.2 Measuring How Well the Competition Meets the Specifications

In step 4, the competitions’ products were compared to customers’ requirements. In this step, they will be measured relative to engineering specifications. This ensures that both knowledge and equipment exist for evaluation of any new products developed in the project. Also, the values obtained by measuring the competition give a basis for establishing the targets. This usually means obtaining actual samples of the competition’s product and making measurements on them in the same way that measurements will be made on the product being designed. Sometimes this is not possible and literature or simulations are used to find values needed here.

The competition values are shown in Fig. 6.6.

6.8.3 Setting Specification Targets

Setting targets early in the design process is important; targets set near the end of the process are easy to meet but have no meaning as they always match what has been designed. However, setting targets too tightly may eliminate new ideas. Some companies refine their targets throughout concept development and then make them firm. The initial targets, set here, may have $\pm 30\%$ tolerance on them.

Most texts on QFD suggest that a single value be set as a target. However, once the design process is underway, often it is not possible to meet these exact values. In fact, a major part of engineering design is making decisions about how to manage targets and the tradeoff meeting them. There are two points to be made here. To make them, we will use a simple example.

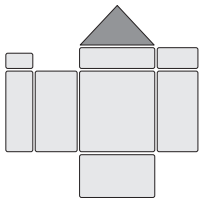
Say you want to buy a new camera. You want to spend less than \$300 and want at least 7.2 megapixels (your only two specifications). You look online and

find a camera with the resolution you want, but it costs \$305. Will you buy it? Probably. What if it costs \$315?—maybe. What about \$400?—probably not. The point here is that most targets are flexible and they may not be met during design. This is not true of all targets. You definitely need to achieve a velocity of 7 m/sec to escape the Earth’s gravitational pull. You cannot say 6.5 is good enough. For those targets that have flexibility, a more robust method for setting targets is to establish the levels at which the customers will be delighted and those where they will be disgusted. The delighted value is the actual target and the disgusted is the threshold beyond which the product is unacceptable. For the camera example, the target cost (delighted) is at \$300 and the threshold (disgusted) between \$315 and \$400, say \$350. For the resolution, delighted may be 7.2 megapixels and disgusted 6.3 megapixels. Note that for cost, less is better and for resolution, more is better.

A second point is that, as a design engineer you often have to trade off one specification against another. Continuing with the camera example, say there are two cameras available, one has 6.3 megapixels and costs \$305 and the other 7.2 megapixels and costs \$330. The question is, how much am I willing to trade off cost for resolution? If the targets were single valued, \$300 and 7.2 megapixels, then neither camera meets the targets. But, by setting the two targets, delighted and disgusted, you can better judge which camera is best.

A final comment on target setting is that if a target is much different than the values achieved by the competition, it should be questioned. Specifically, what do you know that the competition does not know? Do you have a new technology, do you know of new concepts, or are you just smarter than your competition? What is possible should fall in the range of the delighted and disgusted targets.

Figure 6.6 shows values for the aisle chair delighted and disgusted targets.



6.9 STEP 8: IDENTIFY RELATIONSHIPS BETWEEN ENGINEERING SPECIFICATIONS: HOW ARE THE HOWS DEPENDENT ON EACH OTHER?

Engineering specifications may be dependent on each other. It is best to realize these dependencies early in the design process. Thus, the roof is added to show that as you work to meet one specification, you may be having a positive or negative affect on others.

In Fig. 6.6, the roof for the aisle chair QFD shows diagonal lines connecting the engineering specifications. If two specifications are dependent, a symbol is noted in the intersection. There are many different styles of symbols used. One is to use the same symbols as in Step 6. The simplest method is to use a “+” to denote that improvement in meeting one of the specifications will improve the other (they are synergistic), and to use a “−” to show that improvement in meeting one may harm the other (a compromise may be forced). Some people use ++ and −− to show a strong dependency.

In building a house of quality on a spreadsheet, a good way to simulate the roof is as shown in Fig. 6.7. Here the specifications are listed in both the

for the passenger to slide” and “force required by the agent.” The lack of clarity is caused by a poor understanding of exactly what force the agent is applying.

6.10 FURTHER COMMENTS ON QFD

The QFD technique ensures that the problem is well understood. It is useful with all types of design problems and results in a clear set of customers’ requirements and associated engineering measures. It may appear to slow the design process, but in actuality it does not, as time spent developing information now is returned in time saved later in the process.

Even though this technique is presented as a method for understanding the design requirements, it forces such in-depth thinking about the problem that many good design solutions develop from it. No matter how hard we try to stay focused on the requirements for the product, product concepts are invariably generated. This is one situation when a design notebook is important. Ideas recorded as brief notes or sketches during the problem understanding phase may be useful later; however, it is important not to lose sight of the goals of the technique and drift off to one favorite design idea.

The QFD technique automatically documents this phase of the design process. Diagrams like those in Figs. 6.5 and 6.6 serve as a design record and also make an excellent communication tool. Specifically, the structure of the house of quality makes explaining this phase to others very easy. In one project, a member of the sponsoring organization was blind. A verbal description of the structure helped him understand the project and recommend the QFD method to other sighted colleagues.

Often, when working to understand and develop a clear set of requirements for the problem, the design team will realize that the problem can be decomposed into a set of loosely related subproblems, each of which may be treated as an individual design problem. Thus, a number of independent houses may be developed.

The QFD technique can also be applied during later phases of the design process. Instead of developing customers’ requirements, we may use it to develop a better measure for functions, assemblies, or components in terms of cost, failure modes, or other characteristics. To accomplish this, review the steps, replacing customers’ requirements with what is to be measured and engineering requirements with any other measuring criteria.

Although QFD seems to imply a waterfall-type development plan, much learning occurs during the design process. The QFD is considered a working document that is reviewed and updated as needed. Thus, it also is important for spirally developed products. The formality and complexity of the technique forces any change to be carefully considered and thus keeps the project moving toward completion. Without a system like QFD, changes in specifications can occur at the whim of a manager or without the design team even realizing it. These changes will lead to a failure to meet the schedule and a potentially poor product.

6.11 SUMMARY

- Understanding the design problem is best accomplished through a technique called Quality Function Deployment (QFD). This method transforms customers' requirements into targets for measurable engineering requirements.
- Important information to be developed at the beginning of the problem includes customers' requirements, competition benchmarks, and engineering specifications complete with measurable benchmarks.
- Time spent completing the QFD is more than recovered later in the design process.
- There are many customers for most design problems.
- Studying the competition during problem understanding gives valuable insight into market opportunities and reasonable targets.

6.12 SOURCES

ANSI standards are available at www.ansi.org

ASTM standards are available at www.astm.org

Cristiano, J. J., J. K. Liker, and C. C. White: "An Investigation into Quality Function Deployment (QFD) Usage in the U.S.," in *Transactions for the 7th Symposium on Quality Function Deployment*, June 1995, American Supplier Institute, Detroit. Statistics on QFD usage were taken from the study in this paper.

Hauser, J. R., and D. Clausing: "The House of Quality," *Harvard Business Review*, May–June 1988, pp. 63–73. A basic paper on the QFD technique.

Index of Federal Specifications and Standards, U.S. Government Printing Office, Washington, D.C. A sourcebook for federal standards.

Krueger, R. A.: *Focus Groups: A Practical Guide for Applied Research*, Sage Publishing, Newbury Park, Calif. 1988. A small book with direct help for getting good information from focus groups.

Roberts, V. L.: *Products Standards Index*, Pergamon, New York, 1986. A sourcebook for standards.

Salant, P., and D. Dillman: *How to Conduct Your Own Survey*, John Wiley & Sons, New York, 1994. A very complete book on how to do surveys to collect opinions.

Software packages

QFD/CAPTURE, <http://www.qfdcapture.com/default.asp>

QFD Designer, IDEACore, <http://www.ideacore.com/v1/Products/QFDDesigner/>

Templates for Excel are at <http://www.qfdonline.com/templates/>

6.13 EXERCISES

- 6.1 For a design problem (Exercise 4.1), develop a house of quality and supporting information for it. This must include the results of each step developed in this chapter. Make sure you have at least three types of customers and three benchmarks. Also, make a list of the ideas for your product that were generated during this exercise.

- 6.2** For the features of the redesign problem (Exercise 4.2) to be changed, develop a QFD matrix to assist in developing the engineering specifications. Use the current design as a benchmark. Are there other benchmarks? Be careful to identify the features needing change before spending too much time on this. The methods in Chap. 7 can be used iteratively to help refine the problem.
- 6.3** Develop a house of quality for these objects.
- The controls on an electric mixer.
 - A seat for an all-terrain bicycle.
 - An attachment for electric drills to cut equilateral-triangle holes in wood. The wood can be up to 50 mm thick, and the holes must be adjustable from 20 mm to 60 mm per side.
 - A tamper-proof fastener as used in public toilet facilities.

6.14 ON THE WEB



A template for the following document is available on the book's website: www.mhhe.com/Ullman4e

- Voice of the Customer

Concept Generation

KEY QUESTIONS

- How can understanding the function help developing form?
- What does flow have to do with function?
- How can patents help generate ideas?
- How can you get the best out of brainstorming and brainwriting?
- How do contradictions lead to new ideas?
- What is a morphology and what does it do?

7.1 INTRODUCTION

In Chap. 6, we went to great lengths to understand the design problem and to develop its specifications and requirements. Now our goal is to use this understanding as a basis for generating concepts that will lead to a quality product. In doing this, we apply a simple philosophy: *Form follows function*. Thus we must first understand the function of a device, before we design its form. Conceptual design focuses on function.

A concept is an idea that is sufficiently developed to evaluate the physical principles that govern its behavior. Confirming that a concept will operate as anticipated and that, with reasonable further development, it will meet the targets set, is a primary goal in concept development. Concepts must also be refined enough to evaluate the technologies needed to realize them, to evaluate their basic architecture (i.e., form), and, to some limited degree, to evaluate their manufacturability. Concepts can be represented in a rough sketch or flow diagram, a proof-of-concept prototype, a set of calculations, or textual notes—an abstraction of what might someday be a product. However a concept is represented, the key point is that enough detail must be developed to model performance so that the functionality of the idea can be ensured.

On the average, industry spends about 15% of design time developing concepts. Based on a comparison of the companies in Fig. 1.5, this should be 20–25%

If you generate one idea, it is probably a poor one. If you generate twenty ideas, you may have a good one.

Or, alternatively

He who spends too much time developing a single concept realizes only that concept.

to minimize changes later. In some companies, however, design begins with a concept to be developed into a product without working to understand the requirements. This is a weak philosophy and generally does not lead to quality products.

Some concepts are naturally generated during the engineering requirements development phase. Since in order to understand the problem, we have to associate it with things we already know (see Chap. 3), there is a great tendency for designers to take their first idea and start to refine it toward a product. This is also a weak methodology best expressed by the aphorisms above. This statement and the methods in this chapter support one of the key features of engineering design: generate multiple concepts. The main goal of this chapter, then, is to present techniques for the generation of many concepts.

The flow of conceptual design is shown in Fig. 7.1. Here, as with all problem solving, the generation of concepts is iterative with their evaluation. Also part of Conceptual Design, as shown in the figure, is the communication of design information and the updating of the plans.

In line with our basic philosophy, the techniques we will look at here for generating design concepts encourage the consideration of the function of the device being designed. These techniques aid in decomposing the problem in a way that affords the greatest understanding of it and the greatest opportunity for creative solutions to it.

We will focus on techniques to help with *functional decomposition* and *concept variant generation* because these important customer requirements are concerned with the functional performance desired in the product. These requirements become the basis for the concept generation techniques. Functional decomposition is designed to further refine the functional requirements; concept variant generation aids in transforming the functions to concepts.

Once the function is understood, there are many methods to help generate concepts to satisfy them. *Concepts are the means for providing function.* Concepts can be represented as verbal or textual descriptions, sketches, paper models, block diagrams, or any other form that gives an indication of how the function can be achieved.

These techniques support a divergent-convergent design philosophy. This philosophy expands a design problem into many solutions before it is narrowed to one final solution. Before continuing, note that the techniques presented here are useful during the development of an entire system and also for each subsystem,

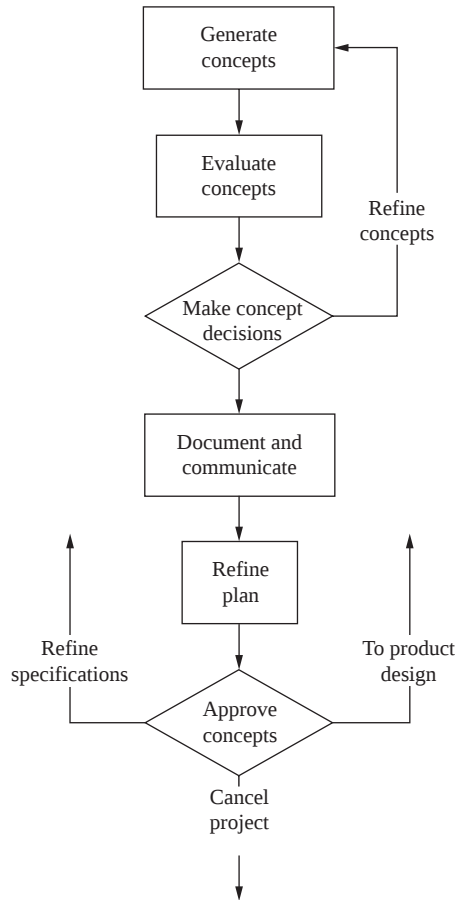


Figure 7.1 The Conceptual Design phase of the design process.

component, and feature. This is not to say that the level of detail presented here needs to be undertaken for each flange, rib, or other detail; however, it helps in thinking about all features and it is especially useful for difficult features.

One example used in this chapter is the redesign of a one-handed bar clamp by Irwin Corporation. This style of clamp was introduced in Chap. 2. By 2004 the Irwin Corp. had sold over \$25 million worth of Quick-Grip one-handed bar clamps. At that time, they decided to develop a new model. A one-handed bar clamp is a simple mechanical device, and although simple, understanding its evolution is very instructive. The following paragraphs describe its early development and the basic theory of operation. The 2004 redesign of the Quick-Grip is used as the basis for an example in the rest of the chapter.

In November 1986, a freelance artist was building an airboat to run on the Platt River in Nebraska. He found he needed a third hand to hold parts together during gluing as he had to hold parts together with one hand and use two hands to apply a clamp. In thinking about how to either grow another hand or work a clamp with one hand, his thoughts went to the common caulking gun (Fig. 7.2). Caulking guns work with one hand. Each time you squeeze the trigger; the rod moves farther into the tube (how energy is transferred from the trigger to the rod will be addressed later). On the end of the rod, a flat disk pushes on a plastic plunger in the tube of caulking, pushing some of the caulking out of the nozzle. What is important here is that when the trigger is fully compressed and the handgrip relaxed, a spring brings the trigger back to its fully extended position, but the rod stays where it was. Holding the rod in position is a jam plate that locks the rod from moving back. (We will explore how this works in a moment.) A jam plate can be clearly seen in Fig. 7.3, the artist's first prototype of the one-handed bar clamp. This prototype was made of some scrap aluminum, pop rivets, and parts from a caulking gun. His idea worked so well he presented his idea to the

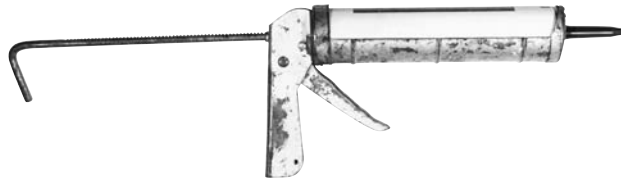


Figure 7.2 A common caulking gun. (Courtesy Arthur S. Aubry/Getty Images.)

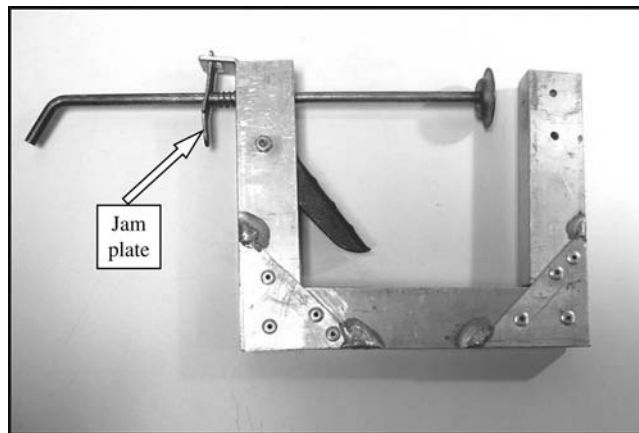


Figure 7.3 The first prototype of a one-handed bar clamp. (Reprinted with permission of Irwin Industrial Tools.)

Most of your best ideas wind up being useless in the final design. Learn to live with the disappointment and take joy in the successes.

American Tool Company. They entered into an agreement with the inventor, hired him, and by March 1989, the sixth prototype looked very much like the product shown in Fig. 7.4. In 2002 Newell Rubbermaid acquired American Tools and changed its name to Irwin.

The operation of all of one-handed clamps is dependent on the use of a jam plate. Figure 7.5 shows a simple schematic of a jam plate with a rectangular rod and a detail of the first prototype showing the jam plate in use. On the prototype, the spring on the rod works to keep the plate in position when not loaded, as will become clear. The operation of this mechanism is due to the height of the hole in the plate, h_p , being slightly more than the height of the rod, h_b . This allows the plate to tilt, $\Theta = 5-10^\circ$, and jam the rod from moving to the left.

On many caulking guns and one-handed clamps there are two jam plates, one for locking the bar in position, as in the diagram, and a second one tilted the other way with the pivot attached to the trigger. Each time the trigger is squeezed, the second plate jams the bar as the trigger is moved. During this motion, the locking jam plate un-tilts sufficiently to allow the bar to move freely and jams when the trigger is released.

This basic introduction to the history and operation of the one-handed clamp will be used later in the chapter.

Before continuing, note that this chapter encourages the development of many ideas. Do be aware that developing ideas is, on one hand, very fulfilling, and on the other hand, disappointing. It is fulfilling in that giving birth to an idea is something that is uniquely your own and you can feel pride and pleasure in being



Figure 7.4 The Irwin Quick-Grip introduced in March 1989.
(Reprinted with permission of Irwin Industrial Tools.)

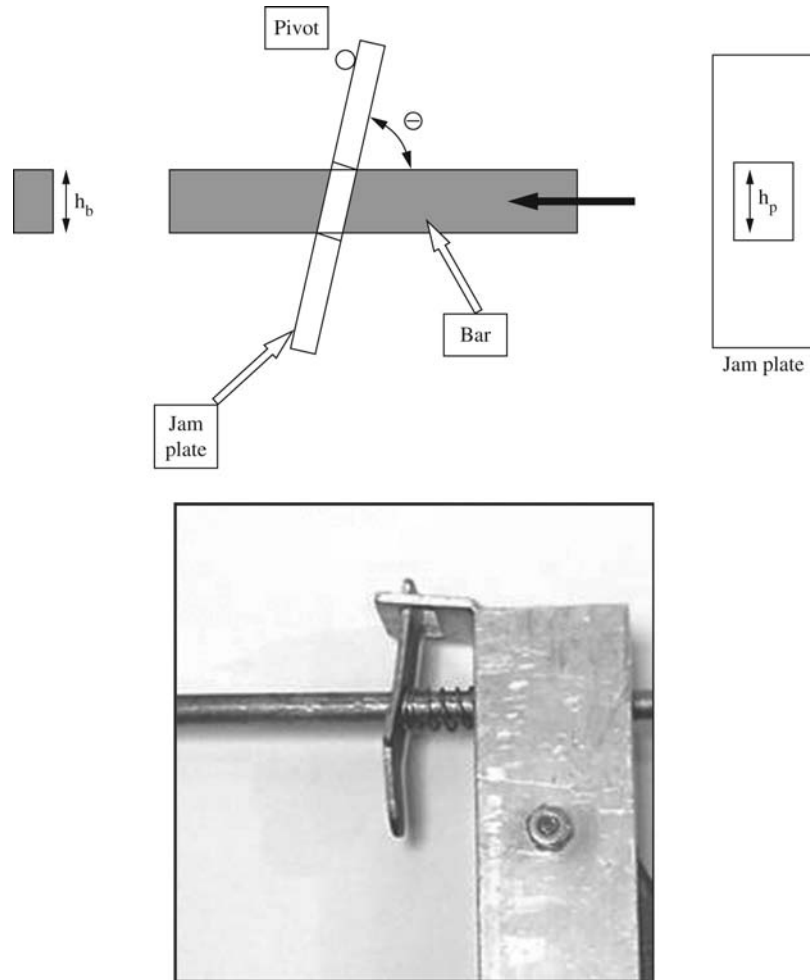


Figure 7.5 Details about how jam plates work. (Reprinted with permission of Irwin Industrial Tools.)

a part of its evolution. However, most ideas never make it to the product stage, as they don't really work, are too complex, or there isn't enough time or money to develop them.

7.2 UNDERSTANDING THE FUNCTION OF EXISTING DEVICES

This section begins with a general discussion of the term “function.” It then focuses on how to decompose existing devices to find their function. We then

turn our attention to the understanding of the function of proposed devices, those described in patents.

7.2.1 Defining “Function”

In reading this section, it is important to remember that *function* tells *what* the product must do, whereas its *form*, or *structure*, conveys *how* the product will do it. The effort in this chapter is to develop the *what* and then map the *how*. This is similar to the QFD in Chap. 6, where *what* the customer required was mapped into *how* the requirements were to be measured. Here we focus on *what* the product must do (its function) and then on *how* to do it (its form).

Function is the logical flow of energy (including static forces), material, or information between objects or the change of state of an object caused by one or more of the flows. For example, in order to attach any component to another, a person must *grasp* the component, *position* it, and *attach* it in place. These functions must be completed in a logical order: grasp, position, and then attach. In undertaking these actions, the human provides information and energy in controlling the movement of the component and in applying force to it. The three flows—energy, material, and information—are rarely independent of each other. For instance, the control and the energy supplied by the human cannot be separated. However, it is important to note that both are occurring and that both are supplied by the human to the component.

The functions associated with the flow of energy can be classified both by the type of energy and by its action in the system. The types of energy normally identified with electromechanical systems are mechanical, electrical, fluid, and thermal. As these types of energies flow through the system, they are *transformed*, *stored*, *transferred (conducted)*, *supplied*, and *dissipated*. These are the “actions” of the components or assemblies in the system. Thus, all terms used to describe the flow of energy are action words; this is characteristic of all descriptions of function. Also, part of the flow of energy is the flow of forces even when they do not result in motion. This concern for force flows is further developed in Section 9.3.4.

The functions associated with the flow of materials can be divided into three main types. *Through-flow*, or material-conserving processes is the first. Material is manipulated to change its position or shape. Some terms normally associated with through-flow are *position*, *lift*, *hold*, *support*, *move*, *translate*, *rotate*, and *guide*. The second type is *diverging flow*, or dividing the material into two or more bodies. Terms that describe diverging flow are *disassemble* and *separate*. *Converging flow*, or assembling or joining materials, is the third. Terms that describe converging flow are *mix*, *attach*, and *position relative to*.

The functions associated with information flow can be in the form of mechanical signals, electrical signals, or software. Generally, the information is used as part of an automatic control system or to interface with a human operator. For example, if you install a component with screws, after you tighten the screws you wiggle the component to see if it is really attached. Effectively you ask the

Function happens primarily at interfaces.

question, Is the component attached? and the simple test confirms that it is. This is a common type of information flow. Software is used to modify information that flows through an electronic circuit—a computer chip—designed to be controlled by the code. Thus, electrical signals transport information to and from the chip and the software transforms the information.

Function can also relate the change of state of an object. If I say that a spring stores energy, then the internal state of stress in the spring is changed from its initial state. The energy that is stored was transferred to (i.e., flowed into) the spring from some other object. Typically, state changes that are important in mechanical design describe transformations of potential or kinetic energy, material properties, form (e.g., shape, configuration, or relative position), or information content.

With this basic understanding of function, we can describe a useful method for reverse engineering an existing product.

7.2.2 Using Reverse Engineering to Understand the Function of Existing Devices

Reverse engineering is a method to understand how a product works. Whereas we used product decomposition in Chap. 2 to understand a product's parts and assemblies, here we will focus on their function. In Chap. 2 we disassembled an Irwin Quick-Grip clamp (Fig. 7.4) and itemized the parts and how they were assembled. Here we will extend this decomposition to understand the function of the clamp—to reverse engineer it. This is more than just taking stuff apart, it is a key part of understanding how others solved the problem.

Reverse Engineering, functional decomposition, or benchmarking is a good practice because many hundreds of engineering hours have been spent developing the features of existing products, and to ignore this work is foolish. The QFD method, featured in Chap. 6, encourages the study of existing products as a basis for finding market opportunities and setting specification targets. Some organizations do not pay attention to products not developed within their walls—a very weak policy. These companies are said to have a case of “NIH” (i.e., Not Invented Here). Dissecting and reverse engineering the products of others helps overcome this policy.

It is a natural tendency to want to understand how things work. Sometimes the operation is obvious and sometimes it is very obscure. The methodology described next is designed to help understand an existing piece of hardware. The primary goal is to find out how the device works—What is its function?

To make sure that the function of a device is understood these steps are suggested. They can be integrated with decomposition or follow on from it. Here it is assumed that the clamp has already been decomposed and the parts named, as in Fig. 2.11.

Step 1: For the Whole Device, Examine Interfaces with Other Objects. Since the function of a device is defined by its effect on the flow of energy, information, and material, a starting place is to examine these flows into and out of the device being examined. Consider the Irwin Quick-Grip clamp shown in Fig. 7.4. Before reading on, identify the energy, information, and material that flow into and out of the clamp.

Energy, information, and materials flow through the clamp. The energy into the clamp is from the *user's hand* squeezing on the hand grip molded into the main body and the trigger and the *parts being clamped* pushing back on the pads that make up the jaw of the clamp. The information flow is back to the *user* to tell her when to stop squeezing. In other words, the user is continuously asking the question "Is the clamp force high enough?" The increase in handgrip force needed to squeeze the parts being clamped plus any change in the look or sound (e.g., something being crushed) answer that question. Finally, even though it does not look like any material is "flowing," it is useful to consider the parts being clamped as material flowing into the clamp and back out again. This forces you to think about the process of aligning the clamp jaw with the work, clamping them, and then removing the parts from the jaws when finished.

There is a second energy flow when the user releases the clamp. We will not explore that here.

Step 2: Remove a Component for More Detailed Study. Remove a single component or an assembly from the device. Note carefully how it was fastened to the rest of the device. Also note any relationships it has to other parts that it may not contact. For example, it may have to have a clearance with some other parts in order to function. It may have to shield other assemblies from view, light, or radiation. It may have to guide some fluid. In fact, the part removed from the assembly may be a fluid, for example, consider the water flowing through a valve in order to study the function of the valve on the water. This step is similar to what was done during product decomposition.

For the clamp, we will focus on the trigger. After you remove the faceplate, you can see the trigger and other internal parts (Fig. 7.6). The part names from the decomposition have been added to the photo in the figure. Now remove the trigger for detailed study. In general, when removing a component for study, note every other part it was in contact with or has to clear (i.e., its interfaces) in order to function. The trigger interfaces with the user, the main body, and the first jam plate, and it has to clear the bar and the faceplate that was removed.

Step 3: Examine Each Interface to Find the Flow of Energy, Information, or Materials. The goal here is to really understand how the functions identified in step 1 are transformed by the device. Additionally, we want to understand how the parts are fastened together, how forces are transformed and flow from one component to another, and the purpose for each component feature.

In looking at each connection, remember that forces may be transferred between components in three directions (x , y , z) and moments transferred about

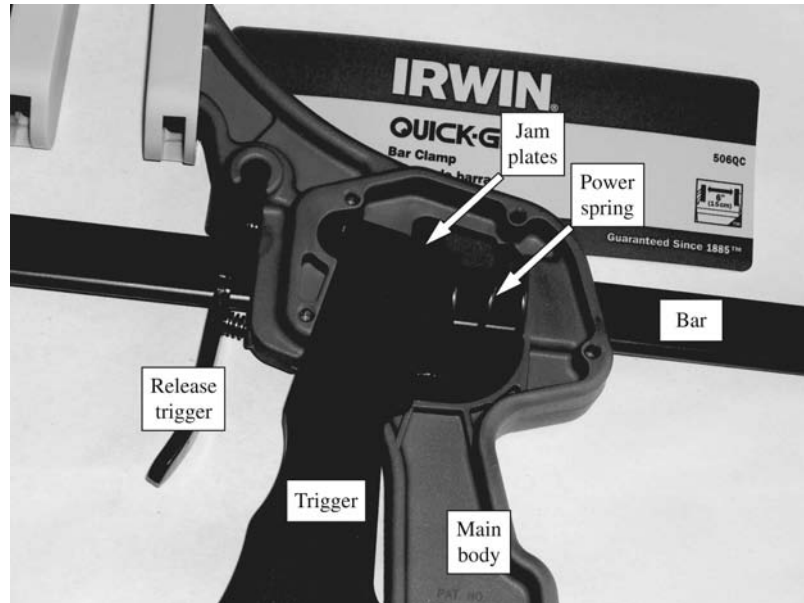


Figure 7.6 The internal parts of the Quick-Grip. (Reprinted with permission of Irwin Industrial Tools.)

three axes. Further, there should be features of each interface that either give a degree of freedom to the force or moment or restrains it.

For the clamp trigger there are three interfaces with other components and the outside world, as shown on the drawing in Fig. 7.7 and the Reverse Engineering Template, Fig. 7.8:

1. The interface between the user's hand and the grip surface, 1a. This force is balanced by the force on the main body, 1b. Energy flows here as described in step 1.
2. The interface to the pivot limits the trigger motion to one degree of freedom—rotation about the circular pivot surface (a virtual axle). Energy flows here as a reaction to clamping force described in item 3. This reaction force is labeled “3” in the figure.
3. The interface to the jam plate. Energy flows between the trigger and the jam plate (2). Moving the jam plate pulls on the bar, closing the jaws and applying a force to the material being clamped.

Also, shown in Fig. 7.7 is the main body. Not including the forces from the release trigger, there are six interfaces as shown. By studying each of these interfaces, the operation of the main body and the design details of it can be understood.

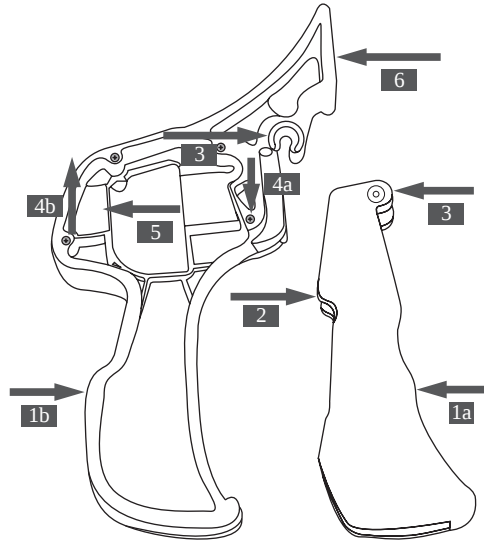


Figure 7.7 Forces on the Quick-Grip main body and trigger.

7.3 A TECHNIQUE FOR DESIGNING WITH FUNCTION

The goal of functional modeling is to decompose the problem in terms of the flow of energy, material, and information. This forces a detailed understanding at the beginning of the design project of *what* the product-to-be is to do. The functional decomposition technique is very useful in the development of new products.

There are four basic steps in applying the technique and several guidelines for successful decomposition. These steps are used iteratively and can be reordered as needed. This technique can be used with QFD to help understand the problem. In this discussion, the usefulness of the technique will be demonstrated with the one-handed bar clamp and with the GE X-ray CT Scanner introduced in Chap 4.

7.3.1 Step 1: Find the Overall Function That Needs to Be Accomplished

This is a good first step toward understanding the function. The goal here is to generate a single statement of the overall function on the basis of the customer requirements. All design problems have one or two “most important” functions. These must be reduced to a simple clause and put in a *black box*. The inputs to this box are all the energy, material, and information that flow into the boundary of the system. The outputs are what flows out of the system.



Reverse Engineering for Function Understanding

Design Organization: Example for the Mechanical Design Process

Date: Dec. 20, 2007

Product Decomposed: Irwin Quick Grip—Pre 2007

Description: This is the Quick-Grip product that has been on the market for many years.

How it works: Squeeze the pistol grip repeatedly to move the jaws closer together and increase the clamping force. Squeeze the release trigger to release the clamping force. The foot (the part on the left in the picture that holds the face that is clamped against) is reversible so the clamping force can be made to push apart rather than squeeze together.

Interfaces with other objects:

Part #	Part Name	Other Object	Energy Flow	Information Flow	Material Flow
1 & 2	Main body and Trigger	User's hand	User squeezes trigger to move jaws closer together and	Squeezing force proportional to jaw force	User's hand grips and releases
8	Pad	Parts being clamped	Clamping force and compressive motion of jaws moving together	None	Parts flow into and out of jaws
Etc.					

Flow of energy, information, and materials:

Part #	Part Name	Interface Part #	Flow of Energy, Information, and Material	Image
1	Trigger	User	Force 1a applied by gripping trigger and main body. Resistance force felt by user proportional to clamping force.	
2	Trigger	1—Main body	Force 3 at pivot—reaction force	
3	Trigger	14—Jam plate	Force 2 pushes on the jam plate to ultimately make the bar move and apply the clamping force.	
4	Etc.			

Links and drawing files:

Team member:	Prepared by:
Team member:	Checked by:
Team member:	Approved by:
Team member:	

The Mechanical Design Process

Designed by Professor David G. Ullman

Copyright 2008, McGraw-Hill

Form # 1.0

Figure 7.8 Reverse Engineering Template sample.

Some guidelines for step 1 are:

Guideline: Energy Must Be Conserved. Whatever energy goes into the system must come out or be stored in the system.

Guideline: Material Must Be Conserved. Materials that pass through the system boundary must, like energy, be conserved.

Guideline: All Interfacing Objects and Known, Fixed Parts of the System Must Be Identified. It is important to list all the objects that interact, or interface, with the system. Objects include all features, components, assemblies, humans, or elements of nature that exchange energy, material, or information with the system being designed. These objects may also constrain the system's size, shape, weight, color, and the like. Further, some objects are part of the system being designed that cannot be changed or modified. These too must be listed at the beginning of the design process.

Guideline: Ask the Question, How Will the Customer Know if the System Is Performing? Answers to this question will help identify information flows that are important.

Guideline: Use Action Verbs to Convey Flow. Action verbs such as those in Table 7.1 can be used to describe function. Obviously, many other verbs beyond those listed tell about the intended action.

Finding the Overall Function: The One-Handed Bar Clamp

For the one-handed bar clamp, the “most important” function is very simple “transform the grip force of one hand to a controllable force capable of clamping common objects together” (Fig. 7.9). This statement is brief, it tells that the goal is to alter the energy flow while sensing the force applied, and that the boundaries of the system are the one hand and the objects being clamped.

Finding the Overall Function: The X-Ray CT Scanner

For the CT Scanner shown in Fig. 7.10 (taken from Fig. 4.2), the top-level function is “convert electrical energy into an image of the organs of a patient.”

Table 7.1 Typical mechanical design functions

Absorb/remove	Dissipate	Release
Actuate	Drive	Rectify
Amplify	Hold or fasten	Rotate
Assemble/disassemble	Increase/decrease	Secure
Change	Interrupt	Shield
Channel or guide	Join/separate	Start/stop
Clear or avoid	Lift	Steer
Collect	Limit	Store
Conduct	Locate	Supply
Control	Move	Support
Convert	Orient	Transform
Couple/interrupt	Position	Translate
Direct	Protect	Verify



Figure 7.9 Top-level function for the one-handed bar clamp.

This statement assumes the boundary considered is the entire CT Scanner and the computer and software that make the image. We could draw the boundary tighter, just around the device shown in the figure, and say “convert electrical energy into a signal that contains information about an image of the organs of a patient.” The difference is small, but indicates the change in boundary.

7.3.2 Step 2: Create Subfunction Descriptions

The goal of this step and step 3 is to decompose the overall function. This step focuses on identifying the subfunctions needed, and the next step concerns their organization.



Figure 7.10 A GE CT Scanner. (Reprinted with permission of GE Medical.)

There are three reasons for decomposing the overall function: First, the resulting decomposition controls the search for solutions to the design problem. Since concepts follow function and products follow concepts, we must fully understand the function before wasting time generating products that solve the wrong problem.

Second, the division into finer functional detail leads to a better understanding of the design problem. Although all this detail work sounds counter to creativity, most good ideas come from fully understanding the functional needs of the design problem. Since it improves understanding, it is useful to begin this process before the QFD process in Chap. 6 is complete and use the functional development to help determine the engineering specifications.

Finally, breaking down the functions of the design may lead to the realization that there are some already existing components that can provide some of the functionality required.

Each subfunction developed will show either

- An object whose state has changed
- or
- An object that has energy, material, or information transferred to it from another object.

The following guidelines are important in accomplishing the decomposition. It will take several iterations to finalize all this information. However, time spent here will save time later when it is realized that the product has intended functions that could have been found and dealt with much earlier. The examples at the end of step 3 will demonstrate the use of the guidelines.

Guideline: Consider *What*, Not *How*. It is imperative that only *what* needs to happen—the function—be considered. Detailed, structure-oriented *how* considerations should be documented for later use as they add detail too soon here. Even though we remember functions by their physical embodiments, it is important that we try to abstract this information. If, in a specific problem solution, it is not possible to proceed without some basic assumptions about the form or structure of the device, then document the assumptions.

Guideline: Use Only Objects Described in the Problem Specification or Overall Function. To ensure that new components do not creep into the product unintentionally, use only nouns previously used (e.g., in the QFD or in step 1) to describe the material flow or interfacing objects. If any other nouns are used during this step, either something is missing in the first step (go back to step 1 and reformulate the overall function), the specifications are incomplete, or a design decision to add another object to the system has been made (consider very carefully). Adding objects is not bad as long as it is done consciously.

Guideline: Break the Function Down as Finely as Possible. This is best done by starting with the overall function of the design and breaking it into the separate functions. Let each function represent a change or transformation in the flow of material, energy, or information. Action verbs often used in this activity are given in Table 7.1.

Guideline: Consider All Operational Sequences. A product may have more than one operating sequence while in use (see Fig. 1.7). The functions of the device may be different during each of these. Additionally, prior to the actual *use* there may be some *preparation* that must be modeled, and similarly, after use there may be some *conclusion*. It is often effective to think of each function in terms of its preparation, use, and conclusion.

Guideline: Use Standard Notation When Possible. For some types of systems, there are well-established methods for building functional block diagrams. Common notation schemes exist for electrical circuits and piping systems, and block diagrams are used to represent transfer functions in system dynamics and control. Use these notation schemes if possible. However, there is no standard notation for general mechanical product design.

7.3.3 Step 3: Order the Subfunctions

The goal is to add order to the functions generated in the previous step. For many redesign problems, this occurs simultaneously with their identification in step 2, but for some material processing systems this is a major step. The goal here is to order the functions found in step 2 to accomplish the overall function in step 1. The guidelines and examples presented next should help with this step.

Guideline: The Flows Must Be in Logical or Temporal Order. The operation of the system being designed must happen in a logical manner or in a time sequence. This sequence can be determined by rearranging the subfunctions. First, arrange them in independent groups (preparation, uses, and conclusion). Then arrange them within each group so that the output of one function is the input of another. This helps complete the understanding of the flows and helps find missing functions.

Guideline: Redundant Functions Must Be Identified and Combined. Often there are many ways to state the same function. If each member of the design team has written his or her subfunctions on self-stick removable notepaper, all the pieces can be put on the wall and grouped by similarity. Those that are similar need to be combined into one subfunction.

Guideline: Functions Not Within the System Boundary Must Be Eliminated. This step helps the team come to mutual agreement on the exact system boundaries; it is often not as simple as it sounds.

Guideline: Energy and Material Must Be Conserved as They Flow Through the System. Match inputs and outputs to the functional decomposition.

Inputs to each function must match the outputs of the previous function. The inputs and outputs represent energy, material, or information. Thus, the flow between functions conveys the energy, material, or information without change or transformation.

Creating a Subfunction Description: The Irwin Quick-Grip Example

A functional decomposition for the one-handed bar clamp is shown in Fig. 7.11. Keep in mind when studying this figure that there is no one right way to do a functional decomposition and that the main reason for doing it is to ensure that the function of the device to be developed is understood. Note that each function statement begins with an action verb from the list in Table 7.1 and then follows with a noun. The boxes are oriented in a logical fashion. Also, note that in this example, the main flow is energy, but there is an information feedback to the user. Would a clamp be as useful, if there were no feedback?

Many functions on this diagram can be further refined. Not shown in the diagram is the release of any locking mechanism, a further refinement of the “hold force on object” box.

Creating Subfunction Description: The CT Scanner

The CT Scanner is a complex device. The functional diagram fills many pages. A partially completed segment, focusing on the X-ray tube, is shown in Fig. 7.12. Here, the function “Convert electrical power to X-rays” is shown

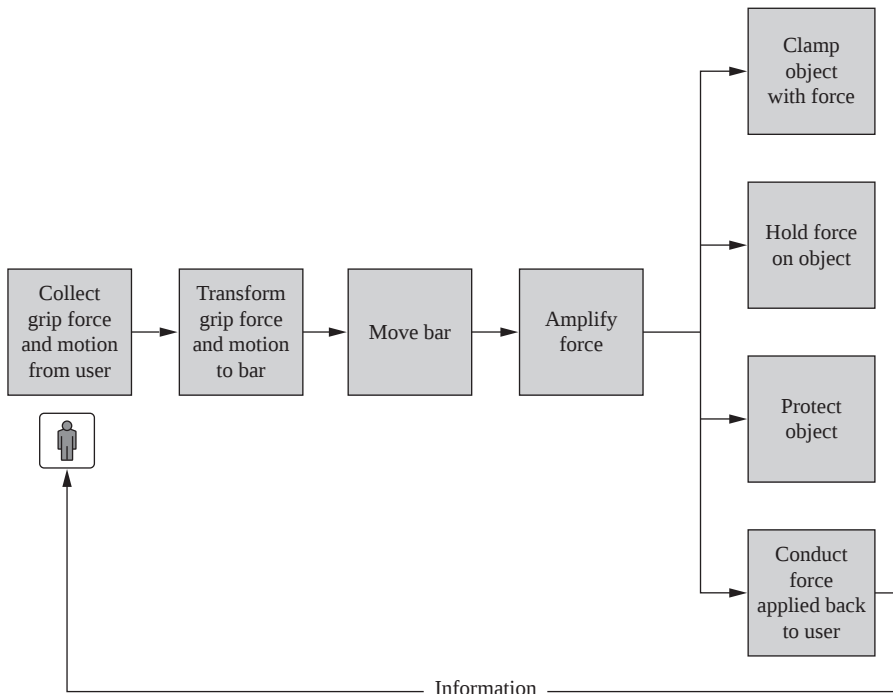


Figure 7.11 Functional decomposition for the one-handed bar clamp.

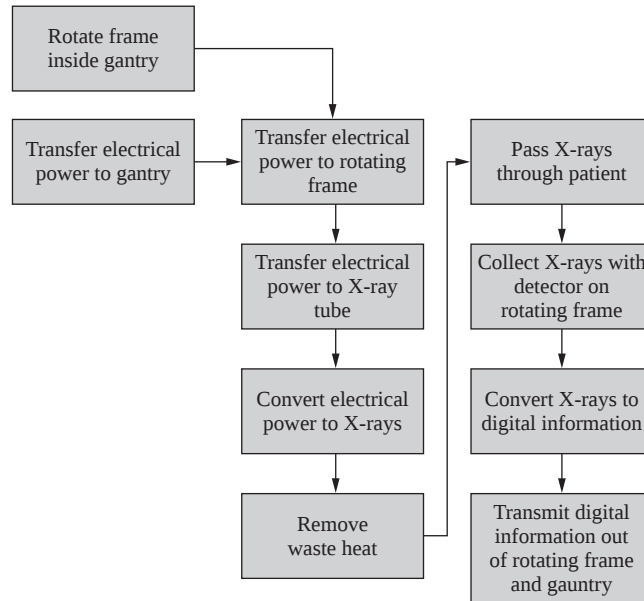


Figure 7.12 Functional decomposition of the CT Scanner.

with many subfunctions yet to be organized. Many of the functions are focused on the transformation of electrical energy. One of them, “Remove waste heat” is especially difficult as only about 1% of the energy is actually converted into X-rays, the other 60+ kW of energy is transformed into waste heat. The removal of this waste heat will be revisited in Chap. 10.

7.3.4 Step 4: Refine Subfunctions

The goal is to decompose the subfunction structure as finely as possible. This means examining each subfunction to see if it can be further divided into sub-subfunctions. This decomposition is continued until one of two things happens: “atomic” functions are developed or new objects are needed for further refinement. The term atomic implies that the function can be fulfilled by existing objects. However, if new objects are needed, then you want to stop refining because new objects require commitment to how the function will be achieved, not refinement of what the function is to be. Each noun used represents an object or a feature of an object.

Further Refining the Subfunctions: The CT Scanner

The function “Convert electrical power to X-rays” can be further decomposed as shown in Fig. 7.13. The function “Maintain vacuum” is shown surrounding all the other functions, as they all must operate within it. The flow of energy in this diagram includes electricity, ions, X-rays, heat, force, and torque.

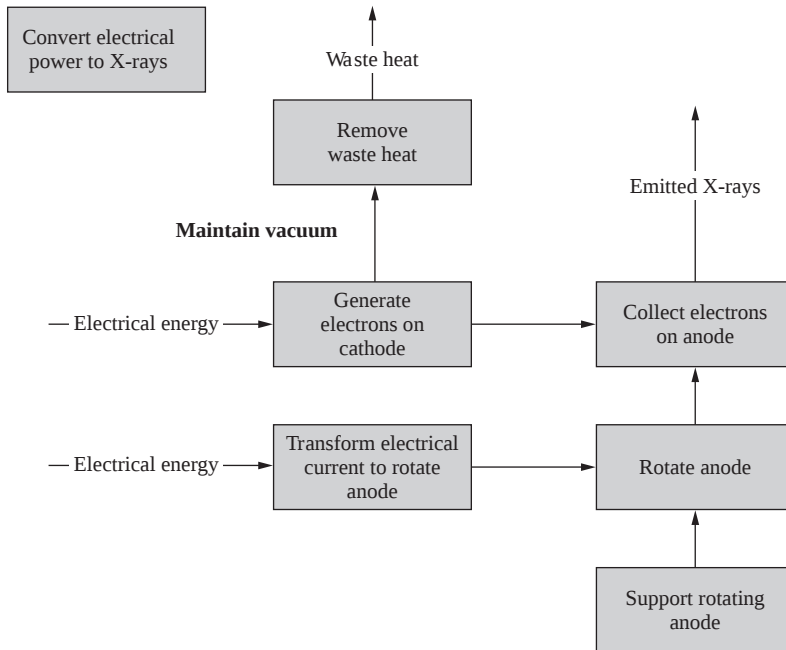


Figure 7.13 Refined functional decomposition for the conversion of electrical power to X-rays.

It must be realized that the function decomposition cannot be generated in one pass and that it is a struggle to develop the suggested diagrams. However, it is a fact that the design can be only as good as the understanding of the functions required by the problem. This exercise is both the first step in developing ideas for solutions and another step in understanding the problem. The functional decomposition diagrams are intended to be updated and refined as the design progresses.

A second goal in refining the functions is to group them. By grouping the functions, chunks of system logic can be isolated and used as building blocks for variant products.

What is important about this four-step decomposition is that concepts must be generated to meet all the functional needs identified. As you read the rest of this chapter note that the methods presented can be focused on entire devices, on collections of subfunctions, or on a single subfunction.

7.4 BASIC METHODS OF GENERATING CONCEPTS

The methods in this section are commonly used to develop concepts. As will be seen, they are based on knowledge of the functions. The methods are presented in

no particular order and can be used together. An experienced designer will jump from one to another to solve a specific problem.

7.4.1 Brainstorming as a Source of Ideas

Brainstorming, initially developed as a group-oriented technique, can also be used by an individual designer. What makes brainstorming especially good for group efforts is that each member of the group contributes ideas from his or her own viewpoint. The rules for brainstorming are quite simple:

1. Record all the ideas generated. Appoint someone as secretary at the beginning; this person should also be a contributor.
2. Generate as many ideas as possible, and then verbalize these ideas.
3. Think wild. Silly, impossible ideas sometimes lead to useful ideas.
4. Do not allow evaluation of the ideas; just the generation of them. This is very important. Ignore any evaluation, judgment, or other comments on the value of an idea and chastise the source.

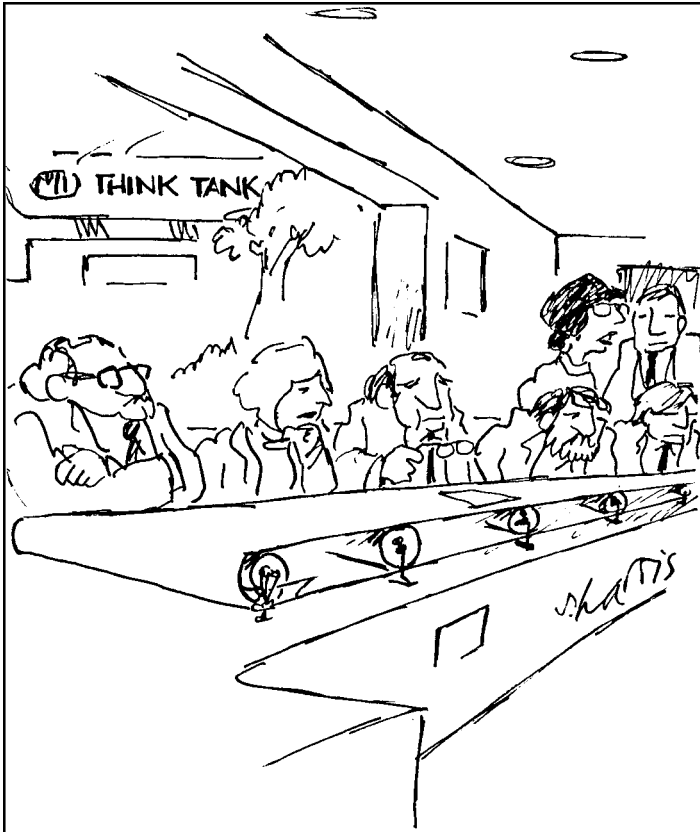
In using this method, there is usually an initial rush of obvious ideas, followed by a period when ideas will come more slowly with periodic rushes. In groups, one member's idea will trigger ideas from the other team members. A brainstorming session should be focused on one specific function and allowed to run through at least three periods during which no ideas are being generated. It is important to encourage humor during brainstorming sessions as even wild, funny ideas can spark useful concepts. This is a proven technique that is useful when new ideas are needed.

7.4.2 Using the 6-3-5 Method as a Source of Ideas

A drawback to brainstorming is that it can be dominated by one or a few team members (see Section 3.3.6). The 6-3-5 method forces equal participation by all. This method is effectively brainstorming on paper and is called *brainwriting* by some. The method is similar to that shown in Fig. 7.14.

To perform the 6-3-5 method, arrange the team members around a table. The optimal number of participants is the "6" in the method's name. In practice, there can be as few as 3 participants or as many as 8. Each takes a clean sheet of paper and divides it into three columns by drawing lines down its length. Next, each team member writes 3 ideas for how to fulfill a specific agreed-upon function, one at the top of each column. The number of ideas is the "3" in the method's name. These ideas can be sketched or written as text. They must be clear enough that others can understand the important aspects of the concept.

After 5 minutes of work on the concepts, the sheets of paper are passed to the right. The time is the "5" in the method's name. The team members now have another 5 minutes to add 3 more ideas to the sheet. This should only be done after studying the previous ideas. They can be built on or ignored as seen fit. As the papers are passed in 5-minute intervals, each team member gets to see the input



"It's our new assembly line. When the person at the end of the line has an idea, he puts it on the conveyor belt, and as it passes each of us, we mull it over and try to add to it."

Figure 7.14 Automated brainwriting. (© 2002 by Sidney Harris. Reprinted with permission from CartoonStock.)

of each of the other members, and the ideas that develop are some amalgam of the best. After the papers have circulated to all the participants, the team can discuss the results to find the best possibilities.

There should be no verbal communication in this technique until the end. This rule forces interpretation of the previous ideas solely from what is on the paper, possibly leading to new insight and eliminating evaluation.

7.4.3 The Use of Analogies in Design

Using analogies can be a powerful aid to generating concepts. The best way to think of analogies is to consider a needed function and then ask, *What else*

provides similar function? An object that provides similar function may trigger ideas for concepts. For example, ideas for the one-handed bar clamp came from a caulking gun (Fig. 7.2).

Many analogies come from nature. For example, engineers are studying the skin of sharks to reduce drag on boats; how ants manage traffic to reduce congestion; and how moths, snakes, and dogs sense odors for bomb detection.

Analogies can also lead to poor ideas. For centuries, people watched birds fly by flapping their wings. By analogy, flapping wings lift birds, so flapping wings should lift people. It wasn't until people began to experiment with fixed wings that the real potential of manned flight became a reality. In fact, what occurred is that by the time of the Wright Brothers in the early 1900s, the problem of manned flight had been divided into four main functions, each solved with some independence of the others: lift, stability, control, and propulsion. The Wright Brothers actually approached each of these in the order listed to achieve controlled, sustained flight.

7.4.4 Finding Ideas in Reference Books and Trade Journals and on the Web

Most reference books give analytical techniques that are not very useful in the early stages of a design project. In some, you will find a few abstract ideas that are useful at this stage—usually in design areas that are quite mature and with ideas so decomposed that their form has specific function. A prime example is the area of linkage design. Even though a linkage is mostly geometric in nature, most linkages can be classified by function. For example, there are many geometries that can be classified by their function of generating a straight line along part of their cycle. (The function is to move in a straight line.) These straight-line mechanisms can be grouped by function. Two such mechanisms are shown in Fig. 7.15.

Many good ideas are published in trade journals that are oriented toward a specific discipline. Some, however, are targeted at designers and thus contain information from many fields. A listing of design-oriented trade journals is given in Sources at the end of this chapter (Section 7.11).

7.4.5 Using Experts to Help Generate Concepts

If designing in a new domain, one in which we are not experienced, we have two choices to gain the knowledge sufficient to generate concepts. We either find someone with expertise in that domain or spend time gaining experience on our own. It is not always easy to find an expert; the domain may even be one that has no experts.

To steal ideas from one person is plagiarism;
to steal from many is research.

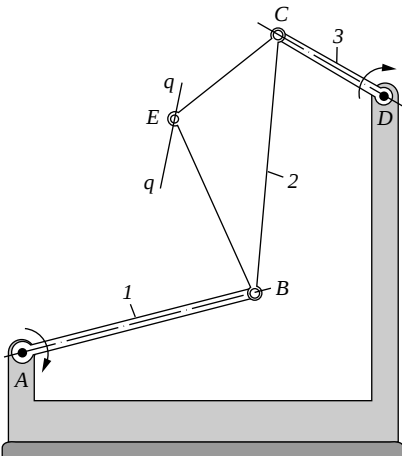
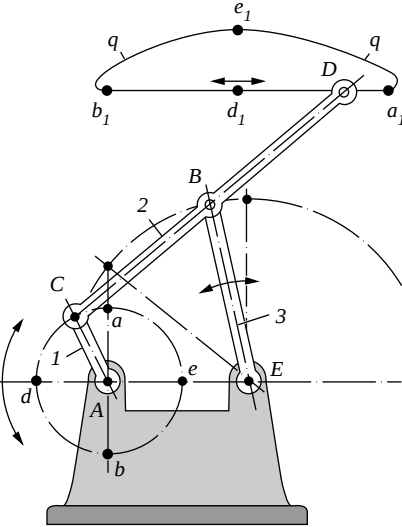
650	WATT FOUR-BAR APPROXIMATE STRAIGHT-LINE MECHANISM	LW GI
 The lengths of the links of four-bar linkage $ABCD$ comply with the conditions: $\overline{AD} = 1.84\overline{AB}$, $\overline{BE} = 0.76\overline{AB}$, $\overline{BC} = 1.03\overline{AB}$, $\overline{EC} = 0.55\overline{AB}$, and $\overline{DC} = 0.52\overline{AB}$. When link 1 turns about fixed axis A, point E of link 2 describes a path of which portion $q-q$ is approximately a straight line.		
651	CHEBYSHEV FOUR-BAR APPROXIMATE STRAIGHT-LINE MECHANISM	LW GI
 The lengths of the links of four-bar linkage $ABCD$ comply with the conditions: $\overline{CB} = \overline{BE} = \overline{BD} = 2.5\overline{AC}$ and $\overline{AE} = 2\overline{AC}$. When link 1 rotates about fixed axis A, point D of link 2 describes path $q-q$. Upon motion of point C along arc $a-d-b$, point D travels along approximately straight line $a_1-d_1-b_1$.		

Figure 7.15 Straight-line mechanisms. (Source: Adapted from I. I. Artobolevsky, *Mechanisms in Modern Engineering Design*, MIR Publishers, Moscow, 1975.)

How do you become an expert in an area that is new or unique? How do you become expert when you cannot find or afford the existing experts? Evidence of expertise can be found in any good designer's office. The best designers work long and hard in a domain, performing many calculations and experiments themselves to find out what works and what does not. Their offices also contain many reference books, periodicals, and sketches of concept ideas.

A good source of information is manufacturers' catalogs and, even better, manufacturers' representatives. A competent designer usually spends a great deal of time on the telephone with these representatives, trying to find sources for specific items or trying to find "another way to do it." One way to find manufacturers is through indexes such as the *Thomas Register*, a gold mine of ideas. All technical libraries subscribe to the 23 annually updated volumes, which list over a million producers of components and systems usable in mechanical design. Beyond a limited selection of reprints of manufacturers' catalogs, the *Thomas Register* does not give information directly but points to manufacturers that can be of assistance. The hard part of using the *Register* is finding the correct heading, which can take as much time as the patent search. The *Thomas Register* is easily searched on the website (see sources in Section 7.11 for the URL).

7.5 PATENTS AS A SOURCE OF IDEAS

Patent literature is a good source of ideas. It is relatively easy to find patents on just about any subject imaginable and many that are not. Problems in using patents are that it is hard to find exactly what you want in the literature; it is easy to find other, interesting, distracting things not related to the problem at hand; and patents are not very easy to read.

There are two main types of patents: *utility patents* and *design patents*. The term *utility* is effectively synonymous with *function*, so the claims in a utility patent are about how an idea operates or is used. Almost all patent numbers you see on products are for utility patents. Design patents cover only the look or form of the idea, so here the term *design* is used in the visual sense. Design patents are not very strong, as a slight change in the form of a device that makes it look different is considered a different product. All design patent numbers begin with the letter "D." Utility patents are very powerful, because they cover how the device works, not how it looks.

There are over 7 million utility patents, each with many diagrams and each having diverse claims. To cull these to a reasonable number, a *patent search* must be performed. That is, all the patents that relate to a certain utility must be found. Any individual can do this, but it is best accomplished by a professional familiar with the literature.

Patent searching changed dramatically in the mid-1990s. Prior to this time, it was necessary to dig through difficult indices and then actually go to one of 50 patent depositories in the United States to see the full text and diagrams. It is now possible to search for patents easily on the Web. Good websites for this are listed in Section 7.11.

Try to not reinvent the wheel.

Before detailing how to best do a patent search, the anatomy of a patent is described. Figure 7.16 is the first page of an early Quick-Grip patent. The heading states that this is a U.S. patent, gives the patent number (since there is not a “D” in front of this number, it is a utility patent), the name of the first inventor, and the date. Important information in the first column is the assignee, the filing (i.e., application) date, its class, and other references cited.

The assignee is the entity which effectively owns the patent, generally the employer of the inventor. Most engineers sign a form on employment that states that the employer owns (is the assignee for) all ideas developed.

The length of time between the filing date and date of the patent is about 15 months in this case. The patent process may take longer depending on revisions (see Section 12.5) and the specific area (e.g., software patents can take three years or longer due to backlog at the patent office).

All patents are organized by their class and subclass numbers. For the example in Fig. 7.16, the primary U.S. class is 81 and subclass is 487. Looking in the *Manual of U.S. Patent Classification*, which can be found in most libraries or at one of the websites, Class 81 is titled “Tools.” Subclass 487 is titled “Hand Held Holder of Having Clamp.” Although the title is not clear, the description is:

Tool comprising either (1) a device adapted to be supported by hand having a work supporting portion or (2) two relatively movable work engaging surfaces for gripping the work of for holding portions of the work in relative position.

Also in the first column of Fig. 7.16 is “references cited.” These are other, earlier patents that are relevant to this patent. Note that in this case, the earliest patent cited is 1932. Referencing a patent this old is often done because all new ideas are based on much older work.

In the second column, after the rest of the references, is the abstract. The abstract is often the first claim of the patent or a paraphrase of it. Often patents have 20 or more claims. Claims are statements about the unique utility (i.e., function) of the device. In patents, subsequent claims are generally built on the first one.

Finally, on the patent front page is a patent drawing. This is usually the first drawing in the patent. As seen in Fig. 7.16, a patent drawing is a stylized line drawing of the device complete with numbers that describe the various parts. In this case, the clamp is shown with the jaws reversed so it can spread rather than clamp. Conversion to this feature is possible with the Irwin product. The remainder of the patent contains a description of the patent, a description of the drawings, the claims, and the drawings.

To use patents as an aid to understanding existing devices, the patent literature can be searched by classification or keyword. If a patent number is known, then use its main class/subclass to search for other similar devices. In the

United States Patent [19]		[11] Patent Number: 5,009,134																																																				
Sorensen et al.		[45] Date of Patent: * Apr. 23, 1991																																																				
[54] QUICK-ACTION BAR CLAMP [75] Inventors: Joseph A. Sorensen; Dwight L. Gatzemeyer, both of Lincoln, Nebr. [73] Assignee: Petersen Manufacturing Co., Inc. [*] Notice: The portion of the term of this patent subsequent to May 22, 2007 has been disclaimed. [21] Appl. No.: 480,283 [22] Filed: Feb. 15, 1990 Related U.S. Application Data [63] Continuation-in-part of Ser. No. 234,173, Aug. 19, 1988, Pat. No. 4,926,722. [51] Int. Cl.⁵ B25B 5/02 [52] U.S. Cl. 81/487; 269/6; 269/166; 269/169; 269/88 [58] Field of Search 81/487, 126; 269/166, 269/167, 170, 169, 165, 6, 203, 204, 88; 29/239 [56] References Cited U.S. PATENT DOCUMENTS <table style="width: 100%; border: none;"> <tr> <td style="width: 30%;">1,878,624</td> <td style="width: 30%;">9/1932</td> <td style="width: 30%;">Estes</td> <td style="width: 10%;">29/239</td> </tr> <tr> <td>3,096,975</td> <td>7/1963</td> <td>Irwin</td> <td>269/169</td> </tr> <tr> <td>4,042,264</td> <td>8/1977</td> <td>Shumer</td> <td>269/203</td> </tr> <tr> <td>4,220,322</td> <td>9/1980</td> <td>Hobday</td> <td></td> </tr> <tr> <td>4,306,710</td> <td>12/1981</td> <td>Vosper</td> <td>269/204</td> </tr> <tr> <td>4,339,113</td> <td>7/1982</td> <td>Vosper</td> <td>269/204</td> </tr> </table> FOREIGN PATENT DOCUMENTS <table style="width: 100%; border: none;"> <tr> <td style="width: 30%;">1408886</td> <td style="width: 30%;">10/1975</td> <td style="width: 30%;">United Kingdom</td> <td style="width: 10%;">.</td> </tr> <tr> <td>1516748</td> <td>7/1978</td> <td>United Kingdom</td> <td>.</td> </tr> <tr> <td>1544156</td> <td>4/1979</td> <td>United Kingdom</td> <td>.</td> </tr> <tr> <td>1555455</td> <td>11/1979</td> <td>United Kingdom</td> <td>.</td> </tr> <tr> <td>2178689</td> <td>2/1987</td> <td>United Kingdom</td> <td>.</td> </tr> <tr> <td>1472278</td> <td>5/1987</td> <td>United Kingdom</td> <td>.</td> </tr> <tr> <td>2204264</td> <td>11/1988</td> <td>United Kingdom</td> <td>.</td> </tr> </table> OTHER PUBLICATIONS Rhombus Pamphlet, Rhombus Tool Limited, Jun. 30, 1989. <i>Primary Examiner</i>—Roscoe V. Parker <i>Attorney, Agent, or Firm</i>—Lackebach Siegel Marzullo & Aronson [57] ABSTRACT A bar clamp having a fixed jaw and a movable jaw which is radially movable over both short and long distances to clamp against a workpiece and is operable using one hand with complete control by the operator at all times. The jaws may either face one another while being mounted on the same side of a handle/grip assembly or face in opposite directions while being mounted on opposite sides of the handle/grip assembly whereby they may be incrementally advanced by the trigger handle/driving lever. 8 Claims, 8 Drawing Sheets			1,878,624	9/1932	Estes	29/239	3,096,975	7/1963	Irwin	269/169	4,042,264	8/1977	Shumer	269/203	4,220,322	9/1980	Hobday		4,306,710	12/1981	Vosper	269/204	4,339,113	7/1982	Vosper	269/204	1408886	10/1975	United Kingdom	.	1516748	7/1978	United Kingdom	.	1544156	4/1979	United Kingdom	.	1555455	11/1979	United Kingdom	.	2178689	2/1987	United Kingdom	.	1472278	5/1987	United Kingdom	.	2204264	11/1988	United Kingdom	.
1,878,624	9/1932	Estes	29/239																																																			
3,096,975	7/1963	Irwin	269/169																																																			
4,042,264	8/1977	Shumer	269/203																																																			
4,220,322	9/1980	Hobday																																																				
4,306,710	12/1981	Vosper	269/204																																																			
4,339,113	7/1982	Vosper	269/204																																																			
1408886	10/1975	United Kingdom	.																																																			
1516748	7/1978	United Kingdom	.																																																			
1544156	4/1979	United Kingdom	.																																																			
1555455	11/1979	United Kingdom	.																																																			
2178689	2/1987	United Kingdom	.																																																			
1472278	5/1987	United Kingdom	.																																																			
2204264	11/1988	United Kingdom	.																																																			

Figure 7.16 A one-handed bar clamp patent front page.

current example, searching under 81/487 yields over 400 recent patents. One problem with patent searches is that usually more information is uncovered than can be reviewed. With each patent found, the claims, drawings, and the reverse engineering methods in the previous section can be used to understand the functionality.

If it is not clear how to start a patent search, then use keywords to search. Prior to the introduction of the Web, keyword searching was not readily possible. Now it is easy to search on the Patent and Trademark Office website for patents issued since 1970 with limited searching back to 1795. Searching “bar” and “clamp” resulted in 1298 patents. Reviewing these showed that many were for concepts for very different applications. However, some seemed to suggest alternative ways for clamping with one hand.

This section has only covered using the patent literature to understand how others have solved similar problems. The process of actually applying for a patent is covered in Section 12.5. Further, over the last few years people have made an effort to organize the patents in other useful ways that help generate concepts. One of these, TRIZ, is discussed in Section 7.7. To make the best use of TRIZ, you first need to understand the concept of contradictions, another idea generation method.

7.6 USING CONTRADICTIONS TO GENERATE IDEAS

Contradictions are engineering “trade-offs.” A contradiction occurs when something gets better, forcing something else to get worse. This means that the ability to fulfill the target for one requirement adversely affects the ability to fulfill another. Some examples are

- Increasing the speed with which squeezing the grip on the one-handed bar clamp moves the jaws together (good) lowers the clamping force (bad).
- The product gets stronger (good) but the weight increases (bad).
- More functions (good) make products larger and heavier (bad).
- An automobile airbag should deploy very fast, to protect the occupant (good), but the faster it deploys, the more likely it is to injure somebody (bad).

Working with contradictions is a powerful method that seems to have evolved in two different fields. The first is as one of the suite of methods used in TRIZ (discussed further in Section 7.7) to generate concepts and as a part of Critical Chain Project Management, a methodology for managing projects (not discussed in this text, but see Sources, Section 7.11, for links that describe it). In project management, using contradictions to generate ideas is called the Evaporating Cloud (EC) method because it helps evaporate the contradiction. The steps developed next help take the amorphous mess of a problem (the cloud), structure it, and then evaporate it by developing better alternative solutions and increasing understanding of the issue.

Figure 7.17 shows the basic EC. The steps in this diagram are

1. Articulate the conflicting positions or functions.
2. Identify the needs forcing the two positions.

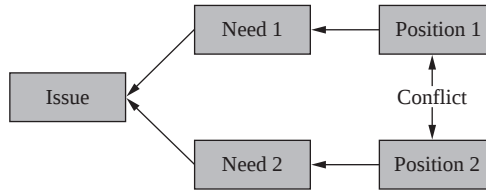


Figure 7.17 Basic structure of the Evaporating Cloud.

3. Identify the issue, the objective of the needs.
4. Generate the assumptions that underlie all of the above.
5. Articulate interjections that can relieve the conflict while meeting the objective.

Let's look at the EC steps through the following example. A company's flagship product was once the market leader but now the competition has caught up. The company can add more functions, but then the product gets heavier and larger. They need to add functions but can't make the product larger and heavier.

1. *Articulate the conflicting positions.* The two positions—initial alternatives—are, “make product smaller and lighter” versus “fit in all the functions.” These are shown in the EC in Fig. 7.18. They represent the basic conflict or dilemma. It is assumed here that many issues start with a basic conflict—the problem that brings the issue to light. These two initial positions are alternative, and mutually exclusive, solutions for the problem. You can't have them both. Another way of formulating the initial positions is to state what you want to improve. This is the first position. Then, identify something else that is preventing you from improving the first position or something that becomes compromised if you do improve it.

The conflict between these two positions is what this method is trying to resolve. Don't get too concerned that there are only two alternative positions; they are merely the starting point, and will evaporate as we progress.

2. *Identify the needs forcing the two positions.* Once the initial positions are identified, the primary “need” or requirement for the position—the “why”—must be discovered. It is the most critical criterion that requires us to choose the position. In this example, we are going to make the product smaller and lighter because we need to make it easier for the customers to move and handle. Similarly, we need the functions to meet the competition. These needs are shown in the diagram in Fig. 7.19. Ideally, we would like to satisfy both of these needs. They are two initial criteria for a good solution to the problem.

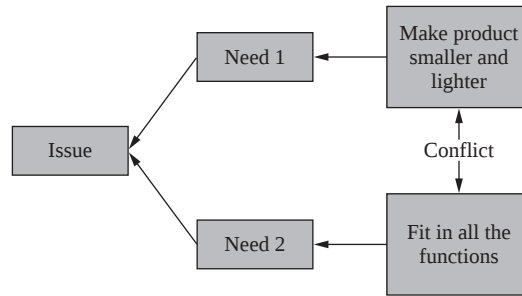


Figure 7.18 The initial positions that cause the conflict.

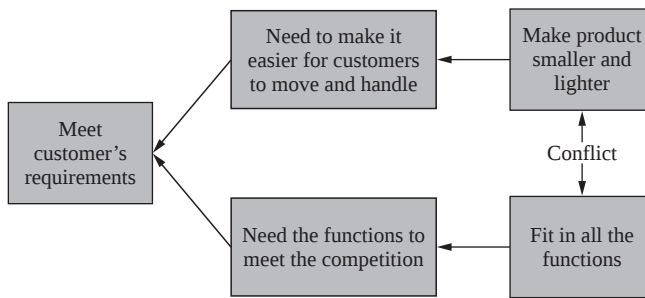


Figure 7.19 The completed initial Evaporating Cloud.

3. *Identify the issue, the objective of the needs.* Based on the needs, you can identify the issue or objective. The issue answers the question, “Why is all this important?” Here, the reason all this is important is that we want to meet the customer’s requirements. Now we can read the entire diagram (Fig. 7.19). Across the top—if we make the product smaller and lighter, we will make it easier for customers to move and handle—some of the customer’s requirements. Across the bottom—if we fit in all the functions, we will meet the competition and customer’s requirements. However, although both lead to the same objective, we have a conflict because we assume we can’t do both with our limited resources.

4. *Generate the assumptions that underlie all of the above.* Now comes the fun part. All of the items in this diagram were predicated on assumptions. These assumptions need to be teased out, as each leads to more criteria and alternatives, and maybe even new issues. To do this, consider each arrow and box, and ask “why”; the “because” answers are the assumptions. There are usually many assumptions. If you find only one per arrow or box, then stretch harder or consider reformulating the cloud.

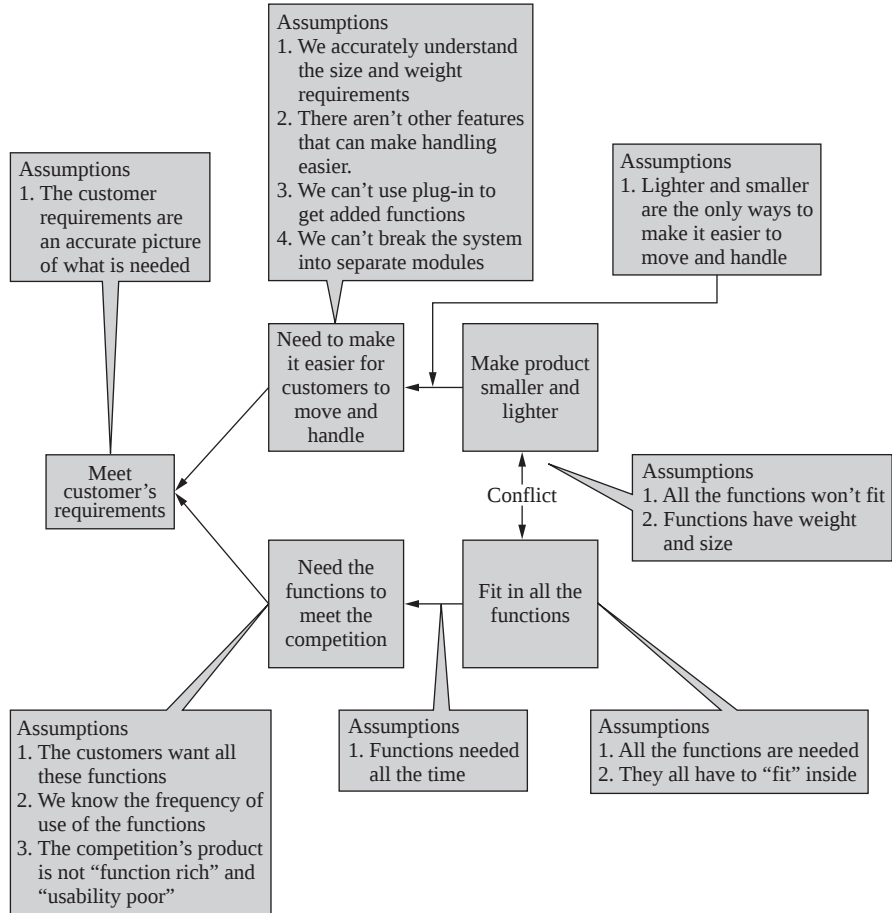


Figure 7.20 The assumptions.

In Fig. 7.20, 14 assumptions have been identified. Some of them may seem obvious, they may overlap, and in some cases, they are trivial. But by noting these assumptions, you can

- Question the diagram for its validity. Some of the assumptions may demand more information (e.g., whether it is true that “the customers are not aware of our product” or “we understand the customers’ desires”). The diagram may need reformulating based on what you now know.
- Note new criteria. Explore how each assumption adds a requirement or constraint to the problem.
- Identify new alternatives. These are called injections and are the focus of the final step.

5. *Articulate injections that can relieve the conflict while meeting the objective.* The final step to evaporate the cloud is to add injections. An injection is a new idea that may help break the conflict. Since virtually all assumptions center on why you can't do something, ask the question, "What can eliminate this assumption?" Answers to this question can help develop directions for further study and new alternatives to consider. In this example, some additional research that might help clarify the situation would be

- Are all the functions on the customers' product used?
- Can we modularize the product?
- Do we really know what the customers want?

Some new ideas that are evident from the EC Fig. 7.20 include:

- Plug ins
- Modules
- Achieving the functions using software (from "Functions have weight and size")

Although the diagram helps tease out much information, the EC mindset is even more important:

- The two alternative views, which seem to conflict, do not conflict in reality if they both support the goal. To meet both needs, we need to fix something that is wrong with our perception (recall the story of the six blind men and the elephant).
- The process brings two sides together to focus on developing a new win-win solution that better meets both needs, thus evaporating the apparent conflict, in which each side defends its position. The win-win solution is not a compromise, which is lose-lose.

7.7 THE THEORY OF INVENTIVE MACHINES, TRIZ

TRIZ (pronounced "trees") is the acronym for the Russian phrase "The Theory of Inventive Machines." TRIZ is based on two ideas:

1. Many of the problems that engineers face contain elements that *have already been solved*, often in a *completely different industry*, for a totally *unrelated situation*, that uses an entirely *different technology* to solve the problem.
2. There are predictable patterns of technological change that can be applied to any situation to determine the most probably successful next steps.

The theory is that with TRIZ we can *systematically innovate*; we don't have to wait for an "inspiration" or use the trial and error common to the other methods

presented earlier. Practitioners of TRIZ have a very high rate of developing new, patentable ideas. To best understand TRIZ, its history is important.

This method was developed by Genrikh (aka Henry) Altshuller, a mechanical engineer, inventor, and Soviet Navy patent investigator. After World War II Altshuller was tasked by the Russian government to study worldwide patents to look for strategic technologies the Soviet Union should know about. He and his team noticed that some of the same principles were used repeatedly by totally different industries, often separated by many years, to solve similar problems.

Altshuller conceived of the idea that inventions could be organized and generalized by function rather than the traditional indexing system discussed in Section 7.5. From his findings, Altshuller began to develop an extensive “knowledge base,” which includes numerous physical, chemical, and geometric effects along with many engineering principles, phenomena, and patterns of evolution. Altshuller wrote a letter to Stalin describing his new approach to improve the rail system along with products the U.S.S.R. produced. The Communist system at the time didn’t value creative, freethinking. His ideas were scorned as insulting, individualistic, and elitist, and as a result of this letter, he was imprisoned in 1948 for these capitalist and “insulting” ideas. He was not released until 1954, after Stalin’s death. From the 1950s until his death in 1998, he published numerous books and technical articles and taught TRIZ to thousands of students in the former Soviet Union. TRIZ has become a best practice worldwide.

Altshuller’s initial research in the late 1940s was conducted on 400,000 patents. Today the patent database has been extended to include over 2.5 million patents. This data has led to many TRIZ methods by both Altshuller and his disciples. The first, contradictions, was developed in Section 7.6. The second, the use of 40 inventive principles, is based on contractions.

TRIZ’s 40 inventive principles, help in generating ideas for overcoming contradictions.¹ The inventive principles were found by Altshuller when researching patents from many different fields of engineering and reducing each to the basic principle used. He found that there are 40 inventive principles underlying all patents. These are proposed “solution pathways” or methods of dealing with or eliminating engineering contradictions between parameters. The entire list of principles and a description of each is on the website. In the list below, the names of the inventive principles are shown organized into seven major categories.

- Organize (6)
 - Segment, Merge, Abstract, Nest
 - Counterweight, Asymmetry

¹ Here, the method has been greatly shortened. In traditional TRIZ practice, the contradictions are used with a large table to find which inventive principles might best be used. The table is too large for inclusion here and simply exploring the 40 principles is not much more time consuming and is more fun than using the table.

- Compose (7)
 - Local Quality, Universality
 - Homogeneity, Composites
 - Spheroids, Thin Films, Cheap Disposables
- Physical (4)
 - Porosity, Additional Dimension, Thermal Expansion, Color Changes
- Chemical (4)
 - Oxidate—Reduce Inertness
 - Transform States, Phase Transition
- Interactions (5)
 - Reduce Mechanical Movement, Bring Fluidity
 - Equipotence, Dynamicity, Vibration
- Process (9)
 - Do It in Reverse, ++ / --, Continued Action, Repeated Action, Skip Through, Negative to Positive
 - Prior Cushioning, Prior Actions, Prior Counteractions
- Service (5)
 - Self-Service, Intermediary, Feedback,
 - Use and Retrieve, Cheap Copies

To see how this works, consider a contradiction in the design of one handed clamp from Section 7.6 “Increasing the speed with which squeezing the grip on the one-handed bar clamp moves the jaws together (good) lowers the clamping force (bad).” Reviewing the list of 40 inventive principles, three ideas were generated. Each inventive principle is listed as a title and clarifying statements followed by the idea generated.

Principle 1. Segmentation

- a. Divide an object into independent parts
- b. Make an object sectional
- c. Increase degree of an object’s segmentation

This leads to the idea of having two mechanisms, one for fast motion with low force and one that gives high force when the motion slows due to clamping pressure. In fact, this two-stage action has been patented by Irwin.

Principle 10. Prior action

- a. Carry out the required action in advance in full, or at least in part
- b. Arrange objects so they can go into action without time loss waiting for action

This leads to the idea of having the clamp automatically move so the jaws come into contact with the work (prior action) and then the grip force is translated into high clamping force with small motion. This is similar to the first idea, but the prior motion is automated.

Principle 17. Moving to a new dimension

- a. Remove problems in moving an object in a line by two-dimensional movement (along a plane)
- b–d. Others are not important here

This leads to the idea of using a linkage to get a more complex motion than purely linear. A linkage is used to get the jaws in contact with the work and then the small motion with high force is action as is typical with a one-handed clamp.

There are many other ideas to be discovered by working through the inventive principles and other TRIZ techniques (see Section 7.11 for TRIZ information sources).

7.8 BUILDING A MORPHOLOGY

The technique presented here uses the functions identified to foster ideas. It is a very powerful method that can be used formally, as presented here, or informally as part of everyday thinking. There are three steps to this technique. The first step is to list the decomposed functions that must be accomplished. The second step is to find as many concepts as possible that can provide each function identified in the decomposition. The third is to combine these individual concepts into overall concepts that meet all the functional requirements. The design engineer's knowledge and creativity are crucial here, as the ideas generated are the basis for the remainder of the design evolution. This technique is often called the "morphological method," and the resulting table a "morphology," which means "a study of form or structure." A partial Morphology for the redesign of the one-handed bar clamp is presented in Figure 7.21. This is highly modified from the morphology done at Irwin to protect their intellectual property. A blank morphology is available as a template.

7.8.1 Step 1: Decompose the Function

The first half of this chapter details this step. For the one-handed clamp example, the function was decomposed in Fig. 7.11. The first four functions in that figure are



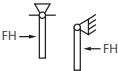
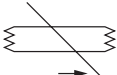
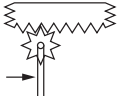
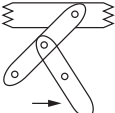
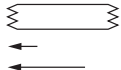
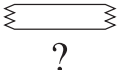
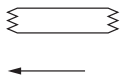
Morphology				
Product: One-handed bar clamp			Organization Name: Irwin Tools	
Subfunctions	Concept 1	Concept 2	Concept 3	Concept 4
Collect grip force and motion from user	One trigger 	Two triggers 		
Transform grip force and motion to bar	Jam plate 	Ratchet 	Rack and pinion 	Linkage 
Move bar	Free sliding 	2 speed system 	> 2 speed system 	
Amplify force	Short stroke 	Long stroke 		
Team Member: DJP	Team Member:		Prepared by: DJP	
Team Member: AKQ	Team Member:		Checked by: AKQ	Approved by:
The Mechanical Design Process Copyright 2008, McGraw-Hill			Designed by Professor David G. Ullman Form #15.0	

Figure 7.21 Example of a morphology.

- Collect grip force and motion from user
- Transform grip force and motion to bar
- Move bar
- Amplify force

These functions were the focus of the new design effort, as Irwin wanted to redesign these to make the clamp more user-friendly. Specifically, the functions “move bar” and “amplify force” are a contradiction. A mechanism that transforms each handgrip cycle (squeeze and release) to move the bar rapidly will result in a lower applied force than one that moves the bar a short distance. As with any other transmission system there is a trade-off between speed and force (or torque in rotational systems). The user would like to be able to move the bar rapidly in the position and then apply a high force. So, this effort focuses on rapidly moving the bar into position and then amplifying the force.

7.8.2 Step 2: Develop Concepts for Each Function

The goal of this second step is to generate as many concepts as possible for each of the functions identified in the decomposition. For the example, there are two ways to collect the grip force and motion from the user, as shown in Fig. 7.21. The first is to use a single trigger as shown in Figs. 7.2, 7.3, and 7.4. This is shown schematically in the morphology with a hand force applied to the trigger and the trigger pivoted someplace in the clamp body. Another option is two triggers, shown as Concept 2 in the morphology. For this concept, both the force on the trigger and the reaction force on the handle are used to enable the clamp. The concepts in the morphology are abstract in that they have no specific geometry. Rough sketches of these concepts and words are both used to describe the concept.

Four ideas were generated to transform the grip. These are not all well thought out, but the morphology is generating ideas, so this is all right. When the project began, discussion centered on a two-speed system, fast to get the clamp in contact with the work and then slow so the force can be amplified during clamping. As can be seen in the “move bar” row, an idea that evolved here is for more than two speeds. Although no immediate ideas were generated, this offered even more possibilities to consider.

If there is a function for which there is only one conceptual idea, this function should be reexamined. There are few functions that can be fulfilled in only one way. The lack of more concepts can be due to

The designer making a fundamental assumption. For example, one function that has to occur in the system is “Collect grip force and motion from user.” It is reasonable to assume that a gripping force will be used to provide motion and clamping force *only if the designer is aware that an assumption has been made.*

The function is directed at how, not what. If one idea gets built into the function, then it should come as no surprise that this is the only idea that gets generated. For example, if “Transform grip force and motion to bar” in Fig. 7.21 had been stated as “use jam plate to transform motion,” then only jam plate ideas are possible. If the function statement has nouns that tell *how* the function is to be accomplished, reconsider the function statement.

The domain knowledge is limited. In this case, help is needed to develop other ideas. (See Sections 7.5, 7.6, or 7.7.)

It is a good idea to keep the concepts as abstract as possible and at the same level of abstraction. Suppose one of the functions is to move some object. Moving requires a force applied in a certain direction. The force can be provided by a hydraulic piston, a linear electric motor, the impact of another object, or magnetic repulsion. The problem with this list of concepts is that they are at different levels of abstraction. The first two refer to fairly refined mechanical components. (They could be even more refined if we had specific dimensions or manufacturers’ model numbers.) The last two are basic physical principles. It is difficult to compare these concepts because of this difference in level of

abstraction. We could begin to correct this situation by abstracting the first item, the hydraulic piston. We could cite instead the use of fluid pressure, a more general concept. Then again, air might be better than hydraulic fluid for the purpose, and we would have to consider the other forms of fluid components that might give more usable forces than a piston. We could refine the “impact of another object” by developing how it will provide the impact force and what the object is that is providing the force. Regardless of what is changed, it is important to try to get all concepts to be equally refined.

7.8.3 Step 3: Combine Concepts

The result of applying the previous step is a list of concepts generated for each of the functions. Now we need to combine the individual concepts into complete conceptual designs. The method here is to select one concept for each function and combine those selected into a single design. So, for example, we may consider combining one trigger with a ratchet as part of a free-sliding system with a short stroke. This configuration frees the bar so that it can be easily pushed into position against the work and then uses the ratchet to apply force to the work. A second system is similar but uses a jam plate. These are both shown in Fig. 7.22 by lines connecting the concepts. In the actual Irwin morphology, six concepts were generated and drawn on their CAD system for evaluation.

There are pitfalls to this method, however. First, if followed literally, this method generates too many ideas. The one-handed clamp morphology, for example, is small, yet there are 48 possible designs ($2 \times 4 \times 3 \times 2$).

The second problem with this method is that it erroneously assumes that each function of the design is independent and that each concept satisfies only one function. Generally, this is not the case. For example, if a two-speed system is used, it has both a long and a short stroke and may not work with a linkage. Nonetheless, breaking the function down this finely helps with understanding and concept development.

Third, the results may not make any sense. Although the method is a technique for generating ideas, it also encourages a coarse ongoing evaluation of the ideas. Still, care must be taken not to eliminate concepts too readily; a good idea could conceivably be prematurely lost in a cursory evaluation. A goal here is to do only a coarse evaluation and generate all the reasonably possible ideas. In Chap. 8, we will evaluate the concepts and decide between them.

Even though the concepts developed here may be quite abstract, this is the time for back-of-the-envelope sketches. Prior to this time, most of the design effort has been in terms of text, not graphics. Now the design is developing to the point that rough sketches must be drawn.

Sketches of even the most abstract concepts are increasingly useful from this point on because (1) as discussed in Chap. 3, we remember functions by their forms; thus our index to function is form; (2) the only way to design an object with any complexity is to use sketches to extend the short-term memory; and



Morphology				
Product: One-handed bar clamp			Organization Name: Irwin Tools	
Subfunctions	Concept 1	Concept 2	Concept 3	Concept 4
Collect grip force and motion from user	One trigger 	Two triggers 		
Transform grip force and motion to bar	Jam plate 	Ratchet 	Rack and pinion 	Linkage
Move bar	Free sliding 	2 speed system 	> 2 speed system ?	
Amplify force	Short stroke 	Long stroke 		
Team Member: D/P		Team Member:	Prepared by: D/P	
Team Member: A/Q		Team Member:	Checked by: A/Q	Approved by:
The Mechanical Design Process Copyright 2008, McGraw-Hill			Designed by Professor David G. Ullman Form #15.0	

Figure 7.22 Combining concepts in a Morphology.

(3) sketches made in the design notebook provide a clear record of the development of the concept and the product.

Keep in mind that the goal is only to develop concepts and that effort must not be wasted worrying about details. Often a single-view sketch is satisfactory; if a three-view drawing is needed, a single isometric view may be sufficient.

7.9 OTHER IMPORTANT CONCERNS DURING CONCEPT GENERATION

The techniques outlined in this chapter have focused on generating potential concepts. In performing these techniques, functional decomposition diagrams, literature and patent search results, function-concept mapping, and sketches of overall concepts are all produced. These are all important documents that can support communication to others and archive the design process.

One of the highest complements that a product designer can receive is “That looks so simple.” It is difficult to find the elegant, simple solutions to complex problems, yet they generally exist. Engineering elegance is the goal of this chapter and thus, keep the following aphorism in mind at all times:

Follow the KISS rule: *Keep It Simple, Stupid.*

Additionally, conceptual design is a good time to review the Hannover Principles introduced in Chap. 1. Questions derived from the Principles that should be asked at this time are

1. Do your concepts enable humanity and nature to coexist in a healthy, supportive, diverse, and sustainable condition?
2. Do you understand the effects of your concepts on other systems, even the distant effects?
3. Are concepts safe and of long-term value?.
4. Do your concepts help eliminate the concept of waste throughout their life cycle?
5. Where possible, do they rely on natural energy flows?

7.10 SUMMARY

- The functional decomposition of existing products is a good method for understanding them.
- Functional decomposition encourages breaking down the needed function of a device as finely as possible, with as few assumptions about the form as possible.
- The patent literature is a good source for ideas.
- Exploring contradictions can lead to ideas.
- Listing concepts for each function helps generate ideas; this list is often called a *morphology*.
- Sources for conceptual ideas come primarily from the designer’s own expertise; this expertise can be enhanced through many basic and logical methods.

7.11 SOURCES

Sources for patent searches

<http://www.uspto.gov/patft/index.html>. The website for the U.S. Patent and Trademark Office. Easy to search but has complete information only on recent patents.

<http://www.delphion.com/home>. IBM originally developed this website. Also, easy to search for recent patents.

<http://gb.espacenet.com/>. Source for European and other foreign patents. Supported by the European Patent Organization, EPO.

Other non-patent sources

Artobolevsky, I. I.: *Mechanisms in Modern Engineering Design*, MIR Publishers, Moscow, 1975. This five-volume set of books is a good source for literally thousands of different mechanisms, many indexed by function.

Chironis, N. P.: *Machine Devices and Instrumentation*, McGraw-Hill, New York, 1966. Similar to Greenwood's *Product Engineering Design Manual*.

Chironis, N. P.: *Mechanism, Linkages and Mechanical Controls*, McGraw-Hill, New York, 1965. Similar to the last entry.

Clausing, D., and V. Fey: *Effective Innovation: The Development of Winning Technologies*, ASME Press 2004. A good overview of recent methods to develop new concepts.

Damon, A., H. W. Stoudt, and R. A. McFarland: *The Human Body in Equipment Design*, Harvard University Press, Cambridge, Mass., 1966. This book has a broad range of anthropometric and biomechanical tables.

Design News, Cahners Publishing, Boston. Similar to *Machine Design*. <http://www.designnews.com/>

Edwards, B.: *Drawing on the Right Side of the Brain*, Tarcher, Los Angeles, 1982. Although not oriented specifically toward mechanical objects, this is the best book available for learning how to sketch.

Greenwood, D. C.: *Engineering Data for Product Design*, McGraw-Hill, New York, 1961. Similar to the above.

Greenwood, D. C.: *Product Engineering Design Manual*, Krieger, Malabar, Fla., 1982. A compendium of concepts for the design of many common items, loosely organized by function.

Human Engineering Design Criteria for Military Systems, Equipment, and Facilities, MILSTD 1472, U.S. Government Printing Office, Washington, D.C. This standard contains 400 pages of human factors information. A reduced version with links to other material is at http://hftag.dtic.mil/hfs_docs.html

Machine Design, Penton Publishing, Cleveland, Ohio. One of the best mechanical design magazines published, it contains a mix of conceptual and product ideas along with technical articles. It is published twice a month. www.machinedesign.com.

Norman, D.: *The Psychology of Everyday Things*, Basic Books, New York, 1988. This book is light reading focused on guidance for designing good human interfaces.

Plastics Design Forum, Advanstar Communications Inc., Cleveland, Ohio. A monthly magazine for designers of plastic products and components.

Product Design and Development, Chilton, Radnor, Pa. Another good design trade journal. www.pddnet.com.

Thomas Register of American Manufacturers, Thomas Publishing, Detroit, Mich. This 23-volume set is an index of manufacturers and is published annually. Best used on the Web at www.thomasregister.com.

TRIZ www.triz-journal.com. The TRIZ Journal is a good source for all things TRIZ.

Functional decomposition or reverse engineering case studies for coffeemaker, bicycle, engine, and other products developed by student of Professor Tim Simpson (Pennsylvania State University) and others: http://gicl.cs.drexel.edu/wiki/Reverse_Engineering_Case_Studies

7.12 EXERCISES

- 7.1 For the original design problem (Exercise 4.1), develop a functional model by
- Stating the overall function.
 - Decomposing the overall function into subfunctions. If assumptions are needed to refine this below the first level, state the assumptions. Are there alternative decompositions that should be considered?
 - Identifying all the objects (nouns) used and defending their inclusion in the functional model.
- 7.2 For the redesign problem (Exercise 4.2), apply items a–c from Exercise 7.1 and also study the existing device(s) to establish answers to these questions.
- Which subfunction(s) must remain unchanged during redesign?
 - Which subfunctions (if any) must be changed to meet new requirements?
 - Which subfunctions may cease to exist?
- 7.3 For the functional decomposition developed in Exercise 7.1,
- Develop a morphology as in Fig. 7.21 to aid in generating concepts.
 - Combine concepts to develop at least 10 complete conceptual designs.
- 7.4 For the redesign problem functions that have changed in Exercise 7.2,
- Generate a morphology of new concepts as in Fig. 7.21.
 - Combine concepts to develop at least five complete conceptual designs.
- 7.5 Find at least five patents that are similar to an idea that you have for
- The original design problem begun in Exercise 4.1.
 - The redesign problem begun in Exercise 4.2.
 - A perpetual motion machine. In recent times the patent office has refused to consider such devices. However, the older patent literature has many machines that violate the basic energy conservation laws.
- 7.6 Use brainstorming to develop at least 25 ideas for
- A way to fasten together loose sheets of paper.
 - A device to keep water off a mountain-bike rider.
 - A way to convert human energy to power a boat.
 - A method to teach the design process.
- 7.7 Use brainwriting to develop at least 25 ideas for
- A device to leap tall buildings in a single bound.
 - A way to fasten a gear to a shaft and transmit 500 watts.
- 7.8 Finish reverse engineering the one-handed bar clamp in Figure 7.7
- 7.9 Choose a relatively simple product and functionally decompose it to find the flow of force, energy and information.

7.13 ON THE WEB

Templates for the following documents are available on the book's website: www.mhhe.com/Ullman4e

- Reverse Engineering
- Morphology



CHAPTER

Concept Evaluation and Selection

KEY QUESTIONS

- How can rough conceptual ideas be evaluated without refining them?
- What is technology readiness?
- What is a Decision Matrix?
- How can I manage risk?
- How can I make robust decisions?

8.1 INTRODUCTION

In Chap. 7, we developed techniques for generating promising conceptual solutions for a design problem. In this chapter, we explore techniques for choosing the best of these concepts for development into products. *The goal is to expend the least amount of resources on deciding which concepts have the highest potential for becoming a quality product.* The difficulty in concept evaluation and decision making is that we must choose which concepts to spend time developing when we still have very limited knowledge and data on which to base this selection.

How can rough conceptual ideas be evaluated? Information about concepts is often incomplete, uncertain, and evolving. Should time be spent refining them, giving them structure, making them measurable so that they can be compared with the engineering targets developed during problem specifications development? Or should the concept that seems like the best one be developed in the hope that it will become a quality product? It is here that we address the question of how soon to narrow down to a single concept.

Ideally, enough information about each concept is known at this point to make a choice and put all resources into developing this one concept. However, it is less risky to refine a number of concepts before committing to one of them. This requires resources spread among many concepts and, possibly, inadequate

development of any one of them. Many companies generate only one concept and then spend time developing it. Others develop many concepts in parallel, eliminating the weaker ones along the way. Designers at Toyota follow what they call a “parallel set narrowing process,” in which they continue parallel development of a number of concepts. As more is learned, they slowly eliminate those concepts that show the least promise. This has proven very successful, as seen by Toyota’s product quality and growth. Every company has its own culture for product development and there is no one “correct” number of concepts to select. Here we try to balance learning about the concepts with limited resources. In this chapter, techniques will be developed that will help in making a knowledgeable decision with limited information.

As shown in Fig. 8.1, after generating concepts, the next step that needs to be accomplished is evaluating them. The term *evaluate*, as used in this text, implies *comparison* between alternative concepts relative to the requirements they must

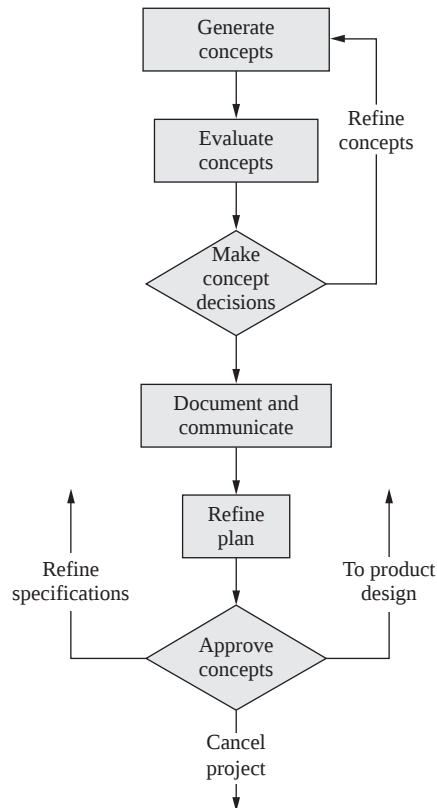


Figure 8.1 The conceptual design phase.

If the horse is dead, get off.

meet. The results of evaluation give the information necessary to *make concept decisions*.

Be ready during concept evaluation to abandon your favorite idea, if you cannot defend it in a rational way. Also, abandon if necessary “the way things have always been done around here.” Reflect on the above aphorism and, if it applies, use it.

Before we get into the details of this chapter, it is worth reflecting on the basic decision-making process introduced in Chap. 4 where we were selecting a project. In Fig. 8.2 (a reprint of Fig. 4.19), the issue is “Select a concept(s) to develop.” We have spent considerable time generating alternatives and criteria. Now we must focus on the remaining steps and decide what to do next. First, we will discuss the types of evaluation information we have available to us, and then we will address different traditional methods for decision making. The criteria importance (step 4) will not really surface until Section 8.5.

The traditional decision-making methods do not do a good job of helping you manage risk and uncertainty. This will be addressed in Section 8.6, and a robust decision-making method, designed for managing uncertainty will be introduced in Section 8.7. Finally, the documentation and communication needs of conceptual design will be detailed.

8.2 CONCEPT EVALUATION INFORMATION

In order to be compared, alternatives and criteria must be in the same language and they must exist at the same level of abstraction. Consider, for example, the spatial requirement that a product fit in a slot 2.000 ± 0.005 in. long. An unrefined concept for this product may be described as “short.” It is impossible to compare “ 2.000 ± 0.005 in.” to “short” because the concepts are in different languages—a number versus a word—and they are at different levels of abstraction—very concrete versus very abstract. It is simply not possible to make a comparison between the “short” concept and the requirement of fitting a 2.000 ± 0.005 in. slot. Either the requirement will have to be abstracted or work must be done on the concept to make “short” less abstract or both.

An additional problem in concept evaluation is that abstract concepts are uncertain; as they are refined, their behavior can differ from that initially anticipated. The greater the knowledge, the less the uncertainty about a concept and the fewer the surprises as it is refined. However, even in a well-known area, as the concept is refined to the product, unanticipated factors arise. Richard Feynman, the Nobel winning physicist said: “If you thought that science was certain—well that is just an error on your part.” A major factor is to manage the uncertain information on which most decisions are based; there is uncertainty in everything.

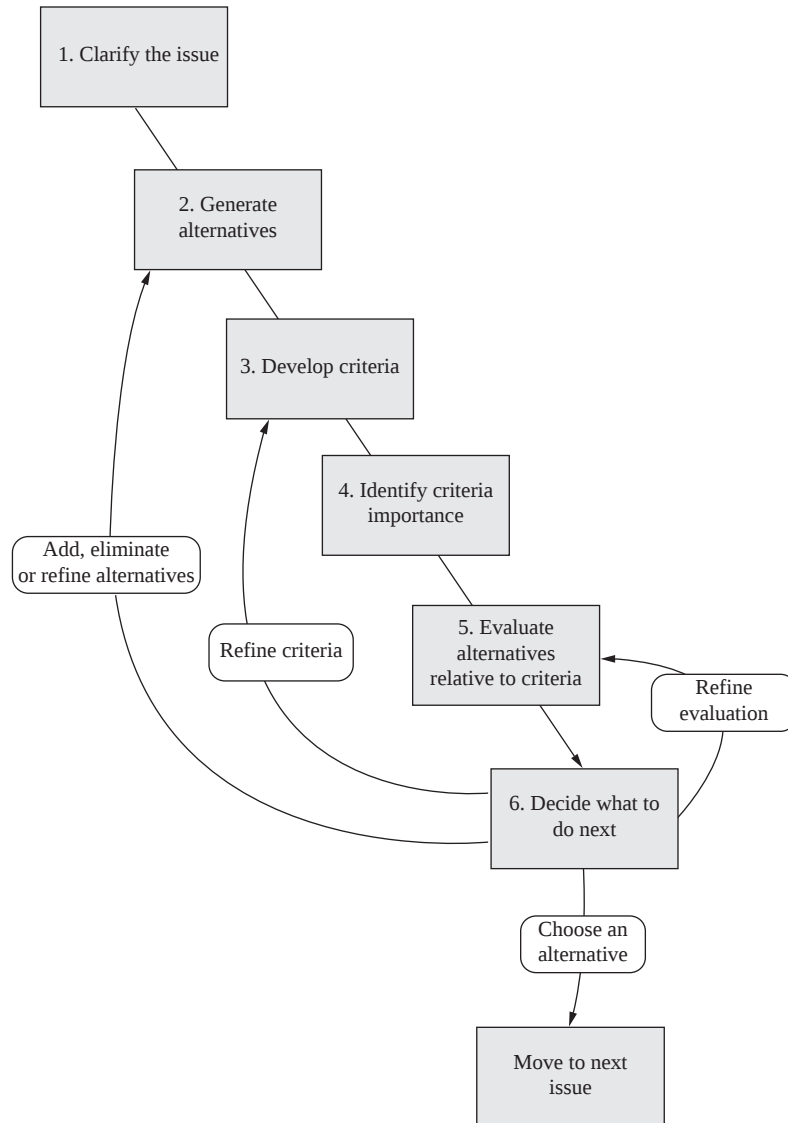


Figure 8.2 The decision-making flow.

When evaluating concepts your information can have a wide range of fidelity (see Section 5.3.3). Back-of-the-envelope calculations are low fidelity, whereas detailed simulations—hopefully—have high fidelity. Experts often run simulations to predict performance and cost. In the early stages of projects, these simulations are usually at low levels of fidelity, and some may be qualitative—just gut feel it. Increasing fidelity requires increased refinement and increased project costs. Increased knowledge generally comes with increased fidelity, but

In planning for the project, we identified the models to be used to represent information during concept development (Table 5.1). Physical models or proof-of-concept prototypes support evaluation by demonstrating the behavior for comparison with the functional requirements or by showing the shape of the design for comparison with form constraints. Sometimes these prototypes are very crude—just cardboard, wire, and other minimal materials thrown together to see if the idea makes sense. Often, when one is designing with new technologies or complex known technologies, building a physical model and testing it is the only approach possible. This *design-build-test cycle* is shown as the inner loop in Fig. 8.3.

The diagram illustrates an iterative design process, divided into three main phases: DESIGN, BUILD, and TEST. The cycle is represented by a circular flow of steps and iterative loops.

- DESIGN Phase:**
 - Starts with **Simulatable technology**.
 - Leads to **Design prototypes**.
 - Includes an **Iterate** loop between **Design prototypes** and **Analytical models and graphical drawings to refine concept and product**.
- BUILD Phase:**
 - Leads from **Design prototypes** to **Build prototypes with each closer to the final product**.
 - Includes an **Iterate** loop between **Build prototypes with each closer to the final product** and **Test physical prototypes**.
- TEST Phase:**
 - Leads from **Build prototypes with each closer to the final product** to **Test physical prototypes**.
 - Includes an **Iterate** loop between **Test physical prototypes** and **Analytical models and graphical drawings to refine concept and product**.
 - Leads from **Test physical prototypes** to **Build final product**.

The process concludes with **Build final product**, which feeds back into the **DESIGN** phase, completing the cycle.

www.EngineeringBooksPdf.com

However, analysis can only be performed on systems that are understood and can be modeled mathematically. New and existing technologies, complex beyond the ability of analytical models, must be explored with physical models.

8.3 FEASIBILITY EVALUATIONS

As a concept is generated, a designer usually has one of three immediate reactions: (1) it is not feasible, it will never work; (2) it might work if something else happens; and (3) it is worth considering. These judgments about a concept's feasibility are based on "gut feel," a comparison made with prior experience stored as design knowledge. The more design experience, the more reliable an engineer's knowledge and the decision at this point. Let us consider the implications of each of the possible initial reactions more closely.

It Is Not Feasible. If a concept seems infeasible, or unworkable, it should be considered briefly from different viewpoints before being rejected. Before an idea is discarded, it is important to ask, Why is it not feasible? There may be many reasons. It may be obviously technologically infeasible. It may not meet the customer's requirements. It may just be that the concept is different from the way things are normally done. Or it may be that because the concept is not an original idea, there is no enthusiasm for it. We will delay discussing the first two reasons until Section 8.4, and we will discuss the latter two here.

As for the judgment that a concept is "different," humans have a natural tendency to prefer tradition to change. Thus, an individual designer or company is more likely to reject new ideas in favor of ones that are already established. This is not all bad, because the traditional concepts have been proven to work. However, this view can block product improvement, and care must be taken to differentiate between a potentially positive change and a poor concept. Part of a company's tradition lies in its standards. Standards must be followed and questioned; they are helpful in giving current engineering practice, and they also may be limiting in that they are based on dated information.

As for the judgment that a concept was "Not Invented Here" (NIH): It is always more ego satisfying to individuals and companies to use their own ideas. Since very few ideas are original, ideas are naturally borrowed from others. In fact, part of the technique presented in Chap. 6 for understanding the design problem involved benchmarking the competition. One of the reasons for doing this was to learn as much as possible about existing products to aid in the development of new products.

A final reason to further consider ideas that at first do not seem feasible is that they may give new insight to the problem. Part of the brainstorming technique introduced in Chap. 7 was to build from the wild ideas that were generated. Before discarding a concept, see if new ideas can be generated from it, effectively iterating from evaluation back to concept generation.

It Is Conditional. The initial reaction might be to judge a concept workable *if something else happens*. Typical of other factors involved are the readiness of

It's hard to make a good product out of a poor concept.

technology, the possibility of obtaining currently unavailable information, or the development of some other part of the product.

It Is Worth Considering. The hardest concept to evaluate is one that is not obviously a good idea or a bad one, but looks worth considering. Engineering knowledge and experience are essential in the evaluation of such a concept. If sufficient knowledge is not immediately available for the evaluation, it must be developed. This is accomplished by developing models or prototypes that are easily evaluated.

8.4 TECHNOLOGY READINESS

One good concept evaluation method is to determine the readiness of its technologies. This technique helps evaluation by forcing a comparison with state-of-the-art capabilities. If a technology is to be used in a product, it must be mature enough that its use is a design issue, not a research issue. The vast majority of technologies used in products are mature, and the measures discussed below are readily met. However, in a competitive environment, there are high incentives to include new technologies in products. Recall from Chap. 1 that a majority of people think that including the latest technology in a product is a sign of quality. Care must be taken to ensure that the technology is *ready* to be included in the product.

Consider the technologies listed in Table 8.1. Each of these technologies required many years from inception to the realization of a physical product. The same holds true for all technologies. Even ones that do not change the world as did the ones in the table. An attempt to design a product before the necessary technologies are ready leads either to a low-quality product or to a project that is canceled before a product reaches the market because it is behind schedule and over cost. How, then, can the maturity of a technology be measured? Six metrics can be applied to determine a technology's maturity:

1. *Are the critical parameters identified?* Every design concept has certain parameters that are critical to its proper operation and use. It is important to know which parameters (e.g., dimensions, material properties, or other features) are critical to the function of the device. It has been estimated that only about 10 to 15% of the dimensions on a finished component are critical to the operation of the product. For a simple cantilever spring, the critical parameters are its length, its moment of inertia about the neutral axis, the distance from the neutral axis to the most highly stressed material, the modulus of elasticity, and the maximum allowable yield stress. These parameters allow for the calculation of the spring stiffness and the failure potential for a given force. The first three parameters are dependent on the geometry; the last two are dependent on the material properties. Say you need a ceramic spring in a

Table 8.1 A time line for technology readiness

Technology	Development time, years
Powered human flight	403 (1500–1903)
Photographic cameras	112 (1727–1839)
Radio	35 (1867–1902)
Television	12 (1922–1934)
Radar	15 (1925–1940)
Xerography	17 (1938–1955)
Atomic bomb	6 (1939–1945)
Transistor	5 (1948–1953)
High-temperature superconductor	? (1987–)

concept. Are the material properties modulus of elasticity and the maximum allowable yield stress the correct material properties to be considering?

Additional critical parameters determine a device's acceptability as a product (e.g., weight, size, and other physical parameters). These too must be identified, but may not be well known at this stage of development.

2. *Are the safe operating latitude and sensitivity of the parameters known?* In refining a concept into a product, the actual values of the parameters may have to be varied to achieve the desired performance or to improve manufacturability. It is essential to know the limits on these parameters and the sensitivity of the product's operation to them. This information is known in only a rough way during the early design phases; during the product evaluation, it will become extremely important.
3. *Have the failure modes been identified?* Every type of system has characteristic failure modes. It is generally useful to continuously evaluate the different ways a product might fail. This is expanded on in Chap. 11.
4. *Can the technology be manufactured with known processes?* If reliable manufacturing processes have not been refined for the technology, then, either the technology should not be used or there must be a separate program for developing the manufacturing capability. There is a risk in the latter alternative, as the separate program could fail, jeopardizing the entire project.
5. *Does hardware exist that demonstrates positive answers to the preceding four questions?* The most crucial measure of a technology's readiness is its prior use in a laboratory model or another product. If the technology has not been demonstrated as mature enough for use in a product, the designer should be very wary of assurances that it will be ready in time for production.
6. *Is the technology controllable throughout the product's life cycle?* This question addresses the later stages of the product's life cycle: its manufacture, use, service, and retirement. It also raises other questions. What manufacturing by-products come from using this technology? Can the by-products be safely disposed of? How will this product be retired? Will it degrade safely? Answers to these questions are the responsibility of the design engineer.



Technology Readiness Assessment				
Design Organization:			Date:	
Technology being evaluated:				
Critical parameters that control function:				
Parameter	Functions Controlled	Operating Latitude	Sensitivity	Failure Modes
Does hardware/software exist that demonstrates the above? (Attach photos or drawings)				
Describe the processes used to manufacture the technology:				
Is the technology controllable throughout the product's life cycle?				
Team member:		Prepared by:		
Team member:		Checked by:		
Team member:		Approved by:		
Team member:				
<i>The Mechanical Design Process</i>		Designed by Professor David G. Ullman		
Copyright 2008, McGraw-Hill		Form # 12.0		

Figure 8.4 Technology readiness assessment.

Often, if these questions are not answered in the positive, a consultant or vendor can be added to the team to help. This is especially true for manufacturing technologies for which the design engineer cannot possibly know all the methods available to manufacture a product. In general, negative answers to these questions may imply that this is a research project not a product development project. This realization may have an impact on the project plan as research takes longer than design. A technology readiness assessment template, Fig. 8.4, can be used for this assessment.

8.5 THE DECISION MATRIX—PUGH'S METHOD

In Chap. 4, we introduced Benjamin Franklin's decision-making method to help choose which projects to undertake. He suggested itemizing the pros and cons

when a choice needs to be made, and then using a process of elimination to decide which way to go. The same methodology can be used here to evaluate concepts one at a time. A big difference here is that we may have many concepts, we have already developed criteria with the QFD, and we may have a mix of qualitative and quantitative evaluations. In this section, a method to handle this additional complexity is developed.

The *decision-matrix method*, or *Pugh’s method*, is fairly simple and has proven effective for comparing alternative concepts. The basic form for the method is shown in Fig. 8.5. In essence, the method provides a means of scoring each alternative concept relative to the others in its ability to meet the criteria. Comparison of the scores in this manner gives insight to the best alternatives and useful information for making decisions. (In actuality, this technique is very flexible and is easily used in other, nondesign situations—such as which job offer to accept, which car to buy, or as in Table 4.2, which project to undertake.)

The decision-matrix method is an iterative evaluation method that tests the completeness and understanding of criteria, rapidly identifies the strongest alternatives, and helps foster new alternatives. This method is most effective if each member of the design team performs it independently and the individual results are then compared. The results of the comparison lead to a repetition of the technique, with the iteration continuing until the team is satisfied with the results. As shown in Fig. 8.5, there are six steps to this method. These steps refine the decision-making steps shown in Fig. 8.2.

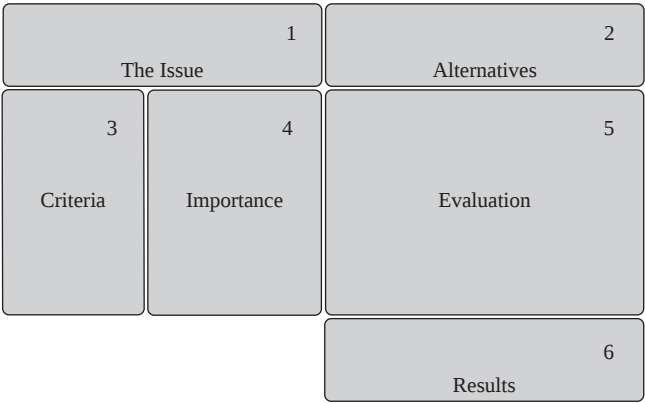
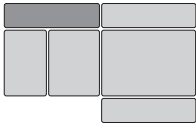


Figure 8.5 The basic structure of a Decision Matrix.

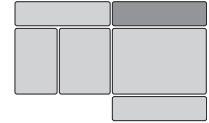
Decision matrices can be easily managed on the computer using a common spreadsheet program. Using a spreadsheet allows for easy iteration and comparison of team members’ evaluations.

The Decision Matrix is completed in six steps.

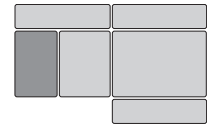
Step 1: State the Issue. The issue is not always obvious, but here it is clearly “Choose a concept for continued development.”



Step 2: Select the Alternatives to Be Compared. The alternatives to be compared are the different ideas developed during concept generation. It is important that all the concepts to be compared be at the same level of abstraction and in the same language. This means it is best to represent all the concepts in the same way. Generally, a simple sketch is best. In making the sketches, ensure that knowledge about the functionality, structure, technologies needed, and manufacturability is at a comparable level in every figure.



Step 3: Choose the Criteria for Comparison. First, it is necessary to know the basis on which the alternatives are to be compared with each other. Using the QFD method in Chap. 6, an effort was made to develop a full set of customer requirements for a design. These were then used to generate a set of engineering requirements and targets that will be used to ensure that the resulting product will meet the customer requirements. However, the concepts developed in Chap. 7 might not be refined enough to compare with the engineering targets for evaluation.



If they are not, we have a mismatch in the level of abstraction and use of the engineering targets must wait until the concept is refined to the point that actual measurements can be made on the product designs. Usually the basis for comparing the design concepts is a mix of customer requirements and engineering specifications, matched to the level of fidelity of the alternatives.

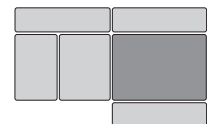
If the customers' requirements have not been developed, then the first step should be to develop criteria for comparison. The methods discussed in Chap. 6 should help with this task.

Additionally, the technology readiness measures can also help with evaluation here. This is especially true if the alternatives are dependent on new technologies.

Step 4: Develop Relative Importance Weightings. In step 3 of the QFD method (Section 6.4) there is a discussion of how to capture the relative importance of the criteria. The methods developed there can be used here to indicate which of the criteria are more important and which are less important. It is often worthwhile to measure the relative importance for different groups of customers, as discussed in Section 6.4.



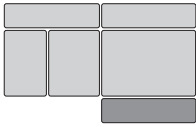
Step 5: Evaluate Alternatives. By this time in the design process, every designer has a favorite alternative; one that he or she thinks is the best of the concepts that have yet to be developed. This concept is used as a *datum*, all other designs being compared with it as measured by each of the customer requirements. If the problem is for the redesign of an existing product, then the existing product, abstracted to the same level as the concepts, can be used as the datum.



For each comparison, the concept being evaluated is judged either better than, about the same as, or worse than the datum. If it is better than the datum, the concept is given a + score. If it is judged to be about the same as the datum or if there is some ambivalence, an S ("same") is used. If the concept does not meet the criterion as well as the datum does, it is given a – score. If the Decision Matrix is on a spreadsheet use +1, 0, –1 for scoring.

Note that if it is impossible to make a comparison to a design requirement, more information must be developed. This may require more analysis, further experimentation, or just better visualization. It may even be necessary to refine the design, through the methods to be described in Chaps. 9–11 and then return to make the comparison. Note that the frailty in doing this step is the topic of Sections 8.6 and 8.7.

In using the Decision Matrix there are two possible types of comparisons. The first type is *absolute* in that each alternative concept is directly (i.e., absolutely) compared with some target set by a criterion. The second type of comparison is *relative* in that alternative concepts are compared with each other using measures defined by the criteria. In choosing to use a datum the comparison is relative. However, many people use the method for absolute comparisons. Absolute comparisons are possible only when there is a target. Relative comparisons can be made only when there is more than one option.



Step 6: Compute the Satisfaction and Decide What to Do Next. After a concept is compared with the datum for each criterion, four scores are generated: the number of plus scores, the number of minus scores, the overall total, and the weighted total. The overall total is the difference between the number of plus scores and the number of minus scores. This is an estimate of the decision-makers' satisfaction with the alternative. The weighted total can also be computed. This is the sum of each score multiplied by the importance weighting, in which an S counts as 0, a + as +1, and a – as –1. Both the weighted and the unweighted scores must not be treated as absolute measures of the concept's value; they are for guidance only. The scores can be interpreted in a number of ways:

- If a concept or group of similar concepts has a good overall total score or a high + total score, it is important to notice what strengths they exhibit, that is, which criteria they meet better than the datum. Likewise, groupings of scores will show which requirements are especially hard to meet.
- If most concepts get the same score on a certain criterion, examine that criterion closely. It may be necessary to develop more knowledge in the area of the criterion in order to generate better concepts. Or it may be that the criterion is ambiguous, is interpreted differently by different members of the team, or is unevenly interpreted from concept to concept. If the criterion has a low importance weighting, then do not spend much time clarifying it. However, if it is an important criterion, effort is needed either to generate better concepts or to clarify the criterion.
- To learn even more, redo the comparisons, with the highest-scoring concept used as the new datum. This iteration should be redone until a clearly "best" concept or concepts emerge.

After each team member has completed this procedure, the entire team should compare each member's individual results. The results can vary widely, since neither the concepts nor the requirements may be refined. Discussion among the members of the group should result in a few concepts to refine. If it does

not, the group should clarify the criteria or generate more concepts for evaluation.

Using the Decision Matrix: The MER Wheel

The Decision Matrix in Fig. 8.7 is completed for the MER wheel, step by step.

Step 1: State the issue. Choose a wheel configuration to develop for the MER.

Step 2: Select the alternatives to be compared. The ideas to be compared are shown in Fig. 8.6.

For this example, the concepts are fairly refined in that wheels were rendered in a CAD system. The same conclusion could have been reached without these solid models, but JPL engineers had the capability to make them and needed the images to present to management. The first wheel is from an earlier concept and was used as the baseline. The *cantilevered beam* design uses eight spokes as cantilever springs. One of the design goals, as described in the next step, is to build a spring into the wheel design. The *hub switchbacks* makes the spring element longer by making the radial section of the wheel a “W” shape—a set of switchbacks. The final idea shown uses spiral spokes to get more length and a better spring rate.

A fifth alternative is included in the Decision Matrix (Fig. 8.7) that is not included in Fig. 8.6, *multipiece*. This idea is to assemble the wheel out of multiple parts. This idea is nowhere near as refined as the others are, and, thus, it is hard to compare to them on the Decision Matrix. This difficulty will be readdressed in Section 8.7.

Step 3: Choose the criteria for comparison. JPL had four basic criteria for choosing a concept:

- Mass efficiency—the estimated weight of the wheel. This was easy to get from the solid model, at least to the accuracy of that model.
- Manufacturability—the ease with which the wheel can be made. This was estimated by a manufacturing expert, but detailed work was needed to get much accuracy here.

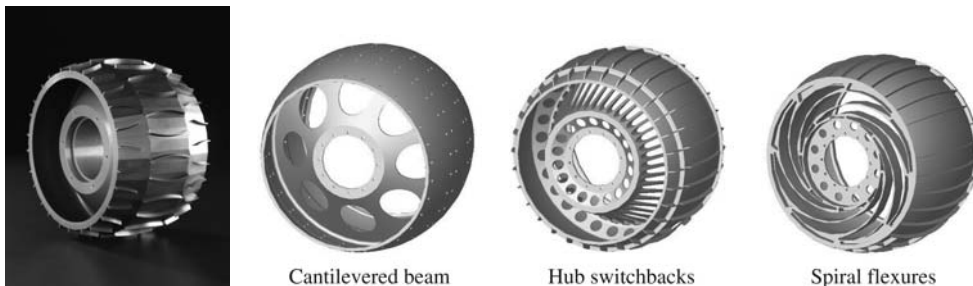


Figure 8.6 All the concepts shown are designed to be milled out of solid blocks of aluminum.

Issue: Choose a MER wheel configuration		Baseline	Cantilevered Beam	Hub Switchbacks	Spiral Flexures	Multipiece
Mass efficiency	35	Datum	0	0	1	?
Manufacturability	10		0	-1	-1	?
Available internal wheel volume	20		1	1	1	?
Stiffness	35		1	1	1	?
Total			2	1	2	?
Weighted total			55	45	80	?

Figure 8.7 MER wheel Decision Matrix.

- Available internal wheel volume—an estimate of the space inside the wheel that can be used for the motor and transmission. This too was easily estimated for the solid model.
- Stiffness 2500 lb/in.—the springiness of the wheel. This was needed to protect the electronic equipment as the Rover went over bumps. It was estimated using strength of materials equations.

Step 4: Develop relative importance weightings. At first, the engineers at JPL assumed all four criteria were equally important. Later they decided that mass efficiency and stiffness were most important. These weights are reflected in Fig. 8.7. The relative weights are shown as percentages totaling 100%.

Step 5: Evaluate alternatives. All the alternatives were compared relative to the datum using the 0 to denote “the same,” 1 equals “better than,” and -1 equals “worse than.”

Step 6: Compute the satisfaction and decide what to do next. From the totals (unweighted results) it is not very clear which configuration is best, but the weighted results show that the Spiral Flexures alternative is best. The matrix suggests that methods to simplify manufacturing should be explored, but this is not as important as the other criteria. The Spiral Flexure case can now be used as a datum if other ideas are developed.

8.6 PRODUCT, PROJECT, AND DECISION RISK

One of the goals when designing a product is to minimize risk. Sometimes this is stated explicitly, other times it goes unsaid. To better manage risk we need to refine exactly what we mean by the term “risk.” There are three types of risk that must be addressed during product development: product risk, project risk, and decision risk. Usually engineers are concerned only with product risk—the risk that the product will fail and potentially hurt someone or something. But, this

Risk is uncertainty falling on you.

view is too narrow. Beyond the risk of the product failing, there is the risk of the project failing to meet its goals, or being behind schedule or over budget. Further, there is the risk, especially during concept development, that a poor decision will be made. In this section, we will address all three types of risks beginning with product safety, liability, and risk.

Before doing so, we need a consistent definition of risk. Formally, *risk is an expected value, a probability that combines the likelihood of something happening times the consequences of it happening*. Thus, risk depends on the answer to three questions:

1. What can go wrong?
2. How likely is it to happen?
3. What are the consequences of it happening?

Keep these three questions in mind in the following sections.

Risk is a direct function of uncertainty. Some uncertainty is just part of nature, and you cannot control it (the weather, material and manufacturing variations, etc). During conceptual design, however, much of the uncertainty is because of a lack of knowledge. If everything is known precisely, then you can design a product with little or no risk. Unfortunately, incomplete knowledge, low-fidelity simulation results, manufacturing and material variations, and unknowable acts of god all contribute to risk. We begin the following sections with a product risk focus and then move to process and decision risk.

Much uncertainty is of no consequence, it has no discernable effect on operation of a product. When it does, then there is a risk. Whether this risk is worthy of design attention is a key determination of product quality.

8.6.1 Product Safety, the Goal of Product Risk Understanding

One area of product understanding that is often overlooked until late in the project is product safety. It is valuable to consider both safety and the engineer's responsibility for it, as safety is an integral part of human-product interaction and greatly affects the perceived quality of the product. Safety is best thought of early in the design process and thus is covered here. Formal failure analysis will be discussed in Chap. 11.

A safe product will not cause injury or loss. Two issues must be considered in designing a safe product. First, who or what is to be protected from injury or loss during the operation of the product? Second, how is the protection actually implemented in the product?

The main consideration in design for safety is the protection of people from injury by the product. Beyond concerns for humans, safety includes concern for

the loss of other property affected by the product and the product's impact on the environment in case of failure. Neglect in ensuring the safety of any of these objects may lead to a dangerous and potentially litigious situation. Concern for affected property means considering the effect the product can have on other devices, either during normal operation or during failure. For example, the manufacturer of a fuse or circuit breaker that fails to cut the current flow to a device may be liable because the fuse did not perform as designed and caused loss of or injury to another product.

There are three ways to establish product safety. The first way is to design safety directly into the product. This means that the device poses no inherent danger during normal operation or in case of failure. If inherent safety is impossible, as it is with most rotating machinery, some electronics, and all vehicles, then the second way to design in safety is to add protective devices to the product. Examples of added safety devices are shields around rotating parts, crash-protective structures (as in automobile body design), and automatic cut-off switches, which automatically turn a device off (or on) if there is no human contact. The third, and weakest, form of design for safety is a warning of the dangers inherent in the use of a product (Fig. 8.8). Typical warnings are labels, loud sounds, or flashing lights.

It is always advisable to design-in safety. It is difficult to design protective shields that are foolproof, and warning labels do not absolve the designer of

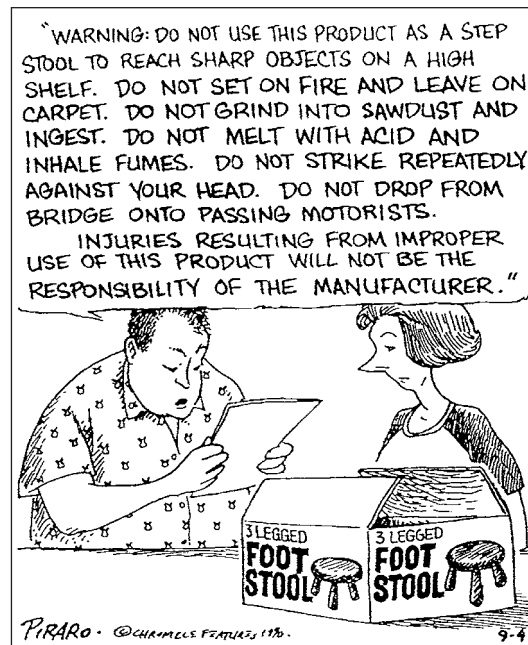


Figure 8.8 Example of the use of a warning label. (Bizarro (New) © Dan Piraro. King Features Syndicate. Reprinted by permission.)

The problem with designing something completely foolproof is to underestimate the ingenuity of a complete fool.
—Douglas Adams

liability in case of an accident. The only truly safe product is one with safety designed into it.

8.6.2 Products Liability, the Result of Poor Risk Understanding

Products liability is the special branch of law dealing with alleged personal injury or property or environmental damage resulting from a defect in a product. It is important that design engineers know the extent of their responsibility in the design of a product. If, for example, a worker is injured while using a device, the designers of the device and the manufacturer may be sued to compensate the worker and the employer for the losses incurred.

A products liability suit is a common legal action. Essentially, there are two sides in such a case, the plaintiff (the party alleging injury and suing to recover damages) and the defense (the party being sued).

Technical experts, professional engineers licensed by the state, are retained by both plaintiff and defense to testify about the operation of the product that allegedly caused the loss. Usually the first testimony developed by the experts is a technical report supplied to the respective attorney. These reports contain the engineer's expert opinion about the operation of the device and the cause of the situation resulting in the lawsuit. The report may be based on an onsite investigation, on computer or laboratory simulations, or on an evaluation of design records. If this report does not support the case of the lawyer who retained the technical expert, the suit may be dropped or settled out of court. If the investigations support the case, a trial will likely ensue and the technical expert may then be called as an expert witness.

During the trial, the plaintiff's attorney will try to show that the design was defective and that the designer and the designer's company were negligent in allowing the product to be put on the market. Conversely, the defense attorney will try to show that the product was safe and was designed and marketed with "reasonable care," as in Fig. 8.8.

Three different charges of negligence can be brought against designers in products liability cases:

The product was defectively designed. One typical charge is the failure to use state-of-the-art design considerations. Other typical charges are that improper calculations were made, poor materials were used, insufficient testing was carried out, and commonly accepted standards were not followed. In order to protect themselves from these charges, designers must

- Keep good records to show all that was considered during the design process. These include records of calculations made, standards considered, results of tests, and all other information that demonstrates how the product evolved.
- Use commonly accepted standards when available. “Standards” are either voluntary or mandatory requirements for the product or the workplace; they often provide significant guidance during the design process.
- Use state-of-the-art evaluation techniques for proving the quality of the design before it goes into production.
- Follow a rational design process (such as that outlined in this book) so that the reasoning behind design decisions can be defended.

The design did not include proper safety devices. As previously discussed, safety is either inherent in the product, added to the product, or provided by some form of warning to the user. The first alternative is definitely the best, the second is sometimes a necessity, and the third is the least advisable. A warning sign is not sufficient in most products liability cases, especially when it is evident that the design could have been made inherently safe or shielding could have been added to the product to make it safe. Thus, it is essential that the design engineers foresee all reasonable safety-compromising aspects of the product during the design process.

The designer did not foresee possible alternative uses of the product. If a man uses his gas-powered lawn mower to trim his hedge and is injured in doing so, is the designer of the mower negligent? Engineering legend claims that a case such as this was found in favor of the plaintiff. If so, was there any way the designer could have foreseen that someone was actually going to pick up a running power mower and turn it on its side for trimming the hedge? Probably not. However, a mower should not continue to run when tilted more than 30° from the horizontal because, even with its four wheels on the ground, it may tip over at that angle. Thus, the fact that a mower continues to run while tilted 90° certainly implies poor design. Additionally, this example also shows us that not all trial results are logical and that products must be “idiot-proof.”

Other charges of negligence that can result in litigation that are not directly under the control of the design engineer are that the product was defectively manufactured, the product was improperly advertised, and instructions for safe use of the product were not given.

8.6.3 Measuring Product Risk

Because safety is such an important concern in military operations, the armed services have a standard—MIL-STD 882D, *Standard Practice for System Safety*—focused specifically on ensuring safety in military equipment and facilities. This document gives a simple method for dealing with any *hazard*, which is defined as a situation that, if not corrected, might result in death, injury, or illness to personnel or damage to, or loss of, equipment (What can go wrong?). MIL-STD

882D defines two measures of a hazard: the likelihood or frequency of its occurrence (How likely is it to happen?) and the consequence if it does occur (What are the consequences of it happening?). Five levels of *mishap probabilities* are given in Table 8.2 ranging from “improbable” to “frequent.” Table 8.3 lists four categories of the *mishap severity*. These categories are based on the results expected if the mishap does occur. Finally, in Table 8.4 frequency and consequence of recurrence are combined in a mishap assessment matrix. By considering the level of the frequency and the category of the consequence, a hazard-risk index is found. This index gives guidance for how to deal with the hazard.

For example, say that during the design of the power lawn mower, the possibility of using the mower as a hedge trimmer was indeed considered. Now, what action should be taken? First, using Table 8.2, we decide that the mishap probability is either remote (D) or improbable (E). Most likely, it is improbable. Next, using Table 8.3, we rate the mishap severity as critical, category II, because severe injury may occur. Then, using the mishap assessment matrix, Table 8.4, we find an index of 10 or 15. This value implies that the risk of this mishap is acceptable, with review. Thus, the possibility of the mishap should not be dismissed without review by others with design responsibility. If the potential for seriousness of injury had been less, the mishap could have been dismissed without further concern. The very fact that the mishap was considered, an analysis was performed according to accepted standards, and the concern was documented might sway the results of a products liability suit.

Paying attention to the risk early is vital. Later, as the product is refined we will make use of this method in a more formal way as part of a Fail Modes and Effects Analysis (FMEA) Section 11.6.1.

Obviously many things can happen that can cause a hazard. It is the job of the designer to foresee these and make decisions that, as best as is possible, eliminates their potential.

8.6.4 Project Risk

Project risk is the effort to identify:

What can happen (What can go wrong?) that will cause the project to
Fall behind schedule, go over budget, or not meet the engineering specifications (What are the consequences of it happening?)
And the probability of it happening (How likely is it to happen?).

Project risks are caused by many factors:

- A technology is not as ready as anticipated—It may take longer than expected to develop the product. The higher the uncertainty in the technology (the lower the technology readiness (Section 8.4), the higher the risk to the project.
- Simulations or tests show unexpected results—The technology was not as well understood as initially thought, really a case of poor estimation of technology readiness.

Table 8.2 The mishap probabilities

Description	Level	Individual item	Inventory
Frequent	A	Likely to occur frequently (probability of occurrence > 10%)	Continuously experienced.
Probable	B	Will occur several times in life of an item (probability of occurrence = 1–10%)	Will occur frequently.
Occasional	C	Likely to occur sometime in life of an item (probability of occurrence = 0.1–1%)	Will occur several times.
Remote	D	Unlikely, but possible to occur in life of an item (probability of occurrence = 0.001–0.1%)	Unlikely, but can reasonably be expected to occur.
Improbable	E	So unlikely that it can be assumed that occurrence may not be experienced (probability of occurrence < 0.0001%)	Unlikely to occur, but possible.

Table 8.3 The mishap severity categories

Description	Category	Mishap definition
Catastrophic	I	Death, system loss, or severe environmental damage
Critical	II	Severe injury, occupational illness, major system damage, or reversible environmental damage
Marginal	III	Minor injury, minor occupational illness, minor system damage, or environmental damage
Negligible	IV	Less than minor injury, occupational illness, system damage, or environmental damage

Table 8.4 The mishap-assessment matrix

Frequency of occurrence	Hazard category			
	I Catastrophic	II Critical	III Marginal	IV Negligible
A. Frequent	1	3	7	13
B. Probable	2	5	9	16
C. Occasional	4	6	11	18
D. Remote	8	10	14	19
E. Improbable	12	15	17	20
Hazard-risk Index		Criterion		
1–5		Unacceptable		
6–9		Undesirable		
10–17		Acceptable with review		
18–20		Acceptable without review		

Source for Tables 8.2–8.4: MIL-STD 882D.

- A material or process is not available—Something that was thought to be usable in the product is not, or at least not at the price and time anticipated.
- Management changes the level of effort or personnel on the project—Fewer or different people are assigned to the project
- A vendor or other project fails to produce as expected—Most projects are dependent on the success of other efforts. If they don't produce on budget, on time, or with the performance expected, it may affect the project.

Of these causes of risk, the design engineer has control of the first three. Poor choices made about the technologies, materials, and process used may be the result of poor decision-making practice.

8.6.5 Decision Risk

Decision-making risks are the chance that choices made will not turn out as expected (What can go wrong?). In business and technology, you only know if you made a bad decision sometime in the future. Since decisions are calls to action and commitment of resources, it's only after the actions are taken that you really know whether the decision was a good one or a bad one.

Decision-making risk is a measure of the probability that a poor decision has been made (How likely is it to happen?) times the consequences of the decision (What are the consequences of it happening?). The goal is to understand the probabilities and consequences during the decision-making process and not have to wait until later, after the action has been taken.

Looking back at the Decision Matrix:

- What can go wrong? = A criterion is not met.
- What are the consequences of it happening? = The customer is not satisfied.
- How likely is it to happen? = It depends on the uncertainty. There is no real measure of uncertainty in the Decision Matrix.

One relatively recent method for managing uncertainty during decision making is called Robust Decision Making. It is introduced in Section 8.7.

8.7 ROBUST DECISION MAKING

The great challenge during conceptual design evaluation is to make good decisions in spite of the fact that the information about the concepts is uncertain, incomplete, and evolving. Recent methods have been developed that are especially designed to manage these types of decision problems. These methods are referred to as robust decision-making methods. The word “robust” will be used again in Chap. 10 to refer to final products that are of high quality because they are insensitive to manufacturing variation, operating temperature, wear, and other uncontrolled factors. Here we use the term “robust” to refer to decisions that are as insensitive as possible to the uncertainty, incompleteness, and evolution of the information that they are based on.

All decisions are based on incomplete, inconsistent,
and conflicting information.

To set the stage for this, reconsider the Decision Matrix. Instead of using the 0, +1, -1 scale, you could refine it by using measureable values. The stiffness of each alternative could be modeled in terms of N/m (lb/in), the mass efficiency in terms of kg, the internal wheel volume in terms of mm³ (in³), and manufacturability in terms of the time to mill each wheel. Then these values could be combined in some fashion (they are all in different units) to generate a measure for each alternative (we will revisit this in Chap. 10). The problem is that it will take significant time to develop these values for each concept.

In fact, many hundreds of hours went into developing the solid models shown in Fig. 8.6. Could JPL have made the decision without refining the wheel ideas to that level? The modeling JPL did was well beyond what most organizations can invest to make concept decisions. So this raises the question, How do you make concept decisions when the information you have is uncertain and incomplete? Or, looking back at the Decision Matrix in Fig. 8.7, How do you include the more abstract idea of a multipiece concept in the Decision Matrix?

To begin we will refine the Decision Matrix a little. The score or total values produced in the Decision Matrix are measures of satisfaction, where *satisfaction* = *belief that an alternative meets the criteria*. Thus, the decision-maker's satisfaction with an alternative is a representation of the belief in how well the alternative meets the criteria being used to measure it. For example, say the criterion for the mass of a MER wheel is 1 kg. You weigh it on a scale you know to be accurate and convert the reading to mass. If you find the mass to be 1 kg, then you would be very satisfied with the object relative to the mass criterion. However, what if the accuracy of the scale was suspect or you were uncertain that the reading was correct? Even though the scale gives you 1 kg, your satisfaction drops because you are uncertain about the accuracy of your reading. Or, what if the concept is only a sketch on a piece of paper and you calculate the mass to be 1 kg. You know this to be uncertain because it was based on incomplete and evolving information, and so your belief that the final object will be 1 kg is not very high. The point here is that regardless of how the evaluation information is developed, it is your belief that is important.

So then, what is "belief?" The dictionary definition of *belief* includes the statement "*a state of mind in which confidence is placed in something.*" A "state of mind" during decision making refers to the decision-maker's *knowledge* and her *confidence* in the result of evaluation of the alternative ("something") compared to the criteria targets. Thus, for our purposes, belief is redefined as

Belief = Confidence placed in an alternative's ability to meet a criterion,
requirement, or specification, based on current knowledge

To further support this concept, if someone hands you an object and says, "I believe that this has 1 kg mass," you might ask, "How do you know?" (a query about

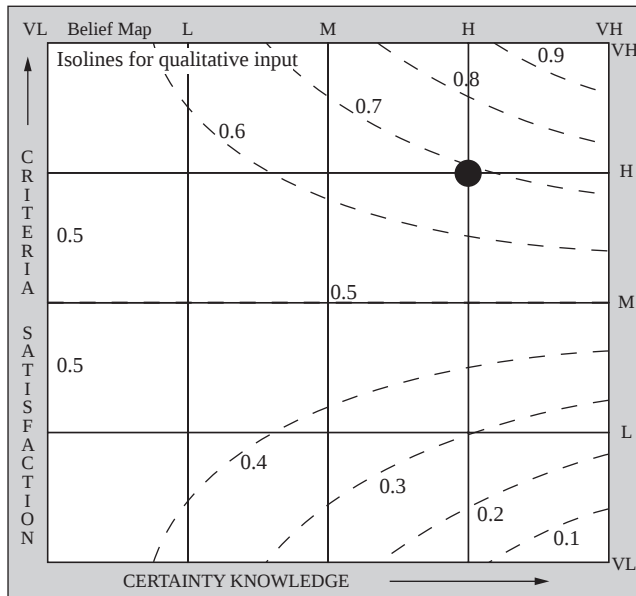


Figure 8.9 A Belief Map.

their knowledge) or “How close to 1 kg is it?” (a query about their confidence in the value).

This virtual sum of knowledge and confidence can be expressed on a *Belief Map*. A Belief Map is a tool to help picture and understand evaluation. A Belief Map organizes the two dimensions of belief: knowledge (or certainty) and criteria satisfaction (Fig. 8.9). For a complete evaluation of an issue, there will be a Belief Map for each alternative/criterion pair corresponding to each cell in a Decision Matrix. By using a Belief Map, the influence of knowledge on the result can be easily found and, as we shall see, the use of Belief Maps can help develop team consensus.

To explain Belief Maps, we will first describe the axes, then the point and finally the lines labeled 0.1–0.9. On the vertical axis of a belief map, we plot the **Level of Criterion Satisfaction**, the probability that the alternative meets the (often unstated) criterion target, or the yes-ness of the alternative. Consider the problem of selecting a MER wheel. Say all we have for the spiral flexure concept is a sketch (Fig. 8.10a) and some rough calculations. The best we can say is that “yes, this concept appears to have high mass efficiency” or “no, it seems to have low manufacturability.” This is similar to what we indicated by the +1 and –1 in the Decision Matrix.

The horizontal axis of the Belief Map is the **Level of Certainty**. This is not commonly measured, yet it is key to understanding belief and decision-making. Think of the Level of Certainty as a probability that ranges from 50% to 100%. The rationale for this is that a certainty of 50% is no better than the flip of a coin, very low—the probability is 50–50 that the evaluation is correct. At the other end

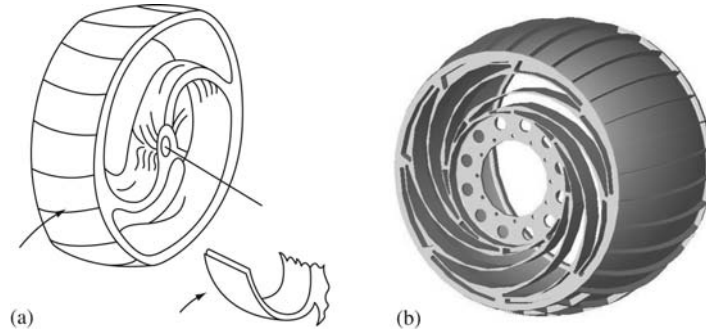


Figure 8.10 Sketch of the MER wheel from Fig. 2.5.

The odds are greatly against your being immensely more knowledgeable than everyone else is.

of the scale, a probability of 100% implies that the evaluation is a sure thing; certainty is very high and the Level of Criterion Satisfaction is a good assessment of the situation.

To better understand Belief Maps, say you are evaluating the manufacturability of the spiral flexure wheel and all you have so far is the above sketch. In making this evaluation you put a point on the Belief Map. If you put your point in the upper right corner as shown in Fig. 8.11, you are claiming that your certainty is very high and you are confident that the Spiral wheel is easy to manufacture [yes, the ability to be manufactured is very high (VH on the belief map)]. Thus, you 100% believe that the Spiral is manufacturable. If you put your point in the lower right corner, at VL on the criterion satisfaction scale, you have high certainty that it is not easy to manufacture. You believe that the Spiral concept has a zero probability of meeting this criterion.

If you put your evaluation point in the upper left corner, you are hopelessly optimistic: “I don’t know anything about this, but I am sure it is easy to manufacture.” This evaluation is no better than flipping a coin, so belief = 50%. If you put your evaluation point in the lower left corner then you believe that the Spiral flexures concept can’t meet the manufacturability criterion, even though you have no knowledge on which to base this belief. This is called the “Eyore corner,” after the character in A.A. Milne’s “Winnie the Pooh,” who thought everything was going to turn out bad no matter how little he knew. This evaluation is also no better than flipping a coin, so belief = 50%. In fact, the entire left border of the Belief Map has belief = 50%, as any point there is based on no certainty or knowledge at all.

If a JPL engineer puts his point anywhere with Level of Criterion Satisfaction = 50%, he is neutral in his evaluation. The Spiral is neither good nor bad in its

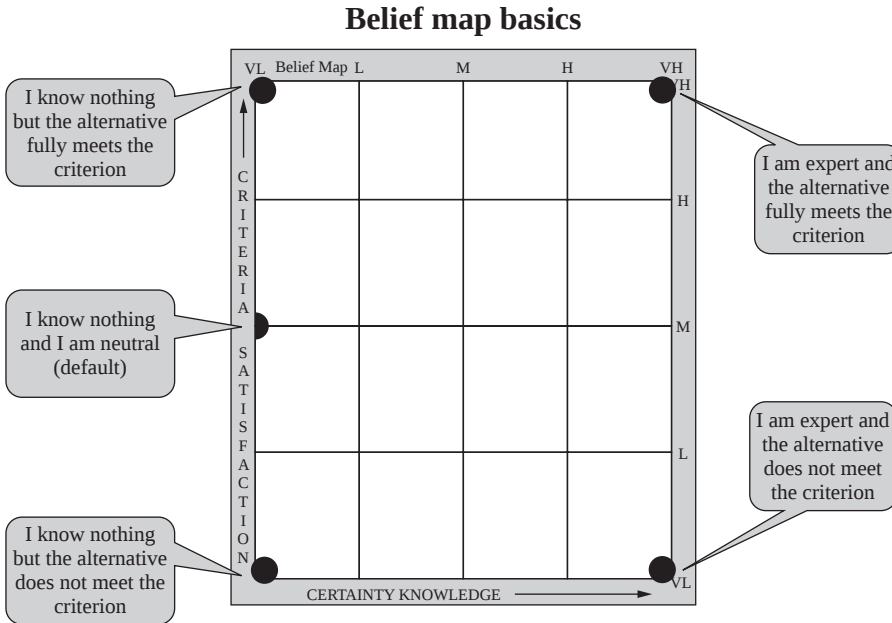


Figure 8.11 The four corners of the belief map.

manufacturability, consequently, regardless of his knowledge or Level of Certainty, his belief is 50%.

Finally, the default position for points on the Belief Map is the center left—you know nothing and you are neutral. A point placed here is the same as not offering any evaluation at all.

The lines on the Belief Map are called *Isolines*. They are belief represented as a probability. Thus, for the point in Fig. 8.9, the belief is 0.69. Note that if the evaluator who put the point on the Belief Map had very high certainty, the point was on the right, then his belief would be 0.75 and if the certainty was very low, Belief = 0.5, all the way over to the right.

The Belief Maps for the five MER wheel options are shown in Fig. 8.12. Assume that no analysis has been done and all the alternatives are sketches like Fig. 8.10a, at best.

The values from the Belief Maps have been entered in a Decision Matrix in Fig. 8.13. To be consistent with the Decision Matrix in Fig. 8.5, the baseline has been assumed 50% satisfactory for each criterion and the other evaluation made relative to it. This is not necessary for using Belief Maps.

The resulting satisfaction values for the alternatives differ from the weighted totals in the decision matrix in Fig. 8.13. This is expected as the evaluation here includes uncertainty. Also, now there is an evaluation for even the multi-piece alternative, but it is highly uncertain as it is only a rough concept. These

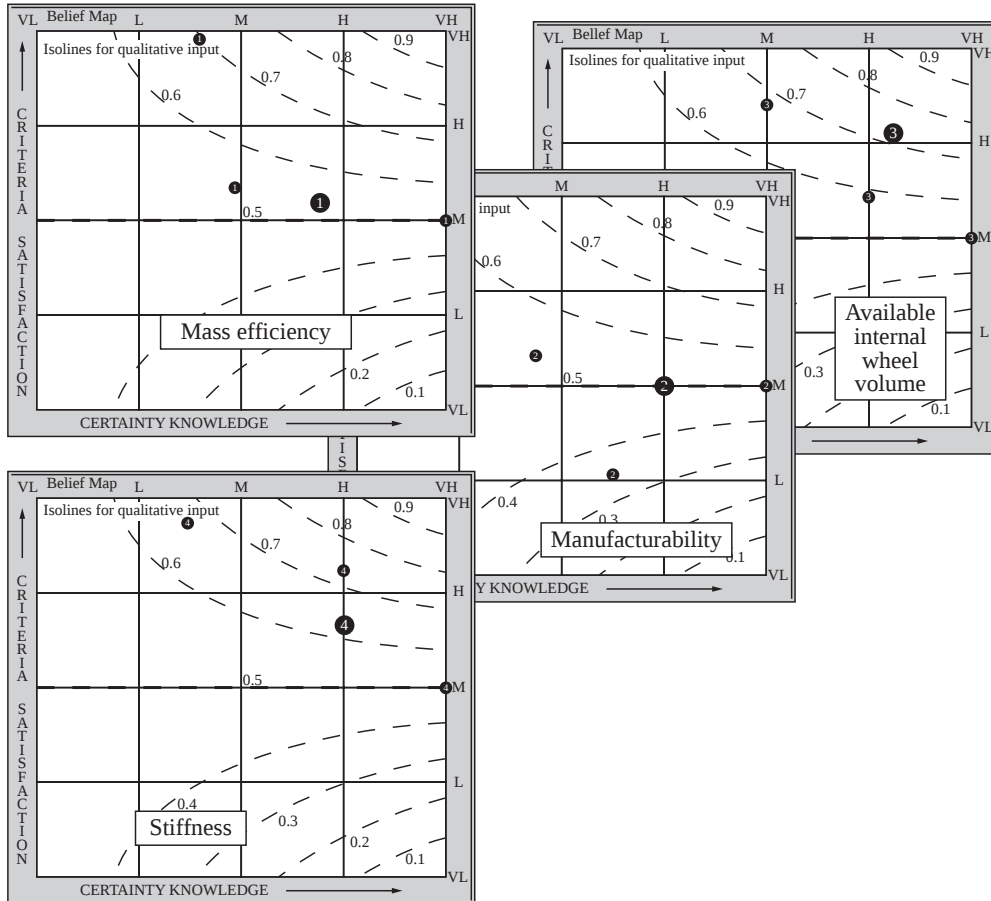


Figure 8.12 Belief Map example for the MER.

Issue:		Baseline	Cantilevered Beam	Hub Switchbacks	Spiral Flexures	Multipiece
Choose a MER wheel configuration						
Mass efficiency	35	0.5	0.55	0.55	0.77	0.71
Manufacturability	10	0.5	0.5	0.35	0.4	0.52
Available internal wheel volume	20	0.5	0.72	0.58	0.84	0.67
Stiffness	35	0.5	0.62	0.74	0.86	0.68
Satisfaction		50	60	60	78	67

Figure 8.13 Decision Matrix with Belief Map results.

satisfaction results still show that the Spiral Flexure alternative is best, but this could have been reached without the time of making a detailed CAD model and doing much analysis. Also, we can now see that the Multipiece alternative may be worth spending time to refine and reevaluate. Its satisfaction is second only to the Spiral.

Another use for Belief Maps is in building team consensus and buy-in. Multiple people putting dots on Belief Maps and comparing them can help ensure that the team is understanding the concepts and criteria in a consistent manner. See links in the Sources, Section 8.9, to learn more about Belief Maps, their use, and software that supports them.

8.8 SUMMARY

- The feasibility of a concept is based on the design engineer's knowledge. Often it is necessary to augment this knowledge with the development of simple models.
- In order for a technology to be used in a product, it must be ready. Six measures of technology readiness can be applied.
- Product safety implies concern for injury to humans and for damage to the device itself, other equipment, or the environment.
- Safety can be designed into a product, added on, or warned against. The first of these is best.
- A mishap assessment is easy to accomplish and gives good guidance.
- The decision-matrix method provides means of comparing and evaluating concepts. The comparison is between each concept and a datum relative to the customers' requirements. The matrix gives insight into strong and weak areas of the concepts. The decision-matrix method can be used for subsystems of the original problem.
- An advanced decision matrix method leads to robust decisions by including the effects of uncertainty in the decision making process.
- Belief maps are a simple yet powerful way to evaluate alternatives and work to gain team consensus.

8.9 SOURCES

- Pugh, S.: *Total Design: Integrated Methods for Successful Product Engineering*, Addison-Wesley, Wokingham, England, 1991. Gives a good overview of the design process and many examples of the use of decision matrices.
- Standard Practice for System Safety*, MIL-STD 882D, U.S. Government Printing Office, Washington, D.C., 2000. The mishap assessment is from this standard. <http://www.core.org.cn/NR/rdonlyres/Aeronautics-and-Astronautics/16-358JSystem-SafetySpring2003/79F4C553-BD79-4A0C-A87E-80F4B520257B/0/882b1.pdf>
- Sunar, D. G.: *The Expert Witness Handbook: A Guide for Engineers*, 2nd edition. Professional Publications, San Carlos, Calif., 1989. A paperback, it has details on being an expert witness for products liability litigation.

Ullman, D. G.: *Making Robust Decisions*, Trafford Publishing, 2006. Details on Belief Maps and robust decision-making. Software that supports the use of belief maps is available from www.robustdecisions.com. Its use is free to students.

8.10 EXERCISES

- 8.1 Assess your knowledge of these technologies by applying the six measures given in Section 8.4.
- Chrome plating
 - Rubber vibration isolators
 - Fastening wood together with nails
 - Laser positioning systems
- 8.2 Use a Decision Matrix or a series of matrices to evaluate the
- Concepts for the original design problem (Exercise 4.1)
 - Concepts for the redesign problem (Exercise 4.2)
 - The alternatives for a new car
 - The alternatives between various girlfriends or boyfriends (real or imagined)
 - The alternatives for a job
- Note that for the last three the difficulty is choosing the criteria for comparison.
- 8.3 Perform a mishap assessment on these items. If you were an engineer on a project to develop each of these items, what would you do in reaction to your assessment? Further, for hazardous items, what has industry or federal regulation done to lower the hazard?
- A manual can opener
 - An automobile (with you driving)
 - A lawn mower
 - A space shuttle rocket engine
 - An elevator drive system



8.11 ON THE WEB

A template for the following document is available on the book's website: www.mhhe.com/Ullman4e

- Technology Readiness

CHAPTER

Product Generation

KEY QUESTIONS

- What are the steps to turn an abstract concept into a quality product?
- What is a BOM?
- In what order should we consider *constraints*, *configuration*, *connections*, and *components* during the design of parts and assemblies?
- How can force flow help in the design of components?
- Who should make the parts you design?

9.1 INTRODUCTION

This chapter and Chaps. 10 and 11 focus on the product design phase, with the goal to refine the concepts into quality products. This transformation process could be called hardware design, shape design, or embodiment design, all of which imply giving flesh to what was the skeleton of an idea. As shown in Fig. 9.1, this refinement is an iterative process of generating products and evaluating them to verify their ability to meet the requirements. Based on the result of the evaluation, the product is patched and refined (further generation), then reevaluated in an iterative loop. Also, as part of the product generation procedure, the evolving product is decomposed into assemblies and individual components. Each of these assemblies and components requires the same evolutionary steps as the overall product. In product design, generation and evaluation are more closely intertwined than in concept design. Thus, the steps suggested for product generation here include some evaluation. In Chaps. 10 and 11, the product designs are evaluated for their performance, quality, and cost. Quality will be measured by the product's ability to meet the engineering requirements and the ease with which it can be manufactured and assembled.

The knowledge gained making the transformation from concept to product can be used to iterate back to the concept phase and possibly generate new concepts. The drawback, of course, is that going back takes time. The natural

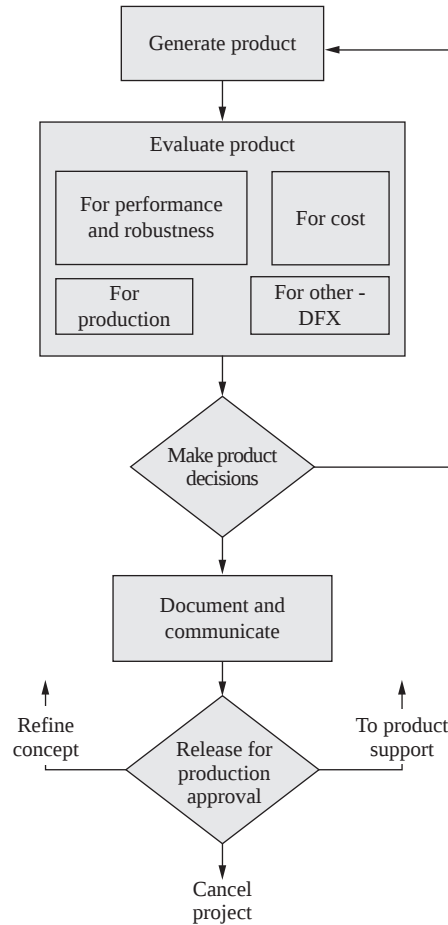
Product Development

Figure 9.1 The product design phase of the design process.

inclination to iterate back and change the concept must be balanced by the schedule established in the design plan.

In two situations design engineers begin at the product design phase in the design process. In the first of these situations, the concept may have been generated in a corporate research lab and then handed off to the design engineers to “productize.” It could be assumed, since the research lab had to develop working models to confirm the readiness of the technology that the design was well on its way to being a successful product at this point. However, the goal of the researchers is to demonstrate the viability of the technology; their working models are generally handcrafted, possibly held together with duct tape and bubble

gum. They probably incorporated very poor product design. The approach of forcing products to be developed from experimental prototypes is very weak. Design engineers, manufacturing engineers, and other stakeholders should have been involved in the process long before the concept was developed to this level of refinement.

In the second situation, the project involves a redesign. Many problems begin with an existing product that needs only to be redesigned to meet some new requirements. Often, only “minor modifications” are required, but these usually lead to unexpected, extensive rework, resulting in poor-quality products.

In either situation, whether the concept comes from a research lab or the project involves only a “simple redesign,” it is wise to ensure that the function and other conceptual design concepts are well understood. In other words, the techniques described in Chaps. 5, 6, 7, and 8 should be applied before the product design phase is ever begun. Only in that way will a good-quality product result.

Before describing the process of refining the concept to hardware, note that only enough detail on materials, manufacturing methods, economics, and the engineering sciences are developed to support techniques and examples of the design process. It is assumed that the reader has the knowledge needed in these areas.

The goal of this and Chaps. 10 and 11 is to transform the concepts developed in Chaps. 7 and 8 into products that perform the desired functions. These concepts may be at different levels of refinement and completeness. Consider the concept examples in Fig. 9.2. The stick-figure representation of a mechanism and a rather complete CAD solid model for a bicycle suspension concept from Marin Bicycles. The sketches are very different levels of abstraction. This is common of concepts and so, the steps for product development must deal with concepts at many varying levels of refinement.

Refining from concept to a manufacturable product requires work on all the elements shown in Fig. 9.3 (a refinement of Fig. 1.1). Central to this figure is the *function* of the product. Surrounding the function, and mutually dependent on each other, are the *form* of the product, the *materials* used to make the product,

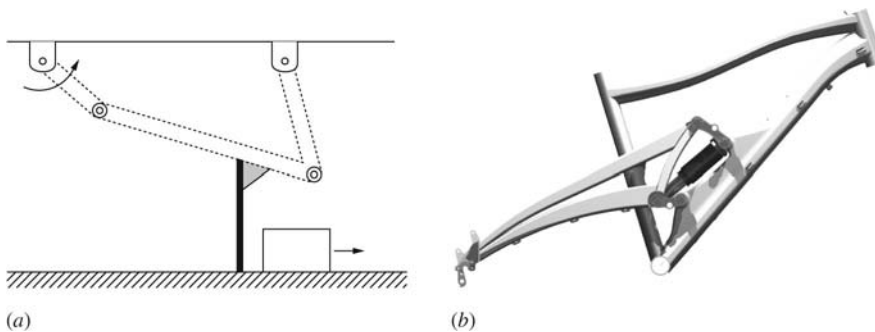


Figure 9.2 Typical concepts. (Reprinted with permission of Marin Bicycles.)

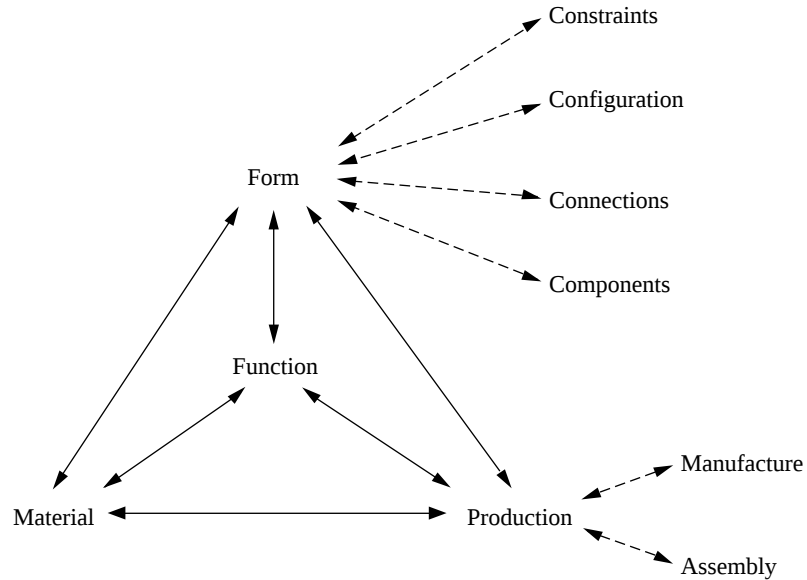


Figure 9.3 Basic elements of product design.

and the *production techniques* used to generate the form from the materials. Although these three may have been considered in conceptual design, the focus there was on developing function. Now, in product design, attention turns to developing producible forms that provide the desired function that are producible with materials that are available and can be controlled.

The form of the product is roughly defined by the spatial *constraints* that provide the envelope in which the product operates. Within this envelope the product is defined as a *configuration of connected components*. In other words, form development is the evolution of components, how they are configured relative to each other and how they are connected to each other. This chapter covers techniques used to generate these characteristics of form.

As shown in Fig. 9.3, decisions on production require development of how the product's components are *manufactured* from the materials and how these components are *assembled*. In general, the term “manufacture” refers to making individual components and “assembly” to putting together manufactured and purchased components. Simultaneous evolution of the product and the processes used to produce it is one of the key features of modern engineering. In this chapter, the interaction of manufacturing and assembly process decisions will affect the generation of the product. Production considerations will become even more important in evaluating the product (Chap. 11).

In the discussion of conceptual design, emphasis was put on developing the function of the product. It is now a reasonable question to ask what should be

worked on next—the form, the materials, or the production? The answer is not easy, because even though we work from function to form, form is hopelessly interdependent on the materials selected and the production processes used. Further, the nature of the interdependency changes with factors such as the number of items to be produced, the availability of equipment, and knowledge about materials and their forming processes. Thus, it is virtually impossible to give a step-by-step process for product design. Figure 9.3 shows all the major considerations in product generation. Sections 9.3–9.5 will begin with form generation and will then cover material and process selection. There is also a section on vendor development, because vendor issues affect product generation. In Chaps. 10 and 11 product evaluation will center on the product's ability to meet the functional requirements, ease of manufacture and assembly, and cost.

Before diving into the development of the product, it is necessary to introduce some basics on how product information is documented and managed.

9.2 BOMs

The *Bill Of Materials (BOM)*, or *parts list*, is like an index to the product. It evolves during this phase of the design process. BOMs are a key part of Product Life-cycle Management (PLM), as introduced in Chap. 1 (Figure 1.8). BOMs are often built on a spreadsheet, which is easy to update (a Word template can also be used). A typical bill of materials is shown in Fig. 9.4. To keep lists to a reasonable length, a separate list is usually kept for each assembly. There are a minimum of six pieces of information on a bill of materials:

1. *The item number or letter.* This is a key to the components on the BOM.
2. *The part number.* This is a number used throughout the purchasing, manufacturing, inventory control, and assembly systems to identify the component. Where the item number is a specific index to the assembly drawing, the part number is an index to the company system. Numbering systems vary greatly from company to company. Some are designed to have context, the part number indicates something about the part's function or assembly. These types of systems are hard to maintain. Most are simply a sequential number assigned to the part. Sometimes, the last digit will be used to indicate the revision number, as in the Fig. 9.4 example.
3. *The quantity needed in the assembly.*
4. *The name or description of the component.* This must be a brief, descriptive title for the component.
5. *The material from which the component is made.* If the item is a subassembly, then this does not appear in the BOM.
6. *The source of the component.* If the component is purchased, the name of the company is listed. If the component is made in-house, this line can be left blank.



Bill of Materials					
Product: Everlast					Date: 03/03/09
Assembly: Shock Assy					
Item #	Part #	Qty	Name	Material	Source
1	63172-2	1	Outer tube	1018 carbon steel	Coyote Steel
2	94563-1	1	Roller bearing		Bearings Inc.
3	..	..	..	..	..
4	..	..	..	..	..
9	74324-2	3	Shaft	304 stainless steel	Coyote Steel
10	44333-8	1	Link rubber	Urethane	Reed Rubber
Team member: Bob			Prepared by: Jan		
Team member: Jan			Checked by: Bob		
Team member:			Approved by: Dr. Roberts		
Team member:					Page 1/4
The Mechanical Design Process			Designed by Professor David G. Ullman		
Copyright 2008, McGraw-Hill			Form # 23.0		

Figure 9.4 Typical bill of materials.

Managing design information such as BOMs, drawings, solid models, simulations, and test results is a major undertaking in a company. In fact, this intellectual property is one of a company's most valuable assets. In past times indexing and finding information in this system was usually difficult and often impossible. As product information has become more computer based, so have methods to manage the information. Generally, BOMs are a part of the PLM system, and thus, the part numbers are linked to drawings, solid models, and other part and assembly information.

9.3 FORM GENERATION

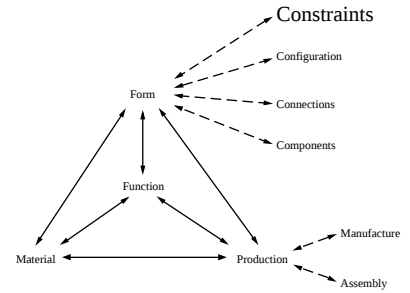
The goal of this section is to give form to the concepts that have been developed. Ideally, form grows from *constraints* with other assemblies and components. After the constraints for components are understood, then the *configuration*, or architecture, can be developed. Next, the *connections* or interfaces with other parts can be developed. These support the functions of the product. Finally, the *components* themselves can be developed. These four steps, although presented sequentially, obviously occur concurrently.

9.3.1 Understand the Spatial Constraints

The spatial constraints are the walls or envelope for the product. Most products must work in relation to other existing, unchangeable objects. The relationships may define actual contact or be for needed clearance. The relationships may be based on the flow of material, energy, or information as well as being physical. For the one-handed clamp the interface with work and the user hand is physical and there is the flow of energy in the form of forces.

Some spatial constraints are for functionally needed space, such as optical paths, or to clear or interfere with the flow of some material such as air or water. Further, most products go through a series of operational steps as they are used. The functional relationships and spatial requirements may change during these. The varying relationships may require the development of a series of layout drawings or solid models.

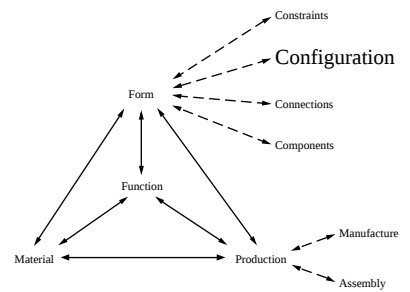
Initially the spatial constraints are for the entire product, system, or assembly; however, as design decisions are made on one assembly or component, other spatial constraints are added. For large products that have independent teams working on different subassemblies, the coordination of the spatial constraint information can be very difficult. PLM and solid modeling systems help in managing the constraints.



9.3.2 Configure Components

Configuration is the architecture, structure, or arrangement of the components and assemblies of components in the product. Developing the architecture or configuration of a product involves decisions that divide the product into individual components and develop the location and orientation of them. Even though the concept sketches probably contain representations of individual components, it is time to question the decomposition represented. There are only six reasons to decompose a product or assembly into separate components:

- Components must be separate if they need to move relative to each other. For example, parts that slide or rotate relative to each other have to be separate components. However, if the relative motion is small, perhaps elasticity can be built into the design to meet the need for motion. This is readily accomplished in plastic components by using elastic hinges, which are thin sections of fatigue-resistant material that act as a one-degree-of-freedom joint.
- Components must be separate if they need to be of different materials for functional purposes. For example, one area of the product may need to conduct heat and another must insulate and both these areas may be served by a single component, were it not for these thermal resistance needs.



- Components must be separate if they need to be moved for accessibility. For example, if the cabinet for a computer is made as one piece, it would not allow access to install and maintain the computer components.
- Components must be separate if they need to accommodate material or production limitations. Sometimes a desired part cannot be manufactured in the shape desired.
- Components must be separate if there are available standard components that can be considered for the product.
- Components must be separate if separate components would minimize costs. Sometimes it is less expensive to manufacture two simple components than it is to manufacture one complex component. This may be true in spite of the added stress concentrations and assembly costs caused by the interface between the two components.

These guidelines for defining the boundaries between components help define only one aspect of the configuration. Equally important during configuration design are the location and orientation of the components relative to each other. *Location* is the measure of components' relative position in x , y , z space. *Orientation* refers to the angular relationship of the components. Usually components can have many different locations and orientations; solid models help with the search for possibilities. Configuration design was introduced in Section 2.4.2 as a problem of location and orientation.

An important consideration in the design of many products is how quickly and cheaply other new products can be developed from them. Designing for use across many products is referred to as modularity or variant design. Where sets of common modules are shared among a product family, cost can be reduced and multiple product variants can be introduced. Consider the design of battery operated power tools or kitchen utensils that all share the same battery. Or, most car and truck manufacturers use common parts across many models.

A module is often defined as a system or assembly that is loosely coupled to the rest of the system. In the ideal world, each module fulfills a single or a small set of related functions as is true with the battery on a laptop computer—where the batteries' function is to store energy. Designing independent modules has many potential advantages:

- They can be used to create product families.
- They provide flexibility so that each product produced can meet the specific customer's needs.
- New technology can be developed without changes to the overall design, modules can be developed independently allowing for overlapped product development.
- They can lead to economies in parts sourcing—the single battery is used for many tools resulting in higher volume and subsequent lower cost.
- Modularity also eases the management of complex product architectures and therefore their development.



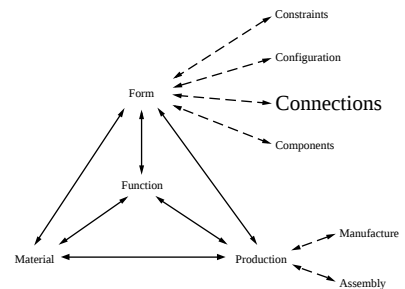
Figure 9.5 An example of integral architecture. (Reprinted with permission of Boeing.)

A pull in the opposite direction from a modular architecture is to design an integral architecture. Integral architectures have fewer parts with all the functions blurred together. An illustration is the Blended Wing Body (BWB) concept developed by Boeing, shown as a test vehicle in Fig. 9.5. In this design, the assignment of functionality between wing, fuselage, and empennage are blended. A traditional aircraft uses wings for lift generation, a fuselage for storage of passengers and cargo, the tail for pitch and yaw control. In the BWB, on the other hand, the integral “blended” body provides all three functions to some extent. This blending when compared to a traditional plane leads to 19% lower weight and 32% less projected fuel burn per passenger per mile flown.

9.3.3 Develop Connections: Create and Refine Interfaces for Functions

This is a key step when embodying a concept because *the connections or interfaces between components support their function and determine their relative positions and locations*. Here are guidelines to help develop and refine the interfaces between components:

- *Interfaces must always reflect force equilibrium and consistent flow of energy, material, and information.* Thus, they are the means through which the product will be designed to meet the functional requirements. Most design effort occurs at the connections between components, and attention to these interfaces and the flows through them, is key to product development. During the redesign of an existing product, it is useful to disassemble it; note the flows of energy, information,



Complexity occurs primarily at interfaces.

and materials at each joint; and develop the functional model one component at a time.

- *After developing interfaces with external objects, consider the interfaces that carry the most critical functions.* Unfortunately it is not always clear which functions are most critical. Generally, they are those functions that seem hardest to achieve (about which the knowledge is the weakest) or those described as most important in the customers' requirements.
- *Try to maintain functional independence in the design of an assembly or component.* This means that the variation in each critical dimension in the assembly or component should affect only one function. If changing a parameter changes multiple functions, then affecting one function without altering others may be impossible.
- *Exercise care when separating the product into separate components.* Complexity arises since one function often occurs across many components or assemblies and since one component may play a role in many functions. For example, a bicycle handlebar (discussed in Section 2.2) enables many functions but does none of them without other components.
- *Creating and refining interfaces may force decompositions that result in new functions or may encourage the refinement of the functional breakdown.*

As the interfaces are refined, new components and assemblies come into existence. One step in the evaluation of each potential embodiment is to determine how each new component changes the functionality of the design.

In order to generate the interface, it may be necessary to treat it as a new design problem and utilize the techniques developed in Chaps. 7 and 8. When developing a connection, classify it as one or more of these types:

- *Fixed, nonadjustable connection.* Generally one of the objects supports the other. Carefully note the force flow through the joint (see Section 9.3.4). These connections are usually fastened with rivets, bolts, screws, adhesives, welds, or by some other permanent method.
- *Adjustable connection.* This type must allow for at least one degree of freedom that can be locked. This connection may be field-adjustable or intended for factory adjustment only. If it is field-adjustable, the function of the adjustment must be clear and accessibility must be given. Clearance for adjustability may add spatial constraints. Generally, adjustable connections are secured with bolts or screws.
- *Separable connection.* If the connection must be separated, the functions associated with it need to be carefully explored.
- *Locator connection.* In many connections, the interface determines the location or orientation of one of the components relative to another. Care

Determine how constrained a component needs to be, and constrain it exactly that amount—no more, nor no less.

must be taken in these connections to account for errors that can accumulate in joints.

- *Hinged or pivoting connection.* Many connections have one or more degrees of freedom. The ability of these to transmit energy and information is usually key to the function of the device. As with the separable connections, the functionality of the joint itself must be carefully considered.

Connections directly determine the degrees of freedom between components and every interface must be thought of as constraining some or all of those degrees of freedom. Fundamentally, every connection between two components has six degrees of freedom—three translations and three rotations. It is the design of the connections that determines how many degrees of freedom the final product will have. Not thinking of connections as constraining degrees of freedom will result in unintended behavior. This discussion on two-dimensional constraints gives a good basis for thinking about connections.

If two components have a planar interface, the degrees of freedom are reduced from six to three, translation in the x and y directions (in both the positive and negative directions) and rotation (in either direction) about the z axis (Fig. 9.6). Putting a single fastener—like a bolt or pin—through component A into component B can only remove the translation degrees of freedom, but leaves rotation. Some novice designers think that tightening the bolt very tight will remove the rotational freedom, but even a slight torque around the z axis will cause A to rotate. Using two fasteners close together may not be sufficient to restrain part A from rotating, especially if the torque is high relative to the strength of the fasteners or the holes in A and B. Even more importantly, most joints need to

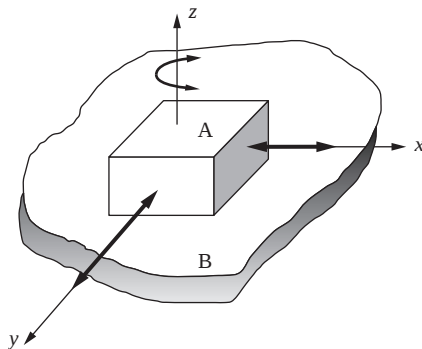


Figure 9.6 Three-degree-of-freedom situation.

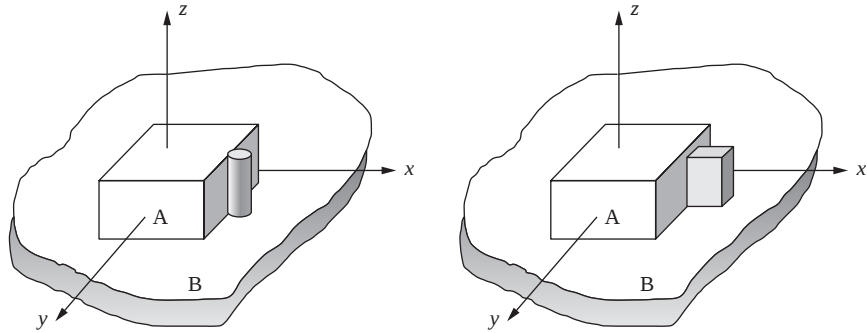


Figure 9.7 Block A restricted by a pin or short wall.

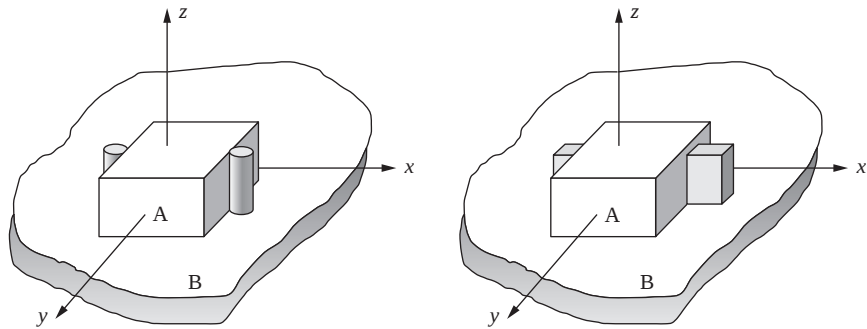


Figure 9.8 Efforts to fully constrain along the x axis.

position parts relative to each other and transmit forces. Thus, it is worthwhile to think in terms of positioning and then force transmission.

Fasteners like bolts and rivets are not good for locating components as the holes for them must be made with some clearance and fasteners are not made with high tolerances. For positioning, first consider a single pin or short wall, as shown in Fig. 9.7. The effect of these will be to only limit the position of A relative to B in the $+x$ direction.

If there is a force always in the positive x direction, then this single constraint fully defines the position on the x axis. Putting a second support on the x axis to limit motion in the negative x direction can have unintended consequences (see Fig. 9.8). Due to manufacturing variations, block A will either be loose or binding. In other words, even though block A looks well constrained in both the $+$ and $-x$ directions, this will be hard to manufacture and to make work like it is drawn. Additionally, the second pin does nothing to constrain the motion in the y direction or rotations about the z axis.

If there are two pins or a long wall positioning the side of the block (see Fig. 9.9), then the x position and angle about the z axis are limited.

If a sufficient force pushes in the $+x$ direction, between the pins, then the block is fully constrained in the x direction and about the z axis. However, if the force has any y component, block A can still move in the y direction.

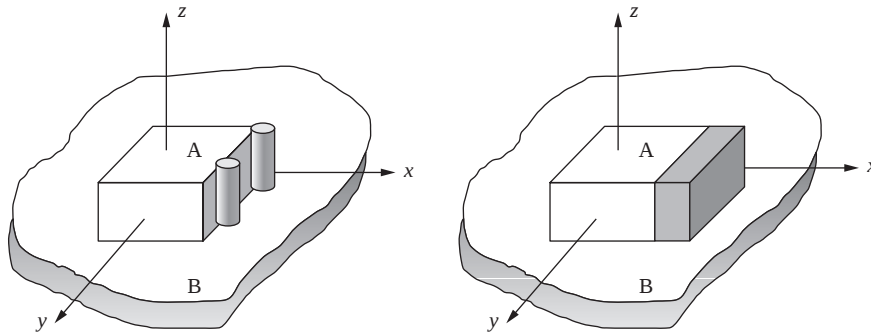


Figure 9.9 Block A restriction in the x direction and z rotation.

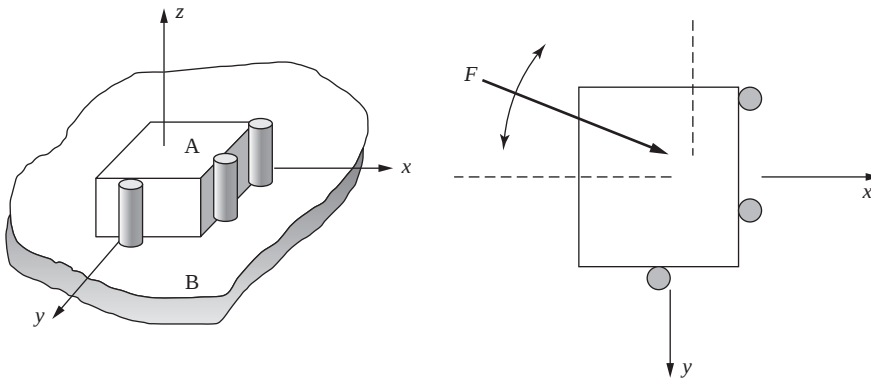


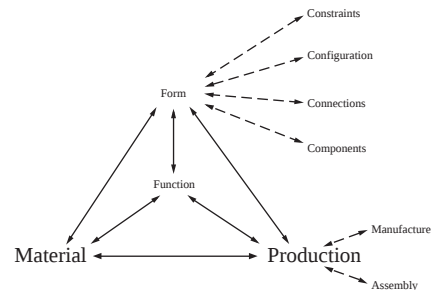
Figure 9.10 Block A fully constrained.

Finally, if a pin is attached to component B so that component A is restrained from moving in the y direction and a force F is directed between the limits shown in Fig. 9.10 (the force has a positive x and y direction), then component A is fully constrained and has no degrees of freedom relative to component B. What is vitally important here is that it takes exactly three points to constrain one component to another.

The three points to constrain component A relative to component B can take many forms. A few of these are shown in Fig. 9.11.

9.3.4 Develop Components

It has been estimated that fewer than 20% of the dimensions on most components in a device are critical to performance. This is because most of the material in a component is there to connect the functional interfaces and therefore is not dimensionally critical. Once the functional interfaces between components have been determined, designing the body of the component is often a sophisticated connect-the-dots problem.



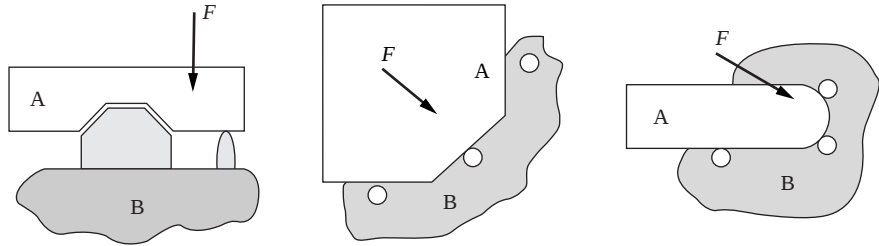


Figure 9.11 Other fully constrained blocks.

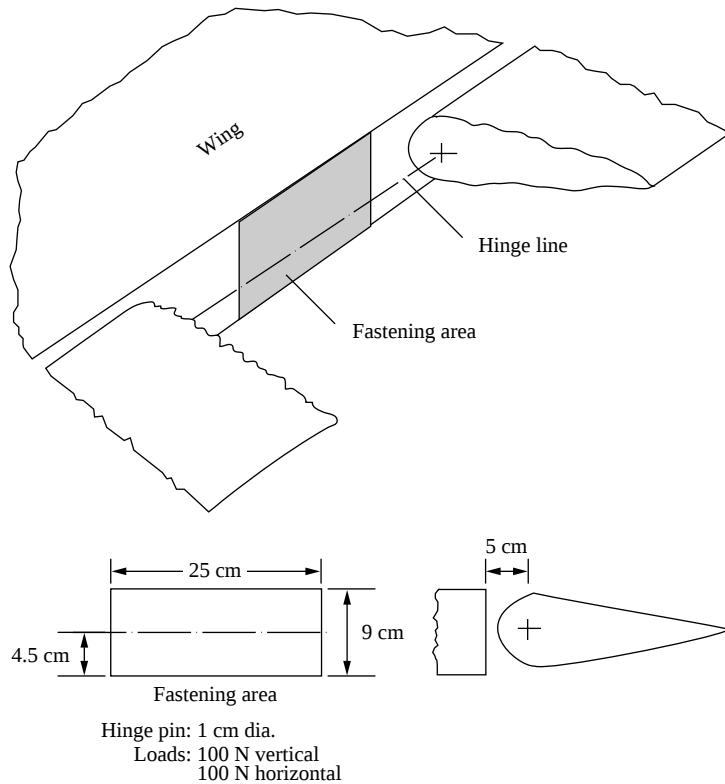


Figure 9.12 Requirements on an aircraft hinge plate.

Consider this example of an aircraft hinge. The spatial constraints for this and its interface points (i.e., fastening area) are shown in Fig. 9.12. The major functions of this individual component are to transfer forces and clear (not interfere with) other components. The load on the component and the geometry of the interfaces are detailed in the figure. The component is a simple structural member that must transfer the load from the hinge line to the fastening area. As shown in Fig. 9.13, there are many solutions to this problem. The solutions in Figs. 9.13a

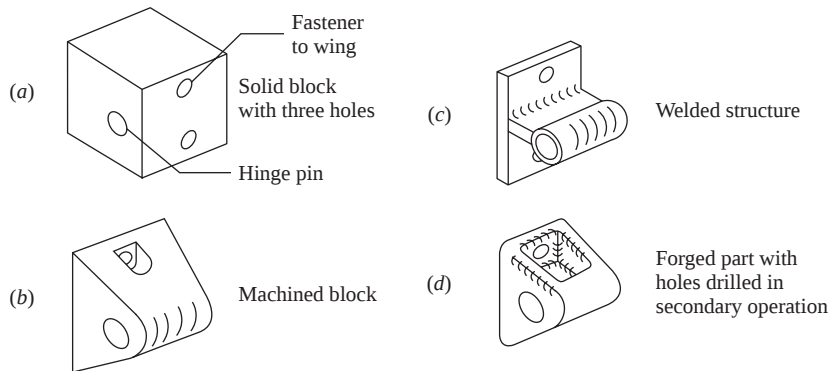


Figure 9.13 Potential solutions for the structure of the aircraft hinge plate.

Components grow primarily from interfaces.

and 9.13b are machined out of a solid block of material. The solution in Fig. 9.13c is made from welded sections of off-the-shelf extruded tubing and plate. These three solutions are good if only a few hinge plates are to be manufactured. If the number to be produced is high, then the forged component in Fig. 9.13d may be a good solution. Note that all four of these components have the same interfaces with adjacent components. One interface is fixed and may need to be removable, and the other has one degree of freedom. The only difference is in the body, the material connecting the interfaces. All of these product designs are potentially acceptable, and it may be difficult to determine exactly which one is best. A decision matrix may help in making this decision.

The material between interfaces generally serves three main purposes: (1) to carry forces or other forms of energy (heat or an electrical current, for instance) between interfaces with sufficient strength and rigidity; (2) to act as an enclosure or guide for other components (guiding airflow, for instance); or (3) to provide appearance surfaces. We have said before that functionality occurs mainly at component interfaces; this is not always true. The exception occurs when the body of a component provides the function—for example, needed mass, stiffness, or strength—in which case, shape can be as important as the interface.

It is best to connect interfaces with strong structural shapes. Strong shapes have material distributed to make the best use of it. Common strong structural shapes are listed next.

- The simplest strong shape is a rod in tension (or compression). If a shape has two interfaces, as shown in Fig. 9.14a, and it needs to transmit a force from one to the other, the strongest shape to use is a rod in tension, Fig. 9.14b. Once away from the ends (interfaces), the forces are distributed as a constant stress

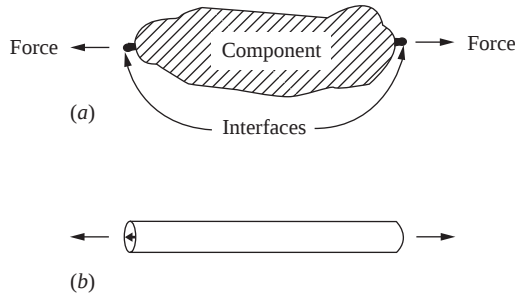


Figure 9.14 A bar in tension.

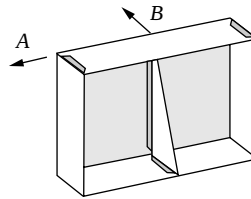


Figure 9.15
A triangulated component.

Triangulate! Unless you have a very good reason not to.

throughout the rod. Thus, this shape provides the most efficient (in terms of amount of material used to transmit the force) shape possible.

- A *truss* carries its entire load as tension or compression. A rule of thumb is *always triangulate the design of shapes*. This is often accomplished by providing shear webs in components to effectively act as triangulating members. The back surface in Fig. 9.15 acts as a shear web to help transmit force *A* to the bottom surface. Take away the back surface and the structure collapses. A rib provides the same function for force *B*.
- A *hollow cylinder*, the most efficient carrier of torque, comes as close as possible to having constant stress throughout all the material. Any closed prismatic shape exhibits the same characteristic. A common example of an approximately closed prismatic shape is an automobile or van body. As the front right wheel of the van shown in Fig. 9.16 goes over a bump, a torque is put on the entire vehicle. Cutting holes in the sides for doors greatly weakens the torque-carrying capability of a van, and it requires additional, heavy structure to make up the difference.
- An *I-beam* is designed to carry bending loads in the most efficient way possible, since most of the material is far away from the neutral bending axis. The principle behind the I-beam is shown in the structural shape of Fig. 9.17.

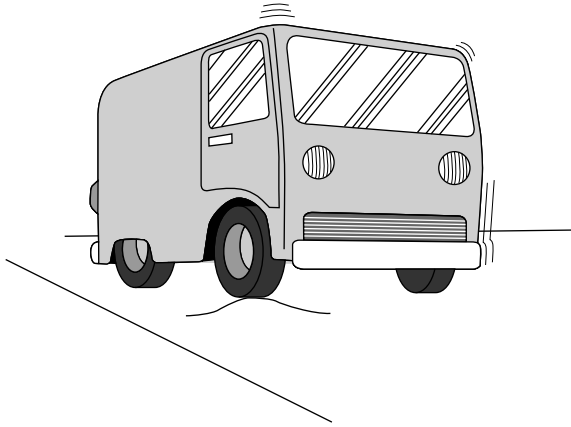


Figure 9.16 Component that efficiently carries torque.

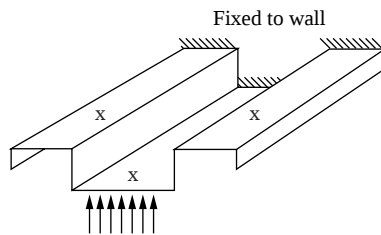


Figure 9.17 Example of an I-beam structure.

Forces flow like water. Failures occur mainly in the rapids.

Although not an I, it behaves much like one, as the majority of the material (labeled “x”) is as far from the neutral bending axis as possible.

Less stress is generally developed if direct force transmission paths are used. A good method for visualizing how forces are transmitted through components and assemblies is to use a technique called *force flow visualization*. These rules explain the method.

1. Treat forces like a fluid that flows in and out of the interfaces and through the component. It makes no difference which way you assume the fluid flows. It is the path that is important.
2. The fluid takes the path of least resistance through the component.
3. Sketch multiple flow lines. The direction of each flow line will represent the maximum principal stress at the location.

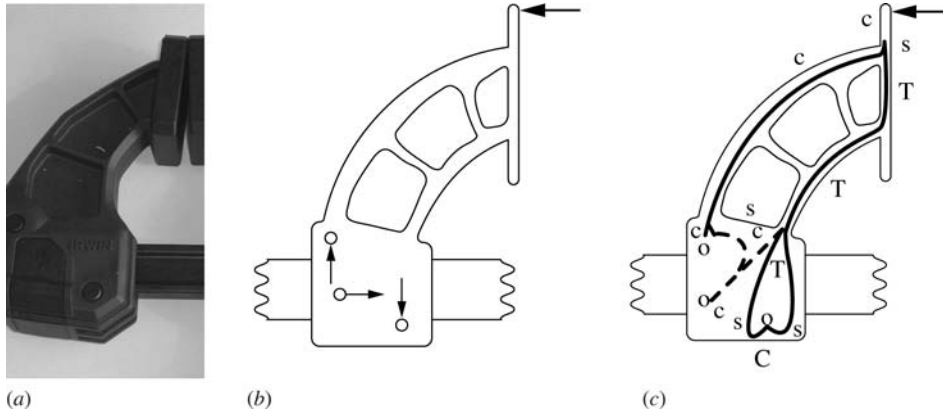


Figure 9.18 Force flow in the tail stock of a clamp. (Reprinted with permission of Irwin Industrial Tools.)

4. Label the flow lines for the major type of stress occurring at the location: tension (T), compression (C), shear (S), or bending (B). Note that bending can be decomposed into tension and compression and that shear must occur between tension and compression on a flow line.
5. Remember that force is transmitted at interfaces primarily by compression. Shear only occurs in adhesive, welded, and friction interfaces.

Two examples clearly illustrate many of the preceding rules. The first is from the tail stock of the Irwin one-handed clamp (Fig. 9.18a). Assume it is loaded at the worst possible condition with a force at the tip as shown. The free-body diagram (Fig. 9.18b) shows the force balanced in the horizontal direction by a pin through the bar. The couple created by these forces is countered by a vertical force couple on the two pins pressing against the bar as shown.

Following the rules just listed, the force flow in the tail stock looks as shown in Fig. 9.18c. The flow enters (leaves) at the tip of the tail stock and leaves (enters) at the compression interface between the tail stock and the three pins. First, consider the bending created by the force on the tip of the tail stock. The middle of the part is like an I-beam, the top is in compression, and the bottom is in tension. Thus, a compression flow line should go from the force on the tip of the tail stock, down the top of the part to the pin. Since the I-beam cross section is in bending, the bottom of the tail stock must be in tension. At some point between the compressive force at the tip and the tensile force in the body there is shear as shown. The tension then flows around the bottom pin to become compressive at the interface with the pin. To visualize this shear take a piece of notebook paper, insert a pencil in one hole, and pull the pencil toward the nearest edge in the plane of the paper. Note that the rip occurs in approximately 45° , signifying a shear failure.

Besides the bending, the force at the tip applies a compressive horizontal load countered by the pin in the center. Depending on the geometry, the entire

part, including the bottom of the I-beam section may be in compression. Also, there will be some shear occurring in order to get the compressive force to the pin. This force flow is shown by a dashed line in Fig. 9.18c.

The tee joint in Fig. 9.19a represents a second example of the use of force flow visualization. Figure 9.19b shows two ways of representing the force flow in the flange. The left side shows the bending stress in the flange labeled B; the right side shows the bending stress decomposed into Tension (T) and Compression (C), which forces consideration of the shear stress. The force flow through the nut and bolt is shown in Figs. 9.19c and 9.19d. The force flow in the entire assembly is shown in Fig. 9.19e.

In summary, force flow helps us visualize the stresses in a component or assembly. It is best if the force paths are short and direct. The more indirect the path, the more potential failure points and stress concentrations. Developing force flows comes with practice and comparison to detailed analyses from finite element programs. With practice, you can learn where to look for failures.

In designing the bodies of components, be aware that stiffness determines the adequate size more frequently than stress. Although component design textbooks emphasize strength, the dominant consideration for many components should be their stiffness. An engineer who used standard stress-based design formulas to analyze a shaft carrying a small torque and virtually no transverse load found that

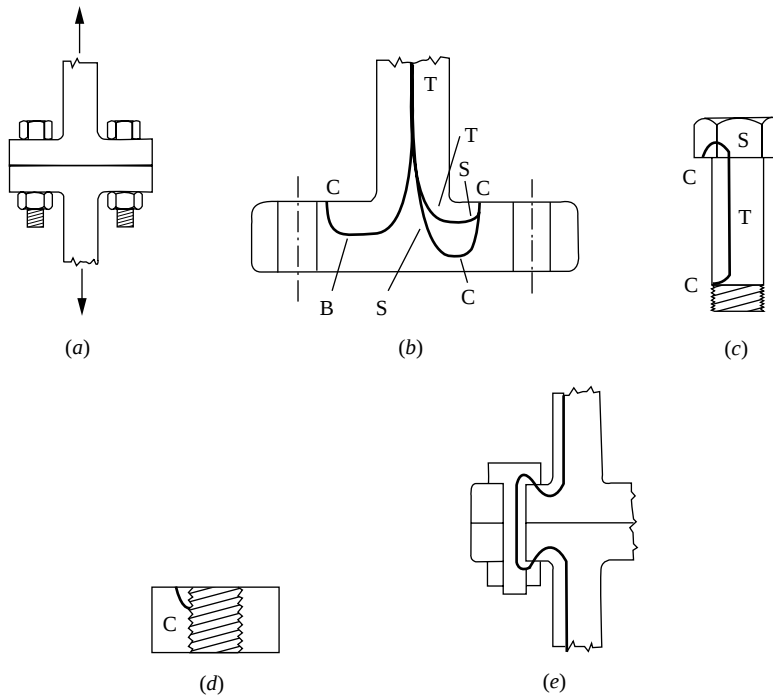


Figure 9.19 Force flow in a tee joint.

it should be 1 mm in diameter. This seemed too small (a gut-feeling evaluation), so the engineer increased the diameter to 2 mm and had the system built. The first time power was put through the shaft, it flexed like a noodle and the whole machine vibrated violently. Redesign based on stiffness and vibration analysis showed that the diameter should have been at least 10 mm to avoid problems.

Finally, in designing components, use standard shapes when possible. Many companies use *group technology* to aid in keeping the number of different components in inventory to a minimum. In group technology, each component is coded with a number that gives basic information about its shape and size. This coding scheme enables a designer to check whether components already exist for use in a new product.

9.3.5 Refine and Patch

Although not shown as a basic element of product design in Fig. 9.3, refining and patching are major parts of product evolution. *Refining*, as described in Section 2.3, is the activity of making an object less abstract (or more concrete). *Patching* is the activity of changing a design without changing its level of abstraction.

The importance and interrelationship of refining and patching the shape can be clearly seen in the following example. A designer was developing a small box to hold three batteries in a series. This subsystem powered the clock/calendar of a personal computer. The designer's notebook sketch of the final assembly is shown in Fig. 9.20. The assembly is composed of a bottom case, a top case, and four contacts. Figure 9.21 shows the evolution of one of the contacts (contact 1), again through the sketches and drawings made by the designer. The number beside each graphic image shows the percentage of the total design effort completed when the representation was made. The designer was simultaneously at work on other components of the product. (The circled letters in Fig. 9.21c were added for this discussion and were not in the original drawings.) The design of the battery contact is one of continued *refinement*. Each figure in the series moves the design closer to the final form of the component. The initial sketch (Fig. 9.21a) shows circles representing contact to the battery and a curved line representing current conduction. The final drawing of the contact (Fig. 9.21e) is a detailed design ready for prototyping. Figure 9.21c is of special interest, as it clearly shows the evolutionary process. The designer began by redrawing the left contact, A, from the earlier sketch (Fig. 9.21b). She also redrew line B, which represents an edge of the structure connecting the contact to the wire. After beginning to draw this line, she realized that, since she last worked on this component, a plastic wall, C, had been added to the product and the contact could no longer continue straight across. At this point, she *patched* the design by tilting the connecting structure B up to position D. The sketch was then completed, with the wire connection still represented by an arc (E). Moments later, the designer further patched the component by combining the wire and the connecting structure, making the structure between contacts all one component (F), and then immediately redrew it, as in Fig. 9.21d.

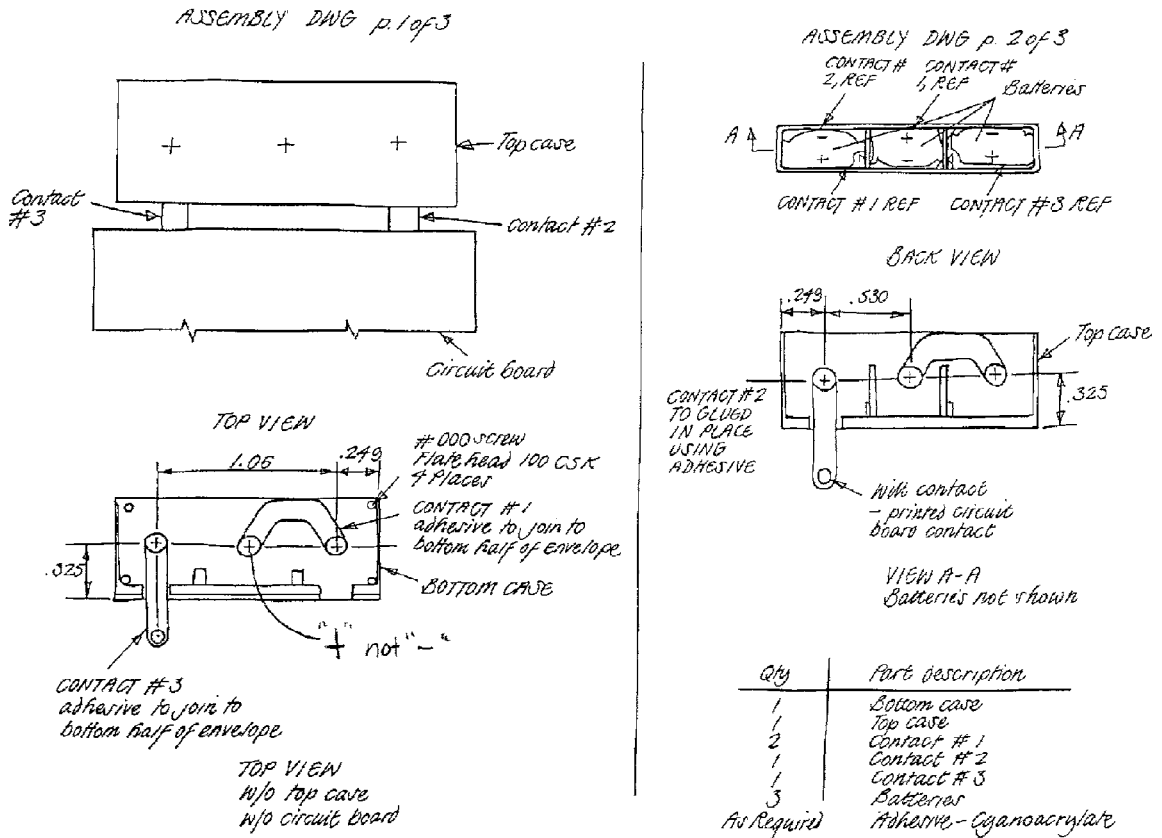


Figure 9.20 Complete layout of battery case.

The elimination of the wire simplified the component; there was no reason for a separate wire in the first place. The battery contact was patched by combining two components. The component was then refined to a fully dimensioned form (Fig. 9.21e).

From this example and others, we can identify many different types of patching:

- **Combining:** Make one component serve multiple functions or replace multiple components. Combining will be strongly encouraged when the product is evaluated for its ease of assembly (Section 11.5).
- **Decomposing:** Break a component into multiple components or assemblies. As new components or assemblies are developed through decomposition, it is always worthwhile to review constraints, configurations, and connections for each one. Because the identification of a new component or assembly establishes a new need, it is even worthwhile to consider returning to the

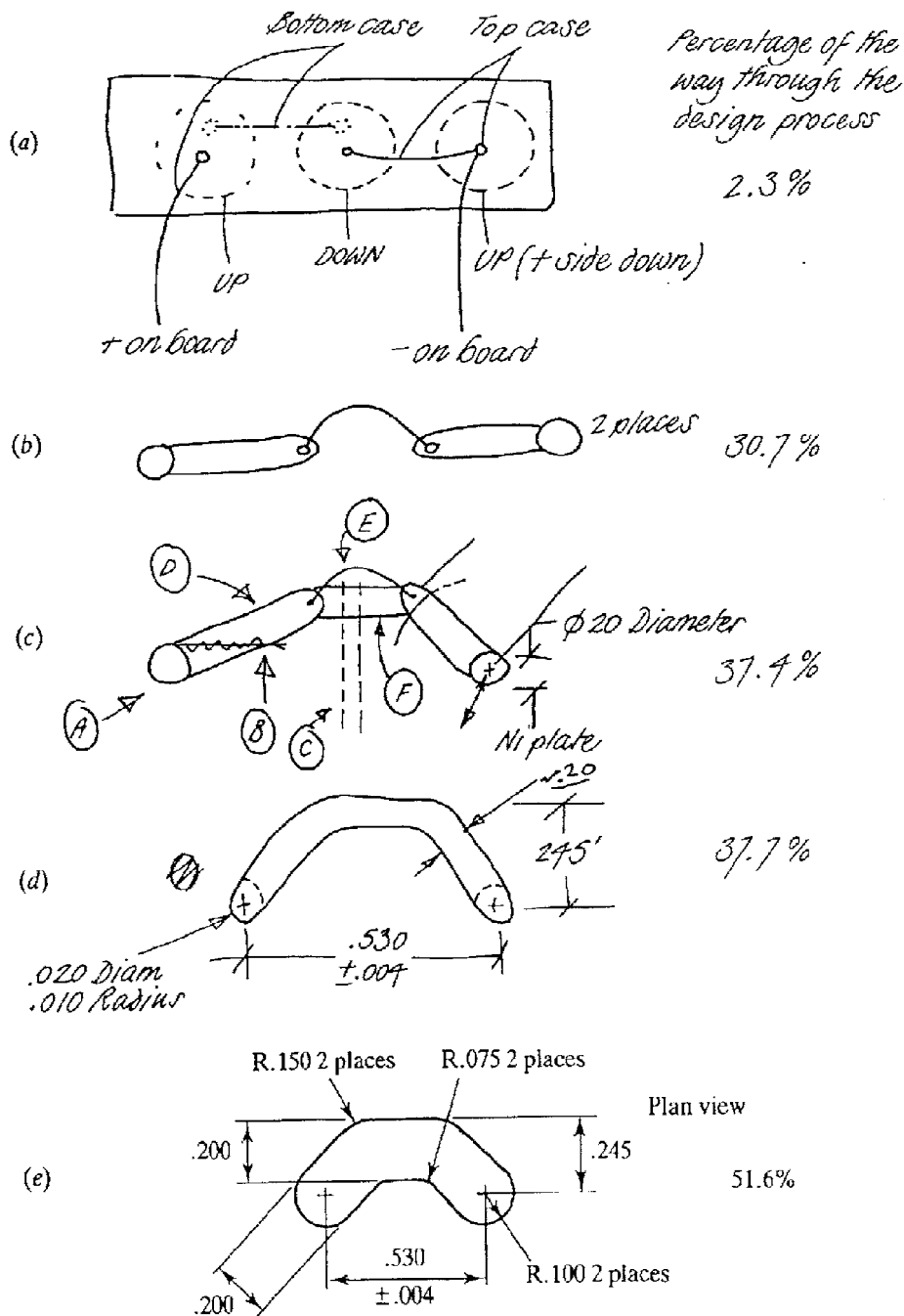


Figure 9.21 Evolution of a battery contact.

Design perfection is achieved not when there is nothing more to add,
but rather when there is nothing more to take away.
— Antoine de Saint-Exupéry

beginning of the design process with it and considering new requirements and functions.

- *Magnifying/Minifying*: Make a component or some feature of it bigger/smaller relative to adjacent items. Exaggerating the size or number of a feature will often increase one's understanding of it. Make one dimension very short or very long. Think about what will happen if it goes to zero or infinity. Try this with multiple dimensions. Sometimes eliminating, streamlining, or condensing a feature will improve the design.
- *Rearranging*: Reconfigure the components or their features. This often leads to new ideas, because the reconfigured shapes force rethinking of how the component fulfills the functions. It may be helpful to rearrange the order of the functions in the functional flow. Take the current order of things and switch them around. Put what is on top, on the bottom; or what is first, last.
- *Reversing*: Transposing or changing the view of the component or feature; it is a subset of rearranging. Try taking what is the inside of something and making it the outside or vice versa.
- *Substituting*: Identify other concepts, components, or features that will work in place of the current idea. Care must be taken because new ideas sometimes carry with them new functions. Sometimes the best approach here is to revert to conceptual design techniques in order to aid in the development of new ideas.
- *Stiffening*: Make something that is rigid, flexible or something that is flexible, rigid.
- *Reshaping*: Make something that is first thought of as straight, curved. Think of it as cooked spaghetti that can be in any form it wants to be and then hardened in that position. Do this with planar objects or surfaces.

A more complete list of ideas for patching can be found in TRIZ's 40 Inventive Principles, discussed in Section 7.7. These principles suggest many ideas for patching products.

The primary goal of patching is to make things work and to make them simpler. The most elegant designs are those that provide the needed functionality but look simple. The quote in the aphorism above states this thought best. Excessive patching implies trouble. If design progress is stuck on one function or component and patching does not seem to be resolving the difficulty, it may be a waste of time to continue the effort. To relieve the problem, apply these three suggestions.

- Return to the techniques in conceptual design; try to develop new concepts based on the functional breakdown and the resources for ideas given in Chap. 7.
- Consider that certain design decisions have altered or added unknowingly to the functions of the component. As products evolve, many design decisions are made; it is easy to unintentionally change the function of a component in the process. It is always worthwhile, when stuck on finding a quality solution, to investigate what functions the component is fulfilling.
- If investigating the changes in functionality does not aid in resolving the problem, the requirements on the design may be too tight. It is possible that the targets based on engineering requirements were unrealistic; the rationale behind them should be reviewed.

The results of efforts to refine or patch any aspect of the product can lead in either of two directions. First, and most often, the refinement or patching is part of the generate/evaluate loop in product design. After each patch or refinement, it is good practice to revisit the decisions that have been made in developing the product to this point before reevaluating. As the product becomes more refined, evaluation usually requires more time and resources; therefore, double-checking can lead to savings. Second, if no satisfactory solution can be found, the result of the refining or patching effort requires a return to an earlier phase of the design process.

9.4 MATERIALS AND PROCESS SELECTION

At the same time form is being developed, it is important to identify materials and production techniques and to be aware of their specific engineering requirements.

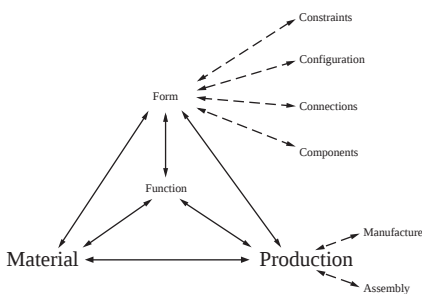
An experienced designer has a short list of materials and processes in mind even with the earliest concepts.

In developing an understanding of the product, we may have set requirements on materials, manufacturing, and assembly. At a minimum we did competitive benchmarking on similar devices, studying them for conceptual ideas and for what they were made of and how they were made. All this information influences the embodiment of the product in several ways:

First, the *quantity of the product to be manufactured* greatly influences the selection of the manufacturing processes to be used.

For a product that will be built only once, it is difficult to justify the use of a process that requires high tooling costs. Such is the case with injection molding, in which the mold cost almost exclusively determines the component cost for low-volume production (see Section 11.2.4). In general, injection-molded plastic components are only cost-effective if the production run is at least 15,000.

A second major influence on the selection of a material and a manufacturing process is *prior-use knowledge* for similar applications. This knowledge can be



When in doubt, make it stout, out of things you know about.

both a blessing and a curse. It can direct selection to reliable choices, yet it may also obscure new and better choices. In general it is best to be conservative, and heed the axiom below.

When studying existing mechanical devices, get into the habit of determining what kind of materials were used for what types of functions. With practice, the identity of many different types of plastics and, to some degree, of the type of steel or aluminum can be determined simply by sight or feel.

Appendix A provides an excellent reference for material selection. It includes two types of information: a compendium of the properties of the 25 materials most often used in mechanical devices and a list of the materials used in common mechanical devices. The 25 most commonly used materials include eight steels and irons, five aluminums, two other metals, five plastics, two ceramics, one wood, and two other composite materials. The properties listed include the standard mechanical properties, along with cost per unit volume and weight. This list is intended to serve as a starting place for material selection. Detailed information on the many thousands of different materials available can be found in the list of references given at the end of Appendix A. Additionally, the appendix contains a list of materials used in common products. Since many different materials can be used in the manufacture of most products, this list gives only those most commonly used.

Knowledge and experience are the third influence on the choice of materials and manufacturing processes. Limited knowledge and experience limit choices. If only available resources can be utilized, then the materials and the processes are limited by these capabilities. However, knowledge can be extended by including on the design team vendors or consultants who have more knowledge of materials and manufacturing processes, so the number of choices can be increased.

Probably the most compelling point in the selection of a material is its *availability*. A product that has a very small production run will probably use off-the-shelf materials. If the design requires structural shapes (I-beams, channels, or L shapes) that must be light in weight, then extruded aluminum shapes could be used. This decision, however, limits the material choices. Aluminum extrusions are readily available in only a few alloy/temper combinations (6061-T6, 6063-T6, and 6063-T52). Other alloys are available on special order. There is a setup charge to obtain these, and a minimum order of a few hundred pounds—a complete run of material—would also apply. If the available alloy/temperatures have properties needed by the product, they can be used. If they do not, the product shape may need to be changed.

During design, the material and production processes selected must evolve as the shape of the product evolves. As a product matures, its layout, details, materials, and production techniques are refined (become less abstract). At the same time a product is refined, changes are sometimes patched with no accompanying

refinement. Suppose the material initially chosen for a component was identified only as “aluminum”; this selection must now be refined and may be patched. For example, the refining/patching history of the selection of material for one component is

“Aluminum” → 2024 → 6061 → 6061-T6.

That is, the selection of “aluminum” was refined to a specific alloy 2024, which was changed (patched) to a different alloy, 6061, which was then refined by identifying its specific heat treatment, T6. This evolution is typical of what occurs as a product is refined toward a final configuration.

Sometimes during the design of a new product, the requirements cannot be met with existing materials or production techniques, no matter how much patching and shape modification occurs. This situation gives rise to the development of new materials and manufacturing processes. Until recently, the thought of designing the materials and processes to meet the product design needs meant postponing the design project so that material or production technology could reach maturity (Section 8.4). However, recent advancements in the knowledge of metal and plastic materials have, to a certain extent, allowed for material and process design on demand.

9.5 VENDOR DEVELOPMENT

When specifying systems, assemblies, or components you either use what is available from vendors, or design new hardware. Mechanical designers seldom design basic mechanical components (e.g., nuts, bolts, gears, or bearings) for each new product, since these components are readily available from vendors. For example, few engineers outside of fastener manufacturing companies design new types of fasteners. Similarly, few designers outside of gear companies design gears. When such basic components are needed in a product, they are usually specified by the designer and purchased from a vendor who specializes in manufacturing them. In general, finding an already existing product that meets the needs in the product is less expensive than designing and manufacturing it, since the companies that specialize in making a specific component have many advantages over an in-house design-and-build effort:

- They have a history of designing and manufacturing the product, so they already have the expertise and machinery to produce a quality product.
- They already know what can go wrong during design and production. A new design effort requires extensive time and experience before reaching the same level of expertise.
- They specialize in the design and manufacture of the component, so they can make it in volumes high enough to keep the cost below what can be achieved through an in-house effort.

Additionally, even if the exact product is not available, most vendors can help develop products or components that are similar to what they already manufacture. Sometimes “design” is specifying Commercial Off The Shelf (COTS) components. This is so common that the terms COTS and the government equivalent, GOTS, are commonly used. COTS and GOTS design is the placement and interfacing between available components.

In past times, it was common for a company to send detailed drawings of components to a number of vendors and select the vendor that quoted the lowest cost. Over the last few years companies have been working with a small number of vendors in the design process from the beginning and including them in the decisions that affect what they will be supplying. In fact, large companies have reduced the number of vendors by an order of magnitude since the mid-1980s. Some companies financially invest in their vendors, and vice versa, to further improve the bond. These tight relationships lead to improved product quality.

Whether to make or buy a component or to choose a component from what is available from vendors, there is need for decision making. For these types of decisions, a good set of criteria are given in the Make/Buy, Vendor Selection template shown in Fig. 9.22. Detailed descriptions of each of the criteria are

- **Low development cost**—How much is it going to cost to develop the component. If it is truly COTS, then there are no development costs. However, if work is needed to change a COTS system or part, or one needs to be developed, then these costs may be significant.
- **Low product cost**—Many decisions are based solely on this criterion. This cost is highly dependent on the volume (the number purchased), delivery costs and many other factors. These will be addressed in Chap. 11 when we discuss DFC, Design For Cost (Section 11.2).
- **High product life cost stability**—Beyond the cost, it is important to consider how the cost may change over time. Cost can be controlled better when you make a component or can be locked in by contract.
- **Low development lead time**—If this and the next criterion are important; they may dominate all the rest and force the purchase of a COTS component. COTS components need no development lead time.
- **Low order lead time**—Even COTS components have an order lead time. Sometimes it can even be longer than the time needed to make the component in house.
- **High product quality**—Sometimes quality must be traded off for cost or time. It is important to understand from the beginning, the level of quality needed to meet the engineering specifications.
- **Good product support**—To address this criterion, two questions must be answered: Who will be responsible for failures and maintenance of the component or product? And, how much support will be needed?
- **Easy to change product**—Sometimes it is necessary to change the product during its lifetime. If it is COTS then you have no control over changes. If



Make/Buy or Vendor Selection

Decision to be made: Make or buy

Date: 09/23/10

Product: Part 234-4B in Espiral

Criterion	Wt.	Vendor 1 Make	Vendor 2 Allied	Vendor 3 Barns	Vendor 4 Crane
Low development cost	5	2	3	2	4
Low product cost	22	4	2	3	4
High product life cost stability	2	5	3	4	4
Low development lead time	7	3	2	4	2
Low order lead time	11	3	2	5	1
High product quality	14	2	3	3	2
Good product support	6	1	4	2	3
Easy to change product	8	3	5	5	4
Strong IP control	18	4	2	4	2
Good control of order volumes	5	4	1	2	4
Good control of supply chain	2	4	4	2	2
Total		35	31	36	32
Weighted total		3.2	2.56	3.47	2.79

Rationale: Choose Barns as it is significantly better than the others in weighted total and has no great weakness.

Team member: Bob

Prepared by: Ivin

Team member: Alvin

Checked by: Becky-Sue

Team member: Becky-Sue

Approved by: Fredrick

Team member:

The Mechanical Design Process

Designed by Professor David G. Ullman

Copyright 2008, McGraw-Hill

Form # 20.0

Figure 9.22 Make/buy or vendor selection example.

this is an important criterion, then it may be best to make the component or have a closely allied vendor make it.

- **Strong IP control**—IP, or Intellectual Property, is a primary asset of a company. IP includes patents, CAD files, drawings, and other documents that give details about the design or production of a product
- **Good control of order volumes**—Sometimes the number of components ordered needs to be flexible. This is generally in response to market changes that can be controlled to some degree through inventory, but that is expensive. So, if order volumes are volatile, then this may be an important criterion.
- **Good control of supply chain**—If you buy a component you can only control the supply chain through your contracts. If this is not sufficient, then this criterion may be important.

These criteria are used in Fig. 9.22 to decide whether to make or buy a component from one of two vendors. This example is a combination of the common make/buy decision and vendor selection decision. Here a simple decision matrix is used to find it. Vendor 3 is the best choice. An online, free robust decision maker is available.

9.6 GENERATING A SUSPENSION DESIGN FOR THE MARIN 2008 MOUNT VISION PRO BICYCLE

The Marin Mount Vision Pro bike was designed for the cross-country mountain bike enthusiast. It is a quality and fairly expensive bicycle (over \$3000USD). The primary demographic for this bicycle is male, 25–50 years old. But, because of its modern look and marketing, it is also designed to attract females and riders of other age groups. It is intended for use on technical trails where there is a mix of uphill and downhill, where light weight and pedaling efficiency are of primary importance. In this section, we will explore how the rear suspension evolved. The story presented here has been tailored for this text, but it does not differ much from the reality of the Marin design process.

9.6.1 Understand the Spatial Constraints for the Mount Vision Bicycle Rear Suspension

For the rear suspension of a mountain bicycle, the spatial constraints are shown in Fig. 9.23. Beyond the obvious need to connect the wheel to the frame, the Marin engineers also wanted to control the path the wheel made relative to the frame as the suspension deflected, the stiffness of the suspension, and the chain length.

Ideally, the wheel of the bicycle should move “nearly” straight up and down as it deflects. If the suspension was designed as a simple bar with a single pivot

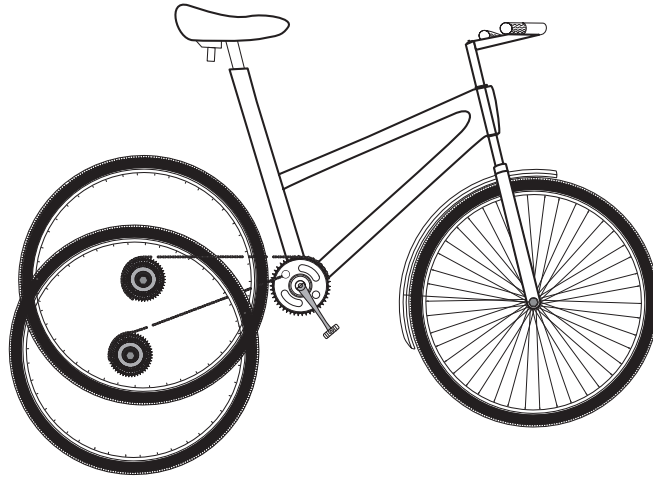


Figure 9.23 Physical constraints for the mount vision.

(see Fig. 9.24), then the wheel would make an arc with it moving closer to the front of the bike as it deflected. This would give the rider the feeling she was falling backward as the wheel deflected. The Marin engineers wanted to control the wheel path to manage the feel transmitted to the rider. As important as the wheel path, is the change in stiffness—flow of energy of the suspension system. The ideal suspension system for any vehicle is soft, has low stiffness, when it goes over small bumps and gets stiffer for large bumps. In other words, the larger the deflection, the stiffer the suspension system should become. This requirement may not seem “spatial” but it constrains how the shock is mounted between the frame and the moving parts as will be seen.

To understand the desire to control the chain length, consider a suspension that was designed so that when the pedals were pressed, the resulting tension in the chain pulled the suspension up (i.e., the frame down). The rider, when feeling the frame drop (flow of information) would then ease off the force and subsequently the frame would rise. Feeling the frame rise, the rider then reapplies the pedal force resulting in a “pogo” motion and a very uncomfortable ride. Thus, an additional constraint is that the motions and accelerations felt by the rider will not lead to poor suspension performance.

Summarizing, the spatial constraints are

1. Wheel and chain must clear frame for all deflections.
2. Wheel should move straight up and down.
3. Low stiffness for small deflections, increasing with deflection.
4. Chain length should not change during deflection.

9.6.2 Configure Components for the Mount Vision Bicycle Rear Suspension

The simplest type of suspension that can be put on a bicycle is a one with a single pivot as shown in Fig. 9.24. On the bike, the pivot is near the center of the crank and every point on the rear triangular structure (called the rear “stay”) rotates around this point. As the wheel deflects, it makes a circular arc and the chain gets shorter, violating two of the spatial constraints. As the wheel moves up, the shock gets shorter. Shocks on bicycles generally have an air or oil damper with a mechanical, coil spring wrapped around it. This spring has a stiffness that remains essentially constant as the wheel deflects. So the spring force increases as the wheel is deflected. Thus, it is clear that this type of suspension will not work for the Marin Mountain Vision Pro.

In 2003, Marin introduced a more sophisticated suspension based on a four-bar linkage and referred to it as their “Quadlink” design. The Quadlink was not the first four-bar suspension used on a mountain bicycle, but it did bring this type of mechanism to a high level of refinement. To understand how Marin configured this suspension, a short refresher on four-bar linkages.

Figure 9.25 shows two simple members, A and B connected by member C. Members A and B, the links, move about fixed points and member C, the “follower,” connects the end points of A and B. Points 1 and 2 move in circular arcs about the fixed points as in Fig. 9.24. For this parallelogram four-bar linkage, member C effectively translates without rotating. This will be clarified in a moment.

To better understand what link C is doing, consider a modification to this basic four-bar where the links are different lengths as shown in Fig. 9.26. The projection of the links intersects at a point called the instant center. The instant center is the point about which link C is rotating when the links are in the configuration shown. The reason for the term “instant” is that the same linkage (all the member’s lengths held constant); with the members in a different position have a different instant center, as can be seen by comparing Figs. 9.26a and 9.26b.

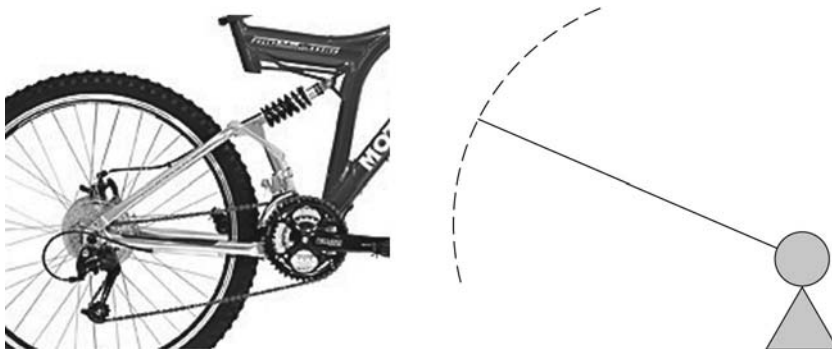


Figure 9.24 A simple, single pivot suspension. (Reprinted with permission of Marin Bicycles.)

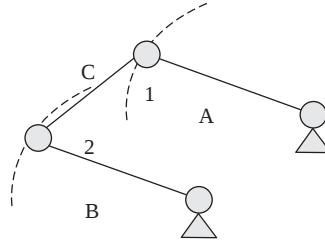


Figure 9.25 A basic four-bar linkage.

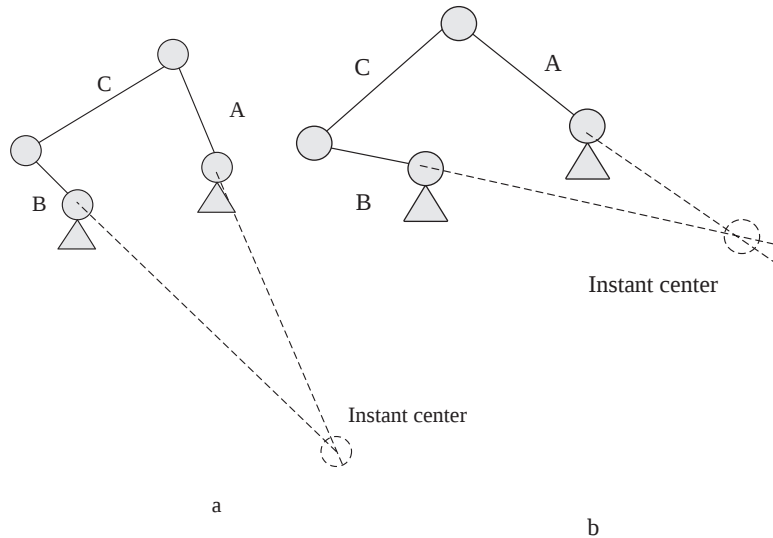


Figure 9.26 A linkage with two of its instant centers.

Thus, as the linkage moves through different positions, the instant center traces a path describing the virtual pivot point for member C. The linkage in Fig. 9.25, the parallelogram, has the instant center always at infinity, thus the link has an infinite radius of rotation—it translates.

One further four-bar concept is needed to understand how the Marin Quadlink was designed. If link C, rather than being a straight member as shown in the figures so far, is a structure as in Fig. 9.27, then every point on this structure or stay is rotating about the instant center. Figure 9.27 is the same linkage as in Fig. 9.26 but with the addition of the stay, CDE.

Point 3 at the left end of member CDE rotates about the instant center and makes a nearly straight line. This is the point where the wheel is mounted. In order to design the shape of the path followed point 3 the engineer specifies the lengths of A, B, C, D, and E and the relative positions for the two fixed pivot

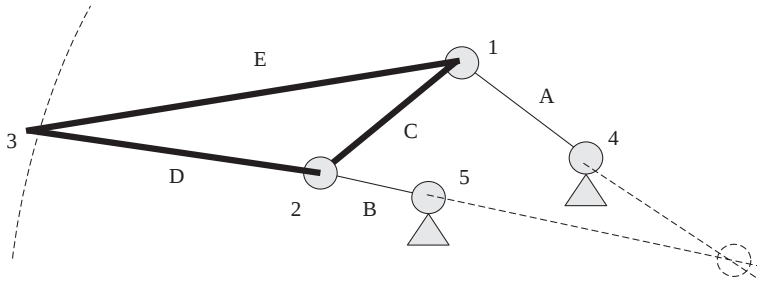


Figure 9.27 A complete four-bar structure.

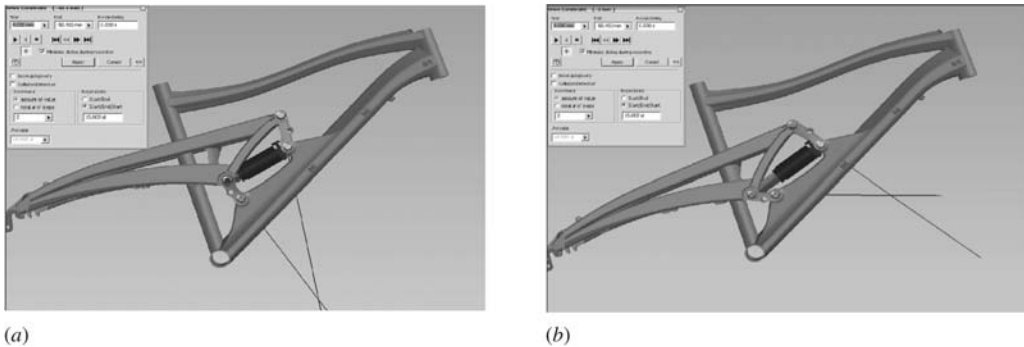


Figure 9.28 Simulation of the Quad link suspension: (a) undeformed and (b) fully deflected. (Marin Bicycles are designed on Autodesk Inventor™. Reprinted with permission of Marin Bicycles.)

points (the distance between them and angle of the line connecting them), for a total of seven variables. There is a lot of design freedom.

The Marin engineers adjusted these variables to meet the spatial constraints. The final design is shown in Fig. 9.28. These solid models were developed in Autodesk Inventor™. This program let the engineers see the motion as the suspension deflected. The block in the upper left corner controls the simulation so the designer can see the motion of the mechanism and instant center.

The Marin designers used the solid modeling software and other computer simulations to determine the best values for the seven variables, one that gave a fairly straight wheel path with near constant chain length. Further, by controlling the location of the virtual center and the positioning of the shock (described in the next section) they were able to achieve low stiffness for small deflections, increasing with deflection. Specifically, when the virtual center is nearly under the crank (Fig. 9.28a) the moment arm of the rear stay is much shorter than when the suspension is deflected (Fig. 9.28b).

9.6.3 Develop Connections: Create and Refine Interfaces for Functions for the Mount Vision Bicycle Rear Suspension

This section focuses on the connections between the components. On the Marin Mount Vision Pro, the connections are those between the links in the four-bar linkage, those connecting the shock to the bike and those that connect the fixed parts together. We will consider these in order. For the four-bar linkage, the connections are the four pivots. These must have one degree of freedom and thus can be either bearings or flexures. For most mountain bikes, either rolling element bearings or bushings are used, but some have used flexures. Considering Fig. 9.27, the shock can be mounted in many different ways. It can be mounted between any two elements that move closer together as the system deflects; for example, element C and the frame, elements A and B, and so on. The addition of the shock adds two more pivots to the assembly making a total of six pivoting connections.

The Marin engineers reduced the number of pivots by mounting the shock between linkage pivots 2 and 4. As the suspension system deflects, pivot 2 moves toward pivot 4. In fact, the engineers, when determining the lengths of all the seven members, took the needed change in length of the shock as an additional constraint. The decision to mount the shock in this manner made the design of linkage more challenging and connections more complex, but the trade-off for fewer pivots made this worthwhile.

Pivots 2 and 4 need to have the link and shock free to rotate about the axel (shown as a centerline in Fig. 9.29). Note in Fig. 9.28, the amount of rotation of these elements is small, only a few degrees in some cases. Bearings that operate primarily in one position and only move a small amount from that position present

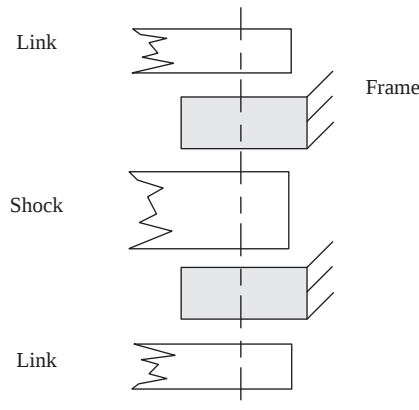


Figure 9.29 The components in pivots 2 and 4.

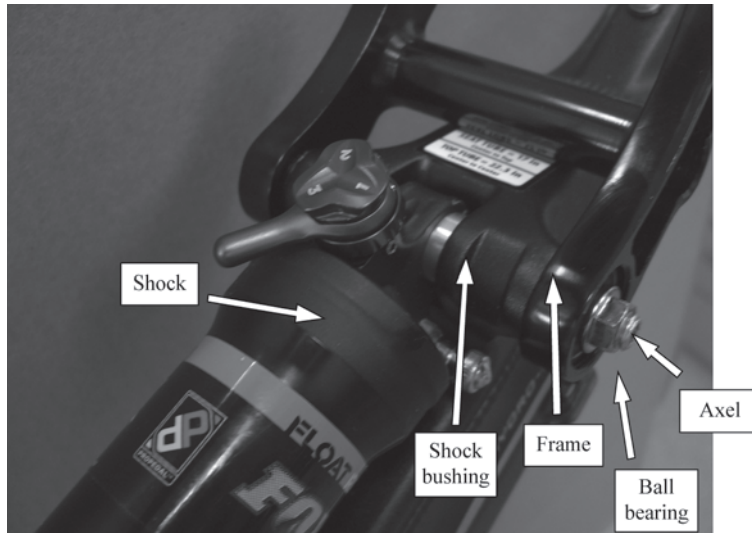


Figure 9.30 Final design of pivot 2. (Reprinted with permission of Fox Racing Shox.)

their own design problems as small deflections do not force the lubricant to flow to all the areas.

The final connection at pivot 4 is shown in Fig. 9.30. Connections between components that are moving relative to each other need to be addressed. They are refined in Section 9.6.4.

9.6.4 Develop Components for the Mount Vision Bicycle Rear Suspension

Finally, the actual components need to be developed. For the Marin engineers these parts needed to be light in weight, manufacturable in volumes that matched the sales projections, and had a look that would attract sales. Thus, these parts were a combination of structure and eye candy. We will discuss three of the components here, link A, the ball bearing in link C, and the lower part of the rear stay.

Link A is a very simple component that needs to connect pivots 1 and 2. The final component, like many on the bike is forged aluminum with the bearing mounting surfaces machined. It is shown in two views in Fig. 9.31.

The bearing between the axel and the link, shown pressed into the link in Fig. 9.31, is a rolling element ball bearing. As mentioned earlier, this bearing does not rotate very much and thus requires special consideration. The final bearing chosen was one that was specially designed for aircraft control systems, another application with small, repetitive motions.

The rear stay components could have been made out of round aluminum tubes welded together as with most aluminum bikes. However, to get a better “look,”



Figure 9.31 Link A. (Reprinted with permission of Marin Bicycles.)

the designers wanted tubes that curved, and to save weight, the engineers wanted tubes that tapered. As shown in Figs. 2.10 and 9.28 these two requirements were met. The manufacturing method used is called hydroforming. To hydroform, a round tube is put in a die and then the tube is filled with high-pressure liquid causing it to deform and be shaped by the die.

9.7 SUMMARY

- A Bill of Materials is a parts list—an index to the product.
- Products must be developed from concepts through concurrent development of *form*, *material*, and *production* methods. This process is driven by the functional decomposition discussed in Chap. 7.
- Form is bound by the geometric *constraints* and defined by the *configuration of connected components*.
- The development of most components and assemblies starts at their interfaces, or connections, since for the most part function occurs at the interfaces between components.
- Product development is an iterative loop that requires the development of new concepts, the decomposition of the product into subassemblies and components, the refinement of the product toward a final configuration, and the patching of features to help find a good product design.
- Vendor selection is an important part of the design process.

9.8 SOURCES

- Ashby, M. F.: *Materials Selection in Mechanical Design*, Pergamon Press, Oxford, U.K., 1992. An excellent text on materials selection. There is a computer program available implementing the approach in this text.
- Blanding, D.: *Exact Constraint: Machine Design Using Kinematic Principles*, ASME Press, 1999. The best reference on the design or connections between components. Written by a design engineer from Kodak.
- Budinski, K.: *Engineering Materials: Properties and Selection*, Reston, Va., 1979. A text on materials written with the engineer in mind.

- Snead, C. S.: *Group Technology: Foundations for Competitive Manufacturing*, Van Nostrand Reinhold, New York, 1989. An overview of group technology for classifying components.
- Tjalve, E.: *A Short Course in Industrial Design*, Newnes-Butterworths, London-Boston, 1979. An excellent book on the development of form.
- Ullman, D. G., S. Wood, and D. Craig: "The Importance of Drawing in the Mechanical Design Process," *Computers and Graphics*, Vol. 14, No. 2 (1990), pp. 263–274. A paper that itemizes the different uses of graphical representations in mechanical design.

Information on modular systems and architecture are from

- Alizon, F., Shooter, S. B. and Simpson, T. W.: "Improving an Existing Product Family based on Commonality/Diversity, Modularity, and Cost," *Design Studies*, 2007 Vol. 28, No. 4, pp. 387–409.
- Qureshi, A., J. T. Murphy, B. Kuchinsky, C. C. Seepersad, K. L. Wood and D. D. Jensen: "Principles of Product Flexibility," *ASME IDETC/CIE Advances in Design Automation Conference*, Philadelphia, Pa., 2006. Paper Number: DETC2006-99583.
- Tripathy, Anshuman, and Steven D. Eppinger: "A System Architecture Approach to Global Product Development," MIT, Sloan School of Management, Working Paper Number 4645-07, March 2007.

9.9 EXERCISES

- 9.1 Develop a bill of materials for
- A stapler
 - A bicycle brake caliper
 - A hole punch
- 9.2 For the original design problem (Exercise 4.1), develop a product layout drawing or solid model by doing these:
- Develop the spatial constraints.
 - Develop a refined house of quality and function diagrams for the most critical interface.
 - Develop connections and components for the product.
 - Show the force flow through the product for its most critical loading.
- 9.3 For the redesign problem (Exercise 4.2):
- Identify the spatial constraints for all important operating sequences.
 - At critical interfaces, identify the energy, information, and material flows.
 - Develop a refined house of quality and function diagrams for the most critical interface.
 - Develop new connections and components for the product.
 - Show the force flow through the product for its most critical loading.
- 9.4 Determine the force flow in
- A bicycle chain.
 - A car door being opened.
 - A paper hole punch.
 - Your body while holding a 5-kg weight straight out in front of you with your left hand.

- 9.5 For a part you designed, decide whether to make it or buy it from a vendor. The cost-estimating templates available on the website for plastic part and machined part cost estimation might be of help. See Sections 11.2.3 and 11.2.4 for discussion about these cost estimators.



9.10 ON THE WEB

A template for the following document is available on the book's website: www.mhhe.com/Ullman4e

- Bill Of Materials
- Make or Buy

CHAPTER 10

Product Evaluation for Performance and the Effects of Variation

KEY QUESTIONS

- Which is best to evaluate the product performance, analytical models or physical testing?
- What is a P-diagram and how does it help identify noise?
- How are trade-offs made?
- What are the three types of noises and how do they affect product quality?
- Why is tolerance stacking important during assembly?
- How is robust design used to ensure quality?

10.1 INTRODUCTION

The primary goal in this chapter is to compare the performance of the product to the engineering specifications developed earlier in the design project. Performance is the measure of behavior, and the behavior of the product results from the design effort to meet the intended function. Thus, part of the goal is to track and ensure understanding of the functional development of the product. If the functional development is not understood, the product may exhibit unintended behaviors.

Another subgoal is to design in quality. Although this chapter is about “evaluation for performance,” it gives another opportunity to be sure that a quality product is developed—that it will always work as it was designed to.

Best practices for product evaluation are listed in Table 10.1, an extension of Table 4.1. The first eight best practices are covered in this chapter. The remainder of the best practices listed in the table are aimed at other, nonperformance product evaluation techniques and are covered in Chap. 11. Although all of these best

Table 10.1 Best practices for product evaluation

-
- Monitoring functional change (Sec. 10.2)
 - Goals of performance evaluation (Sec. 10.3)
 - Trade-off management (Sec. 10.4)
 - Accuracy, variation, and noise (Sec. 10.5)
 - Modeling for performance evaluation (Sec. 10.6)
 - Tolerance analysis (Sec. 10.7)
 - Sensitivity analysis (Sec. 10.8)
 - Robust design (Secs. 10.9 and 10.10)
 - Design for cost (DFC) (Sec. 11.2)
 - Value engineering (Sec. 11.3)
 - Design for manufacture (DFM) (Sec. 11.4)
 - Design for assembly (DFA) (Sec. 11.5)
 - Design for reliability (DFR) (Sec. 11.6)
 - Design for test and maintenance (Sec. 11.7)
 - Design for the environment (Sec. 11.8)
-

practices are discussed as techniques for product evaluation, they all contribute to the generation of the product as part of the iterative generate/evaluate cycle.

10.2 MONITORING FUNCTIONAL CHANGE

Although the main goal of evaluation is comparing product performance with engineering targets, it is equally important to track changes made in the function of the product. Conceptual designs were developed first by functionally modeling the problem and then, on the basis of that model, developing potential concepts to fulfill these functions. This transformation from function to concept does not end the usefulness of the functional modeling tool. As the form is refined from concept to product, new functions are added.

An obvious question about this process arises: What benefit is there in refining the function model as the form is evolving? The answer is that by updating the functional breakdown, the functions that the product must accomplish can be kept very clear. Nearly every decision about the form of an object adds something, either desirable or undesirable to the function of the object. It is important not to add functions that are counter to those desired. For example, in the design of the Marin Mount Vision suspension, the decision to use the air shock necessitated an interface between the user and shock to add air and to adjust the dampening. The final shock chosen, the Fox Float RP23 (Fig 10.1), shows the air valve and adjustment handle near the top of the unit. The exact steps a user must go through to add air to the shock were made clear by refining the function occurring at the interface between the user and the air valve on the shock. Besides tracking the functional

Every feature added brings with it new intended functionality.
It is unintended functionality that can hurt you.



Figure 10.1 Fox Float RP23 used on the Marin Mount Vision. (Reprinted with permission of Fox Racing Shox.)

evolution of the product, the refinement of the functional decomposition also aids in the evaluation of potential failure modes (covered in Chap. 11).

Finally, tracking the evolution of function means continuously updating the flow models of energy, information, and materials. It is these flows that determine the performance of the product. As the product matures, the intended function and actual behavior merge and so what was, in conceptual design, concern for “the desired” now turns to measuring “the reality.”

10.3 THE GOALS OF PERFORMANCE EVALUATION

In Chap. 6, we developed a set of engineering requirements based on the needs of the customer. For each of these requirements, a specific target was set. The goal now is to evaluate the product design relative to these targets. Since the targets are represented as numerical values, the evaluation can only occur after the product is refined to the point that numerical engineering measures can be made. In Chap. 8, the concepts developed were not yet refined enough to compare with the targets and were thus compared with measures that are more abstract. Evaluation can now be based on comparison with the engineering requirements. Beyond comparison to the requirements, effective evaluation procedures should

Evaluation always requires a clear head
and twice the time you estimated.

clearly show what should be altered (patched) in order to make deficient products meet the requirements, and they should demonstrate the product's insensitivity to variation in the manufacturing processes, aging, and operating environment. Restated, the evaluation of product performance must support these factors:

1. Evaluation must result in *numerical measures* of the product for comparison with the engineering requirement targets developed during problem understanding. These measurements must be of sufficient accuracy and precision for the comparison to be valid.
2. Evaluation should give some indication of *which features of the product to modify*, and by how much, in order to bring the performance on target.
3. Evaluation procedures must include the influence of *variations* due to manufacturing, aging, and environmental changes. Insensitivity to these “noises” while meeting the engineering requirement targets results in a robust, quality product.

Where traditionally engineering evaluation has focused on only the first of these three points, this chapter covers all three. Much emphasis is placed on the third point, the consideration of variation because of its direct relationship with product quality.

This chapter is built around Fig. 10.2, the P-diagram. This diagram will be referenced and added to throughout this chapter. In the P-diagram, the letter “P” stands for either product or process and can represent the entire product or some system, subsystem, or process within it. The product or process being evaluated is dependent on the values of many parameters. These parameters may be physical dimensions, material properties, forces from other systems, or forces and motions from humans controlling the system. They may be the temperature of

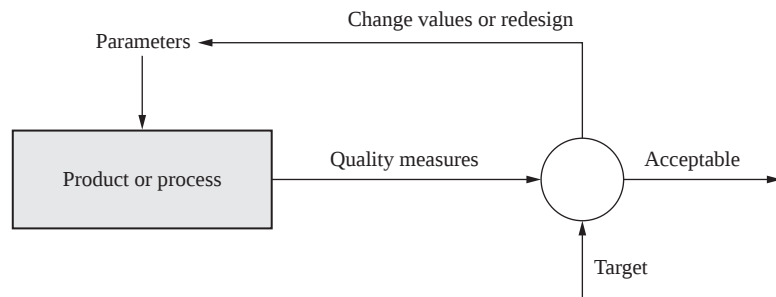


Figure 10.2 The basic P-diagram.

Know how to control what you can, make your product insensitive to what you cannot, and be wise enough to know the difference.

the environment, the humidity, or the amount of dirt on the system. The parameters are all the factors on which the product or process depends and the values of these parameters determine the resulting performance, ease of assembly, quality, and other features of the product or process.

To evaluate the system we need to assess quality measures. These are measures that communicate quality to the customer. To evaluate the product or process these quality measures must be compared to the targets set by the engineering specifications (Chap. 6). If the quality measures compare well to the targets, then we have a quality product. If they do not, then we have to change the values of the parameters or redesign the system—changing the parameters themselves.

One addition to the P-diagram is necessary when considering dynamic function, the product or process may be responding to input signals, as is shown in Fig. 10.3. In this case, the quality measures include system performance. Examples of systems with and without input signals will be given in the chapter.

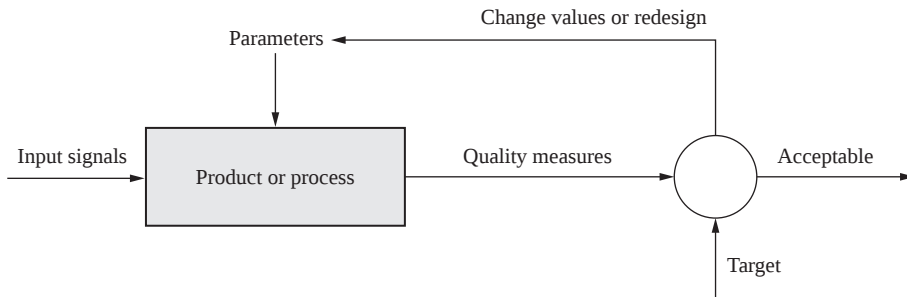


Figure 10.3 The P-diagram with input signal.

Meeting the product performance evaluation goals requires more than throwing together a prototype or running a computer simulation and seeing if it will work. Meeting the goals requires an understanding of concepts such as *optimization*, *trade studies*, *accuracy*, *tolerances*, *sensitivity analysis*, and *robust design*. The remainder of this chapter is focused on these techniques. This phase of the design process is the last chance to design quality into the product.

Consider the design of a tank to hold liquid. Conceptual design of the tank has resulted in a cylindrical shape with an internal radius r and an internal length l . Thus, the volume of the tank V can be written as

$$V = \pi r^2 l$$

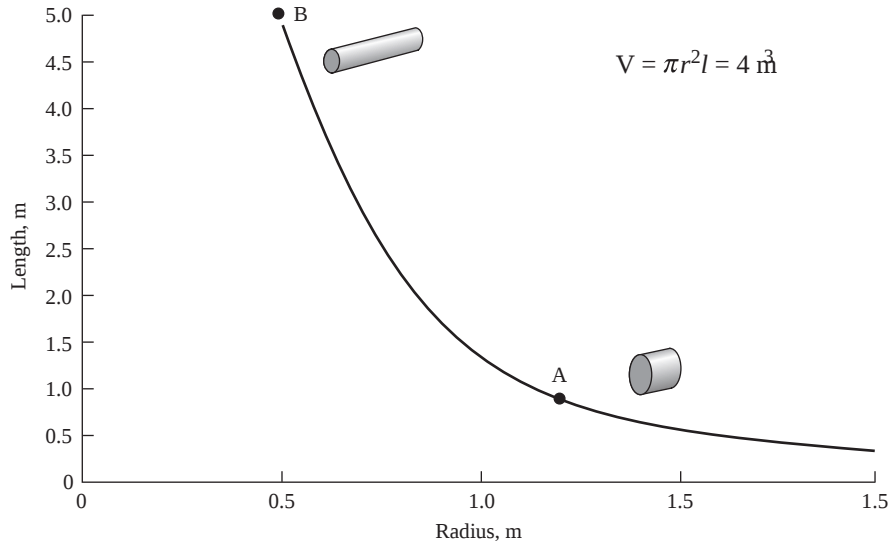


Figure 10.4 Potential solutions for the tank problem.

An inaccurate model is inaccurate no matter how small the variation.

Additionally, a customer's requirement is to design the "best" tank to hold "exactly" 4 m^3 of liquid. This seems a simple enough problem with r and l as the parameters, V as the quality measure, and its target response as

$$V = 4 \text{ m}^3$$

Then

$$r^2 l = 1.27 \text{ m}^3$$

As can be seen from Fig. 10.4, there are an infinite number of solutions to the problem. The tank at point A, for example, is short and fat, and the one at point B is long and thin. It is not clear which point on the curve might be "best" in terms of holding "exactly" 4 m^3 of liquid. Obviously, some more thought on what is meant by the terms "best" and "exactly" is necessary.

There may be other quality measures for this tank. It may have weight, manufacturability, and size targets. These may limit the potential r and l values or may force the design of a tank that is noncylindrical. As with the P-diagram, this sample problem will be used throughout the chapter to clarify the methods described.

10.4 TRADE-OFF MANAGEMENT

With the increasing demand for complex and interrelated systems comes the challenge of managing the decisions that balance the pieces of the puzzle, all

vying for their share of the scarce resources. For example, on the Marin Mount Vision, there was a continuous trade-off between weight, cost, and additional features. Although the bicycle has a great suspension, to gain this required a trade-off for 29 lb (13.2 kg) and price \$3100. In early-stage design, the trade-off process is especially challenging, as there is limited knowledge, uncertainties are high, and the decisions made have far-reaching effects on the directions pursued thereafter, and hence the affordability, reliability/safety, and effectiveness of the final product. It is clearly more viable and less expensive to refine a design at the time that it is being conceived. Therefore, efforts toward making good decisions at this stage have high payoffs.

A *trade study* is the activity of a multidisciplinary team to identify the most balanced technical solutions among a set of proposed viable solutions. These viable solutions are judged by their satisfaction of a series of measures or cost functions. These measures describe the desirable characteristics of a solution. They may be conflicting or even mutually exclusive. Trade studies, often called trade-off studies, are commonly used in the design of aerospace and automotive vehicles and the software selection process to find the configuration that best meets conflicting performance requirements.

The measures are dependent on variables that characterize the different potential solutions. If the system can be characterized by a set of equations, we can write the definition of the trade study problem as:

Find the set of variables, x_i , that give the best overall satisfaction to the measures:

$$T_1 = f(x_1, x_2, x_3, \dots)$$

$$T_2 = f(x_1, x_2, x_3, \dots)$$

$$T_3 = f(x_1, x_2, x_3, \dots)$$

$$\vdots$$

$$T_N = f(x_1, x_2, x_3, \dots)$$

Here T_j is a target value and $f(\dots)$ denotes some functional relationship among the variables. Generally, one or more of the targets is not fixed at a specific value, and it is desired to make these T values as large or small as possible (e.g., weight and cost should be small). These are generally referred to as cost functions, and the other measures are treated as constraints.

If you can write these equations with sufficient fidelity, formal optimization methods can be used to find the optimal trade-off. However, what makes design trade studies most challenging is that much of the critical information is often uncertain, evolving, and may be lacking in fidelity. Further, with team members from many disciplines and with different values about what is important, information may be conflicting.

There are two types of uncertainty in these equations. The first, variability, is the result of the fact that a system can behave in random ways. The weather will change, material properties are variable, or a part that is to be made 10 mm in diameter may be a little larger or smaller than that. In general, even though some

portion of variation can be controlled (e.g., insulation from weather changes, tighter manufacturing tolerances) there is always variation that is either uncontrollable or too expensive or difficult to warrant controlling.

The second type of uncertainty results from the lack of knowledge about a system (i.e., subjective uncertainty or state of knowledge uncertainty). It is a property of the team members' cumulative experience and the amount of time they have spent on the current or similar concepts. Both types of uncertainty are direct causes of risk. In a world with no variability and perfect knowledge, there would be no risk.

Trade studies are essentially decision-making exercises—choose an optional concept or course of action from a discrete or continuous set of viable alternatives. The decision analysis matrix (aka, Pugh's method) or the robust decision methods in Section 8.7 can be used when optimization is not possible. But, before using these methods, a better understanding of variability is needed.

10.5 ACCURACY, VARIATION, AND NOISE

In Section 5.3 we discussed modeling using physical prototypes, analysis, and graphical representation. Regardless of the type of model, the goal of modeling is to find the easiest method by which to evaluate the product for comparison with the engineering targets using available resources. To compare the product under development with the engineering targets means that numerical values must be produced; even a rough value is better than no value at all.

In any model (regardless of the level of fidelity), two kinds of errors may occur: *errors due to inaccuracy* and *errors due to variation*. Accuracy is the correctness or truth of the model's estimate. If there is a distribution of results (each time we measure the performance, we get a different number), then the estimate is the mean value of the distribution. With an accurate model, the best estimate will be a good predictor of product performance; with an inaccurate model, it will be a poor one. The variation in the results obtained from the model refers to the statistical variation of the results about the mean value—where accuracy tells “how much,” distribution tells “how sure.” In Fig. 10.5 the inaccurate estimate is shown with a small variation and the accurate estimate with a large one. The obvious goal in modeling is to develop an accurate model with a small variation. The next best model is accurate with a large variation.

Why so much concern about variation? Every parameter that defines a product or process has variation and so each may vary greatly from the desired mean. During production, not all samples of the product are exactly the same size, are made of exactly the same material, or behave in exactly the same way. For example, the actual dimensions of components that were specified on a drawing to be 38.1 ± 0.06 mm are shown in Fig. 10.6. The target for this dimension was 38.1 mm. However, during manufacturing, the values ranged from 0.03 mm below the target to 0.07 mm above. It was only during inspection that the tool cutting the component was found to be worn, thus causing the 0.07-mm deviation from the mean. The tool was replaced.

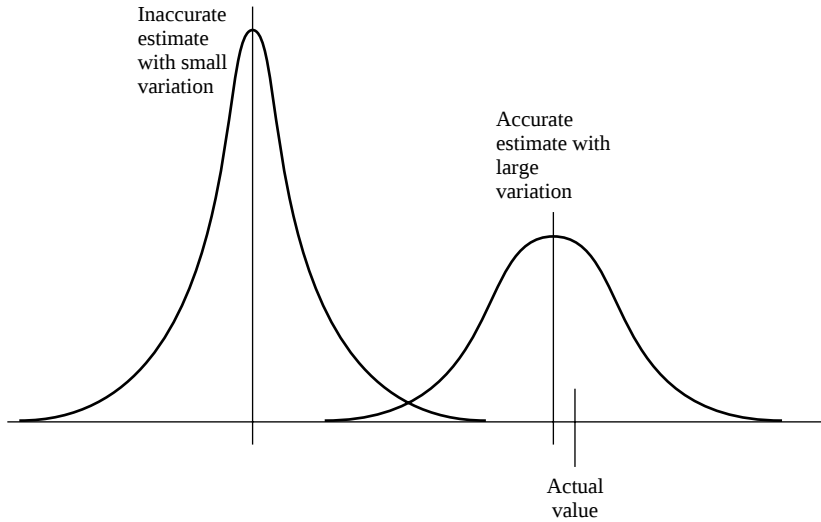


Figure 10.5 Relation of accuracy and resolution of error in modeling.

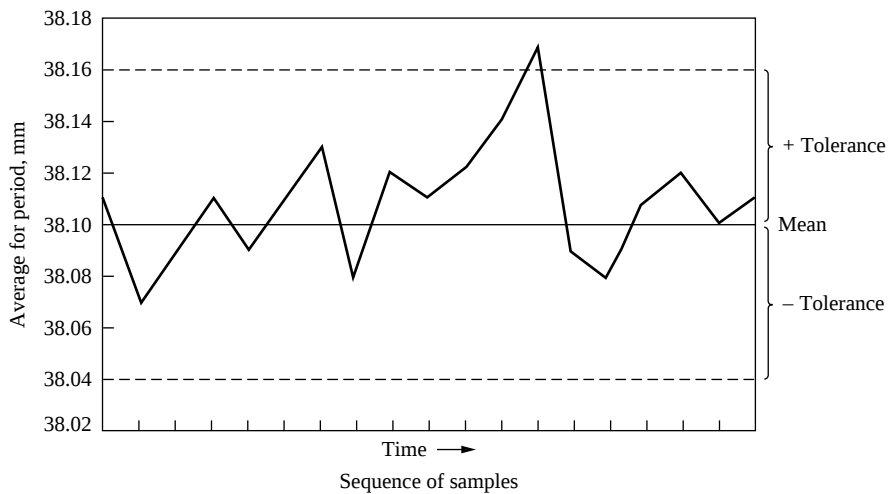


Figure 10.6 Manufactured component distribution relative to design specification.

Nothing is deterministic, everything uncertain.

Another example of variation's importance during design is shown in the data in Fig. 10.7. These data represent the tensile strength for 913 samples of 1035 hot-rolled steel. The data have been grouped to the closest 1 kpsi. The tensile strength varies by as much as 10 kpsi from the mean. These data are replotted

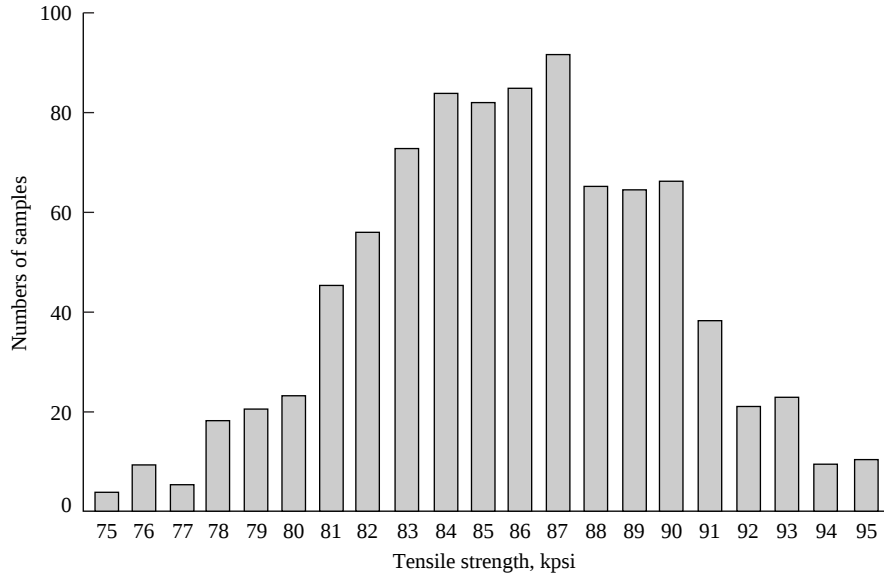


Figure 10.7 Distribution of tensile strength of 1035 steel.

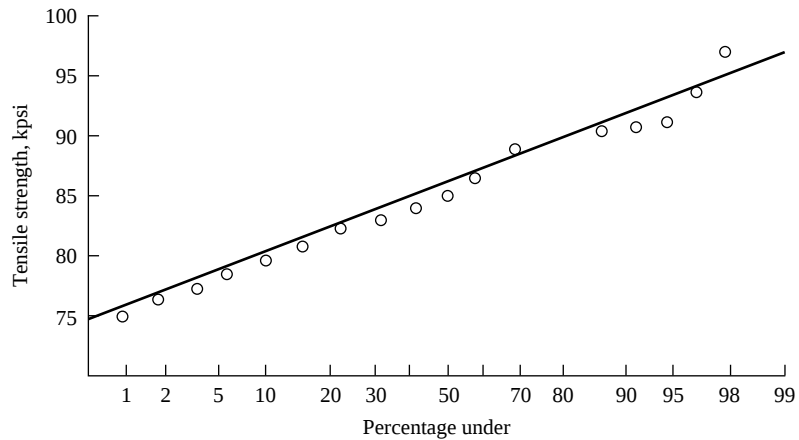


Figure 10.8 Steel data plotted on normal-distribution paper.

in Fig. 10.8 on normal-distribution paper. Since a straight line fits the data in this figure, the tensile strength of the sample material is normally distributed. From Fig. 10.8, the mean strength is 86.2 kpsi (the 50% point) and the standard deviation is 3.9 kpsi. (Details on normal distributions are in App. B.)

The data for the steel in Figs. 10.7 and 10.8 were fit by a normal distribution. The data in Fig. 10.6 for the dimension of a component, although skewed

when the tool wore, were also close to being normally distributed. For most design parameters, variations in value are considered as normal distributions fully characterized by the mean and variance or standard deviation.

However, most analytical models are *deterministic*—that is, each variable is represented by a single value. Since *all* parameters are really distributions, this single value is generally assumed to be the mean. Calculations performed with only mean value information may or may not give accurate estimates. Regardless of accuracy, these models give no information on the variation of the estimated value. There are, however, *nondeterministic*, or *stochastic*, analytical methods that account for both the mean and the variation by using methods from probability and statistics.

10.5.1 The Effect of Variation on Product Quality

In Table 1.1, we listed the results of a customer survey about what determines quality. Based on this survey, the most essential factors in a quality product are “works as it should,” “lasts a long time,” and “is easy to maintain.” The first of these implies that not only does the product match its targets, but that it also stays on them regardless of variations in operating conditions or age of the product, and that all samples of the product work the same. The second quality factor says that the product’s operation and looks should not vary with time. The third says that its operation should not vary or need adjustment or other attention as it ages or is used in different situations. We can reduce all of this to one statement that defines product quality:

A product is considered to be of high quality if its quality measures stay on target regardless of parameter variation due to manufacturing, aging, or the environment.

“Quality measures” are those engineering requirement targets identified in the House Of Quality and result in customer satisfaction. The product quality definition is very important. In fact, designers go to great length to control some parameters so that they won’t have an effect on the quality measures. For example,

- Controlling the temperature of food so it won’t spoil regardless of room temperature
- Controlling the feel of power steering so the driver’s steering experience stays constant regardless of road conditions
- Controlling the dimensions of a part so they will fit with other parts regardless of manufacturing, temperature, or aging

However, some parameters are impossible to control or can be controlled only at great cost. These parameters will be separated out from those that can be controlled and are called noise parameters, as shown in refined P-diagram in Fig. 10.9.

The variations in the dimensions (Fig. 10.6) and in the material properties (Fig. 10.7) are examples of types of noise that affect the performance of a product. These are uncontrollable or can only be controlled at great cost and thus are

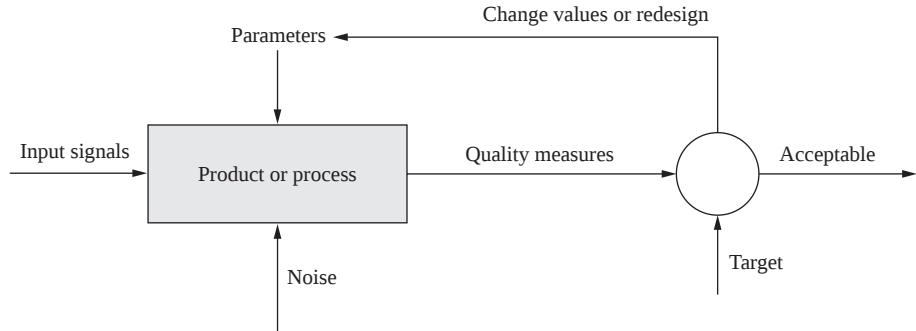


Figure 10.9 The P-diagram with noise and control parameters.

Hand fitting parts is fun when making a prototype,
a disaster on the assembly line.

treated as uncontrollable. Noises affecting the design parameters are generally classified as

- *Manufacturing, or unit-to-unit, variations*, including dimensional variations, variations in material and other properties, and process variations such as those in manufacturing and assembly.
- *Aging, or deterioration, effects*, including etching, corrosion, wear, and other surface effects, along with material property or shape (creep) changes over time.
- *Environmental, or external, conditions*, including all effects of the operating environment on the product. Some environmental conditions, such as temperature or humidity variations, affect the material properties; others, such as the amount of paper in the tray of a paper feeder or the amount of load on a walkway, affect the operating stresses, strains, or positions.

All of these types of noises are inherent in the final product. They all affect the variation in the product's performance. *A quality product is one that is insensitive to noise and thus has a small variation in performance as factors vary.*

Noises that affect strength are often accounted for by using a *factor of safety*. Two methods for calculating the factor of safety are given in App. C. In both of these, noises are caused by uncertainties in knowledge about the material properties, the load causing the stress, unit-to-unit variations, and the ability to analyze failure.

As an example of noises and their effects, we revisit the simple tank problem; namely, what radius and length should a tank be to hold 4 m^3 of liquid? Because the length and radius of the tank must be manufactured, they are not exact values, but have manufacturing variations. Thus, the control parameters, r and l , are actually

distributions about nominal values. If r and l are distributions, then the quality measure, the volume V , must also be a distribution and thus cannot be “exact.” The problem is now reduced to determining the dependence of the distribution of V on r and l and finding the values of r and l that make V as exact as is possible.

Making this problem even more difficult, the liquid that will be stored in the tank is corrosive and over time will etch the inside of the tank, increasing the values of r and l . Additionally, the tank will be installed on Mars and thus operate at a wide range of temperatures, so r and l will vary. Even with the effects of manufacturing variance, the aging effects due to etching, and the environmental effects of the temperature variation, it is still our goal to keep the volume as close to 4 m^3 as possible. Thus, we want to find the values for l and r that make V the least sensitive to noise, that is, manufacturing, aging, and environmental variations.

Consider a second example, the Marin Mount Vision suspension system. Figure 10.10 shows a P-diagram for a bicycle suspension system. During the design of the Mount Vision, a key goal of the team was to ensure that the suspension gave quality performance. During the designers’ effort to understand the problem, they developed the QFD diagram with engineering specifications that defined a quality product. Three of these specifications were for vertical accelerations during different riding conditions:

1. Maximum acceleration on a standard street
2. Maximum acceleration on a 2.5-cm standard pothole
3. Maximum acceleration on a 5-cm standard pothole

Translating these specifications to a P-diagram, the street surface or pothole is the input signal, the maximum acceleration is the quality measure, and the targets are as shown. Also shown in the P-diagram are the control and noise parameters.

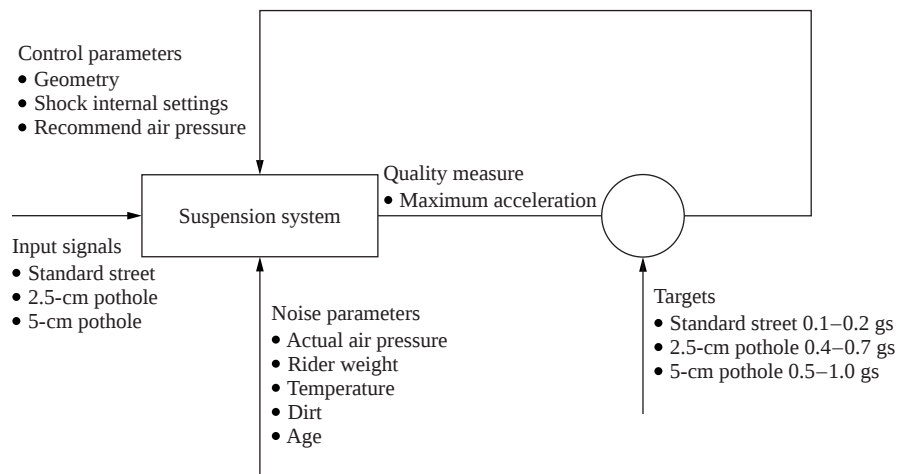


Figure 10.10 P-diagram for bicycle suspension performance.

The design team had control over the dimensions of the suspension system, some of the internal settings in the air shock, and the recommended air pressure for the shock. What they did not have control over was

- The actual air pressure in the shock
- The weight of the rider
- The temperature
- The dirt buildup on the shock
- The age of the shock

These parameters are all noises. Marin's riders will consider the Mount Vision a quality product if it meets the quality measures and is insensitive to these noises.

In general, there are four ways to deal with noises. The first is to keep them small by tightening manufacturing variations (generally expensive). The second is to add active controls that compensate for the variations (generally complex and expensive). The third is shielding the product from aging and environmental effects (sometimes difficult and maybe impossible). The fourth is to make the product insensitive to the noises. A product that is insensitive to manufacturing, aging, and environmental noises is considered *robust* and will be perceived as a high-quality product. If robustness is accomplished, the product will assemble as designed and will be reliable once in operation. Thus, the key philosophy of robust design is to:

Determine values for the parameters based on easy-to-manufacture tolerances and default protection from aging and environmental effects so that the best performance is achieved. The term "best performance" implies that the engineering requirement targets are met and the product is insensitive to noise. If noise insensitivity cannot be met by adjusting the parameters, then tolerances must be tightened or the product shielded from the effects of aging and environment.

With such a philosophy, quality can be designed into a product. For example, in 1981 Xerox had a line fallout that was 30 components per thousand, in other words, 1 out of every 33 components did not fit into the product during assembly. This failure to fit was discovered either during inspection or by the inability of the assembly personnel or machine to mate the components to the product. This high rate led to great expense in reworking components or disposing of them. By 1995, using the robust design philosophy, Xerox had reduced the line fallout to about 30 components per million, 1 out of every 33,000.

Designing robustness into a product is the topic of Sections 10.8 and 10.9. First, a background in modeling, and tolerance and sensitivity analysis is needed.

10.6 MODELING FOR PERFORMANCE EVALUATION

Concern for accuracy and variation is reflected in choosing the best modeling technique for each situation. Although each modeling challenge is different, the

steps discussed here give order to the considerations taken into account during evaluation. The discussion is centered on analytical and physical modeling. A 2008 survey about the use of virtual (analytical) simulations versus physical prototypes in mechatronics found that top companies perform an average of 25 simulations and 6 physical prototypes. Companies that struggle to meet time and cost goals, on the other hand, average 5 simulations and 8 prototypes. Further, simulations are generally less expensive on systems that are sufficiently understood to model. The following steps can help in making the simulation/physical prototype decision.

10.6.1 Step 1: Identify the Output Responses (i.e., the Critical or Quality Parameters) That Need to Be Measured

Often the goal in evaluation is to see if a new idea is feasible. Even with this ill-defined goal, the important critical parameters, those that determine the performance, must be clearly identified. In developing engineering requirements and targets during the specification development phase of the design process, many parameters of interest are identified. As the product is refined, other important requirements and targets arise. Thus, throughout the development of the product, the parameters that demonstrate the performance of the product are identified and measured during product evaluation.

10.6.2 Step 2: Note the Needed Fidelity

Early in the product refinement, it may be sufficient to find only the order of magnitude of some parameters. Back-of-the-envelope calculations may be sufficient indicators of performance for relative comparisons. As the product is refined, the accuracy of the evaluation modeling must be increased to enable comparison with the target values. It is important to realize the degree of fidelity needed before beginning the evaluation. Effort spent on a finite-element model is wasted if a rough calculation using classical strength-of-materials techniques or a simple laboratory test of a piece of actual material is sufficient. Getting this wrong can lead to “paralysis by analysis”—overanalyzing to the point that progress is stifled.

10.6.3 Step 3: Identify the Input Signal, the Control Parameters and Their Limits, and Noises

It is important, before beginning to model a system, that a P-diagram is drawn and the factors affecting the output be at least initially identified and classified. Input signals are the energy, information, and materials modified by the product or process. Usually these signals are important; however, they may be secondary to the control parameters and ignored in many design situations.

Many noises can be controlled if desired; however, controlling them when a design can be made robust is needlessly expensive. Thus, list as noises all

unit-to-unit, aging, and environmental variations that can be identified. Then decide which may have an impact on the output (this may be dependent on the outcome of evaluation).

Control parameters are sometimes difficult to identify, and it is not until a model (either analytical or physical) is built and tested that some dependencies are discovered. One may build a model only to find that the variables thought to be important are not and other, more important variables have been left out of consideration.

It is important to list the control parameters and their upper and lower limits. Considering these limits helps in understanding the design and aids in the development of the layout drawing. The physical limits on these parameters give the limits on patching the design during iteration. Knowledge about limits is one measure of technology readiness discussed in Section 8.4.

10.6.4 Step 4: Understand Analytical Modeling Capabilities

Generally, analytical methods are less expensive and faster to implement than physical modeling methods. However, the applicability of analytical methods depends on the level of accuracy needed and on the availability of sufficient methods. For example, a rough estimate of the stiffness of a diving board can be made using methods from strength of materials. In this analysis, the board is assumed to be a cantilever beam, made of one piece of material, of constant prismatic cross section, and with known moment of inertia. Further, the load of a diver bouncing on the end of the board is estimated to be a constant point load. With this analysis, the important dependent variables—the energy storage properties of the board, its deflection, and the maximum stress—can be estimated.

Using more sophisticated and advanced strength of materials modeling techniques, the fidelity of the model is improved. For example, the taper of the diving board, the distributed nature of the diver in both time and space, and the structure of the board can be modeled. The dependent variables remain unchanged. More parameters that are independent can now be utilized in a more laborious and more accurate evaluation.

Finally, using finite-element methods, even more accuracy can be achieved, though at a higher cost in terms of time, expertise, and equipment. If the diving board is made of a composite material, it may even be that no finite-element methods are yet available to allow for sufficiently accurate evaluation.

Each of these three analytical evaluation methods (basic strength of materials, advanced strength of materials, and finite-element analysis) provides a single answer for the stiffness of the diving board. Although the fidelity of the models and accuracy of the results vary, none gives an indication of the variance in the stiffness; all three are deterministic and provide single answers insensitive to the variations in the dependent variables. Methods are presented later in this chapter for using deterministic analytical models that give stochastic information and thus data on the variance of the dependent parameters.

In this discussion on analytical modeling, a number of issues were raised:

- What level of accuracy is needed? Analytical models can be used instead of physical models only when there is a high degree of confidence in their fidelity.
- Are analytical models available of sufficient fidelity to give the needed accuracy? If not, then physical models are required. Often it is valuable to do both to confirm one's understanding of the product.
- Are deterministic solutions sufficient? They probably are in the early evaluation efforts. However, as the product is finalized, they are not sufficient, as knowledge of the effect of noises on the dependent parameters is essential in developing a quality product.
- If no analytical techniques are available, can new techniques be developed? In developing a new technology, part of the effort is often devoted to generating analytical techniques to model performance. During a design effort, there is usually no time to develop very sophisticated analytical capabilities.
- Can the analysis be performed within the resource limitations of time, money, knowledge, and equipment? As discussed in Chap. 1, time and money are two measures of the design process. They are usually in limited supply and greatly influence the choice of the modeling technique used. Limitations in time and money can often overwhelm the availability of knowledge and equipment.

10.6.5 Step 5: Understand the Physical Modeling Capabilities

Physical models, or prototypes, are hardware representations of all or part of the final product. Most design engineers would like to see and touch physical realizations of their concepts all the way through the design process. However, time, money, equipment, and knowledge—the same resource limitations that affect analytical modeling—control the ability to develop physical models. Generally, the fact that physical models are expensive and take time to produce, controls their use.

However, the ability to develop physical prototypes of complex components has improved greatly since the mid-1980s. During this period, *rapid prototyping* methods were developed. These systems use solid models of components to deposit materials or laser-harden polymers to rapidly make a physical model. The components made by some of the methods are actually usable in tests; others are only visual and usable to test fit and interference.

You can't BS hardware.

10.6.6 Step 6: Select the Most Appropriate Modeling Method

There is nothing as satisfying in engineering as modeling a system both analytically and physically and having the results agree! However, resources rarely allow both modeling methods to be pursued. Thus, the method that yields the needed accuracy with the fewest resources must be selected.

10.6.7 Step 7: Perform the Analysis or Experiments and Verify the Results

Document that the targets have been met or that the model has given a clear indication of what parameters to alter, which direction to alter them in, and how much to alter them. In evaluating models, not only are the results as important as in scientific experimentation, but since the results of the modeling are used to patch or refine the product, the model must also give an indication of what to change and by how much. In analytical modeling, this is possible through sensitivity analysis, as will be discussed in Section 10.7. This is more difficult with physical models. Unless the model itself is designed to allow easily changed parameters, it may be difficult to learn what to do next.

For the Marin suspension system, steps 1–3 are included in the P-diagram developed in Section 10.5 (Fig. 10.10). The goal of step 4 is to understand the analytical modeling capabilities. The engineers at Marin had some simulation capability, but this was only sufficient to ensure that the performance was in the range of the targets. They felt that the best results could be found with physical hardware (steps 5 and 6). Thus, they built a test bike and instrumented it for measuring acceleration. They also set up a test track with 2.5- and 5-cm potholes. Tests were performed with riders of differing weights and with pressures different than those recommended. They also experimented with dirt on the shock and with heating and chilling it. Their goal was to find the best configuration of the parameters they could control and be insensitive to the noises.

10.7 TOLERANCE ANALYSIS

This section focuses on manufacturing variations and tolerances. We begin with a discussion of the relationship between tolerances and manufacturing variations. This is followed by a general discussion about tolerancing, how it affects manufacturing costs and how it is sometimes confused with quality. We then develop two methods for analyzing how tolerances stack up, how the tolerances on a number of components affect the total assembly. These analyses also serve as a basis for the sensitivity analysis in Section 10.8 and robust design in Section 10.9.

Costs generally increase exponentially with tighter tolerances.

10.7.1 The Difference Between Manufacturing Variations and Tolerance

The data in Fig. 10.6 show the manufacturing variation in components that are all supposed to have the same dimension. Also shown are lines ± 0.06 mm from the mean value. These represent the tolerance the designer specified for the dimension. The manufacturing engineer used this value to determine which manufacturing machine to use in making the component. The machinist or the quality-control inspectors used it for determining when to change the tool. Theoretically, the tolerance is assumed to represent ± 3 standard deviations about the mean value. This implies that 99.68% of all the samples should fall within the tolerance range. In actuality, the variation is controlled by a combination of tool control and inspection, as shown in Fig. 10.6.

Recently, the best practice has been to manufacture to “6-sigma.” This term implies that six standard deviations of manufactured product are within tolerance. As stated in Chapter 1, Six Sigma is a quality-oriented best practice that uses the five-step DMAIC process (Define, Measure, Analyze, Improve, and Control). The “measure” in this process is the generation of data, as in Fig. 10.6. Now the focus turns to “analyze.”

10.7.2 General Tolerancing Considerations

Concern about tolerances on dimensions and other variables (i.e., material properties) that affect the product is the focus of *tolerance design*. If the nominal tolerances do not give sufficient performance of the quality measures, then tolerances need to be changed to meet the targets. A drawing of a component or an assembly to be manufactured is incomplete without tolerances on all the dimensions. These tolerances act as bounds on the manufacturing variations such as shown in Fig. 10.6. However, studies have shown that only a fraction of the tolerances on a typical component actually affect its function. The remainder of the dimensions on a typical product could be outside the range set by their tolerances and it would still operate satisfactorily. Thus, when specifying tolerances for noncritical dimensions, always use those that are nominal for the manufacturing process specified to make the component. For example, as shown in Fig. 10.11, operations have *nominal tolerances*. These values for steel reflect the expected variation if standard practices are followed. If tighter tolerances are specified, the cost will increase, as shown in the figure. Most companies have specifications for tolerances that are within the nominal variations for each process. Specifying tolerances tighter than these may require special approval.

Once the designer has specified a tolerance on a drawing, what does it mean downstream in the product life cycle? First, it communicates information to manufacturing that is essential in helping to determine the manufacturing processes that

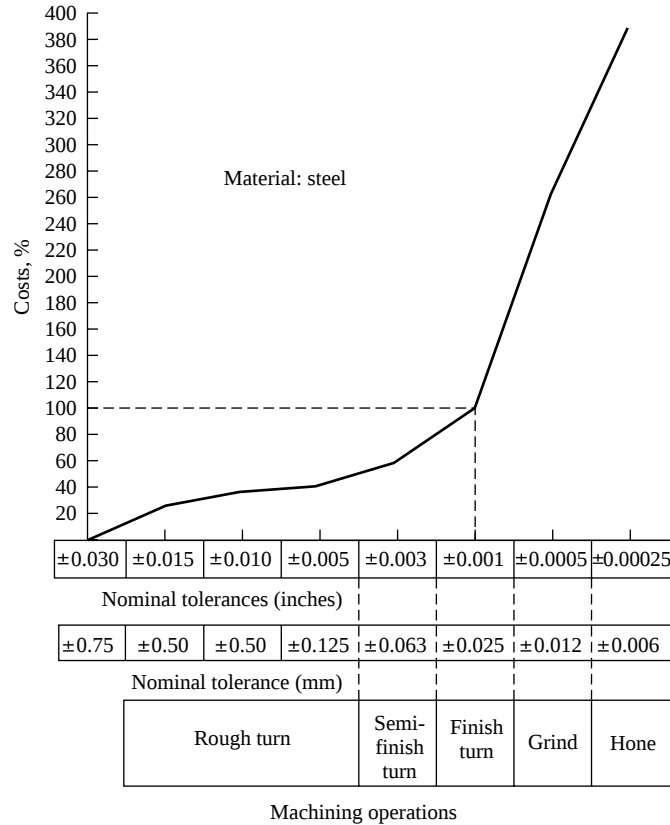


Figure 10.11 Tolerance versus manufacturing process.

will be used. Second, tolerance information is used to establish *quality-control* guidelines, as shown in Fig. 10.6. Quality maintained by comparing the manufactured components to the dimensions and tolerances specified on drawings is called conformance quality. It is a weak form of quality control as it is only as good as what is specified on the drawing.

In the 1920s, when mass production was instituted on a broad scale, quality control by inspection was also begun. This type of quality assurance is often called “on-line,” as it occurs on the production line. Most production facilities have quality-control inspectors whose job it is to verify that produced products are within specified tolerances. This effort to increase product quality through inspection is not very robust, because poor manufacturing process control and poor design can make quality inspection very difficult.

In the 1940s, much effort was expended on designing the production facilities to turn out more uniform components. This moved some of the responsibility of quality control from on-line inspection to off-line design of production

processes. To keep manufactured components within their specified tolerances, many statistical methods were developed for manufacturing process control. However, even if a production process can keep a manufactured component within the specified tolerances, there is no assurance of a robust, quality product. Thus, it was realized in the 1980s that quality control is really a design issue. If robustness is designed in, the burden of quality control is taken off production and inspection.

10.7.3 Additive Tolerance Stack-up

To introduce tolerance stack-up consider the joint in Figs. 9.29, 10.12 and 10.13 for Marin's connection of the air shock to the forward pivot. This is a fairly complex pivot developed in Chap. 9. It consists of two bushings and shock body held between the two fingers of the frame. A shaft goes through the fingers and bushings, holding the assembly together and transferring the forces between the air shock and the frame. The air shock pivots on the bushings, which are clamped between the fingers of the frame. The problem addressed here is how big to make the spacing between the frame fingers and the length of the bushings so the parts all assemble easily. If the spacing between the fingers is too narrow, it will be difficult to get the bushings between them. If the spacing is too wide, then either the shock will rattle or, if the nuts on the end of the shaft are tightened sufficiently, the fingers will be flexed, adding unneeded stress. So, the questions are: What dimension to make the spacing between the fingers? And, How do the tolerances affect the assembly?

Since the bushing is 20 mm long and each flange is 2 mm thick, in the ideal, deterministic world, the spacing should be $20 + 2 * 2 = 24$ mm. However, all the components have variation and it is important to understand how the tolerances

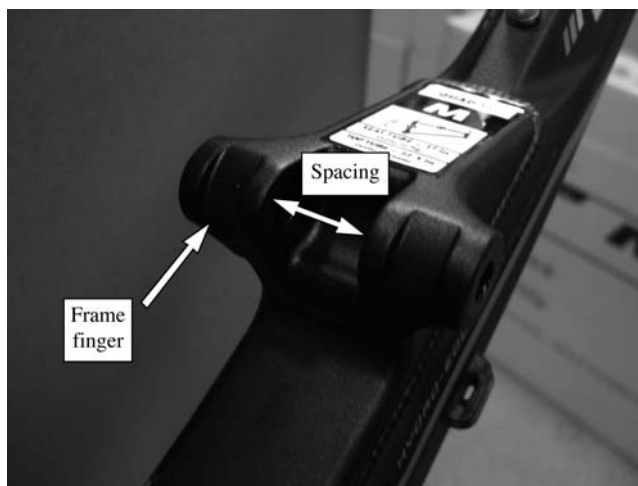


Figure 10.12 Shock-frame connection.

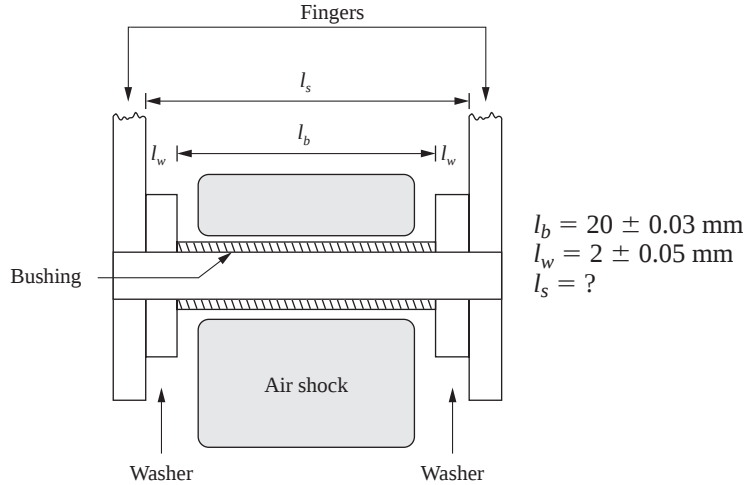


Figure 10.13 Details of connection.

on them add together, or stack up. Analysis of tolerance stack-up is the most common form of tolerance analysis. For this analysis, the notation is

l = dimension

$\bar{l}$ = mean dimension

t = tolerance on dimension

s = standard deviation of dimension

The subscripts refer to

b = bushing length

w = washer thickness

s = distance between fingers

g = gap (+ = clearance, - = interference)

When the joint is assembled, the bushing and washers, if smaller than the distance between the fingers will leave a clearance. If the bushing and two washers are larger than the distance between the fingers, the gap will be negative, an interference. In general,

$$l_g = l_s - (l_b + 2 \times l_w) \quad [10.1]$$

Assume, as an example, that a randomly selected bushing is within the tolerances and at the minimum length ($l_b = 19.97$ mm). Similarly, the washers are also the minimum ($l_w = 1.95$ mm). Finally, say the distance between the fingers is specified to be 24 ± 0.1 mm and the sample randomly chosen happens to be the maximum allowable, $l_s = 24.1$ mm. This situation could happen if the components are selected at random, yet the odds of selecting the thinnest joint components

and widest spacing are low. Suppose it did occur; then, using Eq. [10.1] the gap would be 0.23 mm. Similarly, if the widest bushing and fattest washers were put in the narrowest spacing, there would be 0.23-mm interference.

This analysis implies that if you want assembly to be easy, no interference, then you should specify $l_s = 24.33 \pm 0.1$ mm (the narrowest possible distance between the fingers will still fit the widest components), then you know all possible combinations of components will fit. Of course, some randomly selected combinations may have a gap as much as 0.56 mm $24.43 - (19.97 + 2 * 1.95)$. Overtightening the bolt may close this gap, but it will also add high stresses to the frame.

The method just followed, one of adding the maximum and minimum dimensions to estimate the stack-up, is called *worst-case analysis*. This technique assumes that the shortest and longest components are as likely to be chosen as some intermediate value. In reality, the odds are that the components will be nearer to the mean than to either of their extreme values. In other words, the probability of the two assemblies in the previous paragraphs occurring from the random selection of components is very small.

A much better method is to use statistical stack-up analysis.

10.7.4 Statistical Stack-Up Analysis

A more accurate estimate of the gap can be found statistically. Consider a stack-up problem composed of n components, each with mean length l_i and tolerance t_i (assumed symmetric about the mean), with $i = 1, \dots, n$ (n is the number of uniaxial dimensions). If one dimension is identified as the dependent parameter (in the suspension example, the gap), then its mean dimension can be found by adding and subtracting the other mean dimensions, as in Eq. [10.1]. In general,

$$\bar{l} = \bar{l}_1 \pm \bar{l}_2 \pm \bar{l}_3 \pm \dots \pm \bar{l}_n \quad [10.2]$$

The sign on each term depends on the structure of the device. Similarly, the standard deviation is

$$s = \left(s_1^2 + s_2^2 + \dots + s_n^2 \right)^{1/2} \quad [10.3]$$

where the signs are always positive. (This basic statistical relation is discussed in App. B.) Generally, “tolerance” is assumed to imply three to six standard deviations about the mean value. More recently, this has, in some high-technology industries even been as high as 9-sigma. For 3-sigma, a tolerance of 0.009 in. means that $s = 0.003$ and that 99.73% of all samples should be within tolerance (i.e., within 3σ). Since $s = t / 3$, Eq. [10.3] can be rewritten as

$$t = \left(t_1^2 + t_2^2 + \dots + t_n^2 \right)^{1/2} \quad [10.4]$$

For the example,

$$\bar{l}_g = \bar{l}_s - (\bar{l}_b + 2 \times \bar{l}_w) \quad [10.5]$$

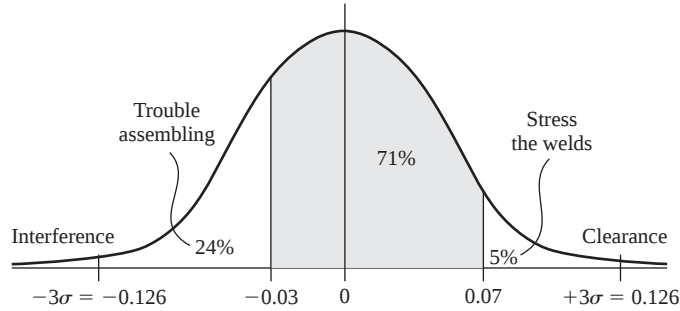


Figure 10.14 Gap distribution.

and

$$t_g = (t_s^2 + t_b^2 + 2 \times t_w^2)^{1/2} \quad [10.6]$$

Say that we make the spacing 24.00 ± 0.10 mm, then the gap and the tolerance on it are

$$\bar{l}_g = 24 - (20 - 2 \times 2) = 0.0 \text{ mm} \quad [10.7]$$

and

$$t_g = (0.10^2 + 0.03^2 + 2 \times 0.05^2)^{1/2} = 0.126 \text{ mm} \quad [10.8]$$

These results show that there is, on the average, no gap and the tolerance on it is 0.126 mm. Say that the fingers can flex up to 0.07 mm inward when bolted without undue stress on the welds to compensate for any clearance. Further, say that assembly personnel can get the parts in between the fingers even if there is a 0.03-mm interference. The question then is, what percentage of the assemblies will meet these requirements?

This situation is plotted in Fig. 10.14. Assuming the tolerance calculated is 3 standard deviations, and using standard normal probability methods (App. B) the shaded area represents 71% of the assemblies. This means that 29% of the time either the assembly people will have trouble assembling the device (24%) or the welds will be overstressed (5%).

Inspecting each joint and reworking those that do not meet the specification or swapping components between joints to meet them could be used to achieve increased quality. Another way to increase the quality is to use the results of the analysis to redesign the joint. This is accomplished through sensitivity analysis.

10.8 SENSITIVITY ANALYSIS

Sensitivity analysis is a technique for evaluating the statistical relationship of control parameters (e.g., dimensions) and their tolerances in a design problem.

In this section, we explore the use of sensitivity analysis for a simple dimensional problem and then apply the method to the problem of the tank volume.

Sensitivity analysis enables the contribution of each parameter to the variation to be easily found. Rewriting Eq. [10.3] in terms of $P_i = s_i^2/s^2$,

$$1 = P_1 + P_2 + \cdots + P_n \quad [10.9]$$

where P_i is the percentage contribution of the i th term to the tolerance (or variance) of the dependent variable. For the current example, these are

$$P_s = \frac{0.10^2}{0.126^2} = 0.63 = 63\%$$

$$P_b = \frac{0.03^2}{0.126^2} = 0.05 = 5\%$$

$$P_w = \frac{0.05^2}{0.126^2} = 0.16 = 16\%$$

$$\text{With two washers} \quad \text{total} = 1.0 = 100\%$$

This result clearly shows that the tolerance on the spacing has the greatest effect on the gap. For one-dimensional tolerance stack-up problems such as this, the results of the sensitivity analysis can be used for *tolerance design*. Since the spacing causes 63% of the noise in the joint, it is the most likely candidate for change.

This technique will work on all one-dimensional problems in which all the parameters are dimensions on the product. To summarize:

- Step 1.** Develop a relationship between the dependent dimension and those it is dependent on, as in Eq. [10.2] or [10.5]. Using each independent dimension's mean value, calculate the mean value of the dependent dimension.
- Step 2.** Calculate the tolerance on the dependent variable using Eq. [10.4] or work in terms of the standard deviations (Eq. [10.3]).
- Step 3.** If the tolerance found is not satisfactory, identify which independent dimension has the greatest effect, using Eq. [10.9], and modify it if possible. Depending on the ease (and expense), it may be necessary to choose a different dimension to modify.

Problems of two or three dimensions are similarly solved, but the equations relating the variables become complex for all but the simplest multidimensional systems.

If the variables are not related in a linear fashion, the equations already given are modified. This is best shown through the tank-volume problem introduced earlier in this chapter. The major difference is that the parameters r (the radius) and l (the length) are not linearly related to the dependent variable, V (the volume), as can be seen in Fig. 10.4. The method shown next is a generalization of the method for the linear problem. It is good for investigating any functional relationship, whether or not the parameters are dimensions.

Consider a general function

$$F = f(x_1, x_2, x_3, \dots, x_n) \quad [10.10]$$

where F is a dependent parameter (dimension, volume, stress, or energy) and the x_i 's are the control parameters (usually dimensions and material properties). Each parameter has a mean $\bar{x}_i$ and a standard deviation s_i . In this more general problem, the mean of the dependent variable is still based on the mean of the independent variables, as in Eq. [10.2]. Thus,

$$\bar{F} = f(\bar{x}_1, \bar{x}_2, \bar{x}_3, \dots, \bar{x}_n) \quad [10.11]$$

Here, however, the standard deviation is more complex:

$$s = \left[\left(\frac{\partial F}{\partial x_1} \right)^2 s_1^2 + \dots + \left(\frac{\partial F}{\partial x_n} \right)^2 s_n^2 \right]^{1/2} \quad [10.12]$$

Note that if $\partial F / \partial x_i = 1$, as it must in a linear equation, then Eq. [10.12] reduces to Eq. [10.3]. Equation [10.12] is only an estimate based on the first terms of a Taylor series approximation of the standard deviation. It is generally sufficient for most design problems.

For the tank problem, the independent parameters are r and l . The mean value of the dependent variable V is thus given by

$$\bar{V} = 3.1416 \bar{r}^2 \bar{l} \quad [10.13]$$

To evaluate this, we must consider specific values of r and l . There is an infinite number of these pairs that meet the requirement that the mean volume be 4 m^3 . For example, consider point A in Fig. 10.15 (which is Fig. 10.4 with added information). With $\bar{r} = 1.21 \text{ m}$ and $\bar{l} = 0.87 \text{ m}$, from Eq. [10.13], $\bar{V} = 4 \text{ m}^3$.

The tolerances on these parameters can be based on what is easy to achieve with nominal manufacturing processes. For example, take $t_r = 0.03 \text{ m}$ ($s_r = 0.01$) and $t_l = 0.15 \text{ m}$ ($s_l = 0.05$). These values are shown in the figure as an ellipse around point A. Using formula [10.12], the standard deviation on this volume is

$$s_v = \left[\left(\frac{\partial V}{\partial l} \right)^2 s_l^2 + \left(\frac{\partial V}{\partial r} \right)^2 s_r^2 \right]^{1/2} \quad [10.14]$$

where

$$\frac{\partial V}{\partial r} = 6.2830rl$$

and

$$\frac{\partial V}{\partial l} = 3.1416r^2$$

For the values in this example, $\partial V / \partial r = 6.61$ and $\partial V / \partial l = 4.60$, so $s_v = 0.239 \text{ m}^3$. Thus, 99.68% (three standard deviations) of all the vessels built will have volumes within 0.717 m^3 (3×0.239) of the target 4 m^3 . Also, the percentage contribution of

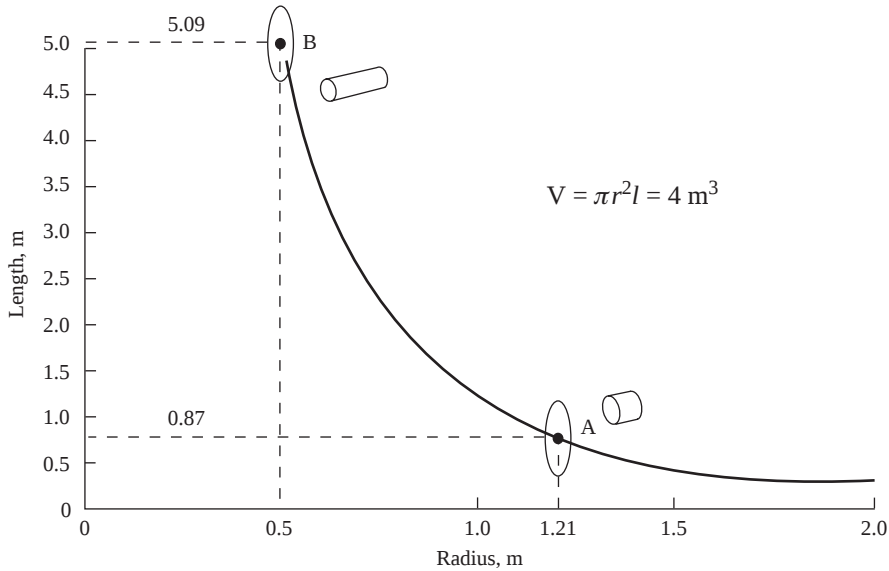


Figure 10.15 Effect of noise on the potential solutions for the tank problem.

She who does not design a robust product
will be cursed with unhappy customers.

each parameter can be found as in Eq. [10.9]. Here the length contributes 92.3% of the variance in volume. However, noting that the tolerance on the length is much larger than that on the radius and considering the shape of the curve in Fig. 10.14, it is evident that a longer vessel with a smaller radius might yield a smaller variance in volume. If the control parameters are taken at $r = 0.50 \text{ m}$ and $l = 5.09 \text{ m}$ (point B in Fig. 10.14), the mean volume is still 4 m^3 . Now $\partial V / \partial r = 16.00$ and $\partial V / \partial l = 0.78$, so $s_v = 0.166 \text{ m}^3$, which is 31% smaller than at point A. Also, now the tolerance on r contributes 94% to the variance in the volume. Note that we achieved *the reduction in variance not by changing the tolerances on the parameters, but by changing only their nominal values*. The second design has higher quality because the volume is always closer to 4 m^3 . If we can find the values of the parameters r and l that give the smallest variance on the volume, then we are employing the philosophy of robust design.

10.9 ROBUST DESIGN BY ANALYSIS

Robust design is often called *Taguchi's method* after Dr. Genichi Taguchi, who popularized the philosophy in the United States and Japan. It must be noted that this philosophy is different from that traditionally used by designers. Traditionally,

A robust design is insensitive to noise. Noise is what the designer cannot control or chooses not to control.

parameter values are determined without regard for tolerances or other noises and the tolerances are added on afterward. These tacked-on tolerances are usually based on company standards. This philosophy does not lead to a robust design and may require tighter tolerances to achieve quality performance.

The implementation of robust design techniques is fairly complex. To ease our explanation of the techniques here, we will make two simplifying assumptions. First, we will consider only noise due to manufacturing variations; second, the only parameters that are considered are dimensions. As a basis for understanding robust design, we will build on dimensional tolerances and sensitivity analysis. Additionally, we first develop robust design analytically so that the philosophy is better understood. The actual methods Taguchi developed are based on experiments rather than analysis and require a background in statistical data reduction beyond the scope of this text. Thus, these experimental methods are only briefly introduced in Section 10.10.

In Section 10.8 we saw that by merely changing the shape of the tank we could improve the quality of the design. The tank with the greater length had less sensitivity to the large tolerance on the length, so the tank's volume varies less. Our goal now is to combine the techniques of sensitivity analysis and optimization to develop a method for determining the most robust values for the parameters. Then we will consider tightening the tolerances to make the best tank possible.

Consider the initial problem: the goal was to have $V = 4 \text{ m}^3$, exactly. This is impossible, as V is dependent on r and l and they are random variables, not exact values. Thus, the "best" we can do is to keep the absolute difference between V and 4 m^3 as small as possible, in other words, minimize the standard deviation of V . We must accomplish this minimization while keeping the mean volume at 4 m^3 . Defining the difference between the mean value ($3.1416\bar{r}^2\bar{l}$) and the target $T (4\text{m}^3)$ as the bias, the objective function to be minimized is

$$C = \text{variance} + \lambda \times \text{bias} \quad [10.15]$$

where λ is a Lagrange multiplier.¹

Using Eq. [10.12], this looks like

$$C = \left[\left(\frac{\partial F}{\partial x_1} \right)^2 s_1^2 + \cdots + \left(\frac{\partial F}{\partial x_n} \right)^2 s_n^2 \right] + \lambda(F - T) \quad [10.16]$$

¹Many different optimization methods could be used. Lagrange's method is well suited to this simple problem.

For the tank,

$$C = (2\pi rl)^2 s_r^2 + (\pi r^2)^2 s_l^2 + \lambda(\pi r^2 l - T)$$

The minimum value of the objective function can now be solved. With known standard deviations on the parameters s_r and s_l (or tolerances t_r and t_l) and a known target T , values for the parameters r and l can be found from the derivatives of the objective function with respect to the parameters and the Lagrange multiplier:

$$\frac{\partial C}{\partial r} = 0 = 2r(2\pi l)^2 s_r^2 + 4r^3 \pi^2 s_l^2 + \lambda 2\pi r l$$

$$\frac{\partial C}{\partial l} = 0 = 2l(2\pi r)^2 s_r^2 + \lambda \pi r^2$$

$$\frac{\partial C}{\partial \lambda} = 0 = \pi r^2 l - 4$$

Solving simultaneously results in

$$r = 1.414l \left(\frac{s_r}{s_l} \right) \quad [10.17]$$

and

$$l = \left[\frac{2}{\pi} \left(\frac{s_l}{s_r} \right)^2 \right]^{1/3} \quad [10.18]$$

Thus, for any ratio of the standard deviations or the tolerances, the parameters are uniquely determined for the best (most robust) design. For the values of $s_r = 0.01$ ($t_r = 0.03$ m) and $s_l = 0.05$ ($t_l = 0.15$ m), these equations result in $r = 0.71$ m and $l = 2.52$ m. Substituting these values into Eq. [10.14], the standard deviation on the volume is $s_v = 0.138$ m³. Comparing this to the results obtained in the sensitivity analysis, 0.239 and 0.165 m³, the improvement in the design quality is evident.

If the radius were harder to manufacture than the length, say $s_r = 0.05$ and $s_l = 0.01$, then, using Eqs. [10.17] and [10.18], the best values for the parameters would be $r = 2.06$ m and $l = 0.29$ m. The resulting standard deviation on the volume would be 0.233 m³.

In summary, the tolerance or standard deviation information on the dependent variables has been used to find the values of the parameters that minimize the variation of the dependent variable. In other words, the resulting configuration is as insensitive to noise as possible and is thus a robust, quality design.

If the standard deviation on the volume is not small enough, then the next step is to tighten the tolerances.

Robust design can be summarized as a three-step method:

Step 1. Establish the relationship between quality characteristics and the control parameters (for example, Eq. [10.10]). Also, define a target for the quality characteristic.

- Step 2.** Based on known tolerances (standard deviations) on the control variables, generate the equation for the standard deviation of the quality characteristic (for example, Eq. [10.12] or [10.14]).
- Step 3.** Solve the equation for the minimum standard deviation of the quality characteristic subject to this variable being kept on target. For the example given, Lagrange's technique was used; other techniques are available, and some are even included in most spreadsheet programs. There are usually other constraints on this optimization problem that limit the values of the parameters to feasible levels. For the example given, there could have been limits on the maximum and minimum values of r and l .

There are some limitations on the method developed here. First, it is only good for design problems that can be represented by an equation. In systems in which the relationships between the variables cannot be represented by equations, experimental methods must be used (Section 10.10). Second, Eq. [10.15] does not allow for the inclusion of constraints in the problem. If the radius, for example, had to be less than 1.0 m because of space limitations, Eq. [10.15] would need additional terms to include this constraint.

10.10 ROBUST DESIGN THROUGH TESTING

It is often impossible to analytically evaluate a proposed design because no mathematical models of the system exist or the fidelity of those that do are too low. In many cases, even when analysis is possible, the analytical model of the system may not allow determination of the effect of the noise on a proposed design. In either case, it is necessary to design and build a physical model for experimental testing. In Chap. 8, physical models were used to verify the concept; now they are needed to refine the product. The material in this section is an introduction to Taguchi's experimental method for the robust design of products. Like many other topics, this subject has entire books devoted to it (see Section 10.12, Sources); however, this material is sufficient to make us appreciate the strength of the method and its complexity and enable us to apply it to the design of the tank as a simple example. We will assume that we do not know the formula $V = \pi r^2 l$ and only know $V = f(r, l)$. To experimentally find dimensions for radius and length, we could begin by building a tank with some best-guess dimensions and then measuring the volume. Then, if the volume was too high, we could build new models, one with a smaller radius and another with a shorter length, and then measure the volumes. Based on these new measurements, we could try to estimate the dependence of the volume on each of the dimensions and iterate (i.e., patch) our way to the target volume. This is the way most experiments are run. This "random walk" toward a solution may require many models, so it is not very efficient. Additionally, the solution found could be anywhere on the curve shown in Figs. 10.4 and 10.15; there is no guarantee that the final design will be the most robust. The following steps can overcome these drawbacks.

10.10.1 Step 1: Identify Signals, Noise, Control, and Quality Factors (i.e., Independent Parameters)

Referring back to the P-diagram in Fig. 10.9, it is necessary to list all the dependent and independent parameters related by the product or system. Then it is necessary to decide which of these are critical to the evaluation of the product. Sometimes this is not easy, and critical parameters or noises may be overlooked. This may not become evident until data are taken and the results are found to have wide distribution, implying that the model is not complete or the experiments have been poorly done. It is essential to take care here to understand the system.

The P-diagram for the tank (Fig. 10.16) shows that the designer has control over the length and radius and that there are many noises that affect the volume of liquid held. The function of the tank is to “hold liquid,” and its performance is measured by how accurately the tank can be held to the target value of 4 m^3 . The noises include the manufacturing variations on the radius and length, and the aging and environmental effects not considered here.

10.10.2 Step 2: For Each Quality Measure (i.e., Output Response) to Be Evaluated, Recall or Determine Its Target Value and the Nature of the Quality Loss Function

During the development of the QFD, target values were determined and the shape of the loss function (see Table 10.2) was identified. If this information has not been previously generated for the parameter being measured, do this before the experiment is developed.

Loss is proportional to the Mean Square Deviation, MSD, the average amount the output response is off the target. This amount is also often referred to as the Signal-to-Noise ratio, or S/N ratio. Generally the S/N ratio is $-10 \log (\text{MSD})$. The minus is included so that the maximum S/N ratio is the minimum quality loss, the 10 is used to get the units to decibels, and the logarithm is used to compress the values.

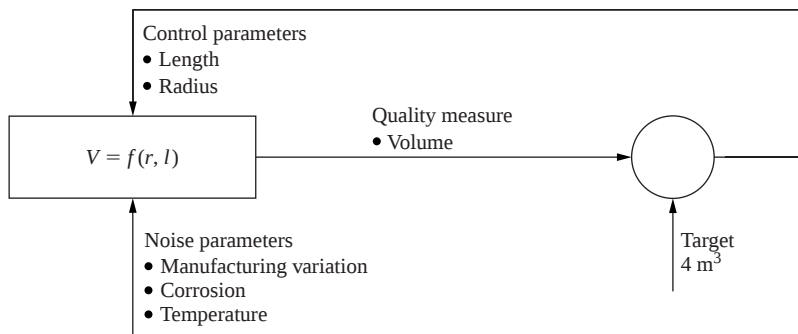


Figure 10.16 P-diagram for tank problem.

Table 10.2 Formulas for means and S/N ratios

Quality loss function	Mean square deviation (MSD)	S/N ratio
Smaller-is-better	$\frac{1}{n} \sum_{i=1}^n y_i^2$	$-10 \log \left(\frac{1}{n} \sum_{i=1}^n y_i^2 \right)$
Larger-is-better	$\frac{1}{n} \sum_{i=1}^n \left(\frac{1}{y_i^2} \right)$	$-10 \log \left(\frac{1}{n} \sum_{i=1}^n \frac{1}{y_i^2} \right)$
Nominal-is-best	$\frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2 + (\bar{y} - m)^2$ $m = \text{target value}$	$-10 \log \frac{1}{n} \sum_{i=1}^n (y_i - \bar{y})^2$

The MSD and S/N for the three most common types of targets identified in Section 6.8 are shown in Table 10.2. For the smaller-is-better target, the larger the value of the output, y , the larger the MSD and the smaller the S/N ratio. In other words, larger values of y are noise, so the signal is weaker relative to that noise. For the larger-is-better case, smaller values of y are seen as noise.

The nominal-is-best target is more complex; there are many ways to calculate the S/N ratio. The most common is shown here. As shown in Table 10.2, the mean square deviation is simply the sum of the variation about the mean and the accuracy about the target. Generally, only the sum of the variation is used in calculating the S/N ratio, as shown in the table.

For the tank problem, 4 m^3 is a nominal-is-best target.

Parameter design is based on maximizing the S/N ratio and then tuning the parameters to bring the design on target. In other words, the goal is to find the conditions that make the product insensitive to noise and then use parameters that do not affect the S/N ratio to bring the quality functions to the desired value. The use of this philosophy will become clear in the example problem.

10.10.3 Step 3: Design the Experiment

The goal is to design an experiment that forces what ever can happen, to happen. It is not sufficient to design a simple experiment in which the model is patched and patched until it works once. This does not lead to a robust design. Instead, the experiment should be designed so that the results give a clear understanding of the effects on the output response of changing control parameters and an understanding of the effects of noise. An ideal experiment will show how to adjust the control parameter to meet the target and show which one to choose so that the resulting system is insensitive to noise.

The physical model of the product or system must be designed so that these can be achieved:

- Control factors can be changed to represent the options available. This may mean designing a number of different physical devices or designing one with

changeable parts or configuration. This model may not be very representative of the final product because its main goal is to support the collection of data.

- Noises can be controlled over the expected range. This may require precision components made to match the upper and lower bounds of tolerances. It may require the use of an environmental chamber capable of temperature, humidity, or other noise control. It may require the components to be artificially aged, corroded, or worn. The noises must be forced to expected extremes so that the effect on output responses can be measured.
- The output responses can be measured accurately. Note that in measuring the output, additional noise is added by the instrumentation. Ensure that this noise is of a lower order of magnitude than the effect of the noise and control variables.

Suppose there are n control factors and data are taken for each at two different settings, there are m noise variables also to be tested at two levels, and, for accuracy, there are k repetitions to be run for each condition. Then there are $k \cdot 2^n \cdot 2^m$ experiments to perform. For example, if there are two control factors, two noises, and three repetitions for each condition, then there are 48 output responses to be recorded. To keep the number of experiments to a reasonable level, on large problems there are statistically based techniques for choosing a subset of experiments to run. These experiments allow the missing data to be inferred (see the book by Taguchi, Chowdhury, and Wu in Sources, Section 10.12).

Table 10.3 shows a layout for an experiment with two control factors, each tested at two levels with two noises also each at two levels. The results for the output response, F , are shown for the 16 experiments. If, for example, there were three repetitions of experiment $F_{21_{12}}$ (control factor 1 at level 2, control factor 2 at level 1, noise 1 at level 1, and noise 2 at level 2), then there would be three $F_{21_{12}}$ values. If all the experiments were run three times, there would be 48 experiments. The mean value and S/N ratio for each control-factor combination are calculated in the last two columns.

For experiments with more than two control factors, with control factors run at more than two levels, or for more than two noises, Table 10.3 is easily extended. Again, for a large number of control factors or noises there are methods of reducing the number of experiments.

Table 10.3 Layout for a two-control-factor experiment

		Noise 1:	Level 1	Level 1	Level 2	Level 2		
		Noise 2:	Level 1	Level 2	Level 1	Level 2		
Control factor 1	Control factor 2						Mean	S/N
Level 1	Level 1		$F_{11_{11}}$	$F_{11_{12}}$	$F_{11_{21}}$	$F_{11_{22}}$	$\overline{F_{11}}$	S/N ₁₁
Level 1	Level 2		$F_{12_{11}}$	$F_{12_{12}}$	$F_{12_{21}}$	$F_{12_{22}}$	$\overline{F_{12}}$	S/N ₁₂
Level 2	Level 1		$F_{21_{11}}$	$F_{21_{12}}$	$F_{21_{21}}$	$F_{21_{22}}$	$\overline{F_{21}}$	S/N ₂₁
Level 2	Level 2		$F_{22_{11}}$	$F_{22_{12}}$	$F_{22_{21}}$	$F_{22_{22}}$	$\overline{F_{22}}$	S/N ₂₂

Table 10.4 Tank experiment results

		$\partial r(\text{m}):$ 0.03 0.03 -0.03 -0.03 $\partial l(\text{m}):$ 0.15 -0.15 -0.15 -0.15					
$r(\text{m})$	$l(\text{m})$					Mean (m^3)	S/N, dB
0.5	0.5	0.57	0.31	0.45	0.244	0.396	3.74
0.5	5.5	5.00	4.76	3.91	3.69	4.34	11.87
1.5	0.5	4.81	2.59	4.39	2.40	3.55	4.40
1.5	5.5	41.89	39.53	38.46	36.13	39.00	19.48

For the tank problem, experimental models are built to enable accurate setting of the length and radius. This may require one model for each experiment, or a model may be designed that allows these values to be changed with sufficient accuracy. In Table 10.4 values of $r = 0.5$ and $r = 1.5$ are chosen as the two levels for the radius. These were chosen as the extreme values of Fig. 10.14 and are only a starting place. Likewise, $l = 0.5$ and 5.5 . The noises are set at the tolerance levels representing the length as harder to manufacture than the radius: $l = \pm 0.15$ and $r = \pm 0.03$. These values are entered into Table 10.4. To find the output response for cell F21₁₂, the experiment needs a tank made as precisely as possible with $r = 1.53$ m and $l = 0.35$ m.

10.10.4 Step 4: Take and Reduce Data

The measured volumes of the tank are shown in Table 10.4 along with the calculated values of the mean and nominal-is-best S/N ratio. Mean values and S/N ratios are calculated for repetitions of each set of control and noise conditions. Two of the mean values are fairly close to the target of 4 m^3 . This was the result of luck in choosing the starting values for r and l . In fact, this result raises the question of which one is best, because they have vastly different values for radius and length.

10.10.5 Step 5: Analyze the Results, and Select New Test Conditions If Needed

The first set of experiments may not yield satisfactory results. The goal is to maximize the S/N ratio and then bring the mean value on target. For analytical problems, we can find the true maximum (Section 10.9); here we can only estimate when we reach that point.

For the tank problem, the experiment with the radius $r = 0.5$ m and the length $l = 5.5$ m gives results near the target and with the highest S/N (11.87 dB). The experiments could be stopped here if a mean value of 4.34 m^3 is close enough. The information in the table could also be used to adjust the control parameters to bring the product closer to target. Since experiments with $l = 5.5$ m resulted in better S/N values, r can be estimated to bring the output to 4 m^3 . However, how much to change r may not be evident from the data. A better idea is to perform

experiments by setting new values for r and l around the values found above and taking new readings. This iteration would eventually lead to a volume $V = 4 \text{ m}^3$ and an S/N ratio of 13.69 at $r = 0.71 \text{ m}$ and $l = 2.52 \text{ m}$, the same values found analytically. Note that the S/N value for this final result is only 1.78 dB higher than the first experimental value found. This implies only a mean square deviation change of 50% [working the S/N equation in Table 10.2 backward, $(V_i - \bar{V})^2 = 10^{(1.78/10)}$].

10.11 SUMMARY

- Product evaluation should be focused on comparison with the engineering requirements and also on the evolution of the function of the product.
- Products should be refined to the degree that their performance can be represented as numerical values in order to be compared with the engineering requirements.
- P-diagrams are useful for identifying and representing the input signals, control parameters, noises, and output response.
- Physical and analytical models allow for comparison with the engineering requirements.
- Concern must be shown for both the accuracy and the variation of the model.
- Parameters are stochastic, not deterministic. They are subject to three types of noises: the effects of aging, of environment change, and of manufacturing variation.
- Robust design takes noise into account during the determination of the parameters that represent the product. Robust design implies minimizing the variation of the critical parameters.
- Tolerance stacking can be evaluated both by the additive method and by statistical means.
- Both analytical and experiment methods exist for finding the most robust design.

10.12 SOURCES

- Barker, T. B.: *Quality by Experimental Design*, 3rd edition, Chapman & Hall, 2005. A very good basic text on experimental design methods.
- Mischke, C. R.: *Mathematical Model Building*, Iowa State University Press, Ames, 1980. An introductory text on the basics of building analytical models.
- Papalambros P., and D. Wilde: *Principles of Optimal Design: Modeling and Computation*, Cambridge University Press, New York, 1988. An upper-level text on the use of optimization in design.
- Rubenstein, M. F.: *Patterns of Problem Solving*, Prentice Hall, Englewood Cliffs, N.J., 1975. An introductory book on analytical modeling.
- Taguchi, G., S. Chowdhury, and Y. Wu: *Taguchi's Quality Engineering Handbook*, Wiley Interscience, 2004, The basic book on robust design.

10.13 EXERCISES

- 10.1** For the original design problem (Exercise 4.1):
- Identify the critical parameters and interfaces for evaluation.
 - Develop a P-diagram for each.
 - Choose whether to build physical models for testing or run an analytical experiment for each.
 - Perform the experiments or analysis and develop the most robust product.
- 10.2** For the redesign problem (Exercise 4.2), repeat the steps in Exercise 10.1.
- 10.3** You have just designed a tennis-ball serving machine. You take it out to the court, turn it on, and quickly run to the other side of the net to wait for the first serve. The first serve is right down the middle, and you return it with brilliance. The second serve is out to the left, the third is long, and the fourth hits the net.
- Does your machine have an accuracy or a variation problem?
 - Itemize some of the potential causes of each type of error. Consider the types of “noise” discussed in Section 10.5.
- 10.4** Convince yourself about the applicability of normal distribution by doing these:
- Measure some feature of at least 20 people and plot the data on normal-distribution paper. Easy measurements to make are weight, height, length of forearm, shoe size, or head circumference.
 - Take a sample of 50 identical washers, bolts, or other small objects and weigh each on a precision scale. Plot the weights on normal-distribution paper and calculate the mean and standard deviation.
- 10.5** For these design problems discuss the trade-offs between using analytical models and using experimental models.
- A new, spring-powered can opener
 - A diving board for your new swimming pool
 - An art nouveau shelf bracket
 - A pogo-stick spring

11

C H A P T E R

Product Evaluation: Design For Cost, Manufacture, Assembly, and Other Measures

KEY QUESTIONS

- What is Design For Cost, DFC, and how can costs be estimated?
- What is Design For Value, DFV, and how is value different from cost?
- How can a product be easy to manufacture (DFM) and assemble (DFA)?
- How do Failure Modes and Effects Analysis (FMEA), Fault Tree Analysis (FTA), and Design For Reliability (DFR) help eliminate failures?
- Can products be designed that are easy to test (DFT) and measure (DFM)?
- What can a designer do to protect the environment (DFE)?

11.1 INTRODUCTION

In Chap. 10 we considered the best practices for evaluating the product design relative to performance, tolerance, and robustness. Also of importance are the evaluations for cost, ease of assembly, reliability, testability and maintainability, and environmental friendliness, all covered in this chapter. These evaluations have come to be known as Design For Cost (DFC), Design For Assembly (DFA), DFR, DFT, and so on, or generically—DFX. This is the TLA (Three Letter Acronym) chapter.

11.2 DFC—DESIGN FOR COST

One of the most difficult and yet important tasks for a design engineer in developing a new product is estimating its production cost. It is important to generate a cost estimate as early in the design process as possible and to compare with the

Eighty percent of the cost is incurred by 20% of the components.

cost requirements. In the conceptual phase or at the beginning of the embodiment phase, a rough estimate of the cost is first generated, and then as the product is refined, the cost estimate is refined as well. For redesign problems, where changes are not extreme, early cost estimates may be fairly accurate, because the current costs are known.

As the design matures, cost estimations converge on the final cost. This often requires price quotes from vendors and the aid of a cost estimation specialist. Many manufacturing companies have a purchasing or cost-estimating department whose responsibility it is to generate estimates for the cost of manufactured and purchased components. However, the designer shares the responsibility, especially when there are many concepts or variations to consider and when the potential components are too abstract for others to cost estimate. Before we describe cost-estimating methods for use by designers, it is important to understand what control the design engineer has over the manufacturing cost and selling price of the product.

Since cost is usually a driving constraint, many companies use the term Design For Cost, DFC, to emphasize its importance. This means keeping an evolving cost estimate current as the product is refined.

11.2.1 Determining the Cost of a Product

The total cost of a product to the customer (i.e., the list price) and its constituent parts are shown in Fig. 11.1. All costs can be lumped into two broad categories,

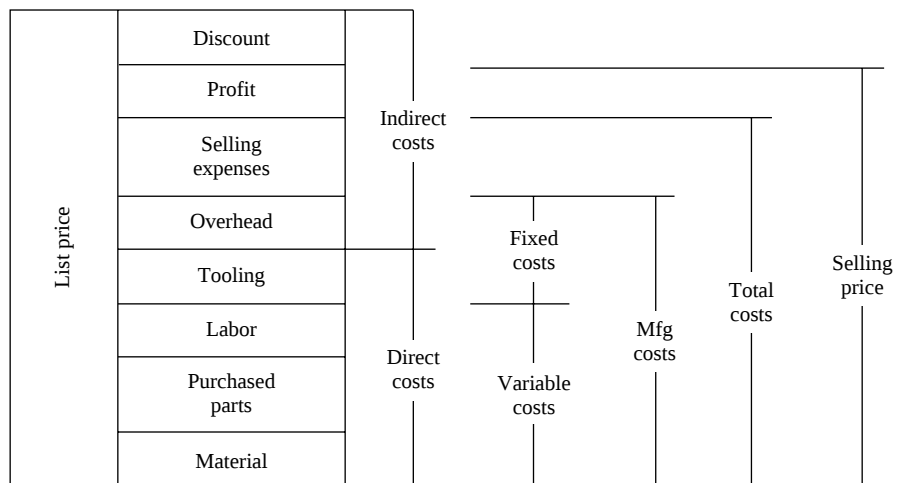


Figure 11.1 Product cost breakdown.

direct costs and indirect costs. *Direct costs* are those that can be traced directly to a specific component, assembly, or product. All other costs are called *indirect costs*. The terminology generally used to describe the costs that contribute to the direct and indirect costs is defined here. Each company has its own method of bookkeeping, so the definitions given here may not match every accounting scheme. However, every company needs to account for all the costs discussed.

A major part of the direct cost is the *material costs*. These include the expenses of all the materials that are purchased for a product, including the expense of the waste caused by scrap and spoilage. Scrap is often an important consideration. For most materials, the scrap can be reclaimed, and the return from the reclamation can be deducted from the material costs. Spoilage includes parts and materials that may not be usable because of manufacturing defects, deterioration, or other damage. Part fallout, those components that cannot be assembled because of poor fit, also contributes to spoilage.

Components that are purchased from vendors and not fabricated in-house are also considered direct costs. At a minimum, this *purchased-parts cost* includes fasteners and the packaging materials used to ship the product. At a maximum, all components may be made outside the company with only the assembly performed in-house. In this case, there are no material costs.

Labor cost is the cost of wages and benefits to the workforce needed to manufacture and assemble the products. This includes the employees' salaries as well as all fringe benefits, including medical insurance, retirement funds, and vacation times. Additionally, some companies include overhead (to be defined shortly) in figuring the direct labor cost. With fringe benefits and overhead included, the labor cost of one worker will be two to three times his or her salary.

The last element of direct costs is the *tooling cost*. This cost includes all jigs, fixtures, molds, and other parts specifically manufactured or purchased for the production of the product. For some products, these costs are minimal; very few items are being made, the components are simple, or the assembly is easy. On the other hand, for products that have injection-molded components, the high cost of manufacturing the mold will be a major portion of the part cost.

Figure 11.1 shows that the sum of the material, labor, purchased parts, and tooling used is the *direct cost*. The *manufacturing cost* is the direct cost plus the *overhead*, which includes all cost for administration, engineering, secretarial work, cleaning, utilities, leases of buildings, and other costs that occur day to day, even if no product rolls out the door. Some companies subdivide the overhead into engineering overhead and administrative overhead, the engineering portion including all expenses associated with research, development, and the design of the product. Many companies subdivide overhead into fixed and variable portions, items such as shop supplies, depreciation on equipment, equipment lease costs, and human resource costs being variable.

The manufacturing cost can be broken down in another important way. The material, labor, and purchased-parts costs are *variable costs*, as they vary directly with the number of units produced. For most high-volume processes, this variation is nearly linear: it costs about twice as much to produce twice as many units.

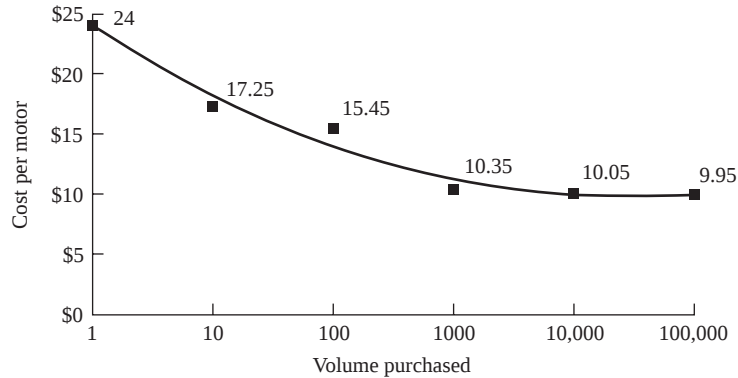


Figure 11.2 Sample of cost per volume purchased for a component.

However, at lower volumes, the costs may change drastically with volume. This is reflected in the price quote made by a vendor for a small electric motor shown in Fig. 11.2.

Other manufacturing costs such as tooling and overhead are *fixed costs*, because they remain the same regardless of the number of units made. Even if production fell to zero, funds spent on tooling and the expenses associated with the facilities and nonproduction labor would remain the same.

In general, the cost of a component, C , can be calculated by:

$$C = C_m + \frac{C_c}{n} + \frac{C_l}{\dot{n}}$$

where C_m is the cost of materials needed for the component (raw materials minus salvage price for scrap), C_c is the capital cost of tooling and a fraction of the cost of the machines and facilities needed, n is the number of components to be made, C_l is the cost of labor per unit time, and $\dot{n}$ is the number of components per unit time. Additionally, if the firm is buying from a vendor, the paperwork and other overhead of selling a small quantity of an item may also appear in C_c . The curve that results from this equation generally looks like that in Fig. 11.2. At low volume, the second and third terms dominate and at high volume the first term, the cost of materials, serves as an asymptote.

The *total cost* of the product is the manufacturing cost plus the selling expenses. It accounts for all the expenses needed to get the product to the point of sale. The actual *selling price* is the total cost plus the *profit*. Finally, if the product has been sold to a distributor or a retail store (anything other than direct sales to the customer), then the actual price to the consumer, the list price, is the selling price plus the *discount*. Thus, the discount is the part of the list price that covers the costs and profits of retail sales. If the design effort is on a manufacturing machine to be used in house, then costs such as discount and selling expenses do not exist. Depending on the bookkeeping practices of the particular company, there may still be profit included in the cost.

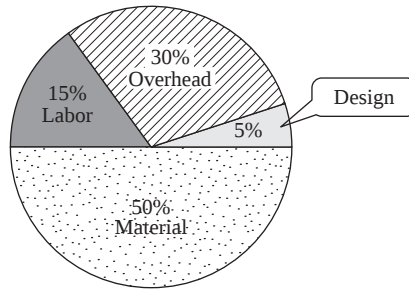


Figure 11.3 Design cost as a fraction of manufacturing cost.

The salaries for the designers, drafters, and engineers and the costs for their equipment and facilities are all part of the overhead. Designers have little control over these fixed expenses, beyond using their time and equipment efficiently. The designer's big impact is on the direct costs: tooling, labor, material, and purchased parts costs. Reconsider Fig. 1.2, reprinted here as Fig. 11.3. These data from Ford show the manufacturing cost, emphasizing the low cost of design activities. If it is assumed that the costs of purchased parts and tooling are included in the material costs, then these account for about 50% of the manufacturing costs. The labor is about 15%, and the overhead, including design expenses, is 35%. As a rule of thumb, for companies whose products are manufactured mainly in house and in high volume, the manufacturing cost is approximately three times the cost of the materials. Also, the selling price is approximately nine times the material cost, or three times the manufacturing cost. This is sometimes called the material-manufacturing-selling 1-3-9 rule. This ratio varies greatly from product to product. The Ford data in Fig. 11.3 show a 1:2 ratio between materials cost and manufacturing cost, less than the rule would predict.

Figure 1.3, reprinted here as Fig. 11.4, shows the influence of design quality on manufacturing cost. As already mentioned, the designer can influence all the direct costs in a product, including the types of materials used, the purchased parts specified, the production methods, and thus the labor hours and the cost of tooling. Management, on the other hand, has much less influence on the manufacturing costs. They can negotiate for lower prices on a material specified by the designer, negotiate lower wages for the workers, or try to trim overhead. With these considerations, it is not surprising that data in Fig. 11.4 show that 50% of the influence on the manufacturing cost is controlled by design.

One final term that should be understood by engineers is *margin*. This is calculated by taking the ratio of profit to selling price. Typically, for product generating companies, a margin of 40–50% will generate a good profit. However, for high-volume production, this may drop to 10%, and for custom production, it may be as high as 60–70%.

To get a feel for these costs, consider a bicycle that has a list price of \$750 (Fig. 11.5). As we can see, only half the list price actually goes into manufacturing

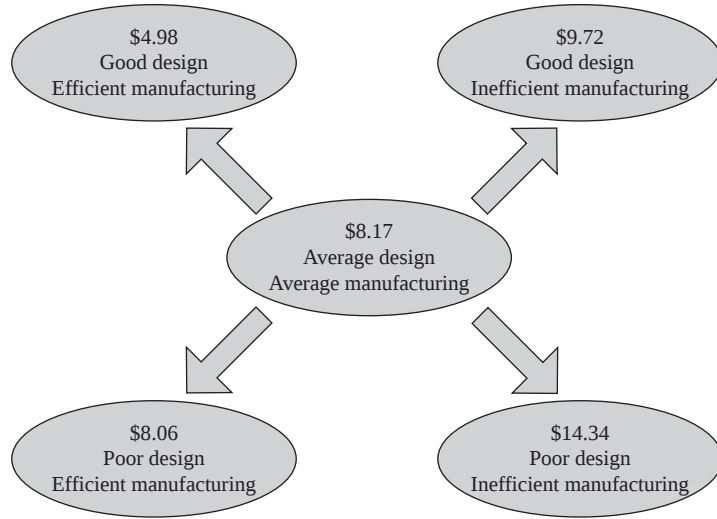


Figure 11.4 The effect of design quality on manufacturing cost.

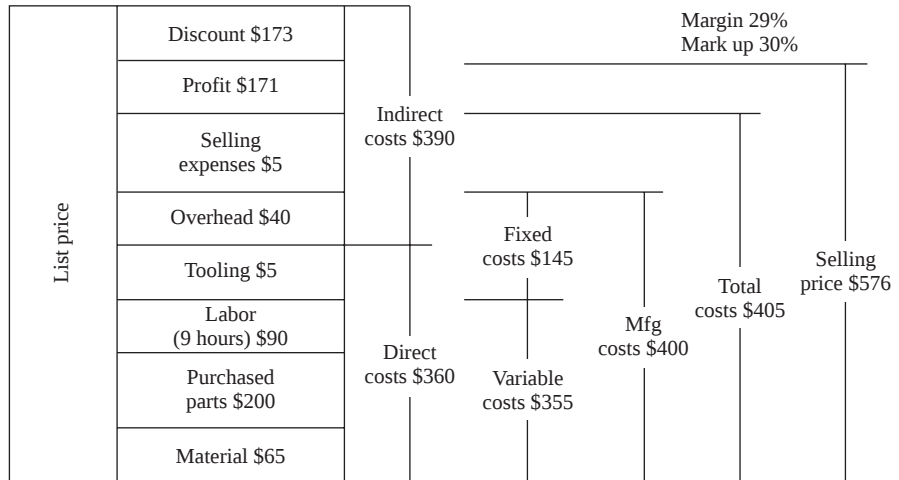


Figure 11.5 Cost breakdown for a \$750 bicycle.

the bicycle (direct costs = \$360). Also, the manufacturing company only makes \$171 profit. Although this seems reasonable, a margin of 29% is just barely high enough to stay in business.

11.2.2 Making a Cost Estimate

It is the responsibility of the engineer to know the manufacturing cost of components designed. The ability to make these estimations comes with experience and with help from experienced team members and vendors. In many companies cost estimating is accomplished by a professional who specializes in determining the

cost of a component whether it is made in house or purchased from a vendor. This person must be as accurate as possible in his or her estimates, as major decisions about the product are based on these costs. Cost estimators need fairly detailed information to perform their job. It is unrealistic for the designer to give the cost estimator 20 conceptual designs in the form of rough sketches and expect any co-operation in return. In most small companies, all cost estimations are done by the engineer.

The first estimations should be made early in the product design phase and be precise enough to be of use in making decisions about which designs to eliminate from consideration and which designs to continue refining. At this stage of the process, cost estimates within 30% of the final direct cost are possible. The goal is to have the accuracy of this estimate improve as the design is refined toward the final product. The more experience one has in estimating similar products, the more accurate the early estimates will be.

The cost-estimating procedure depends on the source of the components in the product. There are three possible options for obtaining the components: purchase finished components from a vendor, have a vendor produce components designed in house, or manufacture components in house.

As discussed in Chap. 9, there are strong incentives to buy existing components from vendors. If the quantity to be purchased is large enough, most vendors will work with the product designer and modify existing components to meet the needs of the new product.

If existing components or modified components are not available off the shelf, then they must be produced, in which case a decision must be made as to whether they should be produced by a vendor or made in house. This is the classic “make or buy” decision, a complex decision that is based on the cost of the component involved as well as the capitalization of equipment, the investment in manufacturing personnel, and plans by the company to use similar manufacturing equipment in the future.

Regardless of whether the component is to be made or bought, cost estimates are vital. We look now at cost estimating for two primary manufacturing processes: machining and injection molding.

11.2.3 The Cost of Machined Components

Machined components are manufactured by removing portions of the material not wanted. Thus, the costs for machining are primarily dependent on the cost and shape of the stock material, the amount and shape of the material that needs to be removed, and how accurately it must be removed. These three areas can be further decomposed into seven significant control factors that determine the cost of a machined component:

1. **From what material is the component to be machined?** The material affects the cost in three ways: the cost of the raw material, the value of the scrap produced, and the ease with which the material can be machined. The first two are direct material costs, and the last affects the amount of labor, the amount of time, and the choice of machines that are used manufacturing the component.

2. **What type of machine is used to manufacture the component?** The type of machine—lathe, horizontal mill, vertical mill, and so on—used in manufacture affects the cost of the component. For each type, there is not only the cost of the machine time itself but also the cost of the tools and fixtures needed.
3. **What are the major dimensions of the component?** This factor helps determine what size of machines of each type will be required to manufacture the component. Each machine in a manufacturing facility has a different cost for use, depending on the initial cost of the machine and its age.
4. **How many machined surfaces are there, and how much material is to be removed?** Just knowing the number of surfaces and the material removal ratio (the ratio of the final component volume to the initial volume) can aid in giving a good estimate for time required to machine the part. Estimates that are more accurate require knowing exactly what machining operations will be used to make each cut.
5. **How many components are made?** The number of components in a batch has a great effect on the cost. For one piece, fixturing is minimal, though long setup and alignment times are required. For a few pieces, simple fixtures are made. For a high volume, the manufacturing process is automated, with extensive fixturing and numerically controlled machining.
6. **What tolerance and surface finishes are required?** The tighter the tolerance and surface finish requirements, the more time and equipment are needed in manufacture.
7. **What is the labor rate for machinists?**

As an example of how these seven factors affect the cost of machined components, consider the component in Fig. 11.6.¹ For this component the seven significant factors affecting cost are

1. The material is 1020 low-carbon steel.
2. The major manufacturing machine is a lathe. Two additional machines need to be used to mill the flat surfaces and drill the hole.
3. The major dimensions are a 57.15-mm diameter and a 100-mm length. The initial raw material must be larger than these dimensions.
4. There are three turned surfaces and seven other surfaces to be made. The final component is approximately 32% the volume of the original.
5. The number of components to be made is discussed in the next paragraph.
6. The tolerance varies over the different surfaces of the component. On most surfaces, it is nominal, but on the diameters, it is a fit tolerance. The surface finish, $.8 \mu\text{m}$ ($32 \mu\text{in.}$), is considered intermediate.
7. The labor rate used is \$35 per hour; this includes overhead and fringe benefits.



¹The cost estimates in this section were made by entering values for these factors on a spreadsheet available as a template that can be used to estimate the cost of any machined part.

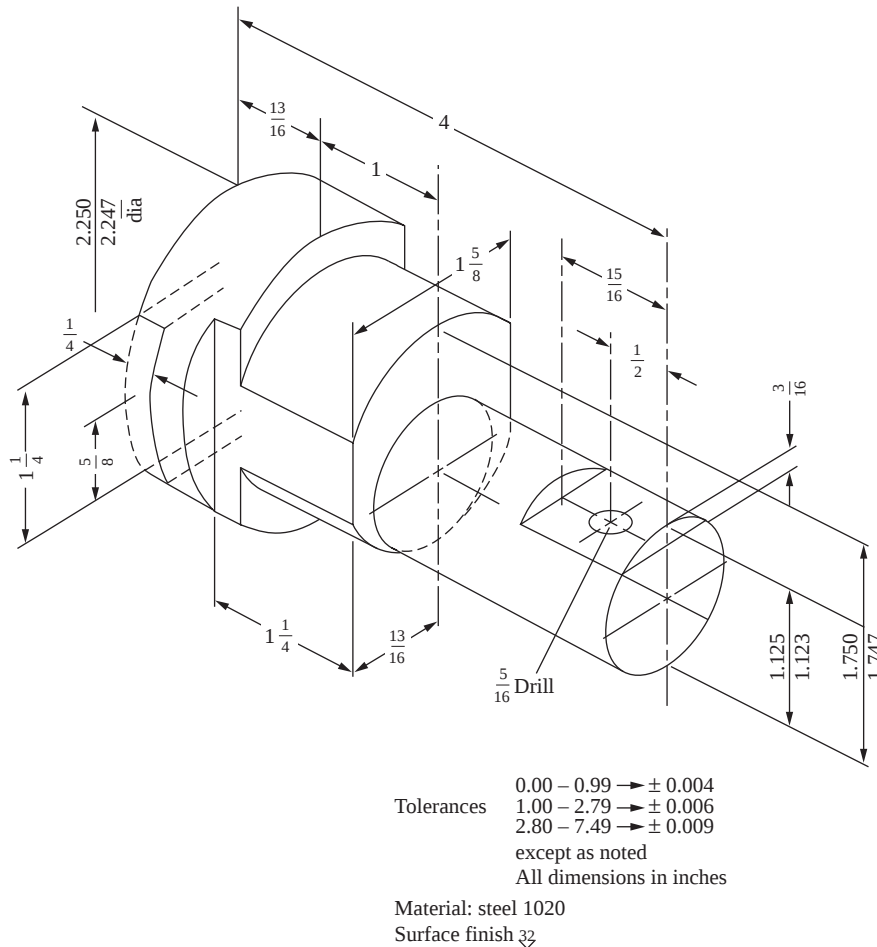


Figure 11.6 Sample component for evaluating machining cost.

Figure 11.7 shows the cost of this component for various manufacturing volumes. The values are the total manufacturing cost per component. The cost of materials per component remains fairly constant at \$1.48, but the labor hours and thus the cost of labor drop with volume. For machined components, the cost dependence on volume is small in quantities above 10 because of the use of Computer-Aided Manufacturing, CAM.

The dependence of the manufacturing cost on other variables is shown in Table 11.1, in which the tolerance, finish, and material are varied. The first three lines show the change with tolerance. A fine tolerance was used for the data in Fig. 11.7 and is shown in line 1. As the tolerance was relaxed to nominal (2) and then to rough (3), the cost dropped.

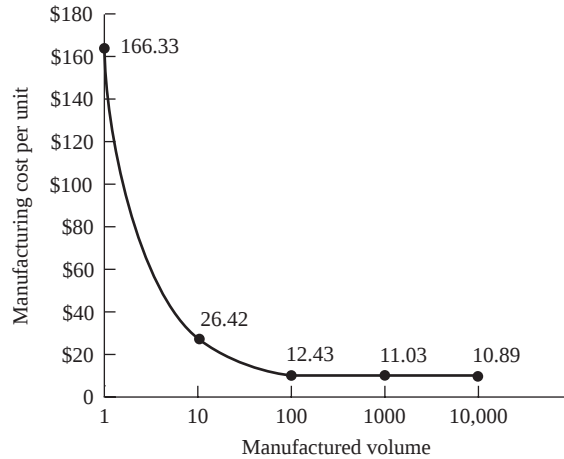


Figure 11.7 Effect of volume on cost.

Table 11.1 Effect of tolerance, finish, and material on cost

Control parameters		
Tolerance	Surface finish	Manufacturing cost
1. Fine	Intermediate	\$11.03
2. Nominal	Intermediate	\$8.83
3. Rough	Intermediate	\$7.36
4. Fine	Polished	\$14.85
5. Fine	As turned	\$8.17
6. High-carbon steel		\$22.45

Note: For 1000 units.

Product cost goes down exponentially with increased production volume.

The fourth and fifth lines show the effect of surface finish on the manufacturing cost. The data in Fig. 11.7 were based on an intermediate surface finish, as specified in the drawing. As this was improved (4), the manufacturing cost rose, and as it was reduced to “as turned” (5), the cost dropped dramatically. Also shown in Table 11.1 is the effect of changing the material from low-carbon steel to high carbon steel (6), which doubles the cost when compared to line 1 in part because of an increase in material cost (+ \$4.00) and an increase in the machining time.

11.2.4 The Cost of Injection-Molded Components

Probably the most popular manufacturing method for high-volume products is plastic injection molding. This method allows for great flexibility in the shape of the components and, for manufacturing volumes over 10,000, is usually cost effective. On a coarse level, all the factors that affect the cost of machined components also affect the cost of injection-molded components. The only differences are that there is only one type of machine, an injection-molding machine, and the questions concerning geometry are modified. Besides the major dimensions of the component, it is important to know the wall thickness and component complexity in order to determine the size of the molding machine needed, the time it will take the components to cool sufficiently for ejection from the machine, the number of cavities in the mold (the number of components molded at one time), and the cost of the mold.

To demonstrate the effect of the factors, we show the cost for a clip, shown in Fig. 11.8.² The significant factors affecting cost are

1. The overall dimensions are 9.46 cm (3.72 in.) by 4.52 cm (1.77 in.) in the mold plane and 4.13 cm (1.6 in.) deep.
2. The wall thickness is 3.2 mm (0.125 in.).
3. The number of components to be manufactured is 1 million.
4. The labor hourly rate is \$35.
5. The tolerance level is intermediate.
6. The surface finish is not critical.

The cost of manufacturing the component in Fig. 11.8 is shown in Fig. 11.9 for varying production volumes. The capital cost of making a mold is high enough to dominate the cost of the component at low volumes. This is why making just 1000 injection-molded plastic parts would be very expensive. A rule of thumb is that if the manufacturing volume is less than 10,000, plastic injection molding may be cost prohibited.

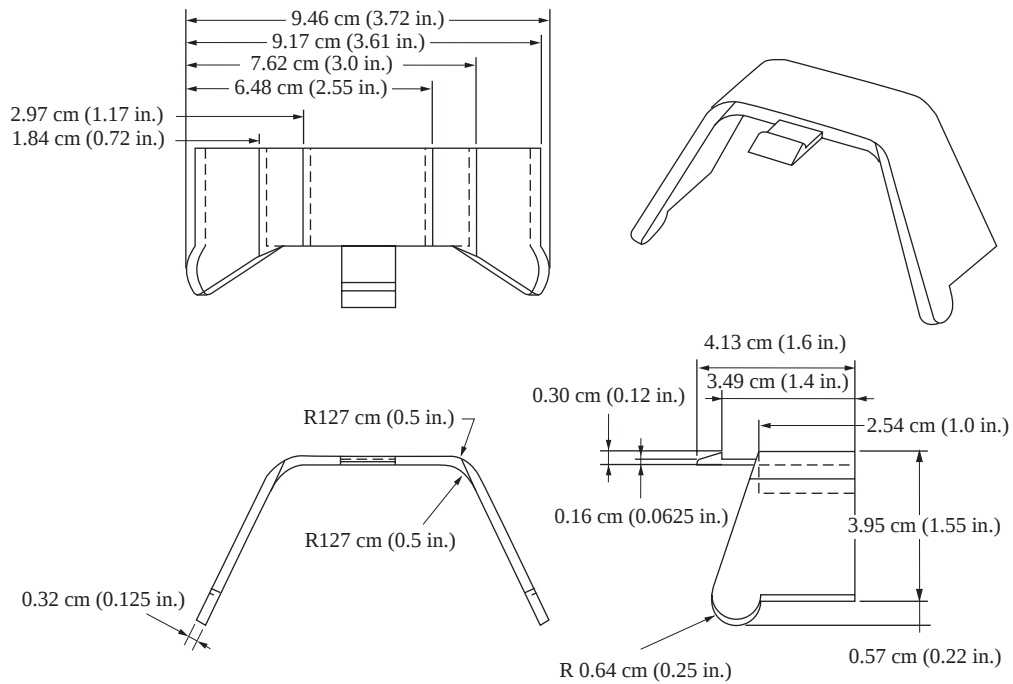
The manufacturing cost can be affected by the wall thickness. In the drawing, the thickness is 3.2 mm. If this is lowered to 2.5 mm, the part cost will drop about 18%. This is primarily because the time needed in the mold for cooling drops from 18 sec to 13 sec, saving cycle time.

11.3 DFV—DESIGN FOR VALUE

The concept of value engineering (also called value analysis) was developed by General Electric in the 1940s and evolved into the 1980s. Value engineering is a customer-oriented approach to the entire design process. It changes the focus from the cost of a component to its value to the customer. The key point of value

²The cost estimates in this section were made by entering values for these factors on a spreadsheet available as a template that can be used to estimate the cost of any machined part.





Brad Tittle Oregon State Univ. December 28, 1990	CLIP	Tol: ± 0.01 cm Approved: <i>#12</i>
--	------	--

Figure 11.8 Component for cost estimation.

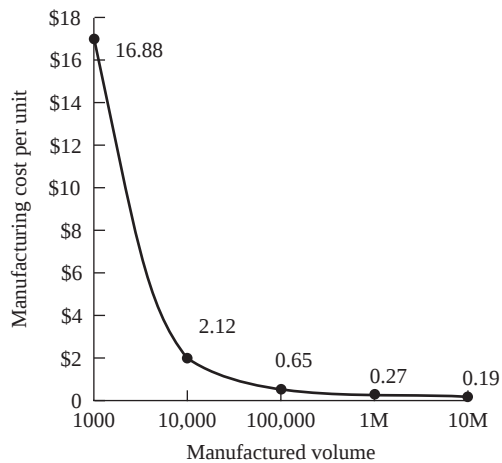


Figure 11.9 The effect of volume on the cost of a plastic part.

engineering is that it is not sufficient to only find cost—it is necessary to find the value of each feature, component, and assembly to be manufactured. Value is defined as

$$\text{Value} = \frac{\text{Worth of a feature, component, or assembly}}{\text{Cost of it}}$$

The worth of a feature of a component, for example, is determined by the functionality it provides to the customer. Thus, a refined definition for value is *function provided per dollar of cost*.

The value formula is used as a theme through the value engineering steps suggested here. These steps are focused on features of components. The method can also be applied to components and assemblies.

Step 1 To ensure that all the functions are known, for each feature of a component ask the question, What does it do? If a feature provides more than one function, this fact must be noted. Features that result from a specific manufacturing operation are at the finest level of granularity that should be considered. For the machined component in Fig. 11.6, each turned diameter and face, each milled surface, and the hole should be considered. For the injection-molded plastic part in Fig. 11.8, the 6.4-mm-radius round feature at the bottom is a good feature to query. This feature provides a number of functions.

Step 2 Identify the life-cycle cost of the feature. This cost should include the manufacturing cost as well as any other downstream costs to the customer. If the feature provides multiple functions, the cost should be divided into cost per function. To do this, consider an equivalent feature that provides only the function in question. Although it is not accurate because of the interdependence of functions, it gives an estimate.

The cost of the round feature ($R = 0.64$ cm) in Fig. 11.8 is not evident. Consultation with tooling and manufacturing engineers revealed that, for a volume of 100,000 components (\$0.65 component cost in Fig. 11.9) \$0.02 was due to this feature. Their logic was that the feature does not contribute to labor cost because the cycle time would not change if the feature were removed. They estimated that, since the feature was hard to machine in the mold, it contributed about 5% to the mold cost. Amortized over the production volume, this gives \$0.017. Finally, the material used for this feature is worth \$0.003. So the feature costs \$0.02 total. It could be argued that the structure of the body of the component should be included because it contributes to the function of the round feature. A decision has to be made as to where to allocate all the costs in the component, one of the challenges of value engineering.

Step 3 Identify the worth of the function to the customer. In an ideal world, we would be able to ask customers how much each function was worth to them. However, this is not realistic. To obtain at least a qualitative indication of function worth, the information developed in the QFD is used. If no formal method was used to develop the customers' requirements and measures of importance,

then the best that can be done is to ask, How important is this feature to the customer?

The feature being used as an example contributes to a number of functions that are very important to the customer. To complicate matters, each of these functions involves other features. The best that can be done is to say that the functions contributed to by the round feature are worth a great deal to the customer. A customer will not pay as much for a product that is hard to attach, so the engineers estimated the worth at \$2.00. Keep in mind that this method compares relative values, and not the values themselves.

Step 4 Compare worth to cost to identify features that have low relative value. If one feature costs more than the others and is worth more—provides important function to the product—then its value may be as high as or higher than the others. On the other hand, if its costs outweigh its worth, then it has low value and should be redesigned.

The round feature contributes to a number of important functions for very low cost and thus is considered to be of high value.

The concept of value is further discussed in Section 11.5, Design for Assembly. In that section, features are added to ease assembly. Even though these features cut assembly time and thus cost, they often raise the manufacturing cost. Whether to use these features is best judged by considering their value.

11.4 DFM—DESIGN FOR MANUFACTURE

The term *Design For Manufacture*, or DFM, is widely used but poorly defined. Manufacturing engineers often use this term to include all or some of the best practices discussed in this book. Others limit the definition to include only design changes that facilitate manufacturing but do not alter the concept and functionality of the product. Here we will define DFM as *establishing the shape of components to allow for efficient, high-quality manufacture*. Notice that the subject of the definition is *component*. In fact, DFM could be called DFCM, Design For Component Manufacture, to differentiate it from Design For Assembly, DFA, the assembly of components covered in the next section.

The key concern of DFM is in specifying the best manufacturing process for the component and ensuring that the component form supports the manufacturing process selected. For any component, many manufacturing processes can be used. For each manufacturing process, there are design guidelines that, if followed, result in consistent components and little waste. For example, the best process to manufacture the clip in Fig. 11.8 is injection molding. Thus, the form of the clip will need to follow design guidelines for plastic injection molding if the product is to be free from sink marks, surface finish blemishes, and other problems causing low-quality results.

Matching the component to the manufacturing process includes concern for tooling and fixturing. Components must be held for machining, released from

If you don't have experience with a manufacturing process you want to use, be sure you consult someone who has—before you commit to using it.

molds, and moved between processes. The design of the component can affect all of these manufacturing issues. Further, the design of the tooling and fixturing should be treated concurrently with the development of the component. The design of tooling and fixturing follows the same process as the design of the component: establish requirements, develop concepts, and then the final product.

In the days of over-the-wall product design processes, design engineers would sometimes release drawings to manufacturing for components that were difficult or impossible to make. The concurrent engineering philosophy, with manufacturing engineers as members of the design team, helps avoid these problems. With thousands of manufacturing methods, it is impossible for a designer to have sufficient knowledge to perform DFM without the assistance of manufacturing experts.

There are far too many manufacturing processes to cover in this text. For details on these, see the *Design for Manufacturability Handbook*.

11.5 DFA—DESIGN-FOR-ASSEMBLY EVALUATION

Design For Assembly, DFA, is the best practice used to measure the ease with which a product can be assembled. Where DFM focuses on making the components, DFA is concerned with putting them together. Since virtually all products are assembled out of many components and assembly takes time (that is, costs money), there is a strong incentive to make products as easy to assemble as possible.

Throughout the 1980s, many methods evolved to measure the assembly efficiency of a design. All of these methods require that the design be a fairly refined product before they can be applied. The technique presented in this section is based on these methods. It is organized around 13 design-for-assembly guidelines, which form the basis for a worksheet (Fig. 11.10). Before we discuss these 13 guidelines, we mention a number of important points about DFA.

Design For Assembly is important only if assembly is a significant part of the product cost.



(DFA) Design For Assembly				
Individual Assembly Evaluation for: Irwin pre 2007 Clamp			Organization Name: Example	
OVERALL ASSEMBLY				
1	Overall part count minimized		Very good	6
2	Minimum use of separate fasteners		Outstanding	8
3	Base part with fixturing features (locating surfaces and holes)		Outstanding	8
4	Repositioning required during assembly sequence		> = 2 Positions	4
5	Assembly sequence efficiency		Very good	6
PART RETRIEVAL				
6	Characteristics that complicate handling (tangling, nesting, flexibility) have been avoided		Most parts	6
7	Parts have been designed for a specific feed approach (bulk, strip, magazine)		Few parts	2
PART HANDLING				
8	Parts with end-to-end symmetry		Some parts	4
9	Parts with symmetry about the axis of insertion		Some parts	4
10	Where symmetry is not possible, parts are clearly asymmetric		Most parts	6
PART MATING				
11	Straight-line motions of assembly		Some parts	4
12	Chamfers and features that facilitate insertion and self-alignment		Some parts	4
13	Maximum part accessibility		All parts	8
Note:	Only for comparison of alternate designs of same assembly	TOTAL SCORE 70		
Team member: Fred Smith		Team member: Jason Peterson	Prepared by: Fred Smith	
Team member: Omhi Ubolu		Team member:	Checked by: Prof Chan	Approved by:
<i>The Mechanical Design Process</i>				
Copyright 2008, McGraw-Hill				
Designed by Professor David G. Ullman				
Form # 21.0				

Figure 11.10 Design for assembly worksheet.

Assembling a product means that a person or a machine must (1) *retrieve* components from storage, (2) *handle* the components to orient them relative to each other, and (3) *mate* them. Thus, the ease of assembly is directly proportional to the number of components that must be retrieved, handled, and mated, and the ease with which they can be moved from their storage to their final, assembled position. Each act of retrieving, handling, and mating a component or repositioning an assembly is called an *assembly operation*.

Retrieval usually starts at some type of component feeder; this can range from a simple bin of loose bulk components to an automatic machine that feeds one component at a time in the proper orientation for a robot to handle.

Component handling is a major consideration in the measure of assembly quality. Handling encompasses maneuvering the retrieved component into position so that it is oriented for assembly. For a bolt to be threaded into a tapped hole, it must first be positioned with its axis aligned with the hole's axis and its threaded end pointed toward the hole. A number of motions may be required in handling the component as it is moved from storage and oriented for mating. If component handling is accomplished by a robot or other machine, each motion must be designed or programmed into the device. If component handling is accomplished by a human, the human factors of the required motions must be considered.

Component mating is the act of bringing components together. Mating may be minimal, like setting one component on the flat surface of another, or it may require threading a fastener into a threaded hole. A term often synonymous with *mating* is *insertion*. During assembly some components are inserted in holes, others are placed on surfaces, and yet others are fitted over pins or shafts. In all these cases, the components are said to be inserted in the assembly, even though nothing may really be inserted, in the traditional sense of the word, but only placed on a surface.

DFA measures a product in terms of the efficiency of its overall assembly and the ease with which components can be retrieved, handled, and mated. A product with high assembly efficiency has a few components that are easy to handle and virtually fall together during assembly. Assembly efficiency can be demonstrated by considering the seat frames designed for a recumbent bicycle (a bicycle ridden in a seated position). Figure 11.11 shows an old frame, which had nine separate components requiring 20 separate operations to put together. These included positioning and welding operations. This frame took 30 min to assemble. In contrast, the new frame (Fig. 11.12) was designed with assembly efficiency as a major engineering requirement. The resulting product has only four components, requiring eight operations and about 8 min to assemble. The savings in labor is obvious. Additionally, there are savings in component inventory, component handling, and dealings with component vendors.

Guidelines similar to those on the worksheet of Fig. 11.10 were used in the design of the new seat frame to make it efficient to assemble. The worksheet is designed to give an assembly efficiency score to each product evaluated. The score ranges from 0 to 104. The higher the score, the better the assembly. This score is



Figure 11.11 Old seat frame.



Figure 11.12 Redesigned seat frame.

A single part costs nothing to assemble.

—M. M. Andreassen

used as a relative measure to compare alternative designs of the same product or similar products; the actual value of the score has no meaning. The design can be patched or changed on the basis of suggestions given in the guidelines and then reevaluated. The difference between the score of the original product and that of the redesign gives an indication of the improvement of assembly efficiency.

Although this technique is only applied late in the design process, when the product is so refined that the individual components and the methods of fastening are determined, its value can be appreciated much earlier in the design process. This is true because, after filling out the worksheet a few times, the designer develops the sense of what makes a product easy to assemble—knowledge that will have an effect on all future products.

Using ease of assembly as an indication of design quality makes sense only for mass-produced products, since the design-for-assembly guidelines encourage a few complex components. These types of components usually require expensive tooling, which can only be justified if spread over a large manufacturing volume.

Finally, the relationship between the cost of assembly and the overall cost of the product must be kept in mind when considering how much to modify a design according to these suggestions. In low-volume electromechanical products, the cost of assembly is only 1 to 5% of the total manufacturing cost. Thus, there is little payback for changing a design for easier assembly; the change will require extra design effort and may raise the cost of manufacturing, with little financial return.

Measures for each of the 13 design-for-assembly guidelines will be discussed in Sections 11.5.1 to 11.5.4; Section 11.5.1 gives guidelines, all concerned with the overall assembly efficiency; Sections 11.5.2 to 11.5.4 give design-for-assembly guidelines oriented toward the retrieval, handling, and mating of the individual components.

11.5.1 Evaluation of the Overall Assembly

Guideline 1: Overall Component Count Should Be Minimized. The first measure of assembly efficiency is based on the number of components or sub-assemblies used in the product. The part count is evaluated by estimating the minimum number of components possible and comparing the design being evaluated to this minimum. The measure for this guideline is estimated in this way:

a. Find the Theoretical Minimum Number of Components. Examine each pair of adjacent components in the design to see if they really should be separate components. Include fastening components such as bolts, nuts, and clips in this accounting. Assuming no production or material limitations: (1) Components must be separate if the design is to operate mechanically. For example, components that must slide or rotate relatively to each other must be separate components. However, if the relative motion is small, then elasticity can be built into the design to meet the need. This is readily accomplished in plastic components by using elastic hinges, thin sections of fatigue-resistant material that act as a one

degree-of-freedom joint. (2) Components must be separate if they must be made of different materials, for example, when one component is an electric or thermal insulator and another, adjacent component is a conductor. (3) Components must be separate if assembly or disassembly is impossible. (Note that the last word is “impossible,” not “inconvenient.”)

Thus, each pair of adjacent components is examined to find if they absolutely need to be separate components. If they do not, then theoretically they can be combined into one component. After reviewing the entire product this way, we develop the theoretical minimum number of components. The seat frame has a minimum of one component. The actual number of components in the redesigned frame (Fig. 11.12) is four.

b. Find the Improvement Potential. To rate any product, we can calculate its improvement potential:

$$\text{Improvement potential} = \frac{\left(\begin{array}{c} \text{Actual number of} \\ \text{components} \end{array} \right) - \left(\begin{array}{c} \text{Theoretical minimum} \\ \text{number of components} \end{array} \right)}{\text{Actual number of components}}$$

c. Rate the Product on the Worksheet (Fig. 11.10).

- If the improvement potential is less than 10%, the current design is *outstanding*.
- If the improvement potential is 11 to 20%, the current design is *very good*.
- If the improvement potential is 20 to 40%, the current design is *good*.
- If the improvement potential is 40 to 60%, the current design is *fair*.
- If the improvement potential is greater than 60%, the current design is *poor*.

The improvement potential of the seat frame in Figure 11.12 is $(4 - 1) / 4 = 75\%$. In this case, design is poor, but the volume is too low to use a method to further reduce the number of components.

As a product is redesigned, keep track of the actual improvement:

Actual improvement

$$= \frac{\left(\begin{array}{c} \text{Number of components} \\ \text{in initial design} \end{array} \right) - \left(\begin{array}{c} \text{Number of components} \\ \text{in redesign} \end{array} \right)}{\text{Number of components in initial design}}$$

Typical improvement in the number of components in the range of 30 to 60% is realized by redesigning the product in order to reduce the component count.

To put this guideline in perspective, compare it with earlier phases of the design process. In the design philosophy of this text, the functionality of the product is broken down as finely as possible as a basis for the development of concepts (Chap. 7). We then used a morphology for developing ideas for each function. This can lead to poor designs, as can the effort to minimize the number of components. Consider the design of the common nail clipper (Fig. 11.13). If the assumption is made that all the functions are independent and that concepts

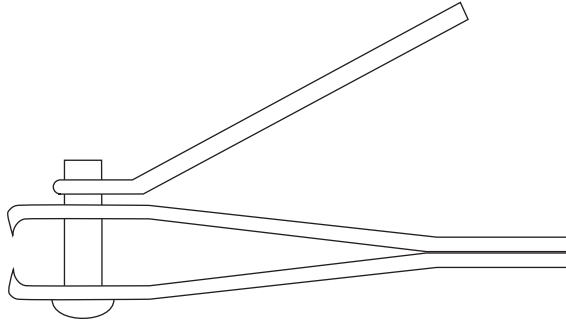


Figure 11.13 Common nail clipper.

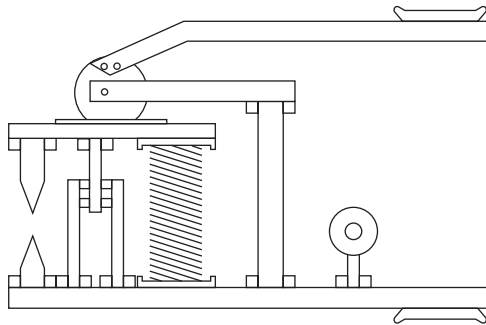


Figure 11.14 Nail clipper with one interface for each function. (Source: Design developed by Karl T. Ulrich, Sloan School of Management, Massachusetts Institute of Technology.)

are generated for each function, then the result, as seen in Fig. 11.14, is a disaster. Note that each function is mapped to one or more interface. At the other extreme, the DFA philosophy leads to the product shown in Fig. 11.15.

Here, in evaluating the product for assembly, this guideline encourages lumping as many functions as possible into each component. This design philosophy, however, also has its problems. The cost of tooling (molds or dies) for the shapes that result from a minimized component count can be high—and that cost is not taken into account here. Additionally, tolerances on complex components may be more critical, and manufacturing variations might affect many functions that are now coupled.

Guideline 2: Make Minimum Use of Separate Fasteners. One way to reduce the component count is to minimize the use of separate fasteners. This is advisable

Every fastener adds costs and reduces strength.

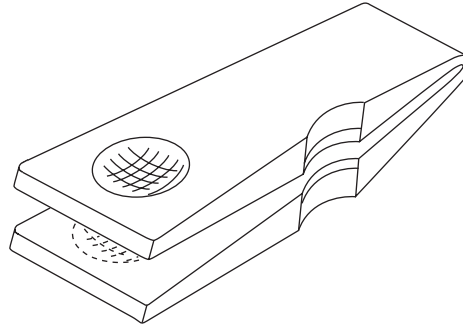


Figure 11.15 A one-piece nail clipper.

for many reasons. First, each fastener used is one more component to handle, and there may be many more than one in the case of a bolt with its accompanying nut, flat washer, and lock washer. Each instance of component handling takes time, typically 10 sec per fastener. Second, the total cost for fasteners is the cost of the components themselves as well as the cost of purchasing, inventorying, accounting for, and quality-controlling them. Third, fasteners are stress concentrators; they are points of potential structural failure in the design. For all these reasons, it is best to eliminate as many fasteners as possible from the design. This is more easily done on high-volume products, for which components can be designed to snap together, than on low-volume products or products utilizing many stock components.

An additional point that should be considered in evaluating a design is how well the use of fasteners has been standardized. A good example of part standardization is the fact that almost everything on the Volkswagen Beetle, a car popular in the 1970s, can be fixed with a set of screwdrivers and a 13-mm wrench.

Finally, if the components fastened together must be taken apart for maintenance, use captured fasteners (fasteners that remain loosely attached to a component even when unfastened). Many varieties of captured fasteners are available, all designed so that they will not be misplaced during assembly or maintenance.

There are no general rules for the quality of a design in terms of the number of separate fasteners. Since the worksheet is just a relative comparison between two designs, an absolute evaluation is not necessary. Obviously, an outstanding design will have few separate fasteners, and those it does have will be standardized and possibly captured. Poor designs, on the other hand, require many different fasteners to assemble. If more than one-third of the components in a product are fasteners, the assembly logic should be questioned.

Figures 11.16 and 11.17 show some ideas for reducing the number of fasteners. In designing with injection-molded plastics, the best way to get rid of fasteners is through the use of snap fits. A typical cantilever snap is shown in Fig. 11.16a. Important considerations when designing snaps are the loads during insertion and when seated. During insertion, the snap acts like a cantilever beam flexed by the amount of the insertion displacement. The major stress during insertion is therefore bending at the root of the beam. Thus, it is important to have

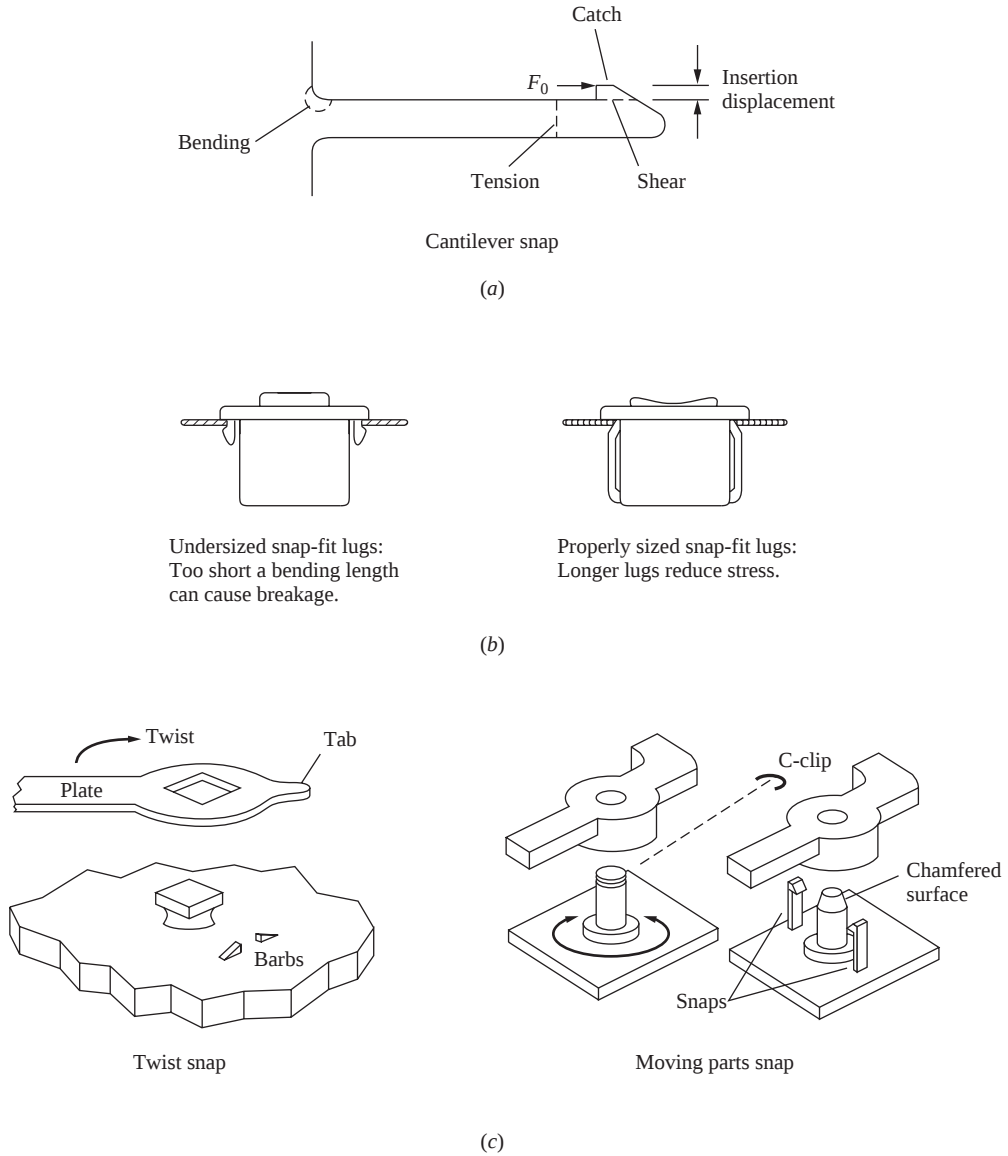


Figure 11.16 Snap-fastener design.

low stress concentrations at that point and to be sure that the snap can flex enough without approaching the elastic limit of the material (Fig. 11.16b). When seated, the snap's main load is the force F_0 , the force holding the components together. It can cause crushing on the face of the catch, shear failure of the catch, and tensile failure of the snap body. (Think of the force flow here.)

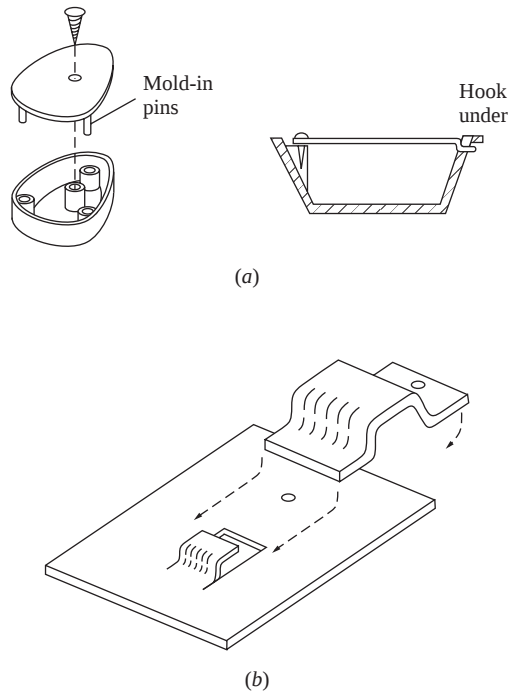


Figure 11.17 Single fastener examples.

Additionally, design consideration must be given to unsnapping. If the device is ever to come apart for maintenance, then consider features that allow a tool or a finger to flex the snap while $F_0 = 0$. Additional snap configurations are shown in Fig. 11.16c. Note that each has one feature that flexes during insertion and another that takes the seated load.

Another way to reduce the number of fasteners is to use only one fastener and either pins, hooks, or other interference to help connect the components. The examples in Fig. 11.17 show both plastic and sheet-metal applications of this idea.

Guideline 3: Design the Product with a Base Component for Locating Other Components. This guideline encourages the use of a single base on which all the other components are assembled. The base in Fig. 11.18 provides a foundation for consistent component location, fixturing, transport, orientation, and strength. The ideal design would be built like a layer cake, with each component or subassembly stacking on top of another one. Without this base to build on, assembly may consist of work on many subassemblies, each with its own fixturing and transport needs and final assembly requiring extensive repositioning and fixturing. The use of a single base component has shortened the length of some assembly lines by a factor of 2.

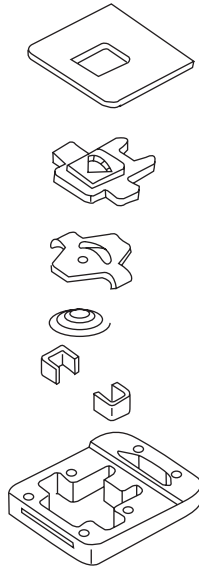


Figure 11.18 Meter assembly.

As with most of these measures, there are no absolute standards for determining an outstanding product and a poor one. Keep in mind that the rating on the worksheet is relative.

Guideline 4: Do Not Require the Base to Be Repositioned During Assembly.

If automatic assembly equipment such as robots or specially designed component placement machines are used during assembly, it is important that the base be positioned precisely. On larger products, repositioning may be time-consuming and costly. An outstanding design would require no repositioning of the base. A product requiring more than two repositionings is considered poor.

Guideline 5: Make the Assembly Sequence Efficient. If there are N components to be assembled, there are potentially $N!$ (N factorial) different possible sequences to assemble them. In reality, some components must be assembled prior to others; thus the number of possible assembly sequences is usually much less than $N!$. An efficient assembly sequence is one that

- Affords assembly with the fewest steps.
- Avoids risk of damaging components.
- Avoids awkward, unstable, or conditionally unstable positions for the product and the assembly personnel and machinery during assembly.
- Avoids creating many disconnected subassemblies to be joined later.

Since even a minor design change can alter the available choices in assembly sequence, it is important to consider the efficiency of the sequence during design. The technique described here will be demonstrated through a simple example, the assembly of a ballpoint pen (Fig. 11.19).

Step 1: List All the Components and Processes Involved in the Assembly Process. Begin with a layout or assembly drawing of the product and a bill of materials. All components for the pen assembly are listed in Fig. 11.19. In some products, the components to be assembled include subassemblies and processes—for example, the component called “ink” in the ballpoint pen includes the process of actually putting the ink in the tube. Additionally, some products require testing during the assembly process. These tests should also be included as components. Finally, fasteners should be lumped with the component they hold in place.

Step 2: List the Connections Between Components and Generate a Connections Diagram. The connection diagram for the ballpoint pen is shown in Fig. 11.20.

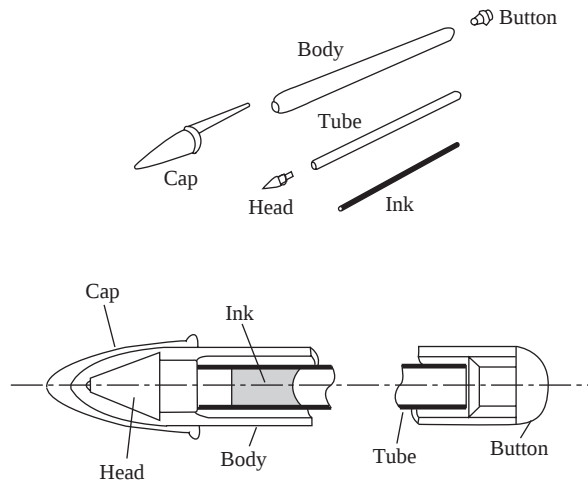


Figure 11.19 Ballpoint pen assembly.

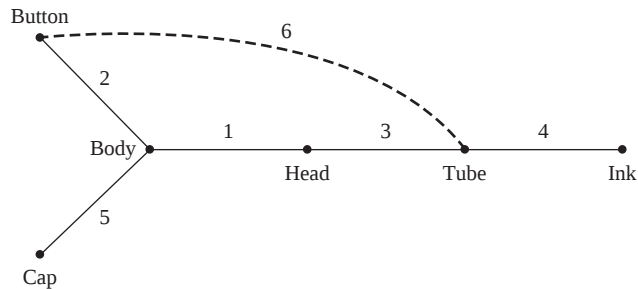


Figure 11.20 Connection diagram for a ballpoint pen.

In this diagram, the nodes represent the components and the links represent the connections. Connection diagrams can have loops. For example, the pen may have the button supporting the end of the tube, creating interface 6, a link between the tube and the button (shown as a dashed line in Fig. 11.20 and assumed not to exist throughout the remainder of this example).

Step 3: Select a Base Component. The base component should be at one end of the connection diagram or be a large component. It should be the component that requires the least subassembly and allows assembly from the fewest directions. For the ballpoint pen, the options are the cap, the button, or the body. The cap requires subassembly of the head in the tube and is thus a poor candidate. The body requires assembly from two directions. The button may be the best base part, but it is hard to hold. Both the body and the button need to be further investigated.

Step 4: Recursively Add the Next Component. Add components to the base using the connection diagram as a guide. It is important to be aware of precedences; for example, the tube must be on the head before the ink is installed. It is useful to list all precedences before starting this step. For the ballpoint pen, the precedences are

Connection 3 must precede connection 4.

Connection 1 must precede connection 5.

Step 5: Identify Subassemblies. Subassemblies can be made of components that have a secure connection with each other, can be reoriented without falling apart, and have a simple connection with the other assembled components. Subassemblies should only be used if they simplify the process. For the pen, the head, tube, and ink form a subassembly that simplifies assembly.

There are many potential assembly sequences for the ballpoint pen. One that is developed using the described procedure is

[2, [3, 4], 1, 5]

or

[button, body, [head, tube, ink], cap]

The first sequence lists the connections, and the second the components, in the order of assembly. The brackets denote subassemblies.

The process given here is very useful in evaluating the assembly sequence and determining the effects of design changes on the sequence. It also measures the efficiency of the assembly sequence. If all connections are made in a logical order, no subassemblies are generated, and no awkward connections made, then the efficiency is rated high; if the connection sequence cannot be accomplished, subassemblies are made, or awkward connections are needed, then the efficiency is low.

11.5.2 Evaluation of Component Retrieval

The measures associated with each guideline for retrieving components range from “all components” to “no components.” If all components achieve the guideline, the quality of the design is high as far as component retrieval is concerned. Those components that do not achieve the guidelines should be reconsidered.

Guideline 6: Avoid Component Characteristics That Complicate Retrieval.

Three component characteristics make retrieval difficult: tangling, nesting, and flexibility. If components of the type shown in Fig. 11.21 column *a* are stored in a box or tray, they will be nearly impossible to pick up individually because they will become tangled. If the components are designed as shown in Fig. 11.21 column *b*, then they cannot tangle.

A second common problem that complicates retrieval is nesting, in which components jam inside each other (Fig. 11.22). There are two simple solutions for this problem: Either change the angle of the interlocking surfaces or add features that prevent jamming.

Finally, flexible components such as gaskets, tubing, and wiring harnesses are exceptionally hard components to retrieve and handle. When possible, make components as few, as short, and as stiff as possible.

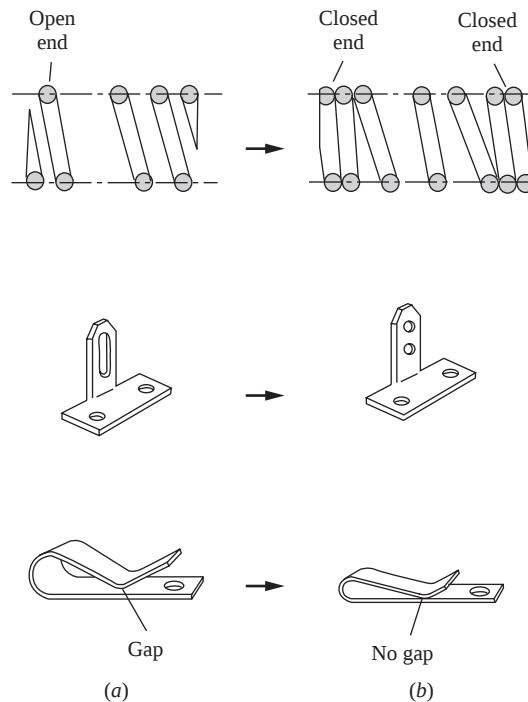


Figure 11.21 Design modifications to avoid component tangling.

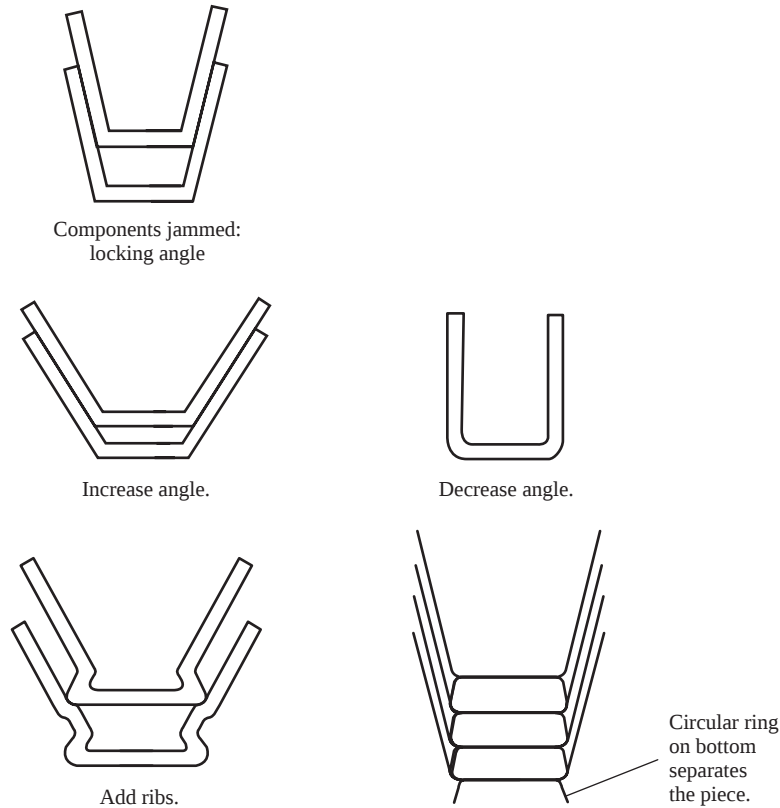


Figure 11.22 Design modifications to avoid jamming.

Guideline 7: Design Components for a Specific Type of Retrieval, Handling, and Mating. Consider the assembly method of each component during design. There are three types of assembly systems: manual assembly, robot assembly, and special-purpose transfer machine assembly. In general, if the volume of the product is less than 250,000 annually, the most economic method of assembly is manual. For products that have a volume of up to 2 million annually, robots are generally best. Special-purpose machines are warranted only if the volume exceeds 2 million. Each of these systems has requirements for component retrieval, handling, and mating. For example, components for manual assembly can be bulk-fed and must have features that make them easy to grasp. Robot grippers, on the other hand, may be fed automatically and can grasp a component externally, like a human; internally, with a suction cup on a flat surface; or with many other end effectors.

11.5.3 Evaluation of Component Handling

The next three design-for-assembly guidelines are all oriented toward the handling of individual components.

Guideline 8: Design All Components for End-to-End Symmetry. If a component can be installed in the assembly only in one way, then it must be oriented and inserted in just that way. The act of orienting and inserting the component takes time and either worker dexterity or assembly machine complexity. If assembly is to be done by a robot, for example, then having only one orientation for insertion may require the robot to be multiaxial. Conversely, if the component is spherical, then its orientation is of no consequence and handling is much easier. Most components in an assembly fall between these two extremes.

There are two measures of symmetry: end-to-end symmetry (symmetry about an axis perpendicular to the axis of insertion) and axis-of-insertion symmetry. (The latter is the focus of guideline 9 and is not discussed here.) End-to-end symmetry means that a component can be inserted in the assembly either end first. Axisymmetric components that are intended to be inserted along their axes are shown in Fig. 11.23. Those in the left-hand column are designed to work in the design only if installed in one way. These same components are shown in the right-hand column modified so that they can be inserted either end first. In each

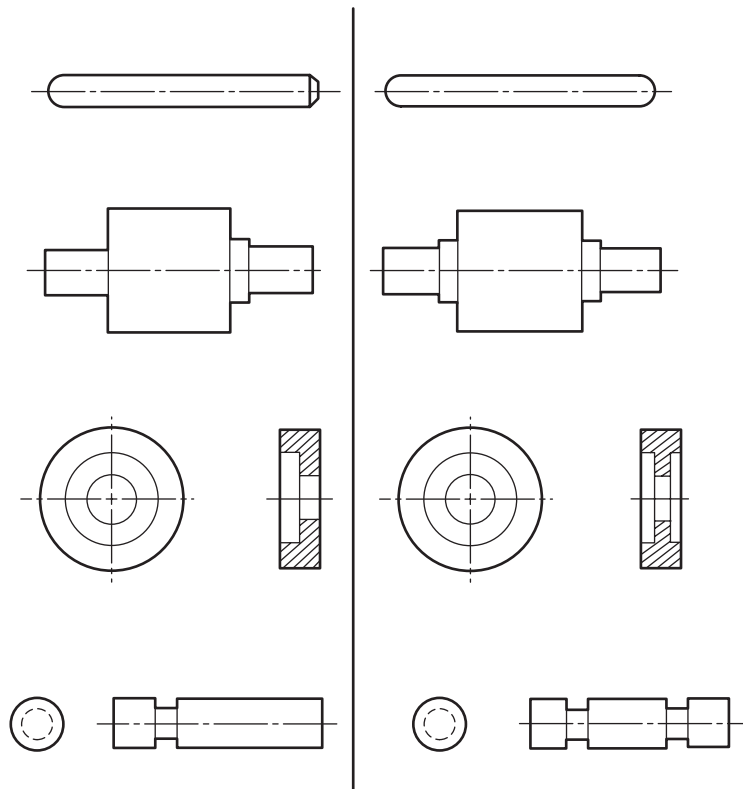


Figure 11.23 Modification of axisymmetric parts for end-to-end symmetry.

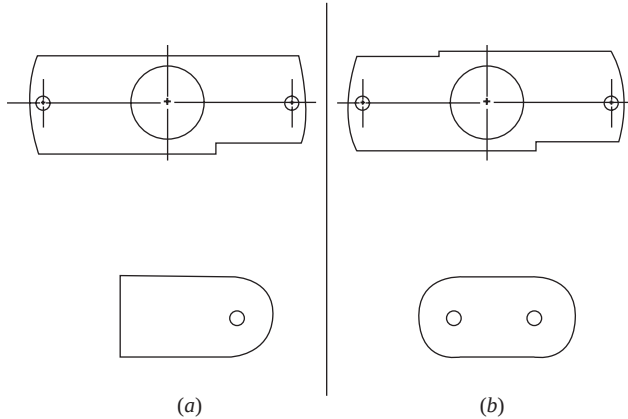


Figure 11.24 Modification of features for symmetry about the axis of insertion.

case, the asymmetrical feature has been replicated to make the component end-to-end symmetrical for ease of assembly.

Before modifying a component to meet this or similar guidelines, it is important to check the value of the modification. The cost of adding a feature may not improve its functionality for the assembler sufficiently to warrant the modification.

Guideline 9: Design All Components for Symmetry About Their Axes of Insertion. Whereas the previous guideline called for end-to-end symmetry, a designer should also strive for rotational symmetry. The components in Fig. 11.23 are all axisymmetric if inserted in the direction of their centerline. In Fig. 11.24 the components in column *a* have only one orientation if they are inserted in the plane of the diagram. However, by adding a functionally useless notch (on the top component) or adding a hole and rounding an end (on the bottom component), we can give the components two orientations for insertion—a decided improvement.

In Fig. 11.25*a*, the original design for the component fits only one way into the assembly. The addition of an opposing finger (Fig. 11.25*b*), which is useless functionally, gives the component two possible insertion orientations. Finally, modifying the component functions (Fig. 11.25*c*) can make the component axisymmetric. It is important to ask if the change in functionality is worth the gained ease of assembly. If not, then the asymmetry should be tolerated.

Guideline 10: Design Components That Are Not Symmetric About Their Axes of Insertion to Be Clearly Asymmetric. The component in Fig. 11.25*a* is clearly asymmetric. If it were not asymmetric, the component could be inserted with the finger pointing the wrong way and, as a result, would not function as it was designed to. In Fig. 11.26 the four component designs of the left-hand column have been modified in the right-hand column to afford easy orientation. The goal of this guideline is to make components that can be inserted only in the way intended.

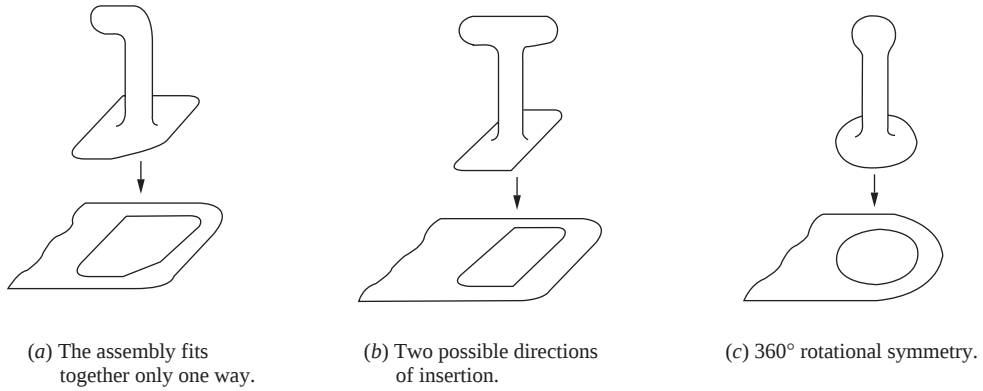


Figure 11.25 Modification of a part for symmetry.

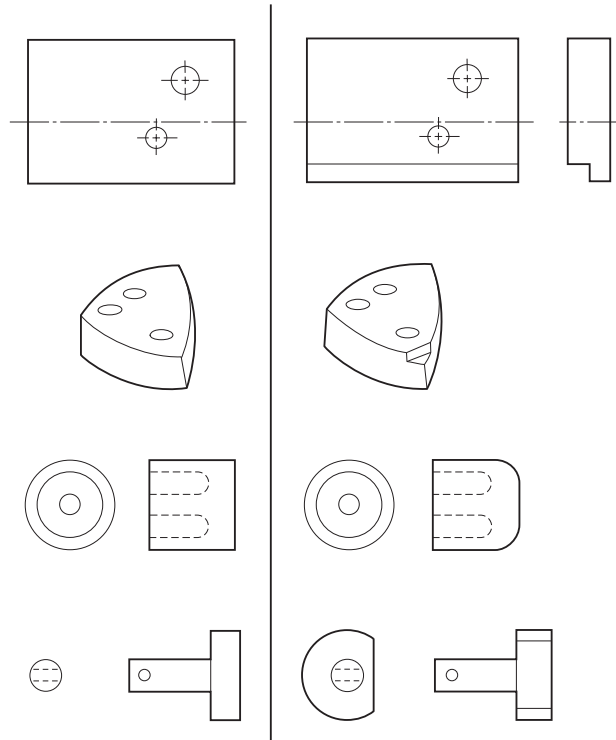


Figure 11.26 Modification of parts to force asymmetry.

11.5.4 Evaluation of Component Mating

Finally, the quality of component mating should be evaluated. Guidelines 11 to 13 offer some design aids for improving assemblability.

Guideline 11: Design Components to Mate Through Straight-Line Assembly, All from the Same Direction. This guideline, intended to minimize the motions of assembly, has two aspects: the components should mate through straight-line motion, and this motion should always be in the same direction. If both of these corollaries are met, the assembly will then fall together from above. Thus, the assembly process will never require reorientation of the base nor any other assembly motion other than straight down. (Down is the preferred single direction, because gravity aids the assembly process.)

The components in Fig. 11.27*a* require three motions for assembly. This number has been reduced in Fig. 11.27*b* by redesigning the interface between the components. Note that the design in Fig. 11.17*b*, although improving the quality in terms of fastener use, has degraded the design in terms of insertion difficulty, again demonstrating that there are always trade-offs to be considered in design.

Guideline 12: Make Use of Chamfers, Leads, and Compliance to Facilitate Insertion and Alignment. To make the actual insertion or mating of a component as easy as possible, each component should guide itself into place. This can be accomplished using three techniques. One common method is to use chamfers, or rounded corners, as shown in Fig. 11.28. Here the four components shown in column *a* are all modified with chamfers in column *b* to ease assembly.

In Fig. 11.29*a* the shaft has chamfers and still the disk is hard to align and press into its final position. This difficulty is alleviated by making part of the shaft a smaller diameter, allowing the disk to mate with the final diameter, as shown in column *b* of the figure. The lead section of the shaft has forced the disk into alignment with the final section. A similar redesign is shown in the lower component, where, in column *b*, by the time the shaft is inserted in the bearing from the right it is aligned properly.

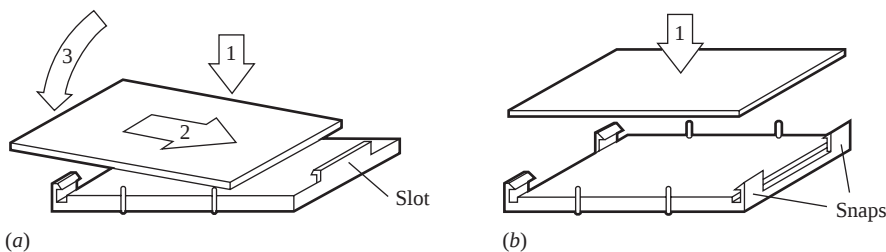


Figure 11.27 Example of one-direction assembly.

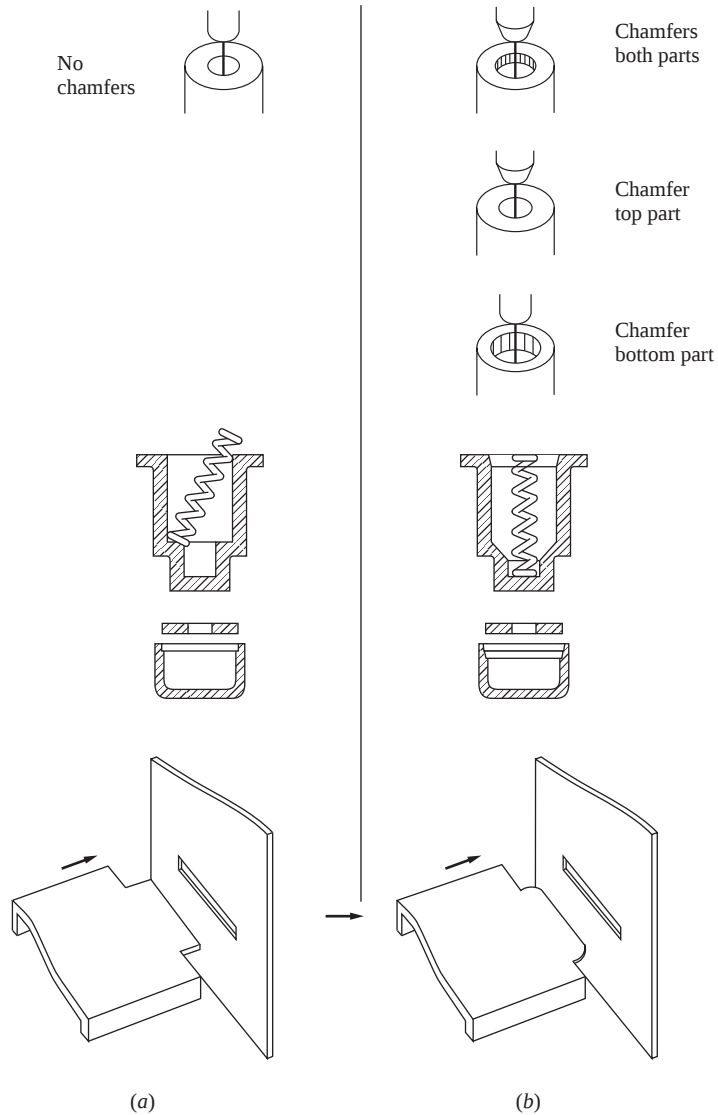


Figure 11.28 Use of chamfers to ease assembly.

Finally, component compliance, or elasticity, is used to ease insertion and also relax tolerances. The component mating scheme in column *b* of Fig. 11.30 need not have high tolerance; even if the post is larger than the hole, the components will snap together.

Guideline 13: Maximize Component Accessibility. Whereas guideline 5 concerned itself with assembly sequence efficiency, this guideline is oriented toward

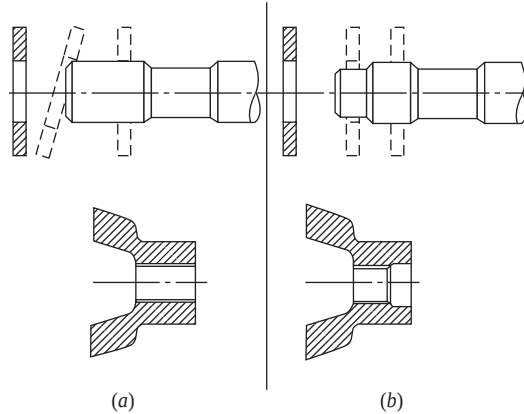


Figure 11.29 Use of leads to ease assembly.

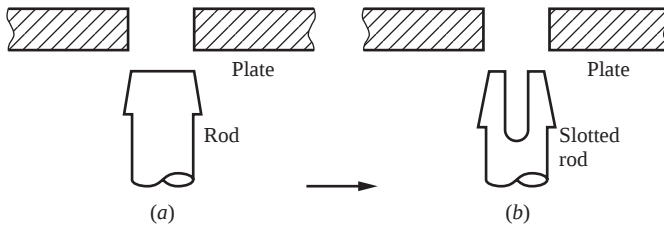


Figure 11.30 Use of compliance to ease assembly.

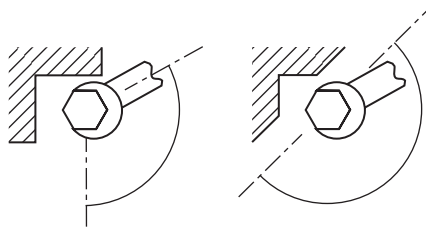


Figure 11.31 Modifications for tool clearance.

sufficient accessibility. Assembly can be difficult if components have no clearance for grasping. Assembly efficiency is also low if a component must be inserted in an awkward spot.

Besides concerns for assembly, there is also maintenance to consider. To replace the fuses in one common computer printer, it is necessary to disassemble the entire machine. In both assembly and maintenance, tools are necessary and room must be allowed for the tools to mate with the components and to be manipulated. As shown in Fig. 11.31, sometimes simple design changes can make tool engagement and motion much easier.

11.6 DFR—DESIGN FOR RELIABILITY

Reliability is a measure of how the quality of a product is maintained over time. Quality here is usually in terms of satisfactory performance under a stated set of operating conditions. Unsatisfactory performance is considered a *failure*, and so in calculating the reliability of a product we use a technique for identifying failure potential called *Failure Modes and Effects Analysis, FMEA*. This best practice is useful as a design evaluation tool and as an aid in hazard assessment, described in Section 8.6.1 (A failure can, but does not necessarily, present a hazard; it presents a hazard only if the consequence of its occurrence is sufficiently severe.) Traditionally, a *mechanical failure* is defined as any change in the size, shape, or material properties of a component, assembly, or system that renders the product incapable of performing its intended function. A failure may be the result of change in the hardware due to aging (for example, wear, material property degradation, or creep) or environmental conditions (for example, overloading, temperature effects, and corrosion). If deterioration or aging noises are taken into account, then the potential for mechanical failure is minimized (see Section 10.7).

To use failure potential as a design aid, it is important to extend the definition of failure to include not only undesirable changes after the product is in service, but also design and manufacturing errors (for example, moving parts interfere, parts do not fit together, or systems do not meet engineering requirements).

Thus, a more general definition is *a mechanical failure is any change or any design or manufacturing error that renders a component, assembly, or system incapable of performing its intended function*. Based on this definition, a failure has two attributes: the function affected and the source of the failure (i.e., the operational change or design or manufacturing error that produced the failure). Typical sources of failure or failure modes are wear, fatigue, yielding, jamming, bonding weakness, property change, buckling, and imbalance.

11.6.1 Failure Modes and Effects Analysis

The Failure Modes and Effects Analysis, FMEA, technique presented here can be used throughout the product development process and refined as the product is refined. The method aids in identifying where redundancy may be needed and in diagnosing failures after they have occurred. FMEA follows these five steps, and can be developed in a simple table, as shown in Figure 11.32:

Step 1: Identify the Function Affected. For each function identified in the evolution of the product, ask, “What if this function fails to occur?” If functional development has paralleled form development, this step is easy; the functions are already identified. However, if detailed functional information is not available, this step can be accomplished by listing all the functions of each component or assembly. For products being redesigned, the functions of a component or assembly are found by examining the connections or component interfaces and identifying the flow of energy, information, or materials through them. Additional considerations come from extending the basic question to read, “What if this



FMEA (Failure Modes and Effects Analysis)									
Product: Mars Rover			Organization Name: Jet Propulsion Lab						
#	Function Affected	Potential Failure Modes	Potential Failure Effects	Potential Causes of Failure	Recommend Actions	Responsible Person	Taken Actions		
1	Propel Rover	No torque to wheel	Wheel stops turning	Motor failure	Ensure motors have high reliability—at least 99.9% reliability for 100 hr	Tim Smithson, Electronics Div.	Vendor required to submit failure test results		
2			Wheel stops turning	Motor failure	Test ability to propel Rover with 1 or 2 drive wheels inoperative	Barb Rojo	Prototype tested with 2 motors off line		
3		Wheel jams against rock	Wheel stops turning	Inability to sense rocks	Develop ability to sense and avoid rocks or feedback torque increase	B. J. Smith	Work in progress		
4			Wheel damages surface	Wheel surface too soft	Specify surface that can stand abrasion	N. Knovo	Hard test developed		
Team member: B. Rojo			Team member:						
Team member: B. J. Smith			Team member:		Prepared by: N. Knovo		Approved by:		
The Mechanical Design Process									
Designed by Professor David G. Ullman									
Form # 22.0									
Copyright 2008, McGraw-Hill									

Figure 11.32 FMEA example for MER.

function fails to occur at the right time?” “What if this function fails to occur in the right sequence?” or “What if this function fails to occur completely?”

Step 2: Identify Failure Modes. For each function, there can be many different failures. The failure mode is a description of the way a failure occurs. It is what is observed, what can be detected when the function fails to occur.

Step 3: Identify the Effect of Failure. What are the consequences on other parts of the system of each failure identified in step 1? In other words, if this failure occurs, what else might happen? These effects may be hard to identify in systems in which the functions are not independent. Many catastrophes result when one system's benign failure overloads another system in an unexpected manner, creating an extreme hazard. If functions have been kept independent, the consequences of each failure should be traceable.

Step 4: Identify the Failure Causes or Errors. List the changes or the design or manufacturing errors that can cause the failure. Organize them into three groups: design errors (D), manufacturing errors (M), and operational changes (O).

Step 5: Identify the Corrective Action. Corrective action requires three parts, what action is recommended, who is responsible, and what was actually done. For each design error listed in step 3, note what redesign action should be taken to ensure that the error does not occur. The same is true for each potential manufacturing error. For each operational change, use the information generated to establish a clear way for the failure mode to be detected. This is important, as it is the basis for the diagnosis of problems when they do occur. For operational changes it may also be important to redesign the device so that the failure mode has a reduced effect on the function. This may include the addition of other devices (for example, fuses or filters) to protect the function under consideration; however, the failure potential of these added devices should also be considered. The use of redundant systems is another way to protect against failures. But redundancy might add other failure modes as well as increase costs.

FMEA is best used as a bottom-up tool. This means focusing on a detailed function and dissecting all its potential failure. Fault Tree Analysis (FTA), Section 11.6.2, is better suited for “top-down” analysis. When used as a “bottom-up” tool, FMEA can augment or complement FTA and identify many more causes and failure modes resulting in top-level symptoms. It is not able to discover complex failure modes involving multiple failures within a subsystem, or to report expected failure intervals of particular failure modes up to the upper level subsystem or system.

An example of an FMEA and its tie to FTA is based on the design of the propulsion system for the Mars Exploration Rover, MER. During its development, the Jet Propulsion Laboratory team made extensive use of FMEA and FTA. The examples in this and the following section are loosely based on their work.

The FMEA analysis in Fig. 11.32 is based on a simple template. The function considered is “propel Rover.” The total analysis for the system may have many

hundreds of failure modes. Only a small part of the analysis is shown in this example. The failure modes identified had to do with one of the six wheels failing to propel the Rover. As can be seen, a failure mode can have multiple effects, causes, or recommended actions.

11.6.2 FTA—Fault Tree Analysis

Fault Tree Analysis (FTA) can help in finding failure modes. FTA evolved in the 1960s during the development of the Minuteman Missile System and has gained in use ever since. The goal of this method is to graphically develop a tree of all the faults that could happen to cause a system failure, and the logical relationships among these faults. Further, there are analytical methods to compute probabilities of faults, but we will only give a basic, usable introduction to the method here.

Fault Trees are built from symbols that signify events and logic. The most basic of these are listed in Table 11.2 and used in an example Fault Tree for the MER (Fig 11.33). This Fault Tree is a partial analysis for the event “Loss of Rover Mobility.” The full Fault Tree had hundreds of events identified. Fault Trees are built from the top down, beginning with an undesired event (loss of Rover mobility) taken as the root (“top event”). The steps for building a Fault Tree are

Step 1: Identify the top event. There should be only one top event.

Step 2: Identify the events (i.e., faults) that can possibly occur to cause the top event. Ask the question “What can go wrong?” repeatedly until all the events that

Table 11.2 Basic Fault Tree symbols

Event block	FTA symbol	Description
Event		An event, something that happens to something and causes a function to fail.
Basic Event		A basic initiating fault or a failure event.
Undeveloped Event		An event that is not further developed.
Logical operation	FTA symbol	Description
AND		The output event occurs if all input events occur.
OR		The output event occurs if at least one of the input events occurs.

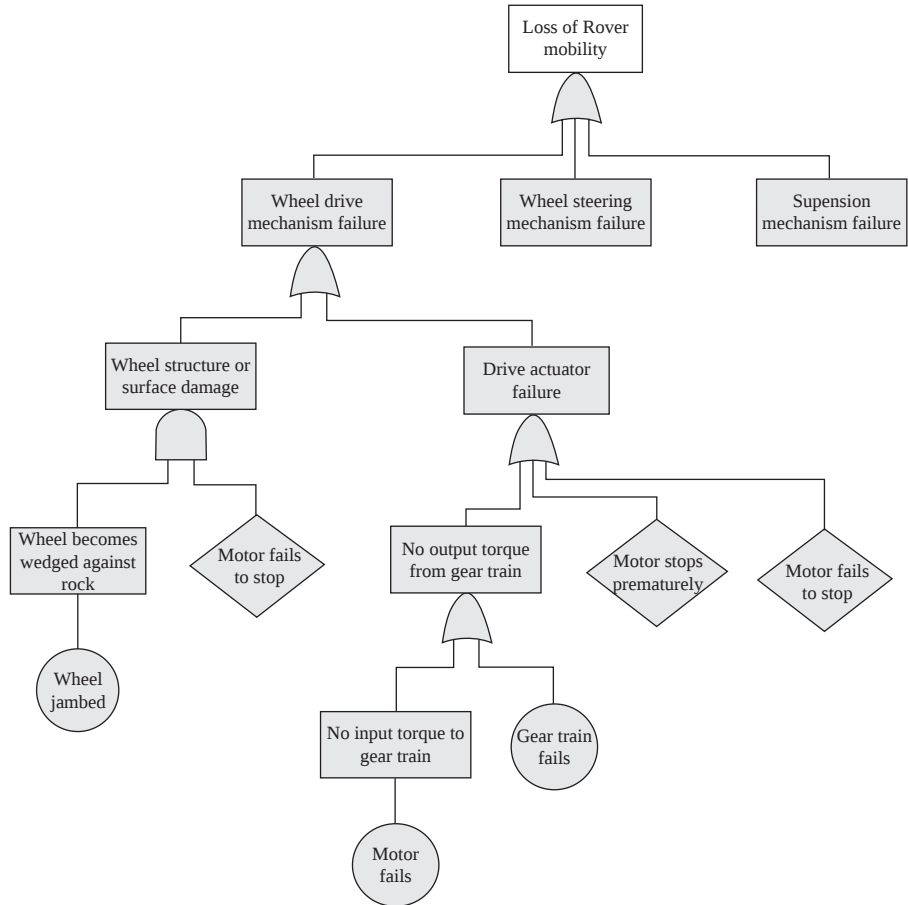


Figure 11.33 Partial Fault Tree for MER Mobility

can occur to cause failure have been identified. Look for hardware and software failures, and human errors. In the MER example, the loss of mobility can be caused by a drive mechanism failure, a steering failure, or a suspension failure.

Step 3: Determine the logical relationship among the events identified in step 2. The goal here is to determine if these newly identified events can happen independently [or], or need to happen together [and]. For example, the loss of rover mobility can be caused by either the drive mechanism failing, **or** the steering mechanism failing, **or** the suspension failing. Thus, on the Fault Tree an “Or” symbol is used to connect the three events to the top event. Farther down in the tree, in order to cause wheel surface damage, the wheel must both be wedged against a rock **and** the motor has to continue to rotate, damaging the wheel surface.

Step 4: Note which events in the tree will not be further developed. These neither need to be considered in the future, considered by others focused on that specific

event, or does not need refinement. For example, “motor fails to stop” can only be caused by a failure of the control system to turn off power to the motor. A separate Fault Tree was developed for the control system by the MER team.

Step 5: Identify the basic events. Each event at the bottom of the tree should end with a basic or initiating event. A basic event is one that cannot be further broken down. In the example Fault Tree “wheel jammed,” “motor fails,” and “gear train fails” cannot be decomposed any further.

11.6.3 Reliability

Once the different potential failures of the product have been identified, the reliability of the system can be found and expressed in units of reliability called *Mean Time Between Failures* (MTBF), or the average elapsed time between failures. MTBF data are generally accumulated by testing a representative sampling of the product. Often these data are collected by service personnel, who record the part number and type of failure for each component they replace or repair.

These data aid in the design of a new product. For example, a manufacturer of ball bearings collected data for many years. The data showed an MTBF of 77,000 hr for a ball bearing operating under manufacturer-specified conditions. On the average, a ball bearing would last 8.8 years $[77,000/(365 \times 24)]$ under normal operating conditions. Of course, a harsh environment or lack of lubrication would greatly reduce this lifetime. Often the MTBF value is expressed as its inverse and called the *failure rate* L , the number of failures per unit time. Failure rates for common machine components are given in Table 11.3, where the failure rate for the ball bearing is $1/77,000$, or 13 failures per 1 million hours.

Table 11.3 Failure rates of common components

Mechanical failures, per 10^6 hr		Electrical failures, per 10^6 hr	
Bearing		Meter	26
Ball	13	Battery	
Roller	200	Lead acid	0.5
Sleeve	23	Mercury	0.7
Brake	13	Circuit board	0.3
Clutch	2	Connector	0.1
Compressor	65	Generator	
Differential	15	AC	2
Fan	6	DC	40
Heat exchanger	4	Heater	4
Gear	0.2	Lamp	
Pump	12	Incandescent	10
Shock absorber	3	Neon	0.5
Spring	5	Motor	
Valve	14	Fractional hp	8
		Large	4
		Solenoid	1
		Switch	6

The actual reliability of a component is determined from the failure rate information. Assuming that the failure rate is constant over the life of the component—which is generally true for all but the initial (infant mortality) and the final (wearout) periods—the reliability is defined as

$$R(t) = e^{-Lt}$$

where R , the reliability, is the probability that the component has not failed. For the ball bearing,

$$R(t) = e^{-0.000013t}$$

with t in hours. Thus,

t , hr	R
0	1.000
100	0.999
1000	0.987
8760 (1 year)	0.892
10,000	0.878
43,800 (5 years)	0.566

If 1000 ball bearings are tested, it would be expected that 892 of them would still be operating a year later within specifications.

What if there are four ball bearings in a product and the product will fail if any one bearing fails? The total reliability of that device is the product of the reliabilities of all its components (this is often called *series reliability*):

$$R_{\text{product}} = R_{\text{bearing 1}} \cdot R_{\text{bearing 2}} \cdot R_{\text{bearing 3}} \cdot R_{\text{bearing 4}}$$

Because of the exponential nature of the definition of reliability, the failure rate for that device would be

$$L_{\text{product}} = L_{\text{bearing 1}} + L_{\text{bearing 2}} + L_{\text{bearing 3}} + L_{\text{bearing 4}}$$

For the product with four bearings, $L = 4 \cdot 0.000013 = 0.000052$. Thus, after one year, $R = 0.634$; about one-third of the products will have had a bearing failure.

There are essentially two ways to increase reliability. First, decrease the failure rate. This is accomplished by lowering the bearing's load or by decreasing its rotation rate. A second way to increase reliability is through redundancy, often called *parallel reliability*. For redundant systems, the failure rate is

$$L = \frac{1}{1/L_1 + 1/L_2 + \dots}$$

Thus, if a ball bearing and a sleeve bearing are designed into the product so that either can carry the applied load, then

$$L = \frac{1}{1/0.000013 + 1/0.000023} = 8.3 \text{ failures}/10^6 \text{ hr}$$

With this technique, reliability evaluations can also be made on complex systems. A model of the failure modes and the MTBF for each of them is needed to accomplish such an evaluation.

11.7 DFT AND DFM—DESIGN FOR TEST AND MAINTENANCE

Testability is the ease with which the performance of critical functions is measured. For instance, in the design of VLSI chips, circuits are included on the chip that allow critical functions to be measured. Measurements can be made during manufacturing to ensure that no errors are built into the chip. Measurements can also be made later in the life of the chip to diagnose failures.

Adding structure in this way, to make testability easier, is often impossible in mechanical products. However, if the technique developed in the previous sections for identifying failures is extended, at least some measure of the testability of the product can be realized. For instance, step 4 of the FMEA technique (Section 11.6.1) required the listing of errors that can cause each failure. An additional step here would address testability:

Step 4A: Is It Possible to Identify the Parameters That Could Cause the Failure? If there are a significant number of cases in which the parameters cannot be measured, there is a lack of testability in the product.

There are no firm guidelines in developing an acceptable level of testability. The designer should ensure, however, that the critical parameters that affect the critical functions can be tested. In this way, the ability to diagnose manufacturing problems and failures when they occur is increased.

The terms *maintainability*, *serviceability*, and *reparability* are often used interchangeably to describe the ease of diagnosing and repairing a product. Since the 1980s, a dominant philosophy has been to design products that are totally disposable or composed of disposable modules that can be removed and replaced. This is in direct conflict to the Hannover Principles introduced in Chap. 1 and DOE—Design For the Environment, in Section 11.8. These modules often contained still-functioning components along with those that had failed. The structure of the module forced replacement of both good and bad components. This philosophy was characteristic of the “throwaway” attitude of the time, and products designed during this period were often easy to replace and hard to repair. A different philosophy is to design products that are easy to diagnose, disassemble, and repair at any level of function. As discussed, designing diagnosability into a mechanical product is possible, but it takes extra effort and may be of questionable value. This also applies to designing a product that is easy to disassemble and

Make it fail where you want. Design in mechanical fuses.

repair. Since the guidelines given for the design-for-assembly technique do not lead to a product that is easy to disassemble, special care must be taken to ensure that, if desired, the snap fits can be unsnapped and that the disassembly sequence has been considered with as much care as the assembly sequence. Further, the ability to disassemble a product is also important if the product is to be recycled at the end of its useful life. This topic is discussed in Section 11.8.

One important feature of design for maintainability is the concept of a “mechanical fuse.” In electrical systems, fuses are used to fail in order to protect the rest of the circuit. The same should be done in mechanical devices. A good use of a mechanical fuse is in high-powered kitchen tabletop mixers. Larger units, those that can mix bread dough, are powerful enough to break fingers and arms. Thus, if something jams these mixers, they stop working. To fix them, you must take a cover off to see that one of the gears has failed. This gear is made of plastic while all the others are of steel. It is designed to break and it is the only gear in the unit that can be purchased at a local appliance repair store.

11.8 DFE—DESIGN FOR THE ENVIRONMENT

Design for the environment is often called green design, environmentally conscious design, life-cycle design, or design for recyclability. Treating environmental concerns as important requirements in the design process began in the 1970s. It was not until the 1990s that it became an important issue in the design community. The major consideration of design for the environment is seen in Fig. 11.34. Here the arrows represent materials that are taken from the Earth or the biosphere and ultimately returned to it. In this figure, all the major green design issues are considered.

When a product’s useful life is over, one of three things happens to its components. They are either disposed of, reused, or recycled. For many products there is no thought given beyond disposal. However, in 1995, 94% of all cars and trucks scrapped in the United States were dismantled and shredded, and 75% of the content by weight was recycled. Whereas, in the 1970s and 1980s, there was design emphasis on disposable products, more and more industries are now trying to design in the ability to recycle or reuse parts of retired products.

For example, even though the single-use camera appears to be disposable after use, Kodak has recycled 41 million of its cameras, or 75% of those sold. Likewise, Xerox reuses or recycles 97% of parts and assemblies from the toner cartridges it manufactures.

You are responsible for the resources used in your products.

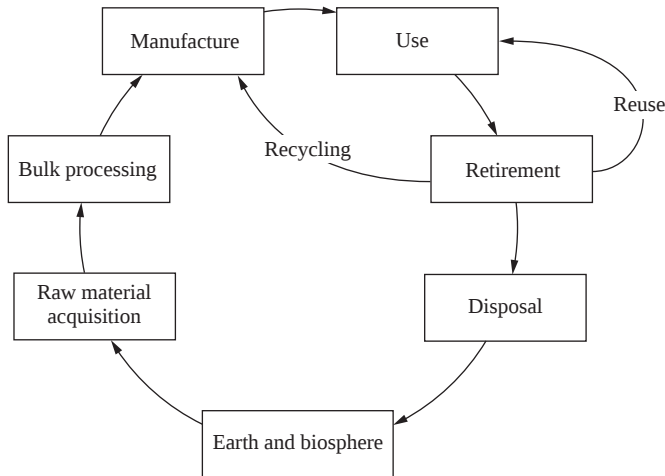


Figure 11.34 Green design life cycle.

This attention to the entire product life cycle is fueled by economics, customer expectation, and government regulation. First, it is becoming less expensive to recycle some materials than it is to pay the expense of processing new raw materials. This is especially true if the product is designed so that it is easily disassembled into components made of a single material. Expense increases if materials are difficult to separate or if one material contaminates another, adversely affecting its material properties. Further, the realization that the resources of raw materials are limited has only recently dawned on many engineers and consumers.

Second, consumers are increasingly more environmentally conscious and aware of the value of recycling. Thus, companies that pollute, generate excessive waste, or produce products that clearly have adverse effects on the environment are looked down on by the public.

Finally, government regulation is forcing attention on the environment. In Germany, manufacturers are responsible for *all* the packaging they create and use. They must collect and recycle it. Further, Mercedes and BMW are designing their new cars so that they, too, can be collected and recycled. European Union laws are forcing this corporate responsibility for the entire life of the product.

In evaluating a product for its “greenness,” the guidelines presented next help ensure that environmental design issues have been addressed. These guidelines are an engineering design refinement of the Hannover Principles introduced in Chap. 1. The guidelines serve to compare two designs as do the Design-For-Assembly, DFA, measures in Section 11.5.

Guideline 1: Be Aware of the Environmental Effects of the Materials Used in Products. In Fig. 11.34, every step requires energy, produces waste products, and may deplete resources. Although it is not realistic for the design engineer to

know the environmental details of every material used in a product, it is important to know about those materials that may have high environmental impact.

Guideline 2: Design the Product with High Separability. The guidelines for design for disassembly are similar to those for design for assembly. Namely, a product is easy to disassemble if fewer components and fasteners are used, if they come apart easily, and if the components are easy to handle. Other aids for high separability are

- Make fasteners accessible and easy to release.
- Avoid laminating dissimilar materials.
- Use adhesives sparingly and make them water soluble if possible.
- Route electrical wiring for easy removal.

One clear measure of separability is the percentage of material that is easily isolated from other materials.

If some of the components are to be reused, the designer must consider disassembly, cleaning, inspection, sorting, upgrading, renewal, and reassembly.

Guideline 3: Design Components That Can Be Reused to Be Recycled. One design goal is to use only recyclable materials. Automobile manufacturers are striving for this goal. In recycling there are five steps: retrieval, separation, identification, reprocessing, and marketing. Of these five, the design engineer can have the most influence on the separation and identification. Separation was just addressed in guideline 2. Identification means to be able to tell after disassembly exactly what material was used in the manufacture of each component. With few exceptions, it is difficult to identify most materials without laboratory testing. Identification is made easier with the use of standard symbols, such as those used on plastics that identify polymer type.

Guideline 4: Be Aware of the Environmental Effects of the Material Not Reused or Recycled. Currently 18% of the solid waste in landfills is plastic and 14% is metal. All of this material is reusable or recyclable. If a product is not designed to be recycled or reused, it should at least be degradable. The designer should be aware of the percentage of degradable material in a product and the time it takes this material to degrade.

11.9 SUMMARY

- Cost estimation is an important part of the product evaluation process.
- Features should be judged on their value—the cost for a function.
- Design for manufacture focuses on the production of components.
- Design for assembly is a method for evaluating the ease of assembly of a product. It is most useful for high-volume products that have molded components. Thirteen guidelines are given for this evaluation technique.

- Functional development gives insight into potential failure modes. The identification of these modes can lead to the design of more reliable and easier-to-maintain products.
- Design for the environment emphasizes concern for energy, pollution, and resource conservation in processing raw materials for products. It also emphasizes concern for recycling, reuse, or disposal of the product after its useful life is over.

11.10 SOURCES

- Boothroyd, G., and P. Dewhurst: *Product Design for Assembly*, Boothroyd and Dewhurst Inc., Wakefield, R.I., 1987. Boothroyd and Dewhurst have popularized the concept of DFA. The range of their tools is much broader than that of those presented here.
- Bralla, J. G.: *Design for Manufacturability Handbook*, 2nd edition, McGraw-Hill, New York, 1998. Over 1300 pages of information about over 100 manufacturing processes written by 60+ domain experts. A good starting place to understand manufacturing.
- Chow, W. W.-L.: *Cost Reduction in Product Design*, Van Nostrand Reinhold, New York, 1978. An excellent book that gives many cost-effective design hints, written before the term *concurrent design* became popular yet still a good text on the subject. The title is misleading; the contents of the book are a gold mine for the designer engineer.
- Lazor, J. D.: "Failure Mode and Effects Analysis (FMEA) and Fault Tree Analysis (FTA)," Chap. 6 in *Handbook of Reliability and Management*, 2nd edition, 1995, <http://books.google.com/books?id=kWa4ahQUPyAC&pg=PT91&lpg=PT91&dq=fault+tree+analysis+fmea&source=web&ots=3WLM58qxy&sig=by3Lbbpi3Uxy8KIMEEnEbsyc9qM&hl=en>
- Life Cycle Design Manual: Environmental Requirements and the Product System*, EPA/600/R-92/226, United States Environmental Protection Agency, Jan. 1992. A good source for design for the environment information.
- Michaels, J. V., and W. P. Wood: *Design to Cost*, Wiley, New York, 1989. A good text on the management of costs during design.
- Nevins, J. L., and D. E. Whitney: *Concurrent Design of Products and Processes*, McGraw-Hill, New York, 1989. This is a good text on concurrent design from the manufacturing viewpoint; a very complete method for evaluating assembly order appears in this text.
- Rivero, A., and E. Kroll: "Derivation of Multiple Assembly Sequences from Exploded Views," *Advances in Design Automation*, ASME DE-Vol. 2, American Society of Mechanical Engineers—Design Engineering, Minneapolis, Minn., 1994, pp. 101–106. More guidance on determining the assembly sequence.
- Trucks, H. E.: *Designing for Economical Production*, 2nd edition, Society of Manufacturing Engineers, Dearborn, Mich., 1987. This is a very concise book on evaluating manufacturing techniques. It gives good cost-sensitivity information.

11.11 EXERCISES

- 11.1 For the product developed in response to the design problem begun in Exercise 4.1, estimate material costs, manufacturing costs, and selling price. How accurate are your estimates?

- 11.2** For the redesign problem begun in Exercise 4.2, estimate the changes in selling price that result from your work.
Exercises 11.3 and 11.4 assume that a cost estimation computer program is available or that a vendor can help with the estimates.
- 11.3** Estimate the manufacturing cost for a simple machined component:
- Compare the costs for manufacturing volumes of 1, 10, 100, 1000, and 10,000 pieces with an intermediate tolerance and surface finish. Explain why there is a great change between 1 and 10 and a small change between 1000 and 10,000 pieces.
 - Compare the costs for fit, intermediate, and rough tolerances with a volume of 100 pieces.
 - Compare the costs of manufacturing the component out of various materials.
- 11.4** Estimate the manufacturing cost for a plastic injection-molded component:
- Compare the costs for manufacturing volumes of 100, 1000, 10,000, and 100,000. The tolerance level is intermediate, and surface finish is not critical.
 - Compare the cost for a change in tolerance.
 - Why does changing the material have virtually no effect on cost at low plastic injection volume (i.e., 100 pieces)?
- 11.5** Perform a design-for-assembly evaluation for one of these devices. Based on the results of your evaluation, propose product changes that will improve the product. Be sure that your proposed changes do not affect the function of the device. For each change proposed, estimate its “value.”
- A simple toy (fewer than 10 parts)
 - An electric iron
 - A kitchen mixing machine or food processor
 - An iPod, cassette, or disk player
 - The product resulting from the design problem (Exercise 4.1) or the redesign problem (Exercise 4.2)
- 11.6** For the device chosen in Exercise 11.5, perform a failure mode and effects analysis.
- 11.7** For one of the products in Exercise 11.5, evaluate it for disassembly, reuse, and recycling.



11.12 ON THE WEB

Templates for the following documents are available on the book's website: www.mhhe.com/Ullman4e

- Machined Part Cost Calculator
- Plastics Part Cost Calculator
- DFA
- FMEA

12

C H A P T E R

Wrapping Up the Design Process and Supporting the Product

KEY QUESTIONS

- What additional documents are needed to launch a product?
- What is important in supporting vendor and customer relationships?
- How are engineering changes managed?
- How can you apply for a patent?
- What does it mean to design for a product's end of life?

12.1 INTRODUCTION

We have come a long way. We began with the need for a product and planning for its development. We then worked our way through product definition and conceptual design. Then we began the hard work of turning this concept into a product that could be manufactured. The diagram shown in Fig. 12.1, a reprint of Fig. 4.1, makes it look easy.

This chapter wraps up the design process and discusses issues that generally occur near its end. Even if the techniques have led to the development of a final design represented by a solid model sufficiently detailed to generate detail and assembly drawings and a bill of materials, the process is not yet complete. We must still finalize all the documentation and pass a final design review before launching the product for production and into the marketplace. Even then, the designer may be involved in product changes and retirement.

Figure 12.2 details the activities necessary for product support. Although all of the best practices we've used in this book have developed documents that trace the evolution of the product, many other documents are still needed. These are detailed in the first section of this chapter.

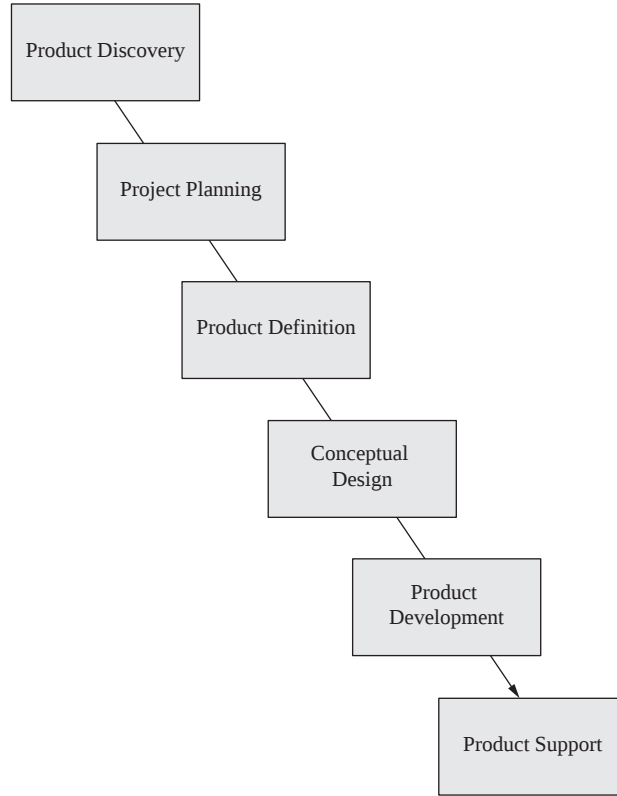


Figure 12.1 The mechanical design process.

A large part of an engineer's activity as a product nears production may be interaction with vendors, manufacturing, and assembly. Without these partners, the product will never reach the customers. If a product is being developed for a specific customer, there may be an extensive interaction with the customer's representatives as the product nears finalization. The nature of an engineer's relationship with the stakeholders will be detailed in this chapter.

In an ideal world, all the correct decisions were made during the development of the product. However, the reality is that manufacturing components will require changes. Managing these changes often becomes a large part of an engineer's responsibilities. In fact, engineers assigned to maintaining existing products may work totally on changes. Managing changes is not easy. Methods to ensure change success will be described.

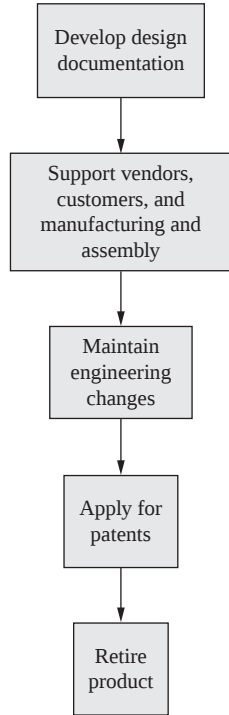


Figure 12.2
Product support details.

If all has gone well, maybe some of the ideas developed are patentable. We used the patent literature as a source of ideas during conceptual design. In this chapter, we will describe how to apply for a patent.

Finally, we will focus on retiring the product. In Chap. 1, we introduced the Hannover Principles, and in Chap. 11, we introduced Design for the Environment. These come together in and are concerned with what happens to the product at the end of its useful life. The final section of this chapter focuses on this concern.

Once all of these activities have been accomplished, there is usually one last design review, made when the product is complete and all of its documentation is in order. Only after passing this review, will the product to be ready to be released for production. Just because a product has been refined to this point does not mean that it will be produced. An engineer can spend many years designing a product, only to have the project canceled prior to production release. However, the probability of a product being approved is greatly increased if the customers' requirements have been met with a quality design in a timely and cost-effective manner. The techniques, or best practices, given in this book have focused on achieving this positive outcome.

Documentation is like the poor crust on a good pie,
you must eat it to clean your plate.

12.2 DESIGN DOCUMENTATION AND COMMUNICATION

In the previous chapters, many design best practices were introduced to aid in the development of a product. The documentation generated by these techniques, along with the personal notebooks of the design team members and the drawings and bill of materials, constitute a record of a product's evolution. Additionally, summaries of the progress for design reviews also exist. All of this information constitutes a complete record of the design process. Most companies archive this information for use as a history of the evolution of the product, or in patent disputes or liability litigation.

Beyond the information generated during the process, there is still much to be done to communicate with those downstream in the product's life. This section briefly describes the types of additional documents that need to be developed and communicated.

12.2.1 Quality Assurance and Quality Control

Even if quality has been a major concern during the design process, there is still a need for Quality Control (QC) inspections. Incoming raw materials and manufactured components and assemblies should be inspected for conformance to the design documentation. The industrial engineers on the design team usually have the responsibility to develop the QC procedures that address the questions, What is to be measured? How will it be measured? How often will it be measured?

Quality Assurance (QA) documentation must be developed if the product is regulated by government standards. For example, medical products are controlled by the Food and Drug Administration (FDA), and manufacturers of medical devices must keep a detailed file of quality assurance information on the types of materials and processes used in their products. FDA inspectors can come on site without prior notification and ask to see this file.

12.2.2 Manufacturing Instructions

A good drawing should have all the information needed to manufacture a component. Nonetheless, each plant has a certain set of manufacturing equipment, jigs and fixtures, and processes to make each component. Industrial engineers have the major responsibility for developing these manufacturing instructions. In very small companies, those with no industrial engineers, manufacturing instructions may become the responsibility of the product designers.

12.2.3 Assembly, Installation, Operating, and Maintenance Instructions

We have all purchased products, opened the box, and seen that there was “some assembly required.” Then, on reading the directions, found that they were unintelligible. Similarly, most software user manuals are impossible to decipher. In smaller organizations, engineers often get to write assembly, installation, operating, and maintenance instructions. In larger organizations, engineers may work with professional writers to create these documents. Either way, it is important to understand what is required to develop a good set of instructions.

For many products, *assembly instructions* are part of the total design package. These instructions spell out, step by step, how to assemble the product. This is necessary whether the assembly is done by hand or by machine. The generation of assembly instructions, while tedious, can be enlightening in that the assembly itself is refined, the assembly sequence (Section 11.5) is refined, and jigs and fixtures for holding the assembly are developed. *Installation instructions* include instructions for unpacking the items and making the necessary connections for power, support, and environmental control. Instructions for initial start-up and testing may also be included. For many systems, these are major parts of the final product package. *Operation instructions* include instructions on how to operate the device over the normal range of activity. Various modes—start-up, standby, emergency operation, and shutdown—may be described. Instructions on how to determine when the equipment is failing may also be included. Finally, all products need maintenance. *Maintenance instructions* may be included with operating instructions. Maintenance can range from something as simplistic as cleaning the surface of the product to total disassembly and inspection.

Although writing instructions may not seem like a task suited for an “engineer,” writing them can help you understand your product in a unique way. It forces you to assume the role of assembler, installer, operator, and maintainer. In fact, writing instructions is helpful to understanding your product if you begin to write them early in the design process.

Some guidelines for writing instructions are

1. Read as many similar instruction manuals as you can. Many companies post their manuals online, or you can obtain one by calling a company’s headquarters and requesting a copy.
2. Organize instructions into sections to make it easy to find answers. Do not write in the order you developed the product, write in the order in which it will be assembled, installed, operated, or maintained. A good way to understand the difference is to walk through assembling, installing, operating, or maintaining the product while pretending you have no knowledge beyond that which you assume the readers of the instructions to have.
3. Recruit members of the user community not familiar with the product to test your instruction manuals. It is best if you hand them the instructions and then watch them assemble, install, operate, or maintain the product using what

you wrote. It is important not to say anything while observing. It is amazing to witness how much you have assumed. You need to observe whether or not the instructions are easy to follow, or if searching, rereading, and interpreting are required? Instructions should consist of short paragraphs explaining the process, plus accompanying numbered or bulleted lists, figures, photographs, or screenshots, and steps for users to follow. Text instructions embedded in long paragraphs are extremely difficult to follow.

4. Make instructions activity centered. Explain the most basic activities and how to accomplish them. Make the explanations short and simple and do not explain every knob and button and menu item.
5. Put legal warnings in an Appendix. When instructions are needed, they are needed right away, and having to work one's way through pages of legal warnings only increases the anxiety level and decreases the pleasure of the product. Moreover, people skip these anyway, so they are ineffective. Consult with a lawyer to make sure you include the right wording to protect your company and employees from potential liability. This is especially important if you have to write instructions for products that may be potentially dangerous.
6. Hire an excellent technical writer. The instruction writers should be a part of the design team. Ideally, instructions are written first, to help understand the voice of the customer.

12.3 SUPPORT

Although not usually thought of as part of the design process, support for downstream activities often takes a sizable portion of engineering time. It has been estimated that about 20 to 30% of all engineering time is spent supporting existing products. Support includes maintaining vendor relationships, interfacing with customers, supporting manufacturing and assembly, and maintaining changes (see Section 12.4).

12.3.1 Vendor Relationships

Very few products are made solely in house. In fact, many companies make no components themselves and only specify, assemble, sell, or distribute what others make. Others only specify and make nothing themselves. Thus, for most companies, relationships with their vendors are crucial. Prior to 1980, many large companies had thousands of vendors, each chosen for its low bid to make a component or assembly. These companies realized, however, that this was a poor way to do business, because the cheapest components were not always of the highest quality even if they met the specifications. Additionally, managing thousands of vendors proved very expensive and difficult.

By the 1990s, the philosophy of including vendors on the design team from the beginning of the project had evolved. This allows for fewer vendors, empowers the vendors as “stakeholders” in the product, and requires that organized design processes be used. Making this shift from low bidder to virtual partner required changing the philosophy of vendor relationships. Most companies reduced their

number of vendors by an order of magnitude. Many now use vendors from only a small, select list. In some cases, the product manufacturing company has a financial interest in the vendor, or vice versa.

Guidelines that can help you build and maintain good vendor relationships include:

1. Know your goals and your vendor's goals. Building a strong vendor relationship means more than cutting product or service costs. It is about improving value provided to the business, reducing the time to deliver solutions, reducing staff effort, and much more. Define the goals and objectives of your department/company and work only with vendors who are aligned with your goals. Vendor's goals may include building a center of excellence, entering new markets, gaining market share within a product line, developing industry verticals, and so on. It is very important for your relationship to understand the vendor's goals and determine how your organization fits into this strategy. When the vendor's goals are aligned with your goals, the relationship will be more successful since you are both working toward the same end results.
2. Define clear relationship guidelines. Meeting with a vendor only when there is a problem with a product is a problem relationship from the beginning. Both organizations lose from this relationship. Clearly defining a regular vendor meeting structure with a defined agenda is the key for both organizations to understand the goals, needs, wants, and actionable items of the other. Both parties must clearly understand each others obligations, who is responsible, and the expected outcomes. Clearly defining this up front is a key success factor.
3. Involve vendors early. When dealing with vendors, you cannot afford delays and extensive alterations. Treat them as your customers early in the product development process, include them on teams, and enlist their expertise as you design the product.
4. Establish relationships. It is important to have vendor partners who understand that the relationship should be win-win for both parties. If you do a lot of business with a particular vendor, he or she will reward you for your loyalty by offering discounts and incentives to you. They will even go out of their way to help you by speeding up the shipment process if you need to quickly ship some orders, for example, or receive a back order. There should be a single point of contact in both organizations and they should get to know each other.
5. Treat vendors with respect. The Golden Rule of any relationship is, "Treat others as you want to be treated, with respect and integrity." Treat your vendors like your customers, and they in turn will treat you like a customer. All successful relationships are built on mutual trust. Only work with vendors that have a good reputation, ones that keep their word. Likewise, be honest and forthcoming in your communications to vendors.
6. Communicate. Put everything in writing—responsibilities, expected sales volume, payment, mode of payment, and so on. Anything you think may cause misunderstanding and strained vendor relationships later must be put

down in writing beforehand. When in doubt, talk it out. What works for interpersonal relationships also serves as a reliable rule of thumb for fostering healthy relationships with your vendors. Poor communication will reduce your relationship to, “It is not in the contract” instead of the response “How can we help you.”

7. Stay professional. Things go wrong in life. When they go wrong in a relationship, the smartest thing to do is to deal with the problem calmly and factually, in order to avoid ruining the relationship.

12.3.2 Customer Relationships

Although many companies isolate their engineers from their customers, others make an effort to close the loop with direct feedback from customers to engineers. Most companies have a product service department that handles day-to-day customer communication and filters information reaching the engineers. This is both necessary and a problem as interruptions slow the development of new products, but some direct contact improves the product developer’s understanding of how the products are being used, their good features, and their bad features.

Other companies, especially those that produce low-volume products, have the engineers work directly with customers. Using methods like quality function deployment keep that communication positive and useful.

12.3.3 Manufacturing and Assembly Relationships

In Chap. 1, the over-the-wall design method showed information flowing from design to production and not back again. Most modern companies try to maintain communication between the two groups so that problems in manufacturing and assembly, those that can lead to changes, are minimized. Methods already discussed like concern for the product life cycle, DFM, DFA, and PLM all help break the over-the-wall way of doing business. For example, quoting from a Neon design manager at Chrysler, “It used to be that the engineers handed off the project to the assembly plant 28 weeks before volume production began . . . now workers began meeting with engineers on the Neon 186 weeks before Job One.” At various stages of Neon development, busloads of engineers traveled en masse to meet with manufacturing and assembly workers to ready the car for production. These meetings focused on designing the product to be easy to manufacture and assemble. This transformation is significant for a company of Chrysler’s size.

12.4 ENGINEERING CHANGES

Although this book encourages change early in the design process, change may still occur after the product is released to production (see Fig. 1.5). Changes are caused by

- Correction of a design error that doesn’t become evident until testing and modeling, or customer use reveals it.

Only a perfect product will never change, and there is no such thing as a perfect product.

- A change in the customers' requirements necessitating the redesign of part of the product.
- A change in material or manufacturing method. This can be caused by a lack of material availability, a change in vendor, or to compensate for a design error.

To make a change in an approved configuration, an *Engineering Change Notice* (ECN), also called an *Engineering Change Order* (ECO), is required. An ECN is an alteration to an approved set of final documents and thus needs approval itself. As shown in the example in Fig. 12.3, an ECN must contain at least this information:

- Identification of what needs to be changed. This should include the part number and name of the component and reference to the drawings that show the component in detail or assembly.
- Reason(s) for the change.
- Description of the change. This includes a drawing of the component before and after the change. Generally, these drawings are only of the detail affected by the change.
- List of documents and departments affected by the change. The most important part of making a change is to see that all pertinent groups are notified and all documents updated.
- Approval of the change. As with the detail and assembly drawings, the changes must be approved by management.
- Instruction about when to introduce the change—immediately (scrapping current inventory), during the next production run, or at some other milestone.

12.5 PATENT APPLICATIONS

In Chap. 7, patent literature was used as a source of conceptual ideas. During the evolution of a product, new ideas, ones not already covered by patents, are sometimes generated and a patent application developed.

Just about any device or process that is new, useful, and not obvious is patentable. In obtaining a patent, the inventor is essentially entering into a contract with the U.S. government, as provided for in the Constitution. The inventor is granted an exclusive right to the idea for 20 years from the filing date.¹ In many cases, the contract is between the inventor's employer—called the

¹Prior to 1995 patents were good for 17 years from the date of issue.



Engineering Change Notice	
Design Organization:	Date:
Subject of Change:	
Reason for Change:	
Description of Change (include drawings as attached pages as needed):	
Impact of change: ☐ Bill of Materials	Team member: Team member: Team member: Team member: Prepared by: Checked by: Approved by:
<i>The Mechanical Design Process</i> Designed by Professor David G. Ullman Copyright 2008, McGraw-Hill Form # 26.0	

Figure 12.3 Engineering change notice.

“assignee”—and the government. Most employers require engineers to sign an agreement that says all ideas belong to the employer, thus most engineers names appear as assignee.

In return for receiving a patent, the inventor must make full public disclosure of the idea through the publication of the patent. This system enables all of society to benefit from the idea and still protects the inventor.

Patents give you bragging rights and a license to litigate.

In reality, a patent is only as good as the inventor's ability to enforce it. In other words, if the holder of a patent is not capable of suing a party that is infringing on that patent, then the patent is virtually useless. Since lawsuits can be extremely expensive, an individual or company holding a patent must decide whether it is better business to allow the infringement or to litigate. Being an assignee to a patent owned by a large organization has a positive benefit when it is necessary to enforce a patent.

Whether applying for a patent or not, it is a good idea to write a *disclosure* whenever a new, potentially patentable idea is developed. A disclosure is description of the idea that is signed, dated, and witnessed. The description of the idea should be in terms of specific claims. Claims describe the new or unique utility or function of the device. A disclosure serves as a legal statement of when the idea was first conceived.

There are essentially two types of patents: design patents and utility patents. *Design patents* relate only to the form or appearance of the article. Thus, there are many design patents for the basic toothbrush, since each toothbrush has a different appearance. *Utility patents*, on the other hand, protect utilitarian or functional aspects of the idea. The patent discussion in Chap. 7 and here pertains only to utility patents. Utility patents are given for processes, machines, manufacturing techniques, or composition of matter.

Applying for a patent is time-consuming but not overly expensive. The procedure outlined in Fig. 12.4 shows the steps involved, whether the application is done by an individual or through a lawyer.

The first step is the preparation of the *specification* or *prospectus* role in the document that becomes the patent. The basic parts of the specification are shown in Fig. 12.5. When preparing a specification, it is best to seek help from a patent lawyer or reference a book with details for writing claims and the proper patent format. Within three months of receiving the patent application, the patent office will acknowledge its submission and accept the application for consideration. The term "patent applied for" can be used after the specification is accepted. (The term "patent pending," often seen on products, has no legal meaning.)

The second phase of the patent process begins after the application is filed. The patent office assigns the application to an examiner familiar with the state of the patent art in the area of the application. This examiner then reviews the application and searches the patent literature. This initial study of the application and the subsequent iteration with the applicant takes months or even years. Seldom does the examiner accept the first application outright. It is more likely that he or she objects to the document itself or to the specification or rejects one (or more) of the claims.

Problems with the document usually stem from not following the rigid style required of patent text and drawings. Problems with the specification concern

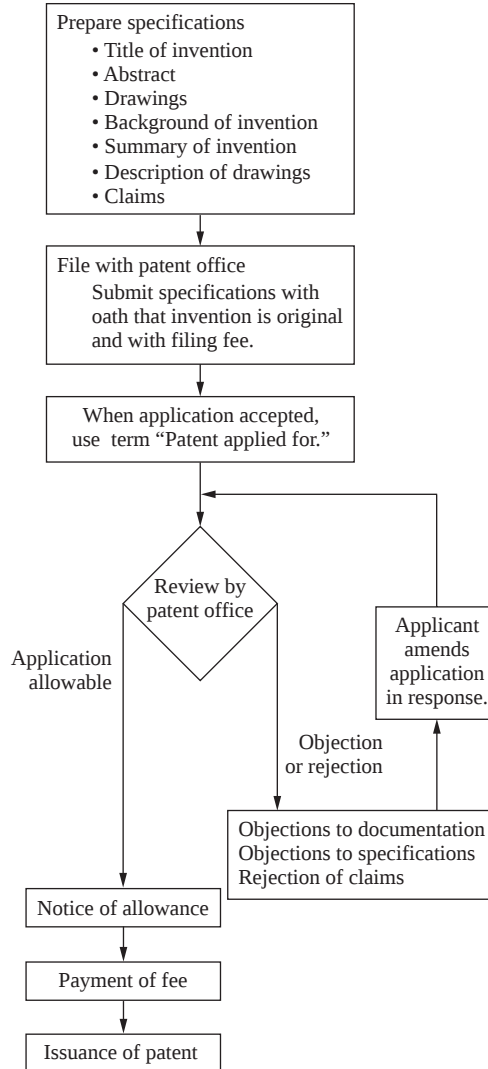


Figure 12.4 Patent application procedure.

the content of the claims. Since the claims define the invention, they are the heart of the patent. All the other information is there simply to support the claims.

After the applicant and the patent examiner agree on the application, the patent office issues a notice of allowance. This means that the patent will be issued on payment of an issue fee. About 65% of the patents applied for are ultimately granted. Virtually all of these have been greatly modified during the process.

Patent Specification	
Design Organization:	Date:
Title of Invention:	
Abstract:	
Background of the Invention:	
Summary of the Invention:	
Description of Drawings:	
Claims:	
Attach drawings as needed	
Notes about filing with the patent office:	Team member:
	Team member:
	Team member:
	Team member:
	Prepared by:
	Checked by:
	Witnessed by:
<i>The Mechanical Design Process</i> Designed by Professor David G. Ullman Copyright 2008, McGraw-Hill Form # 27.0	



Figure 12.5 Patent Specification.

12.6 DESIGN FOR END OF LIFE

Part of Design For the Environment, DFE, is concern for what to do with a product at the end of its life. This is especially so in the automotive industry, because there is so much material tied up in 250 million cars and light trucks

Table 12.1 Typical car composition by weight

Material	Percentage	Changes in last 15 years
Metals		
Ferrous	65%	Down 7%
Aluminum	8%	Up 4%
Other	4%	
Plastics	9%	Up 3%
Rubber	6%	
Glass	3%	
Misc	5%	

registered in the United States (and an equal number in Europe). Thus, it is worth looking at efforts to recycle End-of-Life Vehicles (ELVs). In the United States, approximately 12.5 million cars and light trucks are recycled each year. The average composition of these vehicles is shown in Table 12.1. This composition is changing. In an effort to reduce weight, more aluminum and plastics are now used than there were 15 years ago. The increase in the use of plastics also reduces manufacturing costs. The basic flow and percentage of materials recovered when recycling an ELV is shown in Figure 12.6. There are four steps:

1. **Dismantling:** This is currently a vehicle-by-vehicle effort at a salvage/scrap yard, where a variety of parts and all vehicle fluids and tires are removed. After removal, the remaining gutted vehicle (“hulk”) is flattened prior to shipment to the shredding facility.
2. **Shredding:** The vehicle hulks are transported to a company that shreds, separates, and processes them. First, a shredding machine takes about a minute to reduce the hulks to fist-sized pieces.
3. **Separation and processing:** The shredded material is separated using magnets to attract ferrous metal (all iron and steel, except stainless steel) away from all the nonferrous materials (both metals and nonmetals). The ferrous material is sent for recycling to steel smelters. The nonferrous material fraction is then typically separated into other metals: aluminum, brass, bronze, copper, lead, magnesium, nickel, stainless steel, and zinc; and into Auto Shredder Residue (ASR) or “fluff.” ASR consists of plastics, glass, rubber, foam, carpeting, textiles, and so on.
4. **Landfill disposal of ASR:** For the most part, ASR is considered nonrecoverable waste material and is sent to landfills for disposal.

As the percentage of plastics used in cars and the cost of oil increases, recovering some of the plastics from the ASR becomes more valuable. What makes this challenging is that there are a wide variety of plastics mixed together in the ASR.

Europe is ahead of the United States in its effort to manage ELVs. In an agreement signed in September 2000, the European Union agreed to

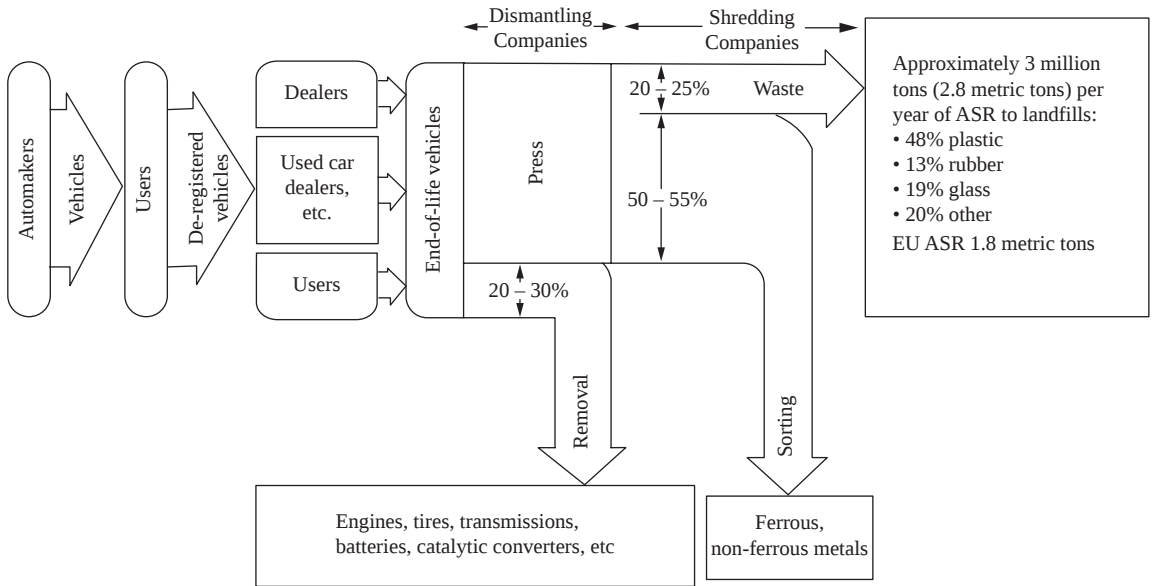


Figure 12.6 The life cycle of a vehicle emphasizing ELVs.

- Establish Extended Producer Responsibility (EPR) for ELV management, requiring manufacturers and importers of autos to pay for the costs of end-of-life management.
- Set increased recycling requirements. By January 1, 2006, reuse and recovery minimums were 85% by weight on average. By January 1, 2015, this needs to be 95% by weight.
- Establish phase-outs in use of certain heavy metals: lead, mercury, cadmium, and hexavalent chromium, except in certain excluded components (e.g., lead in lead-acid batteries; hexavalent chromium as a corrosion preventative coating; lead containing alloys of steel, aluminum, and copper; lead as a coating inside fuel tanks; and mercury in headlamps).
- Encourage Design For the Environment (DFE) practices.
- Code components and materials to facilitate product identification for material reuse and recovery.
- Provide dismantling information for every vehicle built.

The U.S. government has done little nationally about ELVs, but states such as California are taking the lead. Additionally, vehicle manufactures are changing both their business model and, to a slower extent, their design practice. Ford is purchasing recycling operations. DaimlerChrysler has goals to reach 95% recyclability, to reduce the types of materials used by 40%, and to increase the recycled content in their vehicles.

Recovery of materials comes at a cost. First, the vehicles must be transported to the dismantling facility. Then the hulks and removed components must be transported to the shredders and other processors. Shredding takes about 30 kWh/vehicle (375 million kWh per year total) and additional energy is needed to process the shredded material.

To design for the end of life, designers of vehicles and other products need to

1. Design for disassembly.
2. Label components for easy material identification.
3. Use fewer different types of materials.
4. Design products with a longer life span.

12.7 SOURCES

Burgess, J. A.: *Design Assurance for Engineers and Managers*, Marcel Dekker, New York, 1984. A very complete and well-written book on the development and control of engineering documentation.

Guide to Filing A Non-Provisional (Utility) Patent Application, U.S. Patent and Trademark Office <http://www.uspto.gov/web/offices/pac/utility/utility.htm>
<http://ec.europa.eu/environment/waste/index.htm> gives some details on the European Union's effort for retiring products

Kivenson, G.: *The Art and Science of Inventing*, Van Nostrand Reinhold, New York, 1977. Good overview of patents and patent applications; however, it is out of date on application details.

Stevens, Ab: "Design for End-of-Life Strategies and Their Implementation," Chapter 23 in *Mechanical Life Cycle Handbook*, by M. S. Hundal, Marcel Dekker, 2001.

Management of End-of Life Vehicles (ELVs) in the United States

Staudinger, Jeff, Gregory A. Keoleian, and Michael S. Flynn: "A Report of the Center for Sustainable Systems," Report No. CSS01-01, University of Michigan, 2001, http://css.snre.umich.edu/css_doc/CSS01-01.pdf

End-of-Life Vehicle Recycling in the European Union

Kanari, N., J.-L. Pineau, and S. Shallari: JOM, August 2003, C:\Documents and Settings\D\My Documents\MDP 4th\Chpt 12 Launch\End-of-Life Vehicle Recycling in the European Union.htm



12.8 ON THE WEB

Templates for the following documents are available on the book's website: www.mhhe.com/Ullman4e

- Engineering Change Notice
- Patent Specification

Properties of 25 Materials Most Commonly Used in Mechanical Design

A.1 INTRODUCTION

There are literally an infinite number of materials available for use in products. In addition, it is now possible to actually design materials for a specific use. There is no way a design engineer can have knowledge of all these materials; however, all design engineers should be familiar with the materials that are the most available and the most commonly used in product design. Because these same materials are representative of a broad spectrum of materials, the design engineer can use his or her knowledge about them to communicate with materials engineers about other, less common materials.

In addition to the important properties of the 25 most used materials, this appendix also contains a list of the specific materials used in many common items. During material selection it is vital to know what materials have been used for similar applications in the past; this list provides a source of such information.

This appendix concludes with an extensive bibliography; the publications listed there are a source for information beyond the basic data presented here.

A.2 PROPERTIES OF THE MOST COMMONLY USED MATERIALS

The following 25 materials are those most commonly used in the design of mechanical products; in themselves they represent the broad range of other materials.

Steel and irons

1. 1020
2. 1040
3. 4140
4. 4340
5. S30400
6. S316
7. 01 tool steel
8. Gray cast iron

Aluminum and copper alloys

9. 2024
10. 3003 or 5005
11. 6061
12. 7075
13. C268

Other metals

14. Titanium 6-4
15. Magnesium AZ63A

Plastics

16. ABS
17. Polycarbonate
18. Nylon 6/6
19. Polypropylene
20. Polystyrene

Ceramics

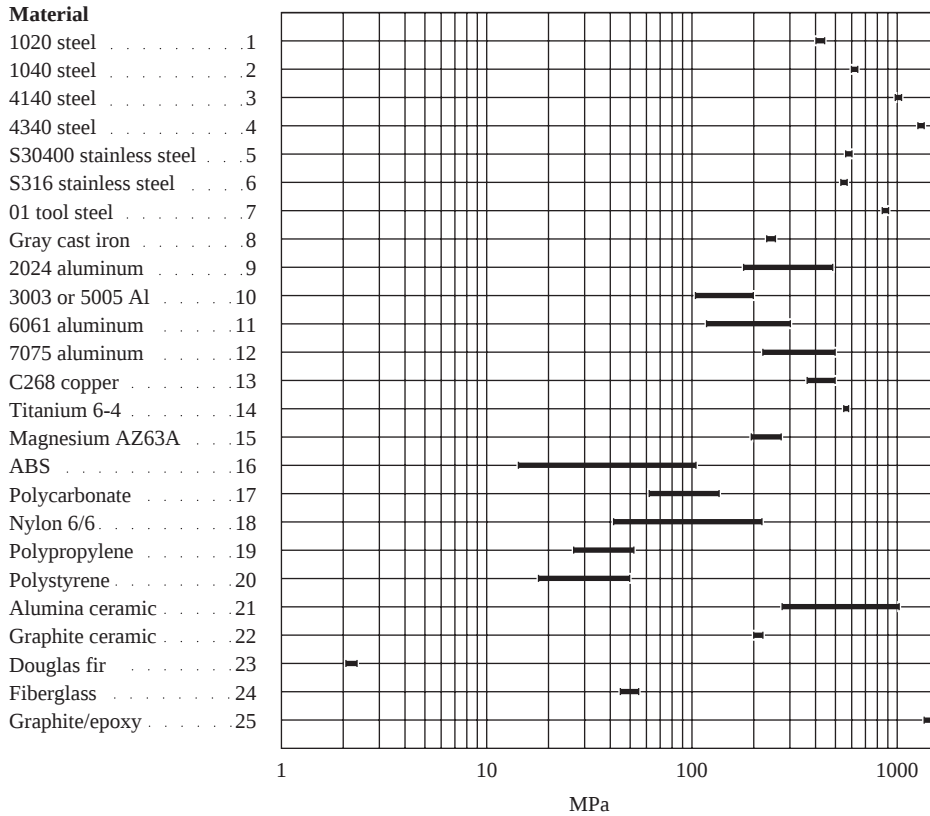
21. Alumina
22. Graphite

Composite materials

23. Douglas fir
24. Fiberglass
25. Graphite/epoxy

The properties of these 25 materials, given in Figs. A.1–A.12, are the properties most commonly needed for design purposes.¹ Other properties can be found in the references at the end of this appendix. The properties are given as ranges, since they will depend on specific heat treatment (metals) and additives (plastics).

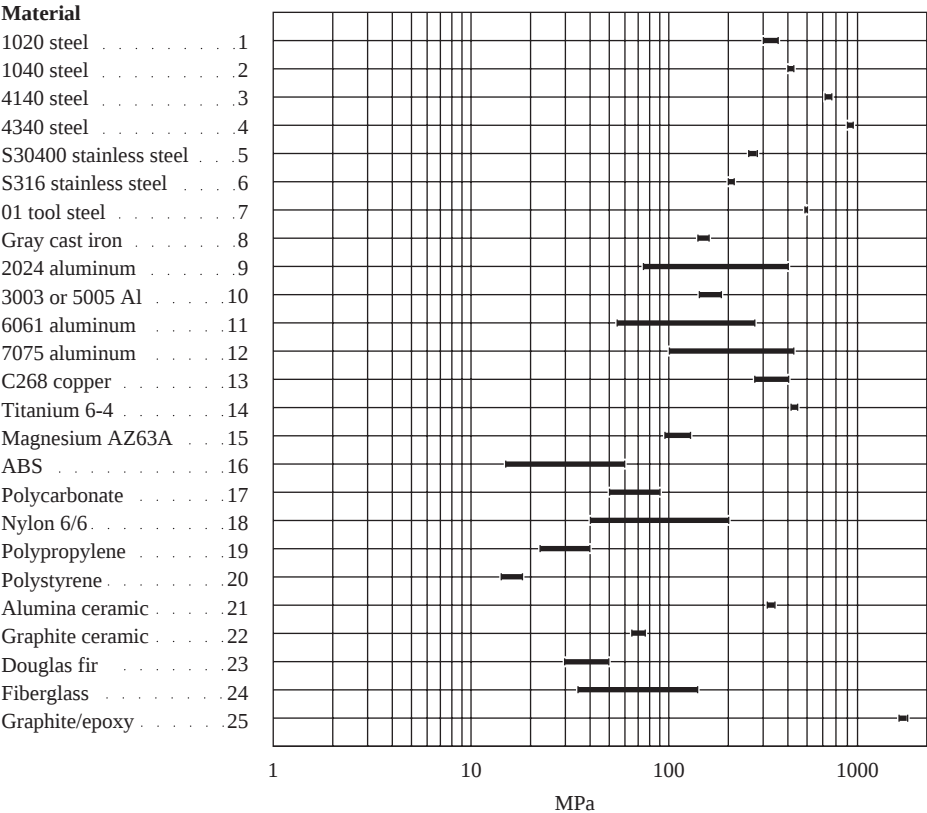
¹An excellent book with more properties presented in this manner and a computer program to support material selection is M. F. Ashby, *Materials Selection in Mechanical Design*, 3rd edition, Butterworth Heinemann, 2005.



Note: Longitudinal value for graphite/epoxy.

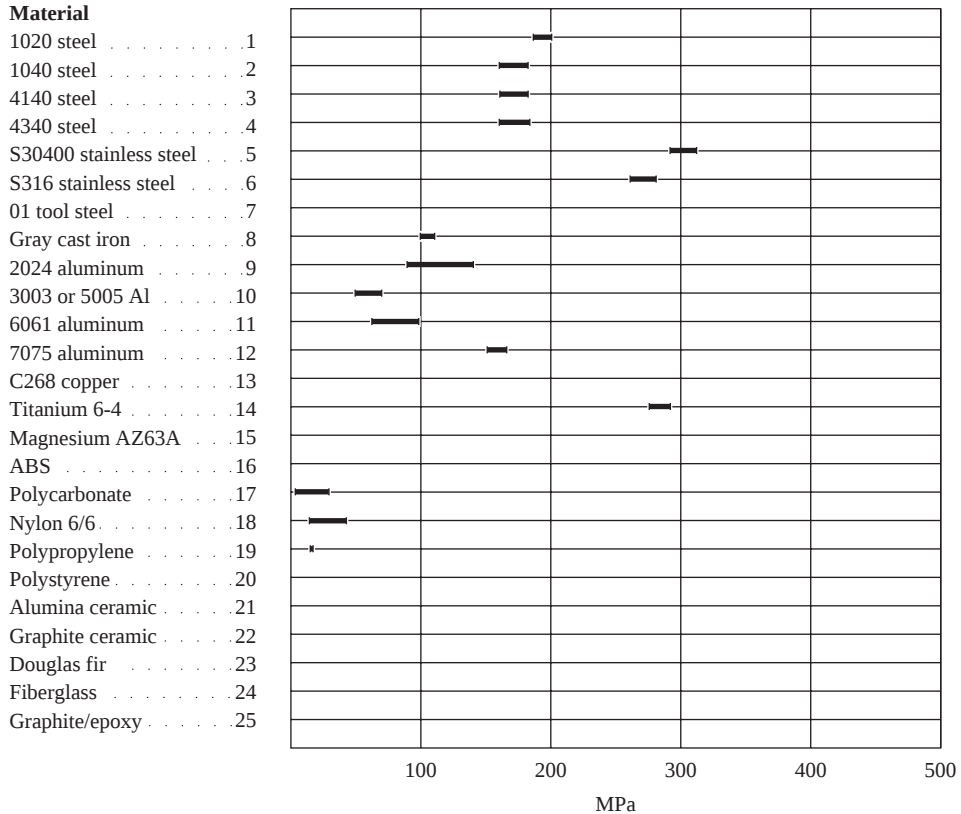
1 MPa = 144.7 psi.

Figure A.1 Tensile strength.



Note: 1 MPa = 144.7 psi.

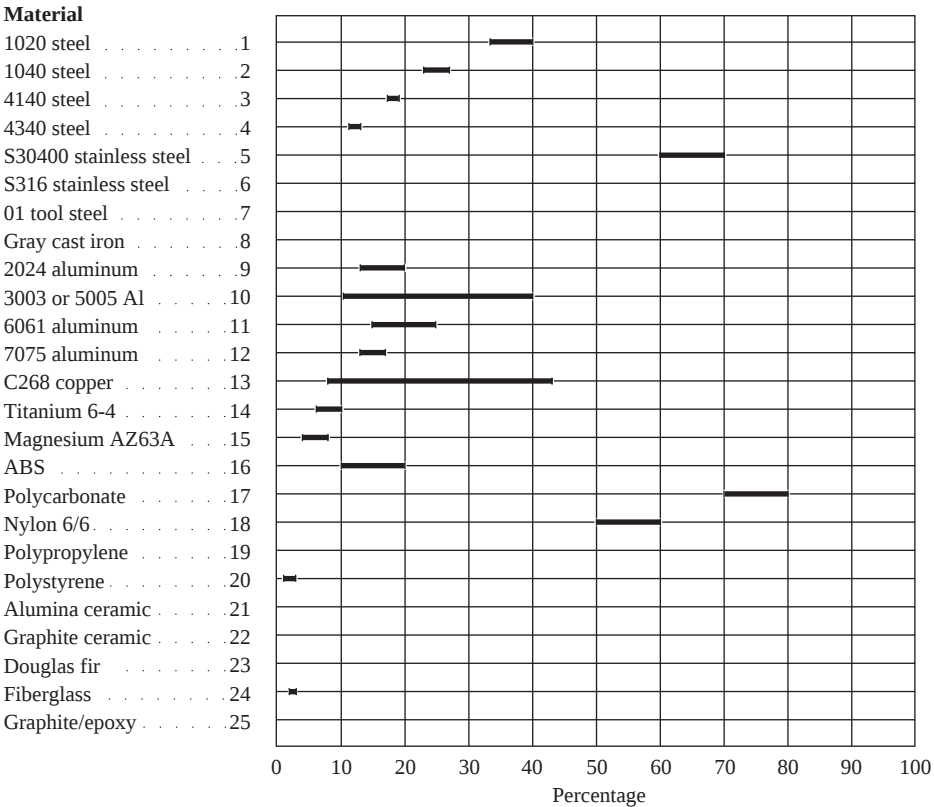
Figure A.2 Yield strength.



Note: Some materials do not have an endurance limit.

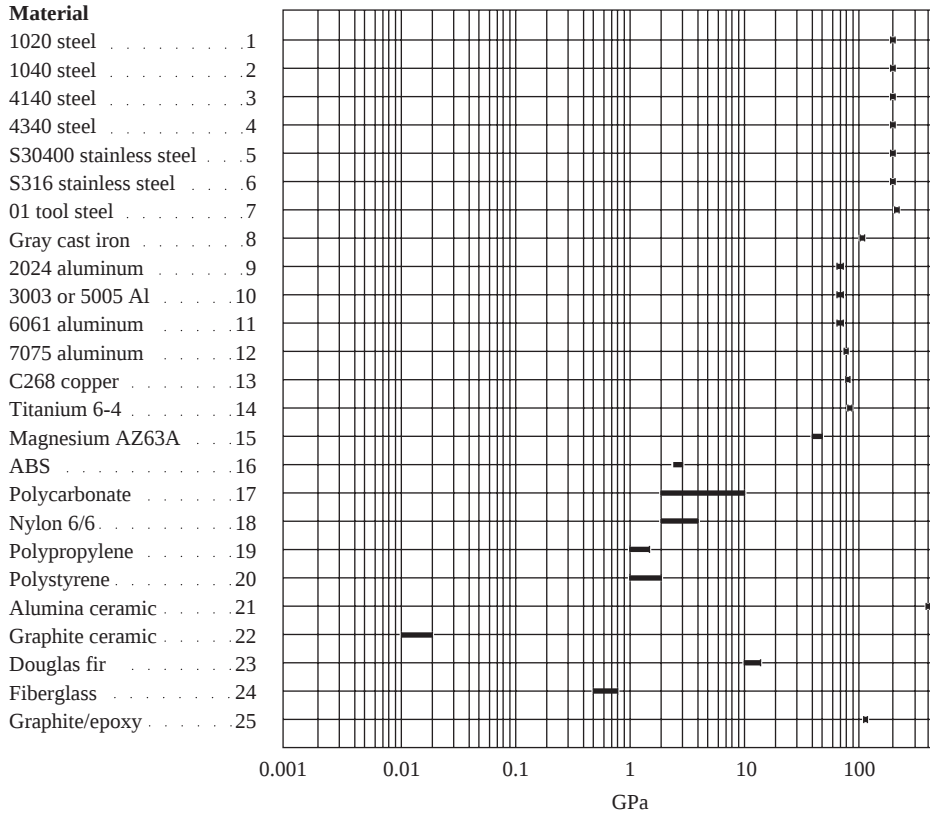
1 MPa = 144.7 psi.

Figure A.3 Endurance limit.



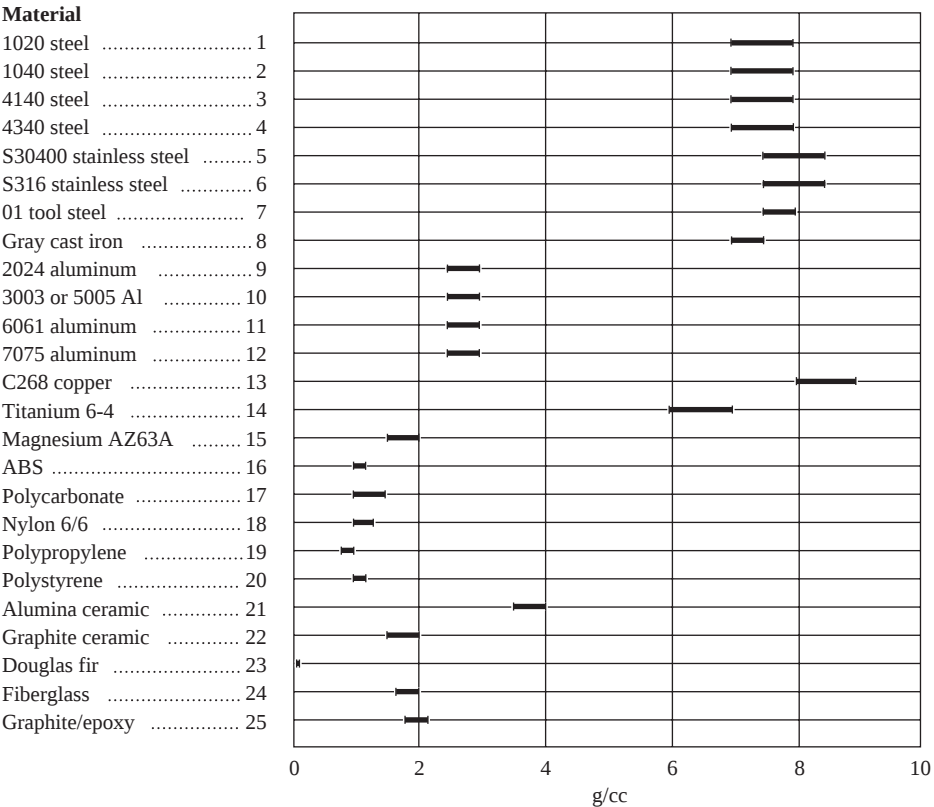
Note: Data unavailable for some materials.
Elongation in plastics depends on filler materials.

Figure A.4 Elongation.



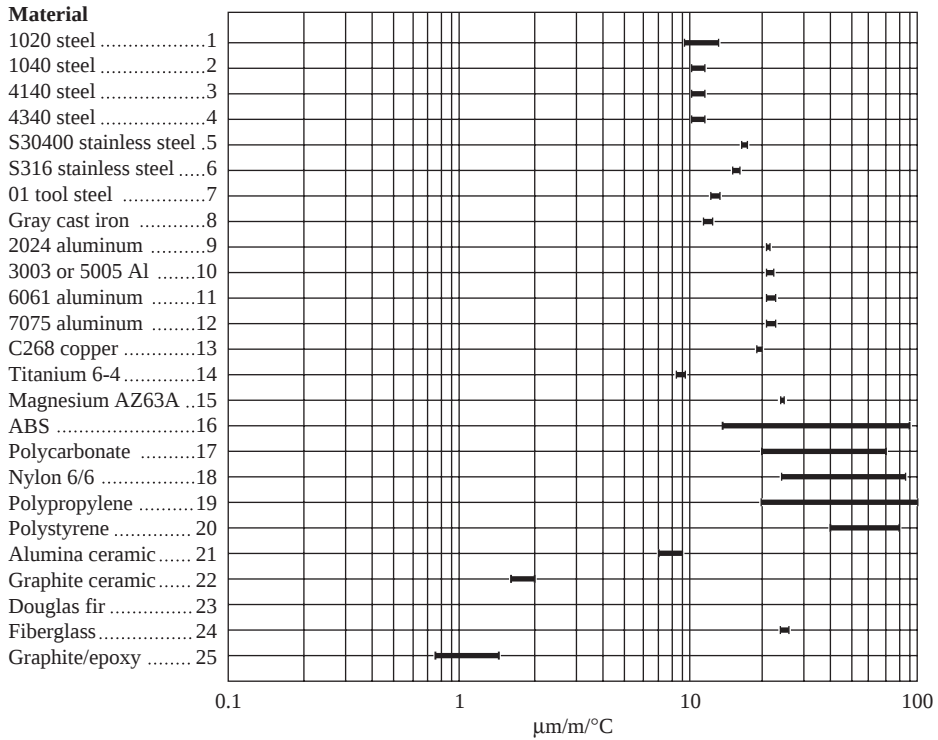
Note: 1 GPa = 144.7 kpsi.

Figure A.5 Modulus of elasticity.



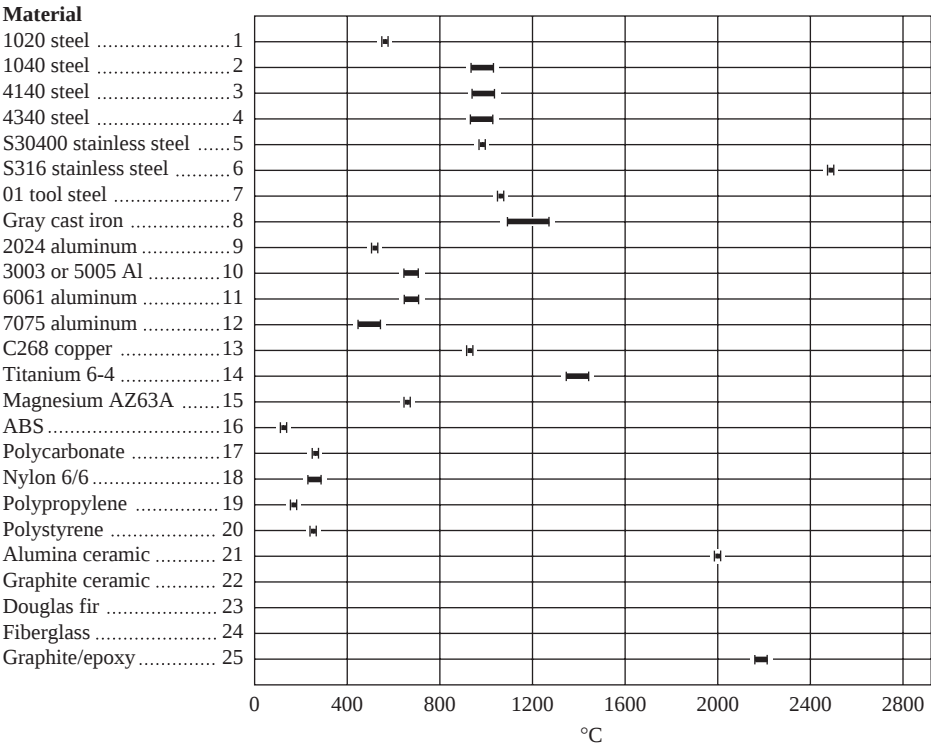
Note: 1g/cc = 0.036 1b/in³.

Figure A.6 Density.



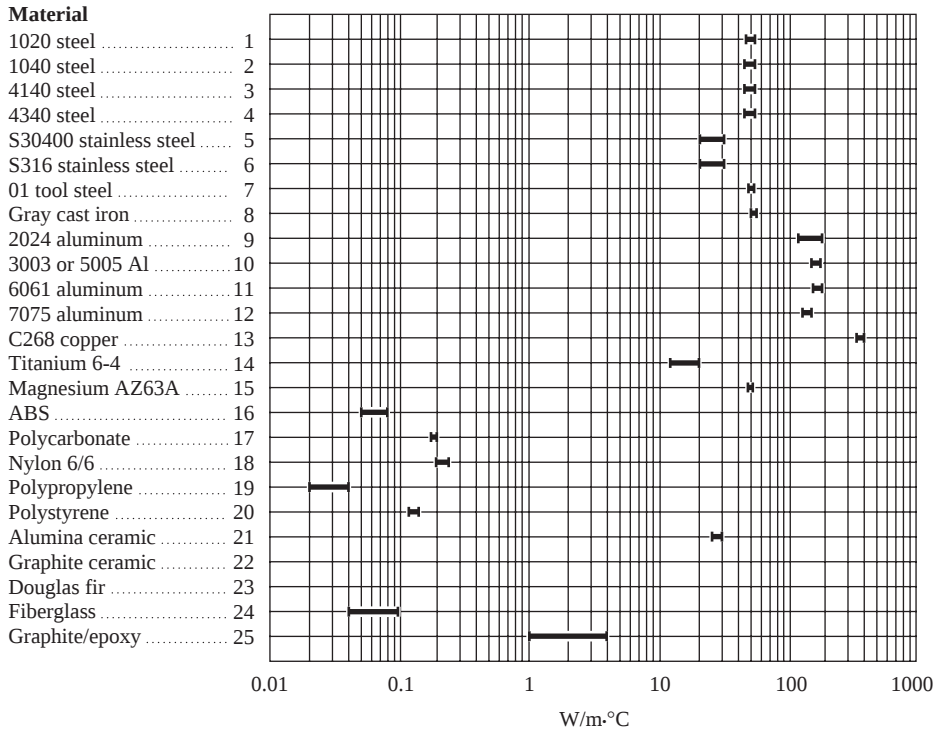
Note: Douglas fir varies greatly.
 $1 \mu\text{m}/\text{m}/^{\circ}\text{C} = 0.55 \mu\text{in}/\text{in}/^{\circ}\text{F}$.

Figure A.7 Coefficient of thermal expansion.



Note: Data unavailable for graphite, Douglas fir, and Fiberglass.
 $^{\circ}\text{F} = 32.2 + (9/5)^{\circ}\text{C}$.

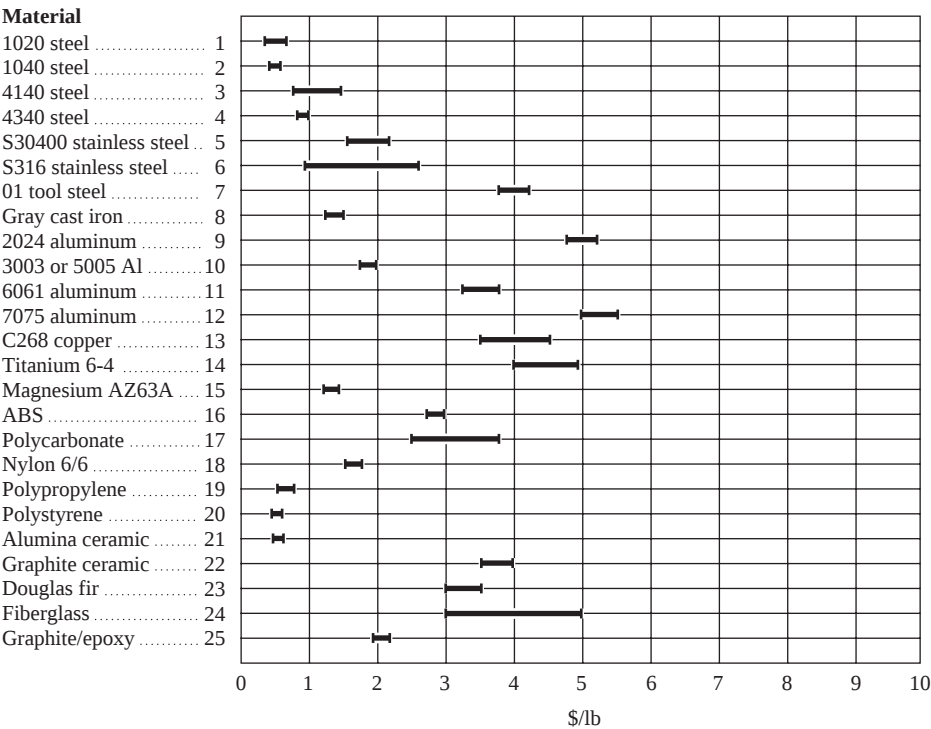
Figure A.8 Melting temperature.



Note: Data unavailable for some materials.

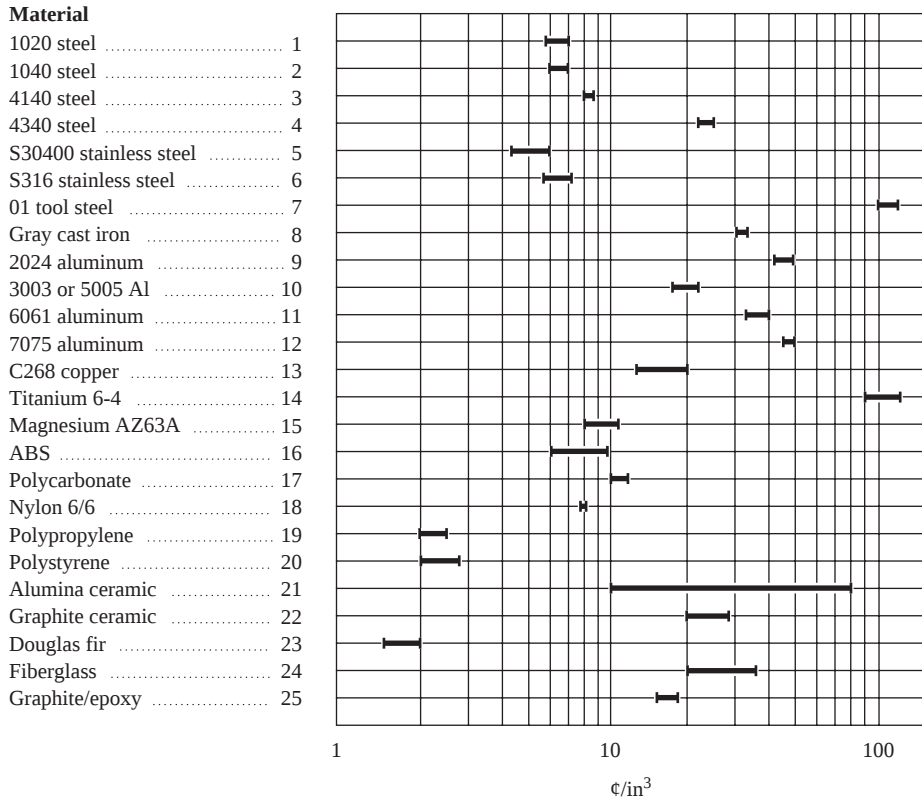
$W/m \cdot ^\circ C = 0.57 \text{ Btu/h} \cdot \text{ft} \cdot ^\circ F$.

Figure A.9 Thermal conductivity.



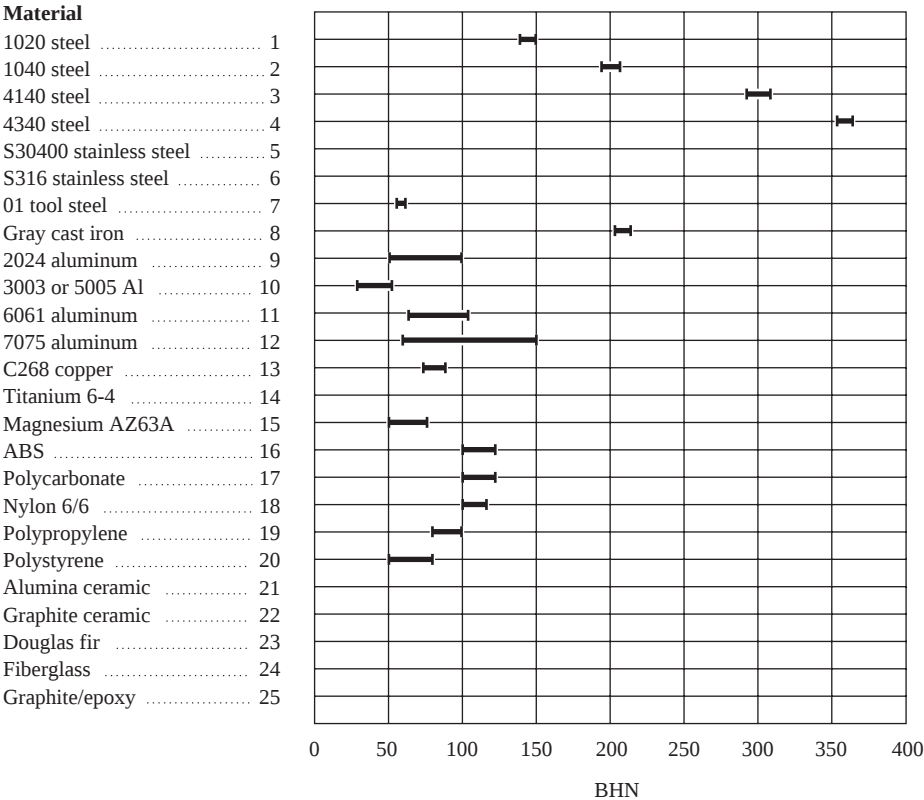
Note: \$ year 2000.

Figure A.10 Cost per pound.



Note: \$ year 2000.

Figure A.11 Cost per unit volume.



Note: Data unavailable for some materials.

Figure A.12 Hardness.

A.3 MATERIALS USED IN COMMON ITEMS

The materials from which many components are commonly made are listed in Table A.1. This list is intended to give ideas for materials to be specified in a new product. In the list, an asterisk denotes that the material is one of the 25 whose properties are given in Figs. A.1–A.12.

Table A.1 Materials used in common components

Component	Materials
Aircraft components	*Aluminum 2024, 5083, *6061, *7075; *titanium 6-4; *steel 4340, *4140; *graphite/epoxy; polyamide-imide, maraging steel
Auto engines	*Gray cast iron
Auto instrument panels	*Polypropylene; modified polyphenylene oxide
Auto interiors	*Polystyrene; polyvinyl chloride; *ABS; *polypropylene
Auto bodies, panels, parts, and coverings	*Steel 1020, *1040; *ABS; *alumina; *nylon; polyphenylene sulfide; aluminum 5083; *stainless steel 316; *fiberglass/epoxy; *polycarbonate
Auto taillight lenses	*Polycarbonate
Automotive trim	*ABS; styrene acrylonitrile; *polypropylene
Bearings and bushings	*Nylon; acetal; bronze; Teflon; beryllium copper; *stainless steel 316
Boat hulls	Styrene acrylonitrile; *aluminum 6061; polyethylene; *fiberglass/epoxy
Bottles	*Polycarbonate; thermoplastic polyester; high-density polyethylene; *polypropylene; polyvinyl chloride; PET
Battery cases	*Polystyrene
Business machine cases	*ABS; *polystyrene; *polycarbonate
Buttons	Melamine; urea
Cams	*Nylon; *aluminum 6061
Cabinets and housings	*ABS; acetal; *polycarbonate; polysulfone; *nylon; *polypropylene; *polystyrene; polyvinyl chloride; *steel 1040
Compact disks	*Polycarbonate
Computer casings	*ABS
Conveyor chains	Acetal; *steel 1040
Cryogenic parts	Boron/aluminum; *graphite/epoxy; aluminum 5086
Dies	*01 tool steel; cemented carbide
Electric connectors	Polysulfone; phenolic; *nylon; polyethersulfone (PES); *ABS
Fan blades	*Nylon; *steel 4340
Forgings	*Steel 1020, *4340; copper alloys; aluminum alloys
Fixtures	*01 and A2 tool steels; epoxy; *aluminum 6061
Gears	Acetal; *gray cast iron; *steel 1020, *4340; *nylon; *polycarbonate; polyamide; filled phenolic; aluminum bronze
Handles and knobs	Phenolic; melamine; urea; nylon*; *ABS; acetal
Helmets	*Polycarbonate; *ABS
Heat exchangers	Stainless steel

(continued)

Table A.1 Materials used in common components (*continued*)

Component	Materials
Hoses	*Nylon; *polycarbonate
House siding and gutters	*Polypropylene
Keys	*Steel 1020, *4140
Levers and linkages	Acetal; *steel 4140, *4340
Machine bases	*Gray cast iron; Steel 1020; ductile iron
Marine parts and instruments	Styrene acrylonitrile; *steel 1020; aluminum 5083, *6061; *polycarbonate; *titanium 6-4
Microwave cookware	Thermoplastic polyester; *polycarbonate; *polypropylene
Molded containers	*ABS; *polypropylene; *polycarbonate; polystyrene; polyvinyl chloride; vitreous graphite
Pipes	Polyvinyl chloride; *copper
Screws and bolts	*Steel 1020, *1040, *4140; acetal; *nylon
Shafts	*Steel 1020, *1040, *4140, *4340
Springs	*Steel 1080, *4140, 6250; stainless steel; beryllium copper; *nylon; titanium 6-4; maraging steel; phosphor bronze; acetal
Structural components	*Cast iron; *Douglas fir; *alumina; *aluminum 2024, *6061, *7075; *steel 1020, *4140
Storage boxes	Polystyrene*
Switches and wire jacketing	Fluoropolymers; thermoplastic polyester; modified polyphenylene oxide; *nylon; copper C11400
Telephone cases	*ABS
Toys (plastic)	*ABS; high- and low-density polyethylene; *polypropylene; polyvinyl chloride
Utensils	*Aluminum 3003; *polycarbonate; polyphenylene sulfide; polytetrafluoroethylene
Valves	Acetal; polyamide-imide; *alumina; *nylon; *aluminum; bronze

A.4 SOURCES

The following books have proven to be good sources of material property data.

Steels

American Society for Metals: *Properties and Selection: Stainless Steels, Tool Steels, and Special Purpose Metals*, vol. 3 of *Metals Handbook*, ASME, Cleveland, Ohio, 1979.

Harvey, P. D.: *Engineering Properties of Steels*, American Society for Metals, Metals Park, Ohio, 1982.

Metals Handbook, Vol. 1, *Properties and Selection: Iron, Steels and High Performance Alloys*, 10th ed., ASM, Cleveland, Ohio, 1990.

Peckner, D., and I. M. Bernstein: *Handbook of Stainless Steels*, McGraw-Hill, New York, 1977.

Other metals

Copper Development Association: *Standards Handbook: Copper-Brass-Bronze*, Copper Development Association, New York, 1988.

Mantell, C. L. (ed.): *Engineering Materials Handbook*, McGraw-Hill, New York, 1958.

Metals Handbook, Vol. 2, *Properties and Selection: Non-ferrous Alloys and Pure Metals*, ASM, Cleveland, Ohio, 1989.

Neale M. J. et al.: *Engineering Materials Selector*, Butterworth-Heinemann, 1998.

Ross, R. B.: *Metallic Materials Specification Handbook*, 4th edition, Chapman & Hall, London, 1992.

Plastics

Flick, E. W.: *Engineering Resins: An Industrial Guide*, Noyes, Park Ridge, N.J., 1988.

Harper, C. A.: *Handbook of Plastics and Elastomers and Composites*, 4th edition, McGraw-Hill, New York, 2002.

Juran, R.: *Modern Plastics Encyclopedia* 96, Vol. 72, No. 12, McGraw-Hill, New York, 1996.

Lubin, G.: *Handbook of Composites*, Van Nostrand Reinhold, New York, 1982.

Pethrick, R. A.: *Polymer Yearbook* 12, Routledge, New York, 1995.

Others

Ashby, M. F.: *Materials Selection in Mechanical Design*, 3rd edition, Butterworth Heinemann, 2005.

Avallone, E. A., and J. T. Baumeister: *Mark's Standard Handbook for Mechanical Engineers*, 11th edition, McGraw-Hill, 2006.

Bever, M. B. (ed.): *Encyclopedia of Materials Science and Engineering*, MIT Press, Cambridge, Mass., 1986.

Budinski, K. G.: *Engineering Materials, Properties and Selection*, 8th edition, Prentice-Hall, Englewood Cliffs, N.J., 2004.

Callister, W. D.: *Materials Science and Engineering*, 7th edition, Wiley, New York, 2006.

Ceramic Source, Vol. 4, American Ceramic Society, Columbus, Ohio, 1989.

Clark, A. F., and R. P. Reed: *Materials at Low Temperatures*, American Society for Metals, Metals Park, Ohio, 1983.

Gere, J. M., and S. P. Timoshenko: *Mechanics of Materials*, 3rd edition, PWS-Kent, Boston, 1990.

Horton, H. L., and E. Oberg: *Machinery's Handbook*, Vol. 25, Industrial Press, New York, 1996.

Kutz, M.: *Mechanical Engineer's Handbook*, Wiley-Interscience, New York, 1986.

Shaw, K.: *Refractories and Their Uses*, Wiley, New York, 1972.

Summitt, R., and A. Sliker: *Handbook of Material Science*, Vol. 4, CRC Press, 1980.

U.S. Forest Products Laboratory: *Wood Handbook: Wood as an Engineering Material*, <http://www.fpl.fs.fed.us/documnts/fplgtr/fplgtr113/fplgtr113.htm>

Normal Probability

B.1 INTRODUCTION

In this book, we consider all data distributions to be normal (Gaussian) distributions. This assumption is fairly accurate for most considerations in mechanical design.

To demonstrate the normal distribution, consider the data in Fig. B.1. This plot is a histogram of the ultimate strength as measured for 913 samples of 1035 steel. The data are plotted to the nearest 1 kpsi. The rounded values are given in this table:

Ultimate strength, kpsi	Number of occurrences	Ultimate strength, kpsi	Number of occurrences
75	4	86	87
76	10	87	93
77	6	88	66
78	19	89	66
79	21	90	67
80	24	91	39
81	46	92	21
82	57	93	24
83	74	94	10
84	85	95	<u>11</u>
85	83	Total	913

These data are but a sample of an entire population. The entire population is all the possible samples of 1035 steel. The object is to use the statistics of this sample to infer something about the statistics of the entire population.

The data plotted in Fig. B.1 are replotted in Fig. B.2 on normal-distribution paper. (Paper for this type of plot is commonly available.) Each ultimate strength value is plotted versus the percentage of the values that are less than it. Consider,

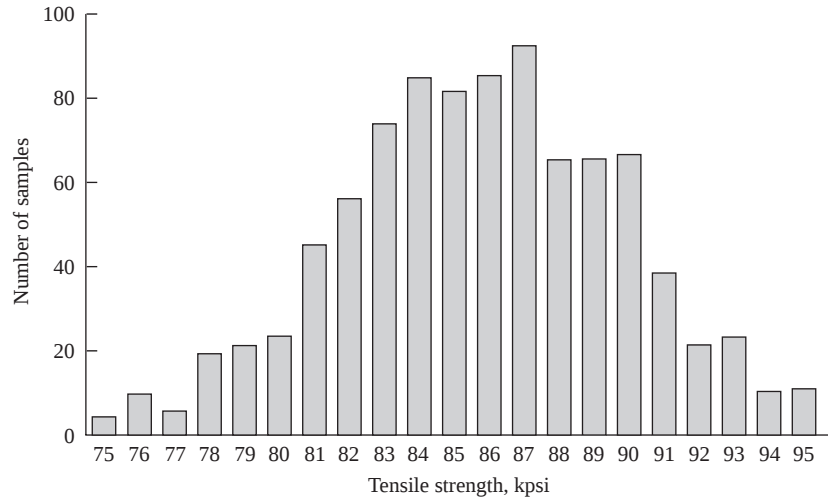


Figure B.1 Distribution of tensile strengths for samples of 1035 steel.

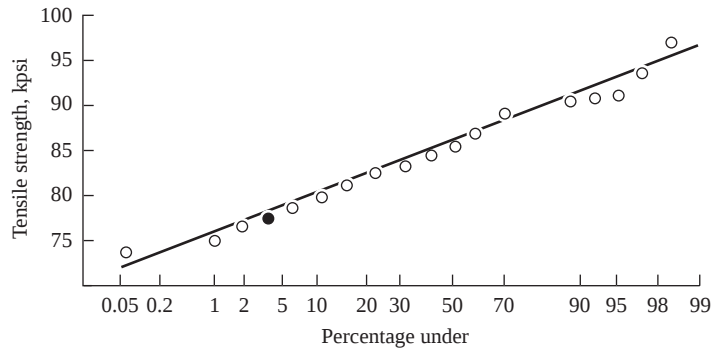


Figure B.2 Ultimate strength test values on normal-distribution paper.

for example, $S_u = 79$ kpsi: there are 39 values (4+10+6+19), or $39/913 = 4.3\%$ of the total samples, that are less than 79 kpsi. This is shown by the shaded point in Fig. B.2. The fact that the data are well fit by a straight line implies that their distribution can be modeled by a normal distribution.

Since any straight line can be characterized by two parameters, these data can be also. These two parameters are the *sample mean*, $\bar{x}$, and the *sample standard deviation*, σ , and they are defined as

$$\bar{x} = \frac{1}{N} \sum_{i=1}^N x_i$$

$$\sigma = \sqrt{\frac{1}{N-1} \sum_{i=1}^N (x_i - \bar{x})^2}$$

where the x_i are the data values S_u and N is the total number of data points (913). For the data here, $\bar{x} = 86$ kpsi and $\sigma = 4$ kpsi (both rounded to the accuracy of the input data). Another term that needs to be defined before we continue is the *sample variance*, which is the square of the standard deviation.

The three statistics just defined give information only about the sample. From the experimental data, much can be learned about the probability of finding the ultimate strength for a specific sample in a range between two values. For example, the probability of the ultimate strength being greater than or equal to 80 kpsi and less than 85 kpsi is

$$\Pr(80 \leq S_u < 85) = \frac{24 + 46 + 57 + 74 + 85}{913} = 31.3 \text{ percent}$$

Also, the probability of being within two standard deviations of the mean value ($86 - 2 \times 4 = 78$ and $86 + 2 \times 4 = 94$) is

$$\Pr(78 \leq S_u < 94) = 95.29\%$$

In both of these examples, the probability was given by

$$\Pr(a \leq x < b) = \frac{1}{N} \sum_{i=a}^{b-1} x_i$$

where the summation limits reflect the inequality or equality of the probability bound.

Leaving the sample data for a moment, we will develop the population normal distribution and then compare it with the sample data. The normal distribution is a continuous distribution, whereas the sample data were treated as discrete (they were rounded to the nearest kpsi). The distribution of the population is based on two parameters: the mean, μ , and the standard deviation, s , and is defined as

$$\Pr(a \leq x < b) = \int_a^b \frac{e^{-(1/2)[(x-\mu)/s]^2} dx}{\sqrt{2\pi}s}$$

The integration implied here is seldom performed, as the variable x is normalized by defining $x' = (x - \mu)/s$. This normalization transforms the mean to 0 and the standard deviation to 1 for the new variable x' . The equation for the normal distribution can then be rewritten as

$$\Pr(a \leq x < b) = \int_{(a-\mu)/s}^{(b-\mu)/s} \frac{e^{-(1/2)x'^2} dx'}{\sqrt{2\pi}}$$

and a table such as Table B.1 used to find the probability. For example, assume that the population mean and standard deviation are the same as in the sample for the 1035 steel, 86 and 4 kpsi. Obviously $\Pr(S_u \leq 86 \text{ kpsi})$ is 50%. This can be found by normalizing the value 86 by subtracting the mean and dividing by the standard deviation, which results in 0. The upper leftmost entry in the table,

Table B.1 Percentiles of the normal distribution: probability of obtaining a value less than or equal to x , where $x = (x - \mu)/s$

x'	.00	.01	.02	.03	.04	.05	.06	.07	.08	.09
0.0	.5000	.5040	.5080	.5120	.5160	.5199	.5239	.5279	.5319	.5359
0.1	.5398	.5438	.5478	.5517	.5557	.5596	.5636	.5675	.5714	.5753
0.2	.5793	.5832	.5871	.5910	.5948	.5987	.6026	.6064	.6103	.6141
0.3	.6179	.6217	.6255	.6293	.6331	.6368	.6406	.6443	.6480	.6517
0.4	.6554	.6591	.6628	.6664	.6700	.6736	.6772	.6808	.6844	.6879
0.5	.6915	.6950	.6985	.7019	.7054	.7088	.7123	.7157	.7190	.7224
0.6	.7257	.7291	.7324	.7357	.7389	.7422	.7454	.7486	.7517	.7549
0.7	.7579	.7611	.7642	.7673	.7704	.7734	.7764	.7794	.7823	.7852
0.8	.7881	.7910	.7939	.7967	.7995	.8023	.8051	.8078	.8106	.8133
0.9	.8159	.8186	.8212	.8238	.8264	.8289	.8315	.8340	.8365	.8389
1.0	.8413	.8438	.8461	.8485	.8508	.8531	.8554	.8577	.8599	.8621
1.1	.8643	.8665	.8686	.8708	.8729	.8749	.8770	.8790	.8810	.8830
1.2	.8849	.8869	.8888	.8907	.8925	.8944	.8962	.8980	.8997	.9015
1.3	.9032	.9049	.9066	.9082	.9099	.9115	.9131	.9146	.9162	.9177
1.4	.9192	.9207	.9222	.9236	.9251	.9265	.9279	.9292	.9306	.9319
1.5	.9332	.9345	.9356	.9370	.9382	.9394	.9406	.9418	.9429	.9441
1.6	.9452	.9463	.9474	.9484	.9495	.9505	.9515	.9525	.9535	.9545
1.7	.9554	.9564	.9573	.9582	.9591	.9599	.9608	.9616	.9625	.9633
1.8	.9641	.9649	.9656	.9664	.9671	.9678	.9689	.9693	.9699	.9706
1.9	.9713	.9719	.9726	.9732	.9738	.9744	.9750	.9756	.9761	.9767
2.0	.9772	.9778	.9783	.9788	.9793	.9798	.9803	.9808	.9811	.9816
2.1	.9820	.9825	.9829	.9833	.9837	.9841	.9845	.9849	.9853	.9856
2.2	.9860	.9863	.9868	.9870	.9874	.9877	.9880	.9883	.9886	.9889
2.3	.9882	.9895	.9897	.9900	.9903	.9905	.9907	.9910	.9912	.9915
2.4	.9917	.9920	.9921	.9924	.9926	.9928	.9929	.9931	.9933	.9935
2.5	.9937	.9938	.9940	.9941	.9943	.9944	.9946	.9947	.9949	.9950
2.6	.9951	.9953	.9954	.9955	.9957	.9958	.9959	.9960	.9961	.9962
2.7	.9963	.9964	.9965	.9966	.9967	.9968	.9969	.9970	.9971	.9971
2.8	.9972	.9973	.9974	.9974	.9975	.9976	.9976	.9977	.9977	.9978
2.9	.9979	.9979	.9980	.9981	.9981	.9982	.9982	.9983	.9983	.9984
3.0	.9987									

$x' = 0$, gives 0.5000 or 50%. If the probability of the ultimate strength being less than the mean plus 1 standard deviation is desired, then the value of x is 90 ($86 + 4$) and so $x = 1.000$, resulting in $\Pr(S_u < 90 \text{ kpsi}) = 84.13\%$, as marked in the table. Obviously, the probability of the ultimate strength being greater than 90 kpsi is $1 - 0.8413$, or $\Pr(S_u > 90 \text{ kpsi}) = 15.87\%$.

The probability of the ultimate strength being greater than 78 kpsi is a little more difficult to find. For $x = 78$, $x' < 2.0$. Since the distribution is symmetrical about the mean value, this can be treated as the same problem as finding the probability that $x' < 2.0$ and subtracting it from 1. From Table B.1, $\Pr(x' < 2) = 97.72\%$; therefore, $\Pr(S_u \leq 78 \text{ kpsi}) = 2.28\%$ ($1 - 97.72$).

The probability that S_u is between ± 2 standard deviations of the mean can be found by taking the probability that $x' \leq 2$, which is 0.9772, and

subtracting the probability that it is less than -2 , so $0.9772 - 0.0228 = 0.9544$, or $\Pr(78 \leq S_u < 94 \text{ kpsi}) = 95.44\%$. This compares well with the value found for the sample of 95.29. For ± 1 standard deviation, the result is 68.26%, and for three standard deviations it is 99.68%. This last value is what is generally assumed to be the limit on dimensional tolerances.

One last example of the use of normal distributions is about dimensions. Say a dimension on a drawing is given as 4.000 ± 0.008 cm. From a statistical viewpoint the dimension is the mean value of all the samples to be manufactured; the tolerance represents ± 3 standard deviations from the mean. This implies that 99.68% of all samples should have dimensions within the tolerance range. Modifying this example, consider the dimension $4.000 + 0.008 - 0.016$ cm. The nominal value is 4.000 cm; however, the mean value is 3.996 cm. Also, the standard deviation of the dimension is $[0.008 - (-0.016)]/6 = 0.004$ cm and the variance is $1.6 \times 10^{-5} \text{ cm}^2$. From normal distribution tables, these data enable the calculation of other statistics. For example, 68.3% of the samples should be between 3.992 and 4.000 cm (± 1 standard deviation), 84.1% should be less than 4.000 cm (less than the mean plus 1 standard deviation), and 95.44% should be between 3.988 and 4.004 cm (± 2 standard deviations).

B.2 OTHER MEASURES

Besides the mean and standard deviation, other measures for the normal distribution are often used. For example, statistics on human size or strength measurements are often given in terms of statistics for the 5th and 95th percentiles. The stature for a 5th-percentile man is 64.1 in., and that for a 95th-percentile man is 73.9 in. From these, the mean, 50th percentile, is 69 in. Additionally, since the 95th percentile is 1.645 standard deviations from the mean (see Table B.1) and the difference between the mean and 73.9 is 4.9, the standard deviation is $4.9/1.645 = 2.97$ in.

Finally, the statistics for the sum and difference of variables with normal distributions are easily found. If x_1 , x_2 , and x_3 are all normally distributed with means μ_1 , μ_2 , and μ_3 and standard deviations s_1 , s_2 , and s_3 and if $y = x_1 + x_2 - x_3$, then the mean value of y is $\bar{y} = \mu_1 + \mu_2 - \mu_3$ and $s_y = (s_1^2 + s_2^2 + s_3^2)^{1/2}$. Note that the mean values are just the sums and differences, whereas all the signs on the standard deviations are positive. These formulas are readily extended to any number of terms, as shown in Eqs. [10.2] and [10.3].

The Factor of Safety as a Design Variable

C.1 INTRODUCTION

The factor of safety is a factor of ignorance. If the stress on a part at a critical location (the applied stress) is known precisely, if the material's strength (the allowable strength) is also known with precision, and the allowable strength is greater than the applied stress, then the part will not fail. However, in the real world, all of the aspects of the design have some degree of uncertainty, and therefore a fudge factor, a factor of safety, is needed. A factor of safety is one way to account for the uncontrollable noises that were discussed in Chap. 10.

In practice the factor of safety is used in one of three ways: (1) It can be used to reduce the allowable strength, such as the yield or ultimate strength of the material, to a lower level for comparison with the applied stress; (2) it can be used to increase the applied stress for comparison with the allowable strength; or (3) it can be used as a comparison for the ratio of the allowable strength to the applied stress. We apply the third definition here, but all three are based on the simple formula

$$FS = \frac{S_{al}}{\sigma_{ap}}$$

Here S_{al} is the allowable strength, σ_{ap} is the applied stress, and FS is the factor of safety. If the material properties are known *precisely* and there is no variation in them—and the same holds for the load and geometry—then the part can be designed with a factor of safety of 1, the applied stress can be equal to the allowable strength, and the resulting design will not fail (just barely). However, not only are these measures not known with precision, they are not constant from sample to sample or use to use. In a statistical sense all these measures have some variance about their mean values (see App. B for the definitions of the mean and variance).

For example, typical material properties, such as ultimate strength, even when measured from the same bar of material, show a distribution of values (a variance) around a nominal mean of about 5%. This distribution is due to inconsistencies in the material itself and in the instrumentation used to take the data. If strength figures are taken from handbook values based on different samples and instrumentations, the variance of the values may be 15% or higher. Thus, the allowable strength must be characterized as a nominal or mean value with some statistical variation about it.

Even more difficult to establish are the statistics of the applied stress. The exact magnitude of the applied stress is a factor of the loading on the part (the forces and moments on the part), the geometry of the part at the critical location, and the accuracy of the analytic method used to determine the stress at the critical point due to the load.

The accuracy of the comparison of the applied stress to the allowable strength is a function of the accuracy and applicability of the failure theory used. If the stress is steady and the failure mode yielding, then accurate failure theories exist and can be used with little error. However, if the stress state is multiaxial and fluctuating (with a nonzero mean stress), there are no directly applicable failure theories and the error incurred in using the best available theory must be taken into account.

Beyond the preceding mechanical considerations, the factor of safety is also a function of the desired reliability for the design. As will be shown in Section C.3, the reliability can be directly linked to the factor of safety.

There are two ways to estimate the value of an acceptable factor of safety: the classical rule-of-thumb method (presented in Section C.2) and the probabilistic, or statistical, method of relating the factor of safety to the desired reliability and to knowledge of the material, loading, and geometric properties (presented in Section C.3).

An additional note on standards. Most established design disciplines and companies have factors of safety used as standards. But often these values are based on lost or outdated material specifications and quality control procedures. At a minimum the following tools will help explore the basis of these standards; at a maximum they can be used to update them.

For example, the Jet Propulsion Laboratory, in its design of the Mars Rover used the factors of safety shown in Table C.1. The factor of safety for both yield and ultimate and for metallic and ceramic materials is given. If the components have not been tested, the required factor of safety is much higher. Note that there is no value for composites yielding—they don’t.

Table C.1 Factors of safety for Mars Rover

	Qualified by testing		Not tested
	Metallic	Composite	
FS _{ultimate}	1.4	1.5	2.00
FS _{yield}	1.25	—	1.60

C.2 THE CLASSICAL RULE-OF-THUMB FACTOR OF SAFETY

The factor of safety can be quickly estimated on the basis of estimated variations of the five measures previously discussed: material properties, stress, geometry, failure analysis, and desired reliability. The better known the material properties and stress, the tighter the tolerances, the more accurate and applicable the failure theory, and the lower the required reliability, the closer the factor of safety should be to 1. The less known about the material, stress, failure analysis, and geometry and the higher the required reliability, the larger the factor of safety.

The simplest way to present this technique is to associate a value greater than 1 with each of the measures and define the factor of safety as the product of these five values:

$$FS = FS_{\text{material}} \cdot FS_{\text{stress}} \cdot FS_{\text{geometry}} \cdot FS_{\text{failure analysis}} \cdot FS_{\text{reliability}}$$

Details on how to estimate these five values are given next. These values have been developed by breaking down the rules given in textbooks and handbooks into the five measures and cross-checking the values with those from the statistical method described in Section C.3.

Estimating the Contribution for the Material

$FS_{\text{material}} = 1.0$	If the properties for the material are well known, if they have been experimentally obtained from tests on a specimen known to be identical to the component being designed and from tests representing the loading to be applied
$FS_{\text{material}} = 1.1$	If the material properties are known from a handbook or are manufacturer's values
$FS_{\text{material}} = 1.2\text{--}1.4$	If the material properties are not well known

Estimating the Contribution for the Load Stress

$FS_{\text{stress}} = 1.0\text{--}1.1$	If the load is well defined as static or fluctuating, if there are no anticipated overloads or shock loads, and if an accurate method of analyzing the stress has been used
$FS_{\text{stress}} = 1.2\text{--}1.3$	If the nature of the load is defined in an average manner, with overloads of 20–50%, and the stress analysis method may result in errors less than 50%
$FS_{\text{stress}} = 1.4\text{--}1.7$	If the load is not well known or the stress analysis method is of doubtful accuracy

Estimating the Contribution for Geometry (Unit-to-Unit)

$FS_{\text{geometry}} = 1.0$	If the manufacturing tolerances are tight and held well
$FS_{\text{geometry}} = 1.0$	If the manufacturing tolerances are average
$FS_{\text{geometry}} = 1.1\text{--}1.2$	If the dimensions are not closely held

Estimating the Contribution for Failure Analysis

$FS_{\text{failure theory}} = 1.0\text{--}1.1$	If the failure analysis to be used is derived for the state of stress, as for uniaxial or multiaxial static stresses, or fully reversed uniaxial fatigue stresses
$FS_{\text{failure theory}} = 1.2$	If the failure analysis to be used is a simple extension of the preceding theories, such as for multiaxial, fully reversed fatigue stresses or uniaxial nonzero mean fatigue stresses
$FS_{\text{failure theory}} = 1.3\text{--}1.5$	If the failure analysis is not well developed, as with cumulative damage or multiaxial nonzero mean fatigue stresses

Estimating the Contribution for Reliability

$FS_{\text{reliability}} = 1.1$	If the reliability for the part need not be high, for instance, less than 90%
$FS_{\text{reliability}} = 1.2\text{--}1.3$	If the reliability is an average of 92–98%
$FS_{\text{reliability}} = 1.4\text{--}1.6$	If the reliability must be high, say, greater than 99%

These values are, at best, estimates based on a verbalization of the factors affecting the design combined and on experience with how these factors affect the design. The stress on a part is fairly insensitive to tolerance variances unless they are abnormally large. This insensitivity will be more evident in the development of the statistical factor of safety.

C.3 THE STATISTICAL, RELIABILITY-BASED, FACTOR OF SAFETY

C.3.1 Introduction

As can be appreciated, the classical approach to establishing factors of safety is not very precise and the tendency is to use it very conservatively. This results in large factors of safety and overdesigned components. Consider now the approach based on statistical measures of the material properties, of the stress developed in the component, of the applicability of the failure theory, and of the reliability required. This technique gives the designer a better feel for just how conservative, or nonconservative, he or she is being.

With this technique, all measures are assumed to have normal distributions (details on normal distributions appear in App. B). This assumption is a reasonable one, though not as accurate as the Weibull distribution for representing material-fatigue properties. What makes the normal distribution an acceptable representation of all the measures is the simple fact that, for most of them, not enough data are available to warrant anything more sophisticated. In addition, the normal distribution is easy to understand and work with. In upcoming sections,

each measure is discussed in terms of the two factors needed to characterize a normal distribution—the mean and the standard deviation (or variance).

The factor of safety is defined as the ratio of the allowable strength, S_{al} , to the applied stress, σ_{ap} . The allowable strength is a measure of the material properties; the applied stress is a measure of the stress (as a function of both the applied load and the stress analysis technique used to find the stress), the geometry, and the failure theory used. Since both of these measures are distributions, the factor of safety is better defined as the ratio of their mean values: $FS = \bar{S}_{al}/\bar{\sigma}_{ap}$.

Figure C.1 shows the distribution about the mean for both the applied stress and the allowable strength. There is an area of overlap between these two curves no matter how large the factor of safety is and no matter how far apart the mean values are. This area of overlap is where the allowable strength has a probability of being smaller than the applied stress; the area of overlap is thus the region of potential failure. Keeping in mind that areas under normal distribution curves represent probabilities, we see that this area of overlap then is the Probability of Failure (PF). The reliability, the probability of no failure, is simply $1 - PF$. Thus, by considering the statistical nature of these curves, the factor of safety is directly related to the reliability.

To develop this relationship more formally, we define a new variable, $z = S_{al} - \sigma_{ap}$, the difference between the allowable strength and the applied stress. If $z > 0$, then the part will not fail. But failure will occur at $z \leq 0$. The distribution of z is also normal (the difference between two normal distributions is also normal), as shown in Fig. C.2. The mean value of z is simply $\bar{z} = \bar{S}_{al} - \bar{\sigma}_{ap}$. If the allowable strength and the applied stress are considered as independent variables (which is the case), then the standard deviation of z is

$$\rho_z = \sqrt{\rho_{al}^2 + \rho_{ap}^2}$$

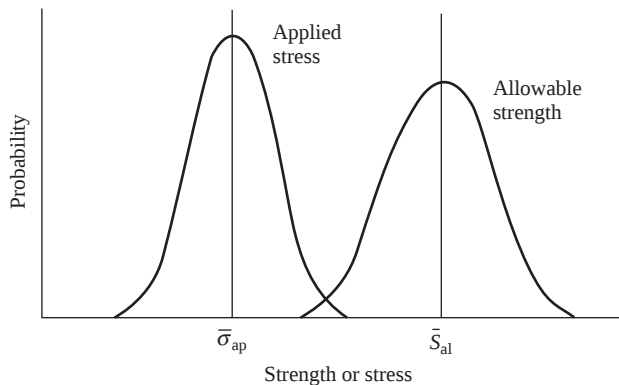


Figure C.1 Distribution of applied stress and allowable strength.

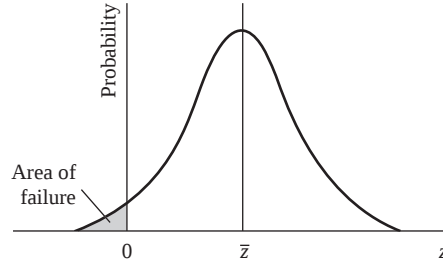


Figure C.2 Distribution of z .

Normalizing any value of z by subtracting the mean and dividing by the standard deviation, we can define the variable t_z as

$$t_z = \frac{(S_{a1} - \sigma_{ap}) - (\bar{S}_{a1} - \bar{\sigma}_{ap})}{\sqrt{\rho_{a1}^2 + \rho_{ap}^2}}$$

The variable t_z has a mean value of 0 and a standard deviation of 1. Since failure will occur when the applied stress is greater than the allowable stress, a critical point to consider is when $z = 0$, $S_{a1} = \sigma_{ap}$. So, for $z = 0$,

$$t_{z=0} = \frac{-(\bar{S}_{a1} - \bar{\sigma}_{ap})}{\sqrt{\rho_{a1}^2 + \rho_{ap}^2}}$$

Thus, any value of t that is calculated to be less than $t_{z=0}$ represents a failure situation. The probability of a failure then is $\Pr(t < t_{z=0})$, which, assuming the normal distribution, can be found directly from a normal distribution table. If the distributions of the applied stress and the allowable strength are known, $t_{z=0}$ can be found from the preceding equation and the probability of failure can be found from normal distribution tables. Finally, the reliability is 1 minus the probability of failure; $R = 1 - \Pr(t_z \leq t_{z=0})$. To make using normal distribution tables (App. B) easier by utilizing the symmetry of the distribution, we can drop the minus sign on the preceding equation and consider values of $t_z > t_{z=0}$ to represent failure. Some values showing the relation of $t_{z=0}$ to reliability are given in Table C.2.

To reduce the equations to a usable form in which the factor of safety is the independent variable, we rewrite the previous equation, dividing by the mean value of the applied stress and using the definition of the factor of safety:

$$t_{z=0} = \frac{FS - 1}{\sqrt{FS^2(\rho_{a1}/\bar{S}_{a1})^2 + (\rho_{ap}/\bar{\sigma}_{ap})^2}}$$

With $t_{z=0}$ directly dependent on the reliability, there are four variables related by this equation: the reliability, the factor of safety, and the coefficients of variation (standard deviation divided by the mean) for the allowable and applied stresses.

Table C.2 Relation of $t_{z=0}$ to R

R	$t_{z=0}$
0.50	0.00
0.90	1.28
0.95	1.64
0.99	2.33

In the development here, the unknown will be the factor of safety. Thus, the final form of the statistical factor of safety equation is

$$FS = 1 + t_{z=0} \frac{\sqrt{(\rho_{a1}/\bar{S}_{a1})^2 + (\rho_{ap}/\bar{\sigma}_{ap})^2 - t_{z=0}^2(\rho_{ap}/\bar{\sigma}_{ap})^2(\rho_{a1}/\bar{S}_{a1})^2}}{1 - t_{z=0}^2(\rho_{a1}/\bar{S}_{a1})^2} \quad [C.1]$$

Before proceeding with details into the development of the applied stress and allowable strength coefficients of variation, let us look at an example of the use of the preceding equations. Say that the allowable strength coefficient of variation (see Sec. C.3.2) is 0.08 (the standard deviation is 8% of the mean value), the applied stress coefficient (see Sec. C.3.3) is 0.20 and the desired reliability is 95%. Using Table C.2, a 95% reliability gives $t_{z=0} = 1.64$. Thus, using Eq. [C.1], the design factor of safety can be computed to be 1.37. If the reliability is increased to 99%, the design factor of safety increases to 1.55. These design factor of safety values are not dependent on the actual values of the material properties or the stresses in the material but only on their statistics and the reliability and applicability of the failure theory. This is a very important point.

C.3.2 The Allowable Strength Coefficient of Variation

Measured material properties, such as the yield strength, the ultimate strength, the endurance strength, and the modulus of elasticity, all have distributions about their means. This is evident in Fig. B.1, which shows the result of static tests on 913 different samples of 1035 steel as hot-rolled, round bars, 1 to 9 in. diameter. Although not perfectly normal, as shown by the fit of the data on the normal distribution paper, an approximation to the straight line is not bad. Not all kinds of data fit this well. Typically, fatigue data and data on ceramic material tend not to be as evenly distributed and are better represented by a skewed distribution, such as the Weibull distribution. (Unfortunately, the four factors needed to represent the Weibull distribution have only been determined for a limited number of materials.) However, the adequacy and simplicity of the normal distribution make it the best choice for representing the material properties here. From Fig. B.2, the mean ultimate strength is 86 kpsi and its standard deviation is 4 kpsi. Note that the standard deviation is 4.6% of the mean, so the coefficient of variation for 1035 hot-rolled steel is 0.046. Unfortunately, it is not always simple to find the standard deviation for the material properties. In looking up the ultimate stress for 1035 HR in standard design books, the following values were found: 72, 85, 72, 82, and 67 kpsi. From this limited sample, the mean value is 75.6 kpsi with a standard

deviation of 6.08 kpsi, resulting in a coefficient of variation of 0.80 (6.08/75.6). If the heat treatment on the material or the exact composition of the material are unknown, the deviation can be much higher.

The allowable stress may be based on the yield, ultimate, or endurance strengths or some combination of them, depending on the failure criteria used. In the preceding formulation, the allowable stress appears only as a ratio of standard deviation to average value. For most materials this ratio, irrespective of which allowable stress is considered, is in the range of 0.05–0.15. It is recommended that the statistics for the strength that best represent the nature of the failure be used. For example, in cases of nonzero-mean fluctuating stresses, the allowable stress coefficient for the endurance limit should be used if the mean is small relative to the amplitude, and the coefficient for the ultimate should be used if the mean is large relative to the amplitude. Any complexity beyond this is not warranted.

C.3.3 The Applied Stress Coefficient of Variation

The applied stress coefficient of variation is somewhat more difficult to develop. The statistics of the geometry and the load obviously affect the statistics of the applied stress, as does the accuracy of the method used to find the stress. Additionally, a measure of the accuracy of the failure analysis method to be used will also be reflected in the statistics of the applied stress. (This measure could be taken into account elsewhere, but it is convenient to consider it as a correction on the applied stress.)

To see how these various factors are combined to form the applied stress coefficient of variation, consider the following example. (The statistics for the stress analysis technique and the failure theory accuracy will be included later in the example.) Consider a round, axially loaded uniform bar. In this bar, the average maximum stress is given by the ratio of the average maximum force divided by the average area:

$$\bar{\sigma}_{\text{ap}} = \frac{\bar{F}}{\pi \bar{r}^2}$$

The standard deviation of the stress is a function of the independent statistics of the geometry and the load. Using standard, normal-distribution relations,

$$\frac{\rho_{\text{ap}}}{\bar{\sigma}_{\text{ap}}} = \sqrt{4 \left(\frac{\rho_r}{r} \right)^2 + \left(\frac{\rho_F}{\bar{F}} \right)^2}$$

Thus, the applied-stress coefficient of variation is written in terms of the coefficients of variation for the geometry ($\rho_r/\bar{r}$) and the loading ($\rho_F/\bar{F}$). In general, the same form can be derived for any loading and shape. No matter whether it is normal or shear, the stress will have the form of force/area, and area always has units of length squared.

In many applications, the magnitude of load forces and moments are well known through either experience or measurement. Essentially, two types of loads

are considered here: static loads and fatigue, or fluctuating, loads. Regardless of which type of loading is considered, the exact magnitude of the forces and moments may have to be estimated. The determination of the statistical factor of safety takes into account the confidence in this estimation. This approach is much like that used in project planning (PERT) and requires the designer to make three estimates of the load: an optimistic estimate o ; a most likely estimate m ; and a pessimistic estimate p . From these three the mean $\bar{m}$, standard deviation ρ , and coefficient of variation can be found:

$$\bar{m} = \frac{1}{6}(o + 4m + p)$$

$$\rho = \frac{1}{6}(p - o)$$

$$\frac{\rho}{\bar{m}} = \frac{p - o}{o + 4m + p}$$

These equations are based on a beta distribution function rather than a normal distribution. However, if the most likely estimate is the mean load, and the optimistic and pessimistic estimates are the mean ± 3 standard deviations, then the beta distribution reduces to the normal distribution. The beauty of this is that an estimate of the important statistics can be made even if the distribution of the estimates is not symmetrical. For example, suppose the maximum load on a bracket is quoted as a force of 25,000 N. This may just be the most likely estimate. There is a possibility that the maximum load may be as low as 15,000 N or, because of light shock loading, the force may be as high as 50,000 N. Thus, from the preceding formulas, the expected value is 27,500 N, the standard deviation is 5833 N, and the coefficient of variation is 0.21. If the optimistic load had been 0, no load at all, the expected load would be 25,000 N and the standard deviation, 8333 N. In this case, the pessimistic and optimistic estimates are ± 3 standard deviations from the expected or mean value. The coefficient of variation is 0.33, reflecting the wider range of estimates. Note again that the load coefficient of variation is independent of the absolute value of the load itself and gives only information on its distribution.

The hardest factor to take into account in failure analysis is the effect of shock loads. In the example just given, the potential maximum load was double the nominal value. Without dynamic modeling there is no way to find the effect of shock loads on the state of stress. These choices are suggested:

- If the load is smoothly applied and released, use a ratio of optimistic:most likely:pessimistic of 1:1:2 or 1:2:4.
- If the load gives moderated shocks, use a ratio of optimistic:most likely:pessimistic of 1:1:4 or 1:4:16. Examples of moderate shock applications are blowers, cranes, reels, and calenders.
- If the load gives heavy shocks, use a ratio of optimistic:most likely:pessimistic of 1:1:10 or 1:10:100. Examples of heavy shock applications are crushers, reciprocating machinery, and mixers.

The geometry of the part is important in that, in combination with the load, the geometry determines the applied stress. Normally, the geometry is given as nominal dimensions with a bilateral tolerance (3.084 ± 0.010 in.). The nominal is the mean value, and the tolerance is usually considered to be three times the standard deviation. This implies that, assuming a normal distribution, 99.74% of all the samples will be within the limits of the tolerance. It is assumed that there is one dimension that is most critical to the stress, and the coefficient of variation for this dimension is used in the analysis. For the example just considered, the coefficient of variation is 0.0011 [$0.010/(3 \cdot 3.084)$], which is an order of magnitude smaller than that for the load. This is typical for most tolerances and loadings.

Using the discussed examples, the applied stress coefficient of variation is

$$\rho_{\text{ap}} = \sqrt{4(0.0011)^2 + 0.21^2} \doteq 0.21$$

Note the lack of sensitivity to the tolerance.

The preceding does not take into account the accuracy of the stress analysis technique used to find the stress state from the loading and geometry or the adequacy of the failure analysis method. To include these factors, the allowable strength needs to be compared with the calculated applied stress corrected for the stress analysis and the failure analysis accuracy. Thus,

$$\sigma_{\text{ap}} = \sigma_{\text{calc}} \times N_{\text{sa}} \times N_{\text{fa}}$$

where N_{sa} is a correction multiplier for the accuracy of the stress analysis technique and N_{fa} is a correction multiplier for the failure analysis accuracy.

If the two corrections are assumed to have normal distributions, they can be represented as coefficients of variation. With the product of normally distributed independent standard deviations being the square root of the sum of the squares (see App. B), we have for the applied stress coefficient of variation

$$\frac{\rho_{\text{ap}}}{\bar{\sigma}_{\text{ap}}} \sqrt{4 \left(\frac{\rho_r}{\bar{r}} \right)^2 + \left(\frac{\rho F}{\bar{F}} \right)^2 + \left(\frac{\rho_{\text{sa}}}{N_{\text{sa}}} \right)^2 + \left(\frac{\rho_{\text{fa}}}{N_{\text{fa}}} \right)^2} \quad [\text{C.2}]$$

This is the same as before, with the addition of the coefficient of variation s for the stress analysis method and for the failure theory.

The coefficient of variation for the stress analysis method can be estimated using the same technique as for estimating the statistics on the loading—namely, estimate an optimistic, pessimistic, and most likely value for the stress, based on the most likely load. Again, consider a load of 25,000 N (the most likely estimate of the maximum load). Assume that at the critical point the normal stress caused by this load is 40.9 kpsi (282.0 MPa), with a stress concentration factor of 3.55. The most likely normal stress is the product of the load and the stress concentration factor, 145 kpsi. However, confidence in the method used to find the nominal stress and the stress concentration factor is not high. In fact, the maximum stress may really be as high as 160 kpsi or as low as 140 kpsi. With these two values as the pessimistic and the optimistic estimates, the coefficient of variation is calculated at 0.023. For strain gauge data or other measured results,

the stress analysis method coefficient of variation will be very small and can, like the geometry statistics, be ignored.

The adequacy of the failure analysis technique, as discussed in the development of the classical factor of safety method, has a marked effect on the design factor of safety. On the basis of experience and the limited data in the references, the coefficient of variations recommended for the different types of loadings are

Static failure theories: 0.02

Fully reversed uniaxial infinite life fatigue failure theory: 0.02

Fully reversed uniaxial finite life fatigue failure theory: 0.05

Nonzero mean uniaxial fatigue failure theory: 0.10

Fully reversed multiaxial fatigue failure theory: 0.20

Nonzero mean multiaxial fatigue failure theory: 0.25

Cumulative damage load history: 0.50

These values imply that for well-defined failure analysis techniques, where the failure mode is identical to that found with the allowable strength material test, the standard deviation is small, namely, 2% of the mean. When the failure theory is comparing a dissimilar applied stress state to an allowable strength, the margin for error increases. The rule used in cumulative damage failure estimation can be off by as much as a factor of 2 and is therefore used with high uncertainty.

C.3.4 Steps for Finding the Reliability-Based Factor of Safety

We can summarize the method discussed in the previous two sections as an eight-step procedure:

Step 1: Select reliability. From Table C.2, find the value of $t_{z=0}$ for the desired reliability.

Step 2: Find the allowable strength coefficient of variation. This can be found experimentally or by following the following rules of thumb. If the material properties are well known, use a coefficient of 0.05; if the material properties are not well known, use a coefficient of 0.01–0.15.

Step 3: Find the critical dimension coefficient of variation. This value is generally small and can be ignored except when the variation in the critical dimensions are large because of manufacturing, environmental, or aging effects.

Step 4: Find the load coefficient of variation. This is an estimate of how well the maximum loading is known. It can be estimated using the PERT method given in Sec. C.3.3.

Step 5: Find the accuracy of the stress analysis coefficient of variation. Even though the variation of the load was taken into account in step 4, knowledge about the effect of the load on the structure is a separate issue. The stress due to a well-known load may be hard to determine because of complex

geometry. Conversely, the stress caused by a poorly known load on a simple structure is no more poorly known than the load itself. This measure then takes into account how well the stress can be found for a known load.

Step 6: Find the failure analysis technique coefficient of variation. Guidance for this is given near the end of Sec. C.3.2.

Step 7: Calculate the applied stress coefficient of variation. This is found using Eq. [C.2].

Step 8: Calculate the factor of safety. This is found using Eq. [C.1].

C.4 SOURCES

Ullman, D. G.: "Less Fudging with Fudge Factors," *Machine Design*, Oct. 9, 1986, pp. 107–111.

Ullman, D. G.: *Mechanical Design Failure Analysis*, Marcel Dekker, New York, 1986.

A P P E N D I X

Human Factors in Design

D.1 INTRODUCTION

Most machines work in coordination with people. Consider the types of interactions you have with a standard gas-powered lawn mower. First, in starting and pushing the mower you *occupy a workspace* around the mower. You have to stoop or bend in this space to reach the starting mechanism, then you have to position yourself while holding your arms at a certain height to push and steer the mower. Second, you *provide a source of power* to the mower to start it and to push it. (Even if it is electrically started, you have to push a button or turn a key.) Additionally, it takes muscle power to steer the mower, whether you are walking behind it or riding on it. Third, you *act as a sensor*, listening to determine if anything is stuck in the mower, seeing where you are going so that you can guide the mower, and feeling with your hands any feedback motion through the steering that might give you information on how well you are guiding the mower. Fourth, based on the information received by the sensory inputs, you *act as a controller*. You determine how much power to provide and in what direction to keep the mower under control.

These four ways a person interacts with the product—as occupant of workspace, as power source, as sensor, and as controller—form the basis for the study of the *human factors* that play a major role in the design of a device. Beyond these four basic types of interactions between person and product, there are further human interaction issues that must be considered during design. First, even those devices that spend their operating life remote from all human interaction, at the bottom of a well or in deep space, for example, must first be assembled. The assembler must interface with the device in the same four ways as described in the lawn-mower example. Second, most devices have to be maintained, which presents yet another situation for the consideration of human interaction in the design of a product. *Human factors must be taken into account for every person who comes into contact with the product, whether during manufacture, operation, maintenance and repair, or disposal.*

Two reasons for this concern with human factors are quality and safety. In Table 1.1 we saw that in determining quality in a product, the most important factor was perceived to be that it “work as it should.” This design requirement relates directly to the four components of human-product interface. Products are perceived to work as they should if they are comfortable to use (there is a good match between the device and the person in the workspace), they are easy to use (minimal power is required), their operating condition is easily sensed, and their control logic is natural, or user-friendly. Of equal importance is the concern for safety. Although not listed as one of the factors in the survey, it is readily assumed that an unsafe design will never be perceived as a quality product. Customers assume that neither they nor others will be injured, and that no property will be destroyed, when a product is in use (obvious exceptions are products that are designed to destroy or injure).

In Sections D.2 through D.4 all of these issues will be further explored, with emphasis on understanding the interactions between humans and machines in order to ensure that quality and safety are designed into the product.

D.2 THE HUMAN IN THE WORKSPACE

It is vital that a product “fit” its intended user; in other words, it must be comfortable for a person to use. The lawn-mower pull starter must be at the right height, or it will be hard to reach and even harder to pull. Likewise, the handle on a push mower must be at a height that is comfortable for a majority of users.

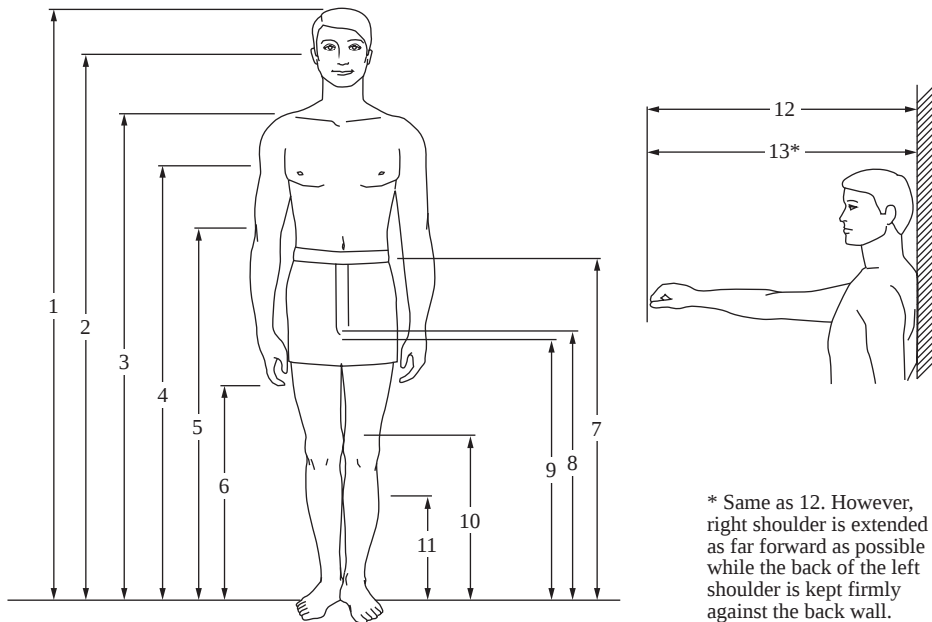


Figure D.1 Anthropometric man (from MIL-STD 1472D).

people or the mower will judged to be of poor quality. The geometric properties of humans—their height, reach, and seating requirements, the size of the holes they can fit through, and so forth—are called *anthropometric data* (literally, “human-measure data”). Many such data have been collected by the armed forces because so many different people must operate military equipment on a day-to-day basis. Typical anthropometric data given in MIL-STD (Military Standard) 1472F are shown in Fig. D.1 and Table D.1. Since people come in a variety of shapes and sizes, it is important that anthropometric data give a range of dimensions. The measures of humans are well represented as normal distributions. (See App. B for details on normal distributions.) Typically, these measures are given for the

Table D.1 Anthropometric data (from MIL-STD 1472F)

	Percentile values in centimeters					
	5th percentile			95th percentile		
	Ground troops	Aviators	Women	Ground troops	Aviators	Women
Weight (kg)	65.5	60.4	46.6	91.6	96.0	74.5
Standing body dimensions						
1. Stature	162.8	164.2	152.4	185.6	187.7	174.1
2. Eye height (standing)	151.1	152.1	140.9	173.3	175.2	162.2
3. Shoulder (acromial) height	133.6	133.3	123.0	154.2	154.8	143.7
4. Chest (nipple) height*	117.9	120.8	109.3	136.5	138.5	127.6
5. Elbow (radiate) height	101.0	104.8	94.9	117.8	120.0	110.7
6. Fingertip (dactylion) height		61.5			73.2	
7. Waist height	96.6	97.6	93.1	115.2	115.1	110.3
8. Crotch height	76.3	74.7	68.1	91.8	92.0	83.9
9. Gluteal furrow height	73.3	74.6	66.4	87.7	88.1	81.0
10. Kneecap height	47.5	46.8	43.8	58.6	57.8	52.5
11. Calf height	31.1	30.9	29.0	40.6	39.3	36.6
12. Functional reach	72.6	73.1	64.0	90.9	87.0	80.4
13. Functional reach, extended	84.2	82.3	73.5	101.2	97.3	92.7
	Percentile values in inches					
Weight (lb)	122.4	133.1	102.3	201.9	211.6	164.3
Standing body dimensions						
1. Stature	64.1	64.6	60.0	73.1	73.9	68.5
2. Eye height (standing)	59.5	59.9	55.5	68.2	69.0	63.9
3. Shoulder (acromial) height	52.6	52.5	48.4	60.7	60.9	56.6
4. Chest (nipple) height*	46.4	47.5	43.0	53.7	54.5	50.3
5. Elbow (radiate) height	39.8	41.3	37.4	46.4	47.2	43.6
6. Fingertip (dactylion) height		24.2			28.8	
7. Waist height	38.0	38.4	36.6	45.3	45.3	43.4
8. Crotch height	30.0	29.4	26.8	36.1	36.2	33.0
9. Gluteal furrow height	28.8	29.4	26.2	34.5	34.7	31.9
10. Kneecap height	18.7	18.4	17.2	23.1	22.8	20.7
11. Calf height	12.2	12.2	11.4	16.0	15.5	14.4
12. Functional reach	28.6	28.8	25.2	35.8	34.3	31.7
13. Functional reach, extended	33.2	32.4	28.9	39.8	38.3	36.5

*Bustpoint height for women.

5th and 95th percentile, as in Table D.1. It is safe to assume that data for civilians do not differ significantly from that for military personnel. Similar data are available for children. (See Sources at the end of this appendix.)

Besides measures for humans standing with their arms at their sides, measures for humans performing various activities are also available. For example, in Fig. D.2, an anthropometric woman is shown standing at a control panel. This

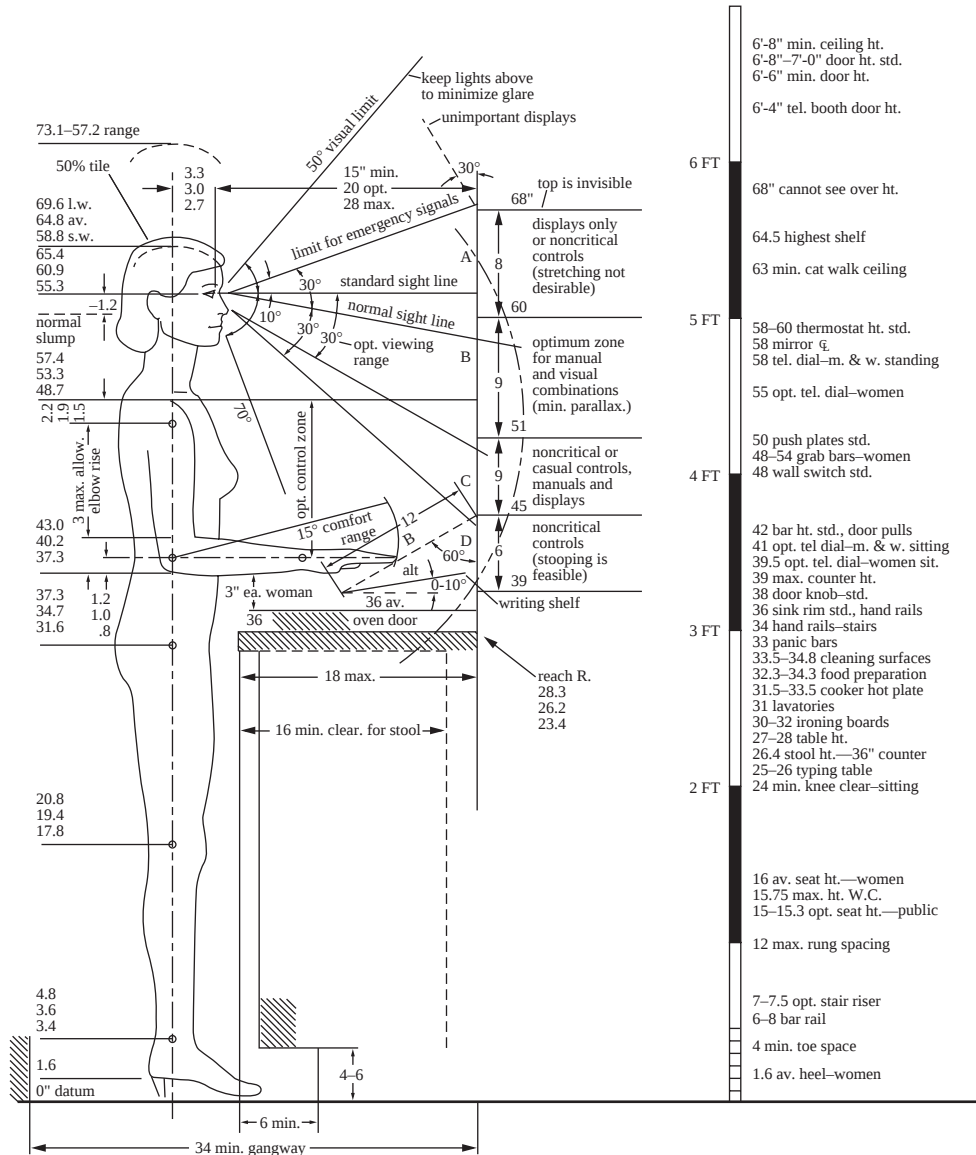


Figure D.2 Anthropometric woman at control panel (Source: Adapted from H. Dreyfuss, *The Measure of Man: Human Factors in Design*, Whitney Library of Design, New York, 1967).

drawing is of a 50th-percentile woman. The dimensions on the control panel are such that a majority of women will feel comfortable looking at the displays and working the controls.

Returning to our lawn mower: the handle should be at about elbow level, height 5 in Fig. D.1 and Table D.1. To fit all men and women between the 5th and 95th percentiles, the handle must be adjustable between 94.9 cm (37.4 in.) for the 5th-percentile woman and 117.8 cm (46.4 in.) for the 95th-percentile man. Anthropometric data from the references also show that the pull starter should be 69 cm (27 in.) off the ground for the average person. For this uncommon position, only an average value is given in the references. For positions even more unique, the engineer may have to develop measurements of a typical user community in order to get the data necessary for quality products.

D.3 THE HUMAN AS SOURCE OF POWER

Humans often have to supply some force to power a product or actuate its controls. The lawn-mower operator must pull on the starter cord and push on the handle or move the steering wheel. Human force-generation data are often included with anthropometric data. This information comes from the study of *biomechanics* (the mechanics of the human body). Listed in Fig. D.3 is the average human strength for differing body positions. In the data for “arm forces standing,” we find that the average pushing force 40 in. off the ground (the average height of the mower handle) is 73 lb, with a note that hand forces of greater than 30 to 40 lb are fatiguing. Although only averages, these values do give some indication of the maximum forces that should be used as design requirements. More detailed information on biomechanics is available in MIL-HDBK (Military Handbook) 759A and *The Human Body in Equipment Design* (see Sources at the end of this appendix).

D.4 THE HUMAN AS SENSOR AND CONTROLLER

Most interfaces between humans and machines require that humans *sense* the state of the device and, based on the data received, *control* it. Thus, products must be designed with important features readily apparent, and they must provide for easy control of these features. Consider the control panel from a clothes dryer (Fig. D.4). The panel has three controls, each of which is intended both to actuate the features and to relate the settings to the person using the dryer. On the left are two toggle switches. The top switch is a three-position switch that controls the temperature setting to either “Low,” “Permanent Press,” or “High.” The bottom switch is a two-position switch that is automatically toggled to off at the end of the cycle or when the dryer door is opened. This switch must be pushed to start the dryer. The dial on the right controls the time for either the no-heat cycle (air dry) on the top half of the dial or the heated cycle on the bottom half.

The dryer controls must communicate two functions to the human: temperature setting and time. Unfortunately, the temperature settings on this panel are

HUMAN STRENGTH

(for short durations)

strength correction factors;

X 0.9 left hand and arm

X 0.84 hand – age 60

X 0.5 arm & leg – age 60

X 0.72 women

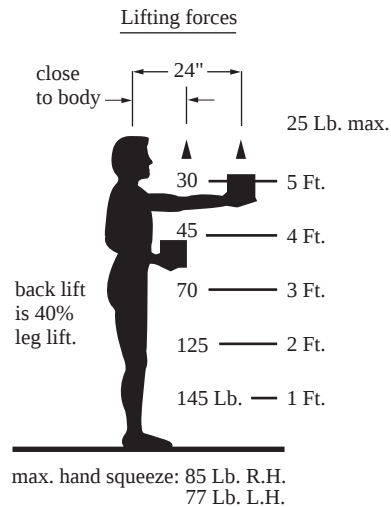
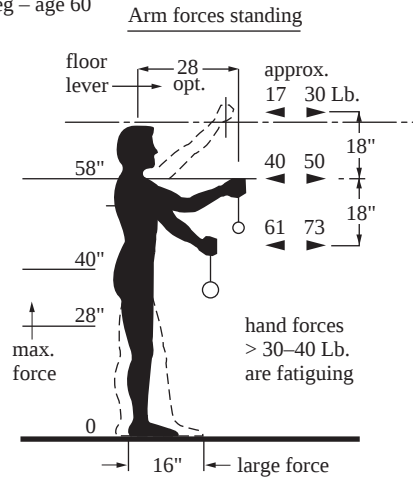
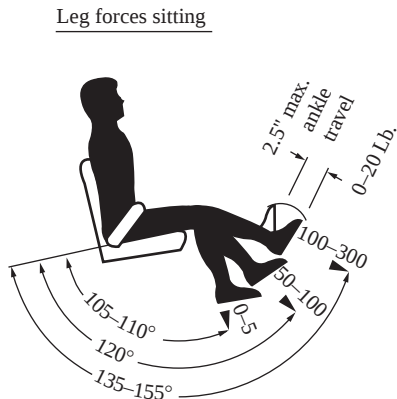
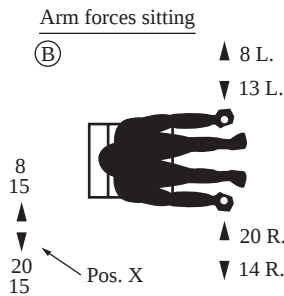
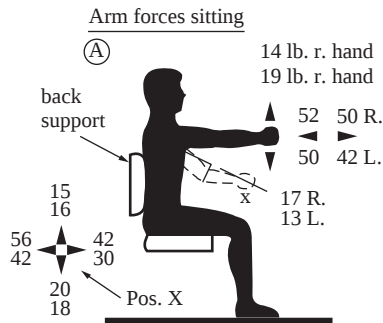


Figure D.3 Average human strength for different tasks (Source: Adapted from H. Dreyfuss, *The Measure of Man: Human Factors in Design*, Whitney Library of Design, New York, 1967).

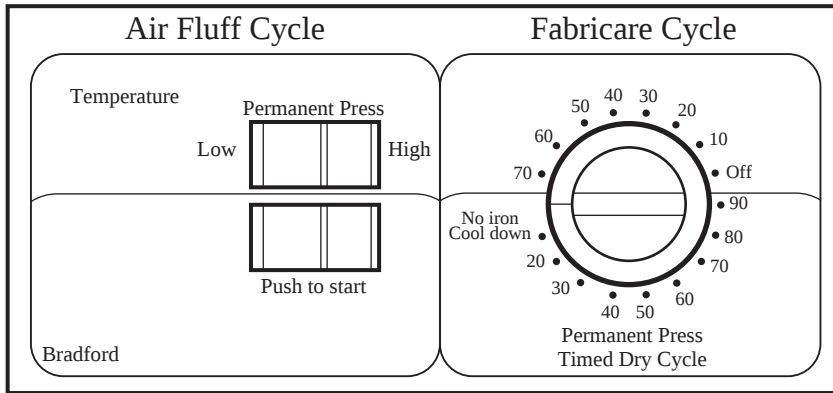


Figure D.4 Clothes dryer control panel (Source: Adapted from J. H. Burgess, *Designing for Humans: The Human Factor in Engineering*, Petrocelli Books, Princeton, N.J., 1986).

hard to sense because the “Temperature” rocker switch does not clearly indicate the status of the setting and the air-dry setting for temperature is on the dial that can override the setting of the “Temperature” switch. There are two communication problems in the time setting also: the difference between the top half of the dial and the bottom half is not clear and the time scale is the reverse of the traditional clockwise dial. The user must not only sense the time and temperature but must regulate them through the controls. Additionally, there must be a control to turn the dryer on. For this dryer, the rocker switch does not appear to be the best choice for this function. Finally, the labeling is confusing.

This control panel is typical of many that are seen every day. The user can figure out what to do and what information is available, but it takes some conjecturing. The more guessing required to understand the information and to control the action of the product, the lower the perceived quality of the product. If the controls and labeling were as unclear on a fire extinguisher, for example, it would be all but useless—and therefore dangerous. There are many ways to communicate the status of a product to a human. Usually the communication is visual; however, it can also be through tactile or audible signals. The basic types of visual displays are shown in Fig. D.5. When choosing which of these displays to use, it is important to consider the type of information that needs to be communicated. Figure D.6 relates five different types of information to the types of displays.

Comparing the clothes-dryer control panel of Fig. D.4 to the information of Fig. D.6, the temperature controls require only discrete settings and the time control a continuous (but not accurate) numerical value. Since toggle switches are not very good at displaying information, the top switch on the panel of Fig. D.4 should be replaced by any of the displays recommended for discrete information. The use of the dial to communicate the time setting seems satisfactory.

To input information into the product, there must be controls that readily interface with the human. Figure D.7 shows 18 common types of controls and

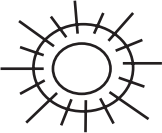
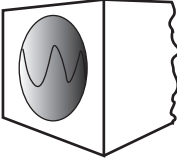
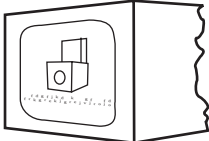
43070 Digital counter	 Icon, symbol display
 Curved dial  Linear dial Fixed pointer on moving scale	 Indicator light
 Circular dial  Linear dial Moving pointer on fixed scale	 Graphical display
 Mechanical indicator	 Pictorial display

Figure D.5 Types of virtual displays.

their use characteristics; it also gives dimensional, force, and recommended use information. Note that the rotary selector switch is recommended for more than two positions and is rated between “acceptable” and “recommended” for precise adjustment. Thus, the rotary switch is a good choice for the time control of the dryer. Also, for rotary switches with diameters between 30 and 70 mm, the torque to rotate them should be in the range from 0.3 to 0.6 N · m. This is important information when one is designing or selecting the timing switch mechanism. In addition, note that for the rocker switch, no more than two positions are recommended. Thus, the top switch on the dryer, Fig. D.4, is not a good choice for the temperature setting.

An alternative design of the dryer control panel is shown in Fig. D.8. The functions of the dryer have been separated, with the temperature control on one rotary switch. The “Start” function, a discrete control action, is now a button, and the timer switch has been given a single scale and made to rotate clockwise. Additionally, the labeling is clear and the model number is displayed for easy reference in service calls.

In general, when designing controls for interface with humans, it is always best to simplify the structure of the tasks required to operate the product. Recall

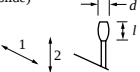
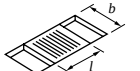
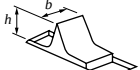
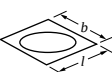
	Exact value	Rate of change	Trend, direction of change	Discrete information	Adjusted to desired value
Digital counter	●	○	○	●	◐
Moving pointer on fixed scale	●	●	●	●	◐
Fixed pointer on moving scale	●	●	○	○	○
Mechanical indicator	○	○	○	●	○
Symbol display	○	○	○	●	○
Indicator light	○	○	○	●	○
Graphical display	◐	◐	●	●	●
Pictorial display	◐	●	●	●	●

 Not suitable
  Acceptable
  Recommended

Figure D.6 Appropriate uses of common visual displays.

the characteristics of the short-term memory discussed in Chap. 3. We learned there that humans can deal with only seven unrelated items at a time. Thus, it is important not to expect the user of any product to remember more than four or five steps. One way to overcome the need for numerous steps is to give the user mental aids. Office reproducing machines often have a clearly numbered sequence (symbol display) marked on the parts to show how to clear a paper jam, for example.

In selecting the type of controller, it is important to make the actions required by the system match the intentions of the human. An obvious example of a mismatch would be to design the steering wheel of a car so that it rotates clockwise for a left turn—opposite to the intention of the driver and inconsistent with the effect on the system. This is an extreme example; the effect of controls is not always so obvious. It is important to make sure that people can easily determine the relationship between the intention and the action and the relationship *between the action and the effect* on the system. *A product must be designed so that when a person interacts with it, there is only one obviously correct thing to do.* If the action required is ambiguous, the person might or might not do the right thing. The odds are that many people will not do what was wanted, will make an error, and, as a result, will have a low opinion of the product.

	Control	Dimension, mm	Force F , N Moment M , N · m		2 positions	>2 positions	Continuous adjustment	Precise adjustment	Quick adjustment	Large force application	Tactile feedback	Setting visible	Accidental actuation
Turning movement	Handwheel 	D : 160–800 d : 30–40	D 160–800 mm 200–250 mm	M 2–40 N · m 4–60 N · m									
	Crank 	Hand (finger) r : <250 (<100) l : 100 (30) d : 32 (16)	r <100 mm 100–250 mm	M 0.6–3 N · m 5–14 N · m									
	Rotary knob 	Hand (finger) D : 25–100 (15–25) h : >20 (>15)	D 15–25 mm 25–100 mm	M 0.02–0.05 N · m 0.3–0.7 N · m									
	Rotary selector switch 	l : 30–70 h : >20 b : 10–25	l 30 mm 30–70 mm	M 0.1–0.3 N · m 0.3–0.6 N · m									
	Thumbwheel 	b : >8	$F = 0.4–5$ N										
	Rollball 	D : 60–120	$F = 0.4–5$ N										
Linear movement	Handle (slide) 	d : 30–40 l : 100–120	$F_1 = 10–200$ N $F_2 = 7–140$ N										
	D-handle 	d : 30–40 b : 110–130	$F = 10–200$ N										
	Push button 	Finger: $d > 15$ Hand: $d > 50$ Foot: $d > 50$	Finger: $F = 1–8$ N Hand: $F = 4–16$ N Foot: $F = 15–90$ N										
	Slide 	l : >15 b : >15	$F = 1–5$ N (Touch grip)										
	Slide 	b : >10 h : >15	$F = 1–10$ N (Thumb-finger grip)										
	Sensor key 	l : >14 b : >14											

*Recessed installation

Figure D.7 Appropriate uses of hand- and foot-operated controls (Source: Adapted from G. Salvendy (ed.), *Handbook of Human Factors*, Wiley, 1987).

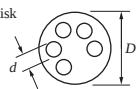
	Control	Dimension, mm	Force F , N Moment M , N-m	2 positions	> 2 positions	Continuous adjustment	Precise adjustment	Quick adjustment	Large force application	Tactile feedback	Setting visible	Accidental actuation
Swiveling movement	Lever 	d : 30–40 l : 100–120	$F = 10\text{--}200\text{ N}$	●	●	●	●	●	●	●	●	○
	Joystick 	s : 20–150 d : 10–20	$F = 5\text{--}50\text{ N}$	●	●	●	●	●	●	●	●	○
	Toggle switch 	b : >10 l : >15	$F = 2\text{--}10\text{ N}$	●	●	○	○	●	○	●	●	○
	Rocker switch 	b : >10 l : >15	$F = 2\text{--}8\text{ N}$	●	○	○	○	●	○	●	●	●
	Rotary disk 	d : 12–15 D : 50–80	$F = 1\text{--}2\text{ N}$	●	●	●	○	●	○	○	○	●
	Pedal 	b : 50–100 l : 200–300 l : 50–100 (forefoot)	Sitting: $F = 16\text{--}100\text{ N}$ Standing: $F = 80\text{--}250\text{ N}$	●	●	●	●	●	●	●	○	○

Figure D.7 (continued).

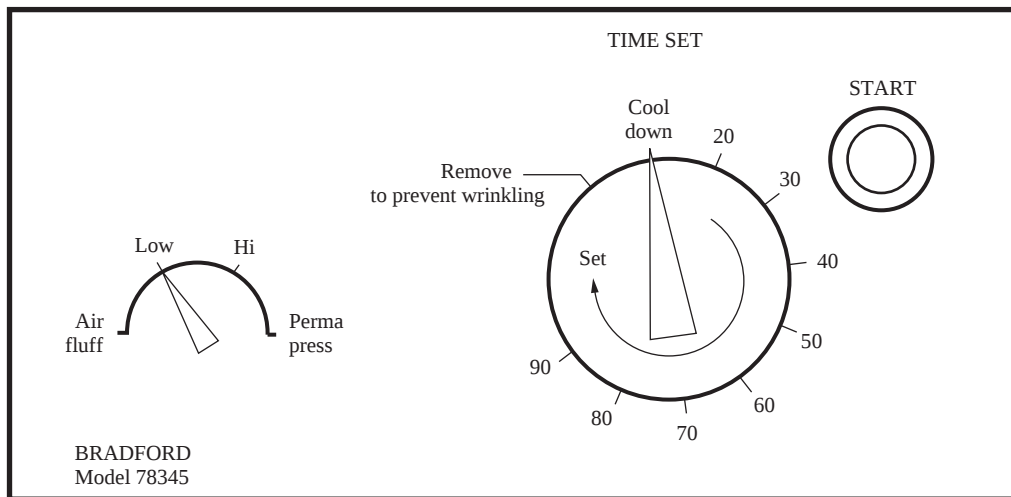


Figure D.8 Redesign of the clothes dryer control panel of Fig. D.4.

D.5 SOURCES

- Burgess, J. H.: *Designing for Humans: The Human Factor in Engineering*, Petrocelli Books, Princeton, N.J., 1986. A good text on human factors written for use by engineers; the dryer example is from this book.
- Damon, A. et al., *The Human Body in Equipment Design*, Harvard University Press, Boston, 1966.
- Dreyfuss, H.: *The Measure of Man: Human Factor in Design*, Whitney Library of Design, New York, 1967. This is a loose-leaf book of 30 anthropometric and biomechanical charts suitable for mounting; two are life-size, showing a 50th-percentile man and woman. A classic.
- Human Engineering Design Criteria for Military Systems, Equipment, and Facilities*, MIL-STD 1472F. <http://hfetag.dtic.mil/docs-hfs/mil-std-1472f.pdf>. Four hundred pages of human factors information.
- Human Engineering Design Data Digest*, Department of Defense Human Factors Engineering Technical Advisory Group, April 2000, http://hfetag.dtic.mil/hfs_docs.html. Excellent online source.
- Human Factors Design Standard (HFDS)*, FAA, <http://hf.tc.faa.gov/hfds/>. Another excellent online source.
- Jones, J. V.: *Engineering Design: Reliability, Maintainability and Testability*, TAB Professional and Reference books, Blue Ridge Summit, Pa., 1988. This book considers engineering design from the view of military procurement, relying strongly on military specifications and handbooks.
- MIL-HDBK-759C, *Human Engineering Design Guidelines*, 1995.
- Norman, D.: *The Psychology of Everyday Things*, Basic Books, New York, 1988. Guidance for designing good interfaces for humans; light reading.
- Moggridge, B.: *Designing Interactions*, <http://www.designinginteractions.com/>. An online book for designing human interfaces for the 21st century.
- Salvendy, G. (ed.): *Handbook of Human Factors*, 3rd edition, Wiley, New York, 2006. Seventeen hundred pages of information on every aspect of human factors.
- System Safety Program Requirements*, MIL-STD 882D. U.S. Government Printing Office, Washington, D.C. <http://safetycenter.navy.mil/instructions/osh/milstd882d.pdf>. The hazard assessment is from this standard.
- Tilly, A. R.: *The Measure of Man and Woman*, Whitney Library of Design, New York, 1993. An updated version of the preceding classic rewritten by one of Dreyfuss's associates.

INDEX

A

- abstraction, levels of, 32, 215
- accuracy, modeling and, 286
- additive tolerance stack-up, 299–301
- aging/deterioration effects, 290
- aisle chair, 147, 158
- analogies, 191–192
- analysis problems, 16
- analytical models, 124–125, 294–295
- assembly
 - drawings, 122–123
 - efficiency, 331, 333
 - instructions, 367
 - manager, 70
 - requirements, in engineering specifications, 162

B

- behavior, human problem-solving, 58–64
- behavior, product, 30
- Belief Map, 235–239
- benchmarking, 157–158
- best practices
 - for product evaluation, 279–280
- bicycle
 - product discovery phase and, 101, 102, 106–109
 - redesign, 37–39
- Bill of Materials (BOM), 15, 245–246
- brainstorming, 190
- brainwriting, 190–191

C

- CAD systems, 118–119, 123–124
- chunks of information, 50, 51, 53
- clamp (see Irwin)
- coefficient of variation, 409–413
- cognitive psychology, 48
- Commercial Off The Shelf (COTS) components, 267
- communication, during design process, 137–141
- competition benchmarking, 157–158
- component
 - assembly, 331
 - development, 253–260
 - handling, 331, 343–346

- mating, 331, 347–349
- retrieval, 331, 342–343
- components, 27
 - configuring, 247–249, 271–273
 - cost of injection-molded, 325
 - cost of machined, 321–324
 - developing, 253–260
 - developing connections/interfaces between, 249–253, 274–275
 - from vendors, 266–269
- Computed Tomography (CT) Scanner, 82–85, 86, 89
- computer-aided design (CAD) systems, 119, 123–124
- computer-generated solid models, 118–119
- concept
 - combining, 207–208
 - defined, 171
 - developed for each function, 206–207
- concept evaluation and selection, 213–239
 - assessing risk and, 226–233
 - decision-matrix method, 221–226
 - feasibility, 218–219
 - level of abstraction and language for, 215–218
 - robust decision-making, 233–239
 - technology readiness, 219–221
- concept generation, 171–209
 - amount of time spent on, 171–172
 - basic methods of, 189–194
 - clamp, 173–176
 - contradictions used for, 197–201
 - functional decomposition technique, 181–189
 - morphological method, 204–208
 - reverse engineering, 178–180
 - Theory of Inventive Machines (TRIZ), 201–204
 - through patent literature, 194–197
 - understanding function of existing designs and, 176–180
- conceptual design, 40, 87–89. *See also* concept evaluation and selection; concept generation
 - phases of, 213–214
 - simplicity and, 208–209
- concurrent engineering, 9
- configuration design, 34–36
- configuration of components, 247–249, 271–273
- conformity, creativity and, 65–66

connections, 249–253
 constraints, design, 40–41
 contradictions, to generate ideas, 197–201
 cost estimates, 320–321
 estimating product development, 133
 of injection-molded components, 325
 of machined components, 321–324
 cost, product
 determining, 316–320
 measuring design process with, 3–6
 cost requirements, 161
 creativity, in designers, 64–66
 “creeping specifications,” 143–144
 Critical Path Method (CPM), 131
 CT Scanner, 82–85, 86, 89
 finding overall function of, 183–184
 subfunction description, 187–188
 customer relationships, 370
 customers
 determining requirements of, 151–155
 evaluating importance of requirements of, 155–157
 identifying, 151
 relating engineering specifications to, 163–164
 satisfaction of, Kano model of, 97–99
 satisfaction with competition, 157–158

D

decision-making
 basics of, 105–106
 choosing a project, 101–109
 concept selection, 216, 233–239
 portfolio decision, 105–109
 risk, 233
 Decision Matrix, 108–109, 221–226, 234
 decisive decision-makers, 62–63
 decomposition, product, 40–44
 functional, 184–188, 204–205
 reverse engineering and, 178–180
 deliverables, 118–124, 128
 design best practices, key features, 10
 design-build-test cycle, 217
 design decisions, 40
 design engineer, 68
 designers. *See also* design teams
 creativity of, 64–66
 generating solutions, 57
 human information processing and, 48–56
 mental processes of, during design process, 56–64
 as part of design team, 69
 problem-solving behaviors by, 58–64
 understanding the design problem, 56–57

design evaluation. *See* concept evaluation and selection
 Design-For-Assembly (DFA), 329–349
 design for cost (DFC), 315–325
 Design for Manufacture (DFM), 12 328–329
 Design for Reliability (DFR), 350–357
 Design for Six Sigma (DFSS), 10
 Design for the Environment (DFE), 20, 358–360, 375–376
 Designing For Sustainability (DFS), 20
 design notebooks, 137–138
 design patents, 373
 design problems. *See also* Quality Function
 Deployment (QFD) technique
 basic actions for solving, 17–19
 configuration design, 34–36
 documentation of, 140
 knowledge and learning during design and, 19–20
 many solutions for, 15–17
 mechanical, 33–40
 mental processes of designers and, 56–57
 original design, 37
 parametric design, 36
 redesign, 37–40
 selection design, 33–34
 solutions for, 15–17
 understanding, 143–144, 143–151
 design process. *See also* designers; mechanical design; product discovery
 communication during, 137–141
 conceptual design phase. *See* concept generation and concept selection
 “creeping specifications” and, 143–144
 defined, 8
 designing quality, 92–95
 documentation and, 363, 366–368
 end of, 363–365
 history of, 8–10
 human factors and, 415–425
 measuring, 3–8
 need for studying, 1–3
 overview of, 81–85
 product definition phase. *See* product generation and product evaluation
 product development phase. *See* product development
 product discovery phase. *See* product discovery
 product support phase, 91–92, 368–370
 project planning phase. *See* project planning
 safety factor in, 403–414
 design report, 139–141
 design reviews, 113, 138–139

Design Structure Matrix (DSM), 132
 design teams
 assessing health of, 76–77
 building performance, 72–73
 characteristics of successful, 72–73
 contract, 73–79
 management of, 71–72
 meeting minutes for, 73, 75, 76
 members of, 68–71
 need for, 66–68
 desktop prototyping, 118
 detail drawings, 121–122
 deterioration/aging effects, 290
 DFA (Design-For-Assembly), 329–349
 DFC (design for cost), 315–325
 DFE (design for the environment). *See* Design for the Environment (DFE)
 DFM (Design for Manufacture), 328–329
 DFR (design for reliability), 350–357
 DFSS (Design for Six Sigma), 10
 DFV (value engineering), 325–328
 disassembly, of product, 13
 disclosure, patent 373
 documentation, communicating final design, 139–141
 domain-specific knowledge, 50
 drawings
 assembly, 122–123
 detail, 121–122
 layout, 120–121
 Dreamliner, Boeing, 146–147

E

efficiency, assembly, 331–333
 end-of-life, product, 13
 End-of-Life Vehicles (ELVs), 376–378
 energy flows, 177, 180
 Engineering Change Notice (ECN), 371
 engineering changes, 370–371
 engineering specifications
 determining importance of, 164–165
 developing, 158–163
 guidelines for good, 162–163
 identifying relationships between, 166–167
 measuring competitors' products, 165
 relating customer requirements to, 163–164
 targets, 165–166
 types of, 160–162
 evaluation. *See also* product evaluation
 of concepts, 88–89
 importance of customer requirements, 155–157
 Evaporating Cloud (EC) method, 197–198
 excitement-level features, 98–99

F

factor of safety, 403–414
 Failure Modes and Effects Analysis (FMEA), 232, 350–353
 failure rate, 355
 fasteners, minimizing use of, 335–338
 Fault Tree Analysis (FTA), 352, 353–355
 Feasibility evaluation, 218–219
 features
 basic, 98
 excitement-level, 98–99
 definition, 27
 performance, 98
 fidelity, 124, 216–217, 293
 flexible decision-makers, 62
 flow of energy, information, and material, 177, 179–180
 focus-group technique, 152, 153, 154
 force flow visualization, 257–259
 form generation, 246–264
 form of the product, 2–3, 29, 243, 244
 Franklin, Benjamin, 102–103
 function, 2–3, 28–40, 243
 behavior and, 30
 defining, 177–178
 developing concepts for each, 206–207
 finding the overall, 181, 183–184
 modeling 181–189
 monitoring change in, 280–281
 using reverse engineering, 178–180
 functional decomposition, 29, 172, 181–194, 204–205
 functional performance requirements, 160
 function diagram, 130

G

Gantt chart, 131, 140
 General Electric CT Scanner. *See* CT Scanner
 generating concepts, 87–88
 graphical models, 118–124
 green design (Design for the Environment), 358–360
 group technology, 260

H

handling, component, 331, 343–346
 Hannover Principles, 20–21, 209, 357
 house of quality. *See* Quality Function Deployment (QFD) Technique
 human factor requirements, 160
 human factors, 415–425
 human information processing, 48–56

I

industrial designer, 70
 information, human memory and, 49–50
 information language, problem-solving behavior and, 61–62
 information processing, human, 48–56
 injection-molded components, costs of, 325
 installation instructions, 367
 installation, product, 13
 instruction manuals, 367–368
 Integrated Product and Process Design (IPPD), 9, 94
 interfaces between components, 249–253
 International Standard Organization's ISO 9000 system, 94–95
 Irwin Quick-Grip clamp, 26, 27
 product decomposition, 41–44, 179
 project planning and, 113–115
 redesign of one-handed bar clamp, 173–176
 reverse engineering, 178–180
 subfunction description, 187
 ISO 9000 quality management system, 94–95

J

Jet Propulsion Laboratory (JPL) (Cal Tech), 26

K

Kano Model of customer satisfaction, 97–99
 Kano, Noriaki, 97
 Key features of design best practice, 10
 knowledge
 increase during design, 19
 types of, 50
 creativity and, 65

L

language
 concept evaluation and, 215–218
 encoding chunks of information, 50
 mechanical design, 30–32
 layout drawings, 120–121
 Lean manufacturing, 9
 level of abstraction, 32, 215
 Level of Certainty, 235, 237
 Level of Criterion Satisfaction, 235, 236
 life cycle, product, 161
 long-term memory, 52–54

M

machined components, costs of, 321–324
 maintainability, 357

maintenance instructions, 367
 manufacturing
 cost, 3–4, 5–6, 317–324
 engineer, 69
 instructions, 366
 processes, 2–3
 requirements, in engineering specifications, 162
 variance, 290, 297
 Marin Mount Vision Pro bike, 39
 product evaluation and, 291–292, 299–300
 product generation for, 269–276
 market pull, 96–97, 99
 Mars Exploration Rover (MER), 26
 Choosing a wheel for
 configuration design, 34–36
 mechanical design language and, 31–32
 planning for, 132
 product support and, 92
 safety factors, 40
 sub-systems, 28
 material costs, 317
 materials,
 properties of the most commonly used, 380–392
 selection of, 264–266
 materials specialists, 69
 mating, component, 331, 347–349
 mature design, 37
 Mean Time Between Failures (MTBF), 355–357
 mean value, 398–399
 measurement of the design process, 3–8
 mechanical failure, 350
 mechanical fuse, 358
 mechatronic devices, 25
 meeting minutes, design team, 73, 75, 76
 memory, human, 48–50
 long-term, 52–54
 short-term, 51–52
 MER. *See* Mars Exploration Rover (MER)
 milestone chart, 131
 MIL-STD 882D (Standard Practice for System Safety), 230–231
 modeling, 117–126, 286, 292–296
 modularity, 248–249
 morphological method, 204–208

N

“NIH” (Not Invented Here) policy, 178, 218
 noise, 290–294
 nominal tolerances, 297
 nonconformity, 65–66
 normal distribution, 397–401

O

objective approach to problem-solving, 61
 observation of customers, 152, 153, 154
 obstructive nonconformists, 66
 operation instructions, 367
 ordering subfunctions, 186–188
 original design, 37
 originality, 60
 overall assembly, evaluation of, 333–341
 overall function, 181, 183–184
 over-the-wall design method, 8–9, 10, 12

P

packaging (configuration) design, 34–36
 parallel tasks, 131–132
 parametric design, 36
 part numbers, 245
 patching, 260–261, 263–264
 patent
 applications, 371–375
 searches, 194–197
 P-diagram, 282–283, 291–292
 Performance and function
 performance evaluation, 281–286, 292–296
 performance features, 98
 PERT (Program Evaluation and Review
 Technique) method, 130–131
 physical models, 117–118, 217, 286, 295–296
 physical requirements, 160
 planning. *See* project planning
 portfolio decision, 105–109
 preproduction run, 118
 Priestly, Joseph, 102–103
 probability, normal, 397–401
 problem-solving behavior, 58–64
 decision closure style, 63–64
 deliberation style, 62–63
 energy source, 58–60
 information language, 61–62
 information management style, 60–61
 pro-con analysis, 102–105
 product change, 96, 99
 Product Data Management (PDM), 14
 product decomposition, 41–44
 product design, 40
 product design engineer, 68
 product development phase, 90–91
 product discovery, 85–86, 95–100
 choosing a project, 101–109
 customer satisfaction and, 97–99
 goal of, 95–96
 market pull and technology push, 96–97

 product maturity and, 97
 product proposal, 99–100
 product evaluation, 279–313, 315–360
 accuracy, variation, and noise, 286–292
 best practices for, 279–280
 Design-For-Assembly (DFA), 329–349
 Design for Cost (DFC), 315–325
 Design for Manufacture (DFM), 328–329
 Design for Reliability (DFR), 350–357
 Design for test and maintenance,
 357–358
 Design for the Environment, 358–360,
 375–376
 goals of performance evaluation, 281–284
 modeling for, 292–296
 monitoring functional change, 280–281
 sensitivity analysis, 302–305
 tolerance analysis, 296–302
 trade-off management, 284–286
 value engineering, 325–328
 product generation, 241–276
 Bill of Materials, 245–246
 developing components, 253–260
 form generation, 246–264
 for Marin Mount Vision Pro bicycle,
 269–276
 materials and process selection, 264–266
 vendor development, 266–269
 Product Life-cycle Management (PLM),
 13–15, 245
 product manager, 69
 product maturity, “S” curve, 97–98
 product proposal, 99–100
 product quality. *See* quality, product
 product risk, 230–233
 product
 function of, 28–29
 liability, 229–230
 life of, 10–15
 safety of, 227–229
 product support, 91–92, 368–370
 project planning, 86, 111–141
 activities of, 111–112
 choosing best models and prototypes for,
 125–126
 design plan examples, 134–137
 goal of, 111
 physical models and prototypes used in,
 117–118
 plan template, 125–133, 128–133
 types of plans, 113–117
 project portfolio management, 101
 project structures, 71

prototypes, 117–118
 choosing, 125
 proof-of-concept prototype, 118
 proof-of-function prototype, 118
 proof-of-process prototype, 118
 proof-of-production prototype, 118
 proof-of-product prototype, 118
 Pugh's method. *See* decision-matrix method
 purchased-parts cost, 317

Q

QFD method. *See* Quality Function Deployment (QFD) technique
 Quality Assurance (QA), 366
 Quality Assurance (QA) specialists, 69, 92
 Quality Control (QC), 366
 Quality Control (QC) specialists, 69, 92
 Quality Function Deployment (QFD) technique, 145–169
 determining what the customers want, 151–155
 developing engineering specifications, 158–163
 evaluating importance of customer requirements, 155–157
 identifying and evaluating the competition, 157–158
 identifying customers, 151
 identifying relationships between engineering specifications, 166–167
 relating customer requirements to engineering specifications, 163–164
 reverse engineering and, 178
 setting engineering specification targets and importance, 164–166
 uses of, 168
 quality, product
 design process and, 92–95
 determinants of, 6
 effect of variation on, 289–292
 measuring design process with, 3, 6
 Quick-grip clamp. *See* Irwin

R

rapid prototyping, 118
 recycling, 13, 359, 360
 redesign, 37–40
 of clamp, 173–176
 QFD method and, 145
 refining products, 260–264
 refining subfunctions, 188–189
 reliability, 161, 350, 355–357

reliability-based factor of safety, 406–414
 reparability, 357
 resource concerns, in engineering specifications, 161–162
 retirement, product, 13
 retrieval, component, 331, 342–343
 reuse, of product, 13
 reverse engineering, 178–180
 risk, 226–233
 decision, 233
 product, 230–233
 project, 232–233
 robust decision making, 233–239
 robust design
 by analysis, 305–308
 through testing, 308–313
 Rover, Mars. *See* Mars Exploration Rover (MER)

S

safety, product, 227–229, 403–414
 sample mean, 398–399
 sample standard deviation, 398–399
 sample variance, 399
 “S” curve, product maturity, 97
 selection design, 33–34
 sensitivity analysis, 302–305
 sequential tasks, 131
 serviceability, 357
 short-term memory, 48, 51–52, 55
 simple design plan, 134–135
 simultaneous engineering, 9
 6-3-5 method, 190–191
 Six Sigma philosophy, 9–10, 297
 sketches, 119
 solid models, 118–119, 123
 spatial constraints, 247, 269–270
 specification, patent, 373–374
 spiral process, 115–117
 standard deviation, 398–399
 Standard Practice for System Safety (MIL-STD 882D), 230–231
 standards, 161–162
 Stage-Gate Process, 113
 statistical stack-up analysis, 301–302
 subfunction
 ordering, 186–188
 refining, 188–189
 descriptions, 184–186
 subjective approach to problem-solving, 61–62
 subsystems, 27
 surveys, 152, 154
 sustainability, design for, 20–21
 SWOT analysis, 101–102, 105

T

Taguchi, Genichi, 305
Taguchi's method, 305–306
tank problem, 283–284
targets, engineering specifications, 165–166
tasks
 planning, 126–128
 sequence, 131–133
teams, design. *See* design teams
technicians, 69
technology push, 96, 99
technology readiness, 219–221
Templates (All available on line
 BOM, 246
 Change order, 372
 Design for Assembly, 330
 Design Report, 139–141
 FMEA, 351
 Machined Part Cost Calculator, 322
 Meeting minutes, 75
 Morphology, 205
 Patent prospects, 375
 Personal Problem Solving Dimensions, 59–63
 Product Decomposition, 42–43
 Project Plan, 127
 Team contract, 74
 Team health inventory, 77
 Plastics Part Cost Calculator, 325
 Product Proposal, 100
 Pro/Con Analysis, 104
 Reverse Engineering, 182
 Swot Analysis, 102
 Technology Readiness, 221
testability, 357
Theory of Inventive Machines (TRIZ), 201–204

time

 product development, 6–8
 project planning and, 128–130
 spent on developing concepts, 171–172
time requirements, in engineering
 specifications, 161
tolerance analysis, 296–302
trade-off management, 284–286
TRIZ. *See* Theory of Inventive Machines

U

UL standards, 162
uncertainties, 285–286
uncoupled tasks, 132
“use” phase of products, 13
utility patents, 373

V

value engineering/analysis, 325–328
variant design, 40
variation, 286–292, 297
vendor development, 266–269
vendor relationships, 368–370
vendor representatives, 70
verbal problem-solvers, 61

W

Waterfall model of project planning, 113
work breakdown structure, 131
worst-case analysis, 301

X

X-Ray CT Scanner. *See* CT Scanner