

Symbolic Logic

Syntax, Semantics, and Proof

By David W. Agler

This textbook is intended solely for personal classroom use.
Copyright © 2025 by David W. Agler

All rights reserved. No part of this publication may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without the prior written permission of the author, except in the case of brief quotations embodied in critical reviews and certain other noncommercial uses permitted by copyright law.

As this book is privately distributed, it does not conflict with my text
Symbolic Logic: Syntax, Semantics and Proof.

Agler, David W., 1982- Symbolic logic : syntax, semantics, and proof / David W. Agler. pages cm Includes bibliographical references and index. ISBN 978-0-9000000-0-0 1. Logic, Symbolic and mathematical. I. Title.

Printed in the United States of America

1 2 3 4 5 6 7 8 9 10

Contents

I Introduction

1	Introduction to Logic	3
1.1	What is Logic?	3
1.2	Evaluating Arguments	15
1.3	Testing for Deductive Validity	24

II Propositional Logic

2	PL Language	41
2.1	PL Symbols	41
2.2	PL Syntax	43
2.3	PL Semantics	58
2.4	PL Translation	69
2.5	End of Chapter Exercises	98
3	PL Truth Tables	101
3.1	Determining the truth of a wff	102
3.2	The Truth-Table Method	105
3.3	Truth-Table Analysis	107
3.4	Limitations of Tables	124
4	PL Truth Trees	131
4.1	Introduction to Trees	131
4.2	Branches and Decomposition	133
4.3	Truth Trees: Decomposition Strategies	148
4.4	Truth Trees: Analysis	152
4.5	Advantages of the Truth-Tree Method	168
4.6	PL Tree Rules	169
5	PL Derivations	171
5.1	Two notions of logical entailment	171
5.2	Syntactic entailment	174
5.3	Setup	175

5.4	The Intelim Derivation Rules	177
5.5	Derivation Strategies	208
5.6	Additional Derivation Rules	215
5.7	Additional Exercises	228
III Predicate Logic		
6	QL Language	235
6.1	QL Symbols	236
6.2	QL Syntax	237
6.3	QL Semantics	249
6.4	QL Translation	261
7	Dummy Chapter	281
8	QL Derivations	283
8.1	QL Derivation Terminology	283
8.2	QL Four Quantifier Rules	284
8.3	Quantifier Negation	303
IV Back Matter: Solutions, Etc.		
	Notes	309
	Bibliography	311
	Solutions to Exercises	313

PREFACE

0

This book is an introduction to symbolic logic. The primary goal is to teach the basics of symbolic logic (propositional and predicate logic). By the end of this book, you will know (1) what an argument is, (2) some criteria for makes a good argument, (3) what it means for a conclusion to "follow from" its premises, (4) how to *test* when a conclusion follows from its premises, and (5) how to *prove* that a conclusion follows from its premises. You will know this with respect to the English language, but also with respect to two different logical languages and (6) you will know how to translate English arguments into these logical languages.

Who is this book for?

In deciding whether this book is worth your time, it is important to know the target audience for this text and how the text presents the material.

First, the primary audience for this book are those who know nothing about logic but want to learn the basics of symbolic logic. As such, it is an ideal textbook for a first course in logic or for those who have taken a logic course but who want to review the basics. I've written the text with undergraduate with a limited mathematical background in mind.

Second, the book is written in a way that (I hope) makes it easier to learn logic outside of the classroom setting. The textbook material is delivered in bite-sized chunks followed by exercises. While this approach is not novel, this text makes a concerted effort to make these chunks smaller than other texts. In other words, exercises are more dispersed throughout the text than in other texts. I'm hoping that this approach gives instructors teaching from the text more flexibility in terms of where to start and stop their courses and a greater opportunity to assess student understanding in the classroom. In addition, I'm hoping that this approach allows students and those learning on their own the opportunity to periodically check their understanding of the material and to take breaks when needed.

Finally, it is worth noting that this book is not intended to be a comprehensive introduction to all aspects of logic as it does not cover topics like informal fallacies, modal logic, metatheory, or axiom systems. As such, this text would not be ideal for a critical thinking course or a course in mathematical logic.

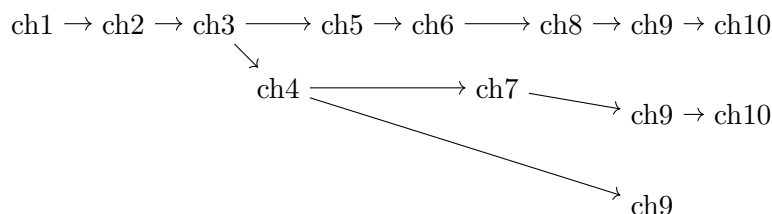
Structure of the book

The book begins with an introductory chapter that orients the reader to the study of logic. Part I contends that the goal of logic is to sort arguments into two types: good arguments and bad arguments. In order to come close to achieving this goal, many hurdles must be overcome. The first hurdle is to define what it means to be an argument. Chapter 1 thus defines what an argument is and distinguishes arguments from other types of discourse. The second hurdle is to define what it means for an argument to be good. Chapter 1 then introduces three criteria for a "good argument" and then focuses on one of these criteria for the remainder of the book: validity. Finally, the third hurdle is that it is one thing to know what a good argument is, but it is another thing to be able to *identify* or *construct* good arguments. The chapter concludes with a discussion of informal methods for identifying valid arguments and limitations in these methods.

From here, the book is divided into two parts. In Part I, the language of propositional logic is introduced. In Chapter 2, the symbols, syntax, and semantics are articulated. It concludes with a discussion of how to translate English sentences into propositional logic. In Chapters 3 and 4, truth tables and truth trees are presented as methods for testing sets of wffs for logical properties. In Chapter 5, the concept of a proof system is introduced and a natural deduction proof system is presented. In my experience, students tend to struggle with proofs and so there are many examples and exercises to help students master this skill.

Part II of the text turns to predicate logic. This part has a similar structure to Part I. In Chapter 6, the formal language of predicate logic is introduced. In Chapter 7, truth trees are articulated and in Chapter 8, the proof system for propositional logic is presented.

Those with limited time or not wishing to read the book cover to cover might still profit from this book by reading chapters in a specific order. Three different courses of reading are provided in the diagram below (although other paths are possible):



The top path skips Chapter 4 and 7, which deal with propositional and predicate logic truth trees. The middle path skips propositional and predicate logic proofs. Finally, the bottom path skips propositional logic proofs and the entire discussion of predicate logic (Chapters 6-8). While the more advanced Chapters 8 and 9 contain elements of both trees and proofs, one may still profit from several sections in these chapters even if one skips discussions of trees and proofs.

About the website

Supplemental material for this book can be found at my website (www.davidagler.com) and on my YouTube Channel (<https://www.youtube.com/LogicPhilosophy>). On the website, you can find course handouts, lecture slides, sample exams, and supplemental materials not included in this text. On my YouTube channel, you can find roughly 75 videos that cover the material in this book. The videos are organized into three playlists. The first covers propositional logic, the second covers predicate logic, and the third covers both propositional and predicate logic. The videos are also useful for students who are learning remotely, absent from class, or want to learn at a different pace than the rest of the class.

Acknowledgments

I owe thanks to more people than I can name here. Thank you to Ayesha Abdullah, Deniz Durmus, Mark Fisher, Emily Grosholz, Cameron O'Mara, and Ryan Pollock for providing valuable feedback on early drafts of this book. I also owe thanks to the students who have used this book and who have provided me with feedback. Some of them include: Christopher Allaman, Meghan Barnett, Kevin Bogle, Ashley Brooks, Isaac Bishop, Aurora Cooper, Maureen Dunn, Elliannies Duran, Ariel Endresen, Nayib Felix, Joy Garcia, Alex Kirk, Edward Lackner, Sarah Mack, Alexander McCormack, Kristin Nuss, Courtney Pruitt. Karantha Parker, Brooke Santkiewicz, Ariel Valdez, Amanda Wise, and Haochen Wang.

I would like to express my gratitude to various individuals and groups who have played a pivotal role in supporting me during the writing process of this book. Among them are "The Faculty Writing Program" at Penn State, the "30/10 Writing Challenge" also at Penn State, and the writing group organized by Aminah Hasan at Emory University. The graduate students in my logic seminar were helpful in providing feedback on early drafts of this book. They include: Cole De Jager, Brooke Hubsch, Madison Phillips, India Rhodes, Corena Smith, and Jamellah Sweeting. In

addition, I want to thank Brooke Hubsch whose interest in and aptitude for logic puzzles inspired me to generate more logic puzzles for this book. I owe a special thanks to students who have served as Learning Assistants (LAs) in the past. These are Rachel Xue and Isabel Newby.

Additionally, I extend my appreciation to those who engage with the logic videos on my YouTube Channel. Over the years, my viewers have provided valuable feedback and support on how I teach logic. The feedback has helped to improve the final product of this text and the support has kept me motivated in writing this work.

The construction of this text benefits from [Brian Dunn](#), the creator of [lwarp](#) and Peter Wilson and Lars Madsen, the creator and maintainer of the [memoir class](#).

In reading this textbook, if you notice errors, have any constructive criticism or suggestions, or you just want to say hello, please feel free to contact me at dwa132@psu.edu or through my YouTube Channel <https://www.youtube.com/davidagler>.

I want to use your book

Sure.

- You can download a pdf version of this text here: [Symbolic Logic – An Introduction](#)
- You can compile one yourself from L^AT_EX.
- This book is distributed under the [GNU AGPLv3 license](#).

Part I

Introduction

1.1 WHAT IS LOGIC?

In our day to day lives, we find ourselves *arguing* with other people. Sometimes we want someone to believe something that they don't currently believe. Other people are no different. They wish to convince us of something and so they use various tactics to get our assent. Sometimes we find their tactics convincing, other times we are unpersuaded, and other times we are unsure what we should think.

Often times we think that the way to get someone to believe what we believe is by giving them *reasons* to accept what we accept. We take ourselves to be reasonable people. At least for many of the propositions we believe, we take these propositions to be true not merely because we believe them but because (i) we have *excellent reasons* for those beliefs and (ii) those excellent reasons support the beliefs in question. And, insofar as we take other people to be reasonable like ourselves, we will often convey those excellent reasons in support of a belief to these people. We do this in the hopes that they recognize that not only the reasons are excellent but also that the argument has some particular quality to it that makes it worth accepting.

In addition to putting forward reasons, we also criticize other people's arguments. When we do this, we either take the reasons they put forward in support of their argument to be *false* or we point out that even if the reasons they have were true, this doesn't mean one ought to rationally accept the conclusion in question. That is, we think that their argument has some particular feature (or features) that make it flawed.

The primary goal of logic is to separate good arguments and bad arguments. Doing this requires a variety of other subgoals, e.g., defining what an argument is, what it means for an argument to be good (or bad), methods for identifying good versus bad arguments, and methods for constructing good arguments.

Definition 1.1.1: Logic

Logic is a science that to separate good and bad arguments.

In the remainder of this chapter, we will focus on three of these subgoals. First, we will define what an argument is and how to identify arguments.

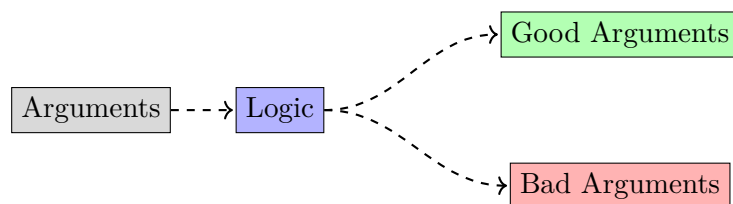


Figure 1.1: Logic aims to identify good arguments

Second, we will develop a preliminary account of what it means for an argument to be *good* or *bad*. Third, we will outline two methods for identifying good and bad arguments.

To begin, in formulating criteria for good and bad arguments, the focus of logic is on *arguments*, not on *arguing*. Arguing and arguments differ in many important ways. Arguing is a broader, more complex phenomena that often, although not always, includes arguments. When we argue, it is often directed at some other person (our friends, family, or strangers), it can involve yelling, the rolling of our eyes in disbelief, or some other gesture. In contrast, an “argument” is a sequence of sentences (called “propositions”) where some proposition (called the “conclusion”) is represented as following from another set of propositions (called the “premises”).

Definition 1.1.2: Argument

An argument is a series of propositions in which a certain proposition (the conclusion) is represented as following from another set of propositions (the premises or assumptions).

One way to think about the concept of an argument is that an argument is an *abstraction* from what goes on when people argue. It is a selection of part of what goes on when people are arguing. That is, part of what people do (or at least aim to do) when they argue is to utter sentences that support some other sentence. In doing this, they put forward arguments.

1.1.1 Propositions

The definition of an argument contains a number of terms that need clarification. First, an argument is a series of propositions. What is a *proposition*?

Definition 1.1.3: Proposition

A proposition is something that is typically expressed by a sentence that is capable of being true or false.

Let's consider some examples of sentences that do express propositions and those that do not. First, consider the following sentences:

1. The sky is blue.
2. Tek is 6'0 tall.
3. If there are three cookies, then one cookie is missing.

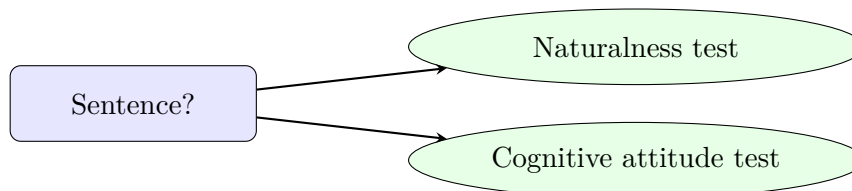
Each of the above sentences express something that can be true or false. In contrast, consider the following items that do not express propositions:

1. A rock on the ground
2. The blueness in Tek's shirt
3. The sentence: do you have any water?

Notice that the rock (which is an *object*) may exist in the world but it is not something that is true or false. If Liz were to say "the rock is on the ground", then this sentence (since it expresses something that can be true or false) is a proposition. Similarly, the *property* of being blue may be in Tek's shirt, but the blueness itself (as a property) does not express a proposition. Finally, not all sentences express propositions since while it may make sense to answer "yes" or "no" to the question of whether you have water, this type of sentence does not express content that can be true or false.

1.1.1.1 Identifying propositions

With a definition of a proposition along with a few examples in place, it would be helpful if there were a method for *identifying* propositions. In this subsection, two informal methods are presented: the naturalness test and the cognitive attitude test.



There are at least two ways to pick out sentences that express propositions. *Naturalness test*
First, there is the "naturalness test". To use the naturalness test, imagine a

scenario where someone utters the sentence S and another person responds "True" or "False". If the response sounds "natural", then there is evidence that S expresses a proposition. While the notion of something "sounding natural" is a vague notion, we can take it to mean that a native speaker of the language in which the sentence is uttered will judge it to be an acceptable response.

Let's consider two examples. First, consider the sentence "the sky is green". Now suppose Tek utters this sentence and Liz responds "False" or "That is false". This response is natural and so there is evidence that "the sky is green". Second, imagine that Tek utters "where is the nearest gas station?" to Liz and Liz responds "False". This is a strange response. Liz would say that it does not make sense to respond in this way. Given Liz's judgment, we have reason to believe that "where is the nearest gas station does not express a proposition?"

Cognitive Attitude Test

The second test is the "cognitive attitude test". There are a variety of mental and emotional attitudes we may have toward things in the world or our minds. Suppose Tek just saw a new movie. Tek may hate the movie or love it or be indifferent about it. A cognitive attitude is a mental or intellectual attitude that someone has toward something. Such attitudes have some relation (explicitly or implicitly) to a person's belief about the world. For example, most explicitly, Tek can believe he saw some specific movie. Alternatively, Tek might *doubt* he saw some specific movie ten years ago. Attitudes like *belief*, *doubt*, *knowing* are all examples of cognitive attitudes.

With cognitive attitudes in mind, suppose we have a sentence S and we want to determine if S expresses a proposition. To test, we place it to the right of (1) a phrase that expresses a cognitive attitude and (2) the word "that". For example,

- Tek believes that S .
- Tek knows that S .
- Tek doubts that S .

If it makes sense to make S the object of a propositional attitude, then S likely expresses a proposition. To use a concrete example, suppose the sentence is "The sky is green".

- Tek believes that *the sky is green*.
- Tek knows that *the sky is green*.
- Tek doubts that *the sky is green*.

Since "the sky is green" can be the object of a propositional attitude, there is good reason to believe that "the sky is green" expresses a proposition.

Finally, it is important to stress that not every sentence expresses a proposition. In general, questions (interrogatives), exclamatory sentences, and commands (imperatives) do not express propositions since the content they express is incapable of being true or false. For example, consider the following sentences:

- Questions: Who has my money? Where am I? Is John tall?
- Exclamations: Woah! Yikes! Congratulations!
- Commands: Close the door. Get out of here.

Some additional propositions include: (1) John is tall, (2) There are 10,234 trees in Paris, (3) Socrates is mortal, and (4) The earth is flat.

In the case of questions, it seems unnatural to respond "True" to the question of "Who has my money?" The right sort of response is to give the name of the person who has the speaker's money or to say that you don't know. Similarly, applying the cognitive attitude test to the question "Who has my money?" does not make sense. Doing so would give us "Tek believes *who has my money*" which is not a natural utterance and is ungrammatical. The same is true for both exclamations and commands. Responding "true" or "false" to these sentences or applying the cognitive attitude test yields a strange result.

Exercise 1.1

Identify Propositions. For the following sentences, state which express propositions and which do not express not propositions. If necessarily, use the naturalness and cognitive attitude test.

1. Be a yardstick of quality.
2. Let's Go Pens!
3. How may I help you?
4. Let the dog out.
5. In a fixed rate par bond, the issuer issues the bond at par value.
6. Brandon has Finance 301 at 11:15PM on Thursdays.
7. Recycling bins are blue.
8. Tek is eating a sandwich.
9. Mike goes to the University of Miami.
10. Billboards are a great way to advertise for your company.
11. Can you pass me the pepper?
12. Isaac Newton discovered gravity when he dropped a piano on his brother's head.

1.1.1.2 Some finer points involving propositions

With the definition of a proposition and some basic examples, let's consider some finer points about propositions.

Propositions are abstract First, we have defined a proposition as something that is capable of being true or false. We have also claimed that it is what is typically expressed by a sentence rather than the sentence itself. A proposition then is something like the *meaning* or *content* typically expressed by certain sentences. As such, a proposition is something *abstract* as it is not the ink on the page or the sound waves created by a person. Some individuals find talk of abstract objects worrisome since they believe that the only things that exist are things located in space and time (concrete things). These individuals might be inclined to define a proposition not as that which is expressed by a sentence but the sentence itself (e.g., the ink on the page).

Questions that express propositions Second, in the previous section, it was noted that questions, exclamations, and commands do not (in general) express propositions. However, the situation is not so simple since there are some contexts where questions do express propositions. For example, suppose you and your friend Tek are having dinner. Your friend Tek tells you that he has joined a cult. Tek begins to describe the belief system of his new cult, but you stop him before he can continue. You utter "are you crazy?" On the surface, this is a question and so, if questions fail to express propositions, then you have not uttered a proposition. However, in this context (and in general) this question is *rhetorical*. You are not asking Tek whether he is crazy but instead telling Tek that he is crazy. Your utterance "are you crazy" is expressing some content that is capable of being true or false. As such, the sentence expresses a proposition.

Declaratory sentences that do not express propositions Third, not only do some questions express propositions, but some sentences that appear to be describing the world (and therefore expressing propositions) are, on closer inspection, not expressing propositions. Here is an example. Suppose you and Tek are having dinner at Tek's home. You insult Tek and are now in a heated argument. Angrily, Tek stands up, points to the door, and says "there is the door." What is Tek saying? On the surface, he is saying that there exists a door in the room and it is the door that he is pointing at. But, as most people would rightly take Tek to be *commanding* you to leave. Therefore, just because a sentence is, on the surface, describing the world does not mean that it is expressing a proposition.

Different sentences but the same proposition Fourth, if the proposition is nothing more than the content expressed

by a sentence, and two different sentences can express the same content (have the same meaning), then two different sentences can express the same proposition. For example, the sentence "Tek loves Liz" and "Liz is loved by Tek" express the same thing: some content that is capable of being true or false. Therefore, they can be understood as expressing the same proposition in two different ways. A similar point might be made when the same content is expressed in two different languages. Take the sentences "Tek quiere la chaqueta" (Tek wants the jacket) and "La chaqueta es buscada por Tek" (The jacket is wanted by Tek). Although these are two different sentences (one expressed in Spanish, the other in English), given that they mean the same thing, they can be understood as expressing the same proposition.

Fifth, different utterances of the same sentence can express different propositions. The meaning of a sentence often depends upon contextual factors (who said it, when it was said, how it was said, etc.). For example, consider how the following three different utterances of the same sentence "I ate breakfast" express three different propositions.

One sentence but different propositions

1. "I ate breakfast" is uttered by John expresses the proposition *John ate breakfast*
2. "I ate breakfast" is uttered by Liz expresses the proposition *Liz ate breakfast*
3. "I ate breakfast" is uttered by Tek expresses the proposition *Tek ate breakfast*

Since the meaning of "I ate breakfast" depends upon who uttered the sentence (the speaker of the sentence), if the speaker changes, then the proposition expressed by the sentence changes.

Sixth, one common mistake with respect to identifying whether something expresses a proposition is thinking that the truth value of a proposition is known. This is not correct for two reasons.

A sentence can express a proposition even if you don't know if it is true

First, there are many sentences that are either true or false even though no one knows whether they are true or false. For example, consider the sentence "there are 250,304 trees in Morelia." Let's suppose that no one knows whether this is true or false. Even if that were the case, presumably there is some fact of the matter about this sentence. That is, presumably there is some precise number of trees in Morelia and this number either makes this sentence true or false. Second, it is important to note that the definition of a proposition says that it is something that is capable of being true or false. What this means is that it is the type of thing that could be assigned a truth value. Some sentences are capable of being true

or false even if (1) no one knows whether they are true or false or (2) there is no fact to the matter as to whether it is true or false. Take a sentence about the future like "Tek will eat a sandwich tomorrow." Suppose that we do not know whether Tek will eat a sandwich tomorrow. In addition, suppose the world is indeterministic and so the present state of affairs along with the physical laws of the universe do not determine whether Tek will or won't eat the sandwich. On our definition, this sentence is capable of being true or false even if (1) we don't know if it is true or false and (2) there are no present facts fixing its truth value.

In other words, to know that a sentence expresses a proposition, you only need to know that it can take a truth value, not what that truth value is.

Propositions expressed by things other than sentences

Seventh, propositions are typically expressed by sentences, but depending upon one's conception of a sentence, it is possible that propositions can be expressed by other means. If a sentence is understood as something that can only be written or spoken, then propositions can also be expressed through sign language. One step further than this is that propositions can be expressed through gestures or other non-linguistic conventions. For example, suppose that an artist has painted a landscape. They place the painting on the wall and next to the painting they place a placard that has an image of a hand with a finger pointing at the artist. Conventions for hanging art along with the juxtaposition of the painting, the placard, and the artist express that the artist is the creator of the painting.

Exercise 1.2

Identify Propositions: For the following sentences, state which express propositions and which do not express not propositions.

1. God does not exist.
2. I know that God exists.
3. You are beautiful.
4. It is morally wrong to eat meat.
5. I bet you five dollars.

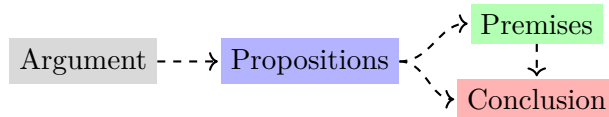
1.1.2 Some key ideas about arguments

There are a number of features of arguments. Let's consider a few of these features.

Arguments have premises and a conclusion

First, an argument is a series of propositions involving premises and a conclusion. The terms "premises" and "conclusion" refer to the *role* certain propositions play in an argument. Intuitively, the "premises" or the

“premise” of an argument plays a supporting role in that they serve to “support” or act as “evidence” or reasons for a conclusion. The “conclusion” plays the role of proposition that is supported.



From the fact that an argument requires premises and a conclusion, a second key point follows. This second point is that not every set of propositions is an argument. For example, consider the following collection of propositions:

Not every set of propositions is an argument

- The sky is blue.
- The sky is not blue.
- The sky is green.

In this example, while there is a set of propositions since the propositions do not take the role of premises or conclusions, there is no argument. It is a mere list of propositions. In other words, a description of events, a work of fiction, a list of items you want from the grocery store, a diatribe of insults all may contain propositions, but none are “arguments” since the propositions do not take the role of premises or conclusions. Let’s illustrate this second point with another example. Consider the following passage that consists of propositions but is not an argument

Yesterday, I saw a little bunny. He was so white and fuzzy and cute. I tried to walk up to him and pet him, but he cowered in fear. Yes, I had just eaten a delicious piece of rabbit meat not too long ago, but the little bunny did not know that.

In the above example, there is a story about a bunny. The story consists of a series of propositions. The propositions are even ordered where one proposition follows from another. However, the order is a temporal order (first this happened, then this, then this) and so none of the propositions take on the role of premise or conclusion. As such, the above passage is not an argument.

A third point to make concerning arguments is that a passage of text might contain an argument but not everything in that passage need be either a premise or a conclusion. To put this crudely, a passage of text, a speech, or the transcription from a debate may contain arguments but not everything within text, speech, or transcription is part of an argument. For example, consider that in everyday life people ask rhetorical questions,

Not everything in the presentation of an argument is part of an argument

utter exclamatory statements, or even issue commands in the course of putting forward an argument. Such sentences, while present in the course of putting forward an argument, are not part of the argument. Let's consider an example.

Does God exist? Now, listen closely! I first started thinking about whether God exists when I was young. Surely, God does not exist. If God did exist, then there would be no suffering in the world. But, as a matter of fact, there is suffering in the world. Therefore, God does not exist.

In the above example, notice that there are several features of the presentation of an argument that are not a part of the argument proper. There is the initial question "Does God exist?", which does not express a proposition. There is the imperative sentence "Now, listen closely!" which does not express a proposition. And, finally, there is a proposition that seemingly does not seem to serve as either a premise or a conclusion: the proposition "I first started thinking about whether God exists when I was young." In sum, not everything in the presentation of an argument is part of an argument.

Argument indicators

A fourth point concerning arguments is that an argument's premises and conclusions are sometimes (but not always) indicated by expressions called "argument indicators". Argument indicators are various words or expressions that indicate that a premise or conclusion is being expressed.

Indicators of Premises

because ... , for, since, for the reason that ... , is supported by ..., may (or can) be deduced from, ... is a reason, ... suggests

Indicators of Conclusions

therefore, hence, in conclusion, so, entails, implies that, indicates, therefore, consequently, we may (can) deduce that, suggests ..., it follows that ..,

Argument indicators are helpful not only for pointing out which propositions are premises and which are conclusions, but also help to indicate that an argument is being expressed. For example, consider the following two propositions:

The sky is blue. The sky is blue.

In this example, it does not appear there is an argument since neither propositions appears to take the role of premise or conclusion. However, contrast this with the following example:

The sky is blue. Therefore, the sky is blue.

In this example, the word "therefore" prefixes the second "the sky is blue". In making this addition, it is now clear (1) that the second instance of "the sky is blue" is a conclusion and (2) that these two propositions together form an argument.

A fifth point concerning arguments is the order in which the premises and conclusion are presented typically does not play a role in determining whether or not something is an argument. That is, an argument might be presented in a passage of text (or in a speech) and the writer (speaker) starts their argument by letting you know the point they are trying to prove (their conclusion). For example, Tek may wish to argue that teaching religious texts in schools funded by tax dollars ought to be illegal. He might say the following:

The order of presenting the premises and the conclusion does not matter

We ought not let religious texts into public schools. I believe this for at least two reasons. First, there are various religions and so selecting one would be prejudicial to one religion over another. Second, schools are funded by taxpayers and taxpayers prefer other subjects take priority over the study of religion (parents want their children to learn math, science, and history).

Notice that the above passage begins with the conclusion and then proceeds to give reasons for that conclusion. However, Tek could have easily began his argument by stating his two premises and concluded with his conclusion. In sum, the order in which the premises and conclusion are presented does not play a role in determining whether or not something is an argument.

A sixth point concerning argument involves a common practice for presenting arguments. It is common practice to take arguments and put them in what is called "argument standard form". Argument standard form is a conventional way of presenting arguments so that (1) the premises of the argument are clearly distinguished from the argument's conclusion and (2) each premise is numbered. To illustrate, let's consider the following passage of text:

Argument Standard Form

Have you heard of this great argument? All women are mortal.
Liz is a woman. Therefore, Liz is mortal.

To put this argument in argument standard form, we identify the propositions of the argument. These are the following:

1. All women are mortal.
2. Liz is a woman.
3. Therefore, Liz is mortal.

The argument above is in argument standard form. Since the initial question is not a proposition, we ignore it as it is not a part of the argument. Next, we label and number the premises using "P" and label the conclusion with "C":

- P1: All women are mortal.
- P2: Liz is a woman.
- C: Therefore, Liz is mortal.

Expressing arguments in argument standard form is helpful for several reasons. First, it makes clear what features of a passage of text are part of the argument and what part are filler. Second, it clearly distinguishes the premises from the conclusion. Third, it can be helpful for better identifying problems with arguments. To illustrate this final point consider the following argument:

- P1: Tek is guilty of fraud.
- P2: Tek signed the document without permission.
- C: Tek is guilty of fraud.

Suppose you were trying to use the above argument to convince someone that Tek is guilty of fraud. In examining the above argument, very few people would be convinced that P1 provides an independent reason to accept C as it is simply a restatement of the conclusion itself!

Exercise 1.3

Identify Arguments: For the following sets of sentences, state which express arguments and which do not express arguments. If a passage of text is an argument, try to put the argument in argument standard form.

1. If Jimmy goes to school, he will get a good grade. His mom would be really happy if Jimmy gets a good grade. Therefore, Jimmy should go to school.
2. I really like elephants. They have super big ears and a really long nose. What a cool animal!
3. Going to the doctor is hard enough but the cost of health care is making it even harder. People got by before all these medical advances. I wish health insurance was more affordable.

4. You should read the review of the new restaurant that was in the paper this morning. It had great information on the types of food available. From the way it sounds, it could be a pretty neat place. It also describes the environment pretty well. Definitely check out the paper when you get a chance.

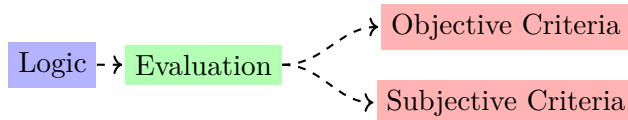
1.2 EVALUATING ARGUMENTS

Recall that logic is a science that aims to separate good arguments from bad arguments.



In previous sections, we have clarified the definition of an argument and various concepts upon which it depends. Now, it is worth considering what makes an argument “good” or “bad”. Saying an argument is good or bad involves an evaluation of the *quality* of an argument. If an argument is good, then it meets certain standards or criteria. If it is bad, then it fails to meet those standards. Much of the business then of logic involves identifying and clarifying these standards.

Let’s begin by distinguishing two different types of criteria: **subjective** and **objective** criteria.



When we evaluate an argument using **subjective criteria** we judge whether the argument is good or bad by appealing to some standard that involves the relation of the argument to some subject. For example, if Tek says an argument is good because he likes it, Tek is evaluating the argument based on how the argument makes him feel (or how he feels about the argument). Here *how Tek feels* is a subjective consideration used to evaluate the argument. Another example is that if we expect arguments to be entertaining, then we evaluate whether the argument produces the right type of response from me or others. Similarly, if we expect an argument to be thought-provoking, then we might evaluate the argument as good if it makes us consider things differently. For example, if an argument against the reality of free will caused us to question whether we were

free, we might evaluate the argument as "good" because we are rethinking whether we are free.

In contrast to subjective criteria, when we use objective criteria to evaluate the quality of an argument, we use judge whether the argument is good or bad using features found in the argument itself. There are three **objective criteria** used in evaluating arguments.

Truth of the propositions First, arguments are evaluated according to whether the propositions (premises and conclusion) that compose the argument are true. If the propositions that compose the argument are true, then the argument is considered "good" in that particular sense. This type of evaluation of arguments is an evaluation independent of the *relation* between the premises and the conclusion. In short, each proposition of the argument is evaluated independently for its truth or falsity. For the most part, we won't consider what makes a proposition true or false.

Relevance Second, arguments are often evaluated with respect to whether the premises of the argument are relevantly related to their conclusion. Take the following argument:

- P1: Tek's dog died.
- C: Therefore, Tek owes Liz five dollars.

Let's suppose that the death of Tek's dog is completely unrelated to Tek owing Liz five dollars. Let's also suppose P1 and C are true. If we only evaluated the quality of an argument based on whether it contained true propositions, then the above argument would be a good argument. However, intuitively, since the truth of P1 has no relation to the truth of C, we can rightly criticize the argument as being bad.

Conclusion follows from the premises Third, arguments are evaluated according to whether the conclusion "follows from" the premises. For example, suppose Tek has a brown dog and that dog is a collie. Now consider the following argument:

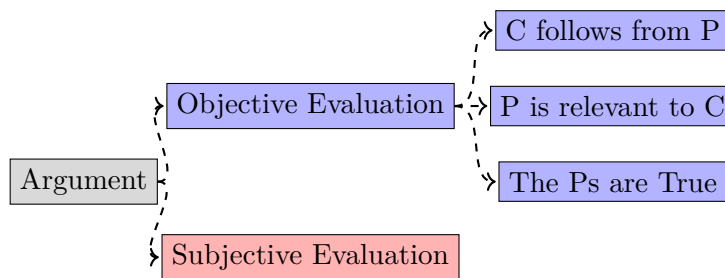
- P1: Tek's dog is brown.
- C: Therefore, Tek's dog is a collie.

Intuitively, while both P1 and C are true, the conclusion does *not* follow from the premises. This is because just because Tek's dog is brown does not mean that Tek's dog is automatically a collie. In contrast, consider the following argument:

- P1: Tek's dog is a collie.
- C: Therefore, Tek has a dog.

In the above example, the conclusion follows from the premises. From the proposition that Tek has a specific type of dog (a collie), it follows that Tek has a dog. While these two examples are straightforward, we might wonder about what exactly it means for a conclusion to "follow from" the premises. One of the major tasks of logic is to clarify the indeterminate idea of a "conclusion following from its premises".

In this section, we considered two general ways of evaluating arguments. First, we considered that arguments may be evaluated subjectively or objectively. Second, we introduced three different ways that arguments might be objectively evaluated. That is, three different objective standards that can be employed when saying an argument is a "good argument" or a "bad argument". The three objective standards are (1) the truth of the propositions, (2) the relevance of the premises to the conclusion, and (3) the conclusion following from the premises.



As a final point, it is worth noting that we have gained some sophistication with respect to saying an argument is "good" or "bad". By introducing three objective criteria for evaluating argument, we might sometimes avoid saying that an argument is "wholly good" or "wholly bad" and instead opt for saying the ways in which the argument is good and the ways in which it is bad. For example, suppose an argument is such that the conclusion follows from the premises. In this way, the argument meets one of the objective standards for evaluating arguments. On this standard, the argument is good. However, if that same argument has false premises, then it fails to meet one of the objective standards for evaluating arguments. On this standard, the argument is bad. In this way, we might say that the argument is "good" in one way and "bad" in another way. Or, suppose an argument has all true propositions and its conclusion follows from the premises, but the conclusion is not relevantly related to the conclusion. In such a case, we might say that the argument is "good" in two ways and "bad" in one way.

1.2.1 *Deductively Valid or Invalid*

In the previous section, three objective standards were cited for evaluating arguments. Let's clarify one of these standards further. When saying that a conclusion "follows from" the premises, it is not entirely clear what is meant. One way to clarify this idea is by saying that a conclusion "follows from" its premises when the argument is *truth-preserving*. Intuitively, an argument is "truth-preserving" if and only if it preserves the truth of its premises in its conclusion. So, in the case of a truth-preserving argument, if the premises are true, the conclusion is true. Let's clarify this idea further in two different ways.

First, an argument can be truth-preserving in that it is **impossible** for the premises to be true and the conclusion to be false. On this sense of truth-preservation, if an argument is truth-preserving, it can never be the case that the premises are true and the conclusion is false. So, necessarily, if the premises are (in fact) true, then the conclusion is true. When an argument is truth-preserving in this sense, the argument is said to be **deductively valid**. In contrast, if an argument fails to be truth-preserving in this sense (that is, it fails to be deductively valid), then the argument is deductively invalid.

Second, an argument can be truth-preserving in that it is **improbable** for the premises to be true and the conclusion to be false. On this sense of truth-preservation, if an argument is truth-preserving, it is unlikely that the premises of the argument are true and the conclusion false. So, it is probably the case that if the premises are true, then the conclusion is true. When an argument is truth-preserving in this sense, the argument is said to be **inductively strong**. In contrast, if an argument fails to be truth-preserving in this sense (that is, it fails to be inductively strong), then the argument is inductively weak.

This text will focus on the former sense of truth preservation (deductive validity) and not the latter (inductive strength). As noted, arguments are "deductively valid" or "deductively invalid". An argument is deductively valid if and only if it is impossible for all of the premises of the argument to be true and the conclusion of the argument to be false.

Definition 1.2.1: Deductive Validity

An argument is deductively valid if and only if it is impossible for all of the premises to be true and the conclusion of the argument to be false.

If an argument is not deductively valid, then it is deductively invalid. An argument is deductively invalid when it is possible for the premises to be true and the conclusion false.

Definition 1.2.2: Deductive Invalidity

An argument is deductively invalid (not valid) if and only if it is not deductively valid.

Before clarifying these definitions further, let's consider some straightforward examples. First, consider the following argument:

- P1: All human beings are mortal.
- P2: Tek is a human being.
- C: Therefore, Tek is mortal.

This argument is deductively valid since it is impossible for the premises (P1 and P2) to be true and the conclusion false. Since we have refined what it means for a conclusion to "follow from" the premises in terms of deductive validity, we can also say that since the argument is deductively valid, the conclusion (C) follows from the premises (P1 and P2). Now let's consider an example of a deductively invalid argument:

- P1: Tek is happy.
- P2: Sal is happy.
- C: Therefore, everyone is happy.

The above argument is deductively invalid. This is because it is *not* impossible for the premises to be true and the conclusion false. That is, it *is* possible for the premises to be true and the conclusion false. To illustrate, it is possible for Tek and Sal to be happy but not everyone to be happy, e.g., suppose Liz is not happy. Similarly, we can say that since the argument is deductively invalid, the conclusion does not follow from the premises.

Exercise 1.4

Consider whether the following arguments are deductively valid or invalid.

1. Some people are friendly. Therefore, all people are friendly.
2. Some plants are not edible. This poison ivy is a plant. Therefore, this poison ivy is edible.
3. All humans are dancers. All dancers are flexible. Therefore,

all humans are flexible.

4. It is not the case that all criminals deserve punishment. Tek is a criminal. Therefore, Tek deserves punishment.
5. Some basketball players are millionaires. Some millionaires drive fancy cars. Therefore, some basketball players drive fancy cars.

1.2.1.1 Further clarifying deductive validity

While the property of deductive validity refines the intuitive idea that an argument's conclusion follows from its premises, it nevertheless lacks clarity in some respects and some individuals find the idea of deductive validity confusing. Let's focus our attention on the latter problem by noting some features of deductively valid arguments and by considering several examples.

Deductive validity is a property of arguments

First, deductive validity is a property of arguments rather than propositions, sentences, or points made in a discussion. In everyday speech, Tek may say to Liz that some proposition or point she raised with respect to a topic is a "valid point". In this case, Tek is using a different sentence of validity. That is, Tek is saying that Liz's point is "a good point" or "a relevant point". Similarly, Liz may tell Tek that a proposition in some argument he raises is a "valid proposition". In this case, Liz does not mean "deductively valid" but instead means "true" or "correct" or "good" proposition. As we have defined the term, "deductive validity" refers to a property of arguments: it concerns the *relation* between the premises and the conclusion. When an argument is truth-preserving in the sense of it being impossible for all of the premises to be true and the conclusion false, then the argument is deductively valid.

Validity concerns whether it is possible for the premises to be true and the conclusion false

Second, except in one case, the *actual* truth and falsity of the propositions that compose an argument are irrelevant to whether or not the argument is valid. What matters instead is whether it is *impossible* or *possible* for the premises to be true and the conclusion false. To illustrate that this is the case, consider that an argument may be valid and its propositions may take any of the following truth values:

1. false premises and a false conclusion
2. true premises and a true conclusion
3. false premises and a true conclusion
4. some true premises, some false premises, and a false conclusion

5. some true premises, some false premises, and a true conclusion

Let's consider some of these. First, an argument can be deductively valid when all of the propositions that compose the argument are false. That is, just because an argument is deductively valid does not mean that its premises are true. Consider the following argument:

Deductively valid argument composed of all false propositions

- P1: All pigs can fly.
- P2: David is a biological pig.
- C: Therefore, David can fly.

In the above argument, P1 and P2 are false. Pigs cannot fly and David (this is me) is not a biological pig. Nevertheless, the argument is valid. Recall that validity concerns the relation between the premises and the conclusion and whether or not it is *possible* for all of the premises to be true and the conclusion is false. Consider what would be the case if the premises (P1 and P2) were true. If the premises were true, then the conclusion would also be true. That is, there is no possibility where the premises are true and the conclusion is false, and so the argument is valid.

Not only are there valid arguments with all false propositions but there are *invalid* arguments consisting of only true propositions. The point to keep in mind is that just because an argument has all true premises does not mean that the argument is valid. To illustrate, suppose it is a hot day. In fact, it is 100 degrees today. Now consider the following argument:

Deductively invalid argument composed of all true propositions

- P1: Red is a color.
- C: Therefore, it is a hundred degrees today.

In the argument above, P1 is true and C is true. However, the argument is invalid. This is because it is possible for P1 to be true and C to be false. That is, it is possible for "red is a color" to be true (because it is a color) and "it is a hundred degrees today" to be false (it could have been a different temperature while "red is a color" is true).

Next, let's consider a valid argument that has false premises but a true conclusion.

- P1: All cats have four legs.
- P2: Danny is my cat.
- C: Therefore, Danny has four legs.

In this example, P1 and P2 are false. Some cats do not have four legs, so P1 is false. In addition, my neighbor's cat (who is named "Danny") is not my cat, so P2 is false. Finally, the conclusion is true: Danny has all four

legs. In determining whether this argument is valid or invalid, notice that it is impossible for P1 and P2 to be true and C to be false. That is, it is impossible for all cats to have four legs, for Danny to be my cat, and for Danny to *not* have four legs. So, the above argument (that has two false premises and a false conclusion) is valid.

Let's consider one more example. Let's consider an argument that has some true premises, some false premises, and a true conclusion.

- P1: Tek is allergic to all cats.
- P2: Danny is my cat.
- C: Therefore, Tek is allergic to my cat.

First, let's suppose that Tek is allergic to every cat. Second, suppose that "Danny is my cat" is false (Danny is my neighbor's cat). Finally, suppose that it is true that Tek is allergic to my cat. In this case, P1 is true, P2 is false, and C is true. In determining whether this argument is valid or invalid, notice that it is impossible for P1 and P2 to be true and C to be false. It must be the case that if Tek is allergic to all cats and Danny is my cat (contrary to what is actually the case), then Tek is allergic to my cat. That is, it is impossible for P1 and P2 to be true and C to be false. Therefore, the above argument is valid.

Exercise 1.5

1. Create a valid argument in standard form.
2. Create an invalid argument in standard form.
3. Create an argument that has all true premises and a true conclusion but is invalid.
4. Create an argument that has all false premises and a false conclusion but is valid.
5. Create an argument that has some true premises, some false premises, and a false conclusion but is valid.
6. Suppose an argument has a conclusion that cannot be false. Is the argument valid or invalid? Explain.
7. Suppose an argument has premises that cannot be false. Is the argument valid or invalid? Explain.

1.2.2 Sound Arguments

Logic is concerned with determining which arguments are good and which arguments are bad. We mentioned that there are three objective criteria

for making this separation. These criteria can be formulated in terms of three questions:

1. Are the propositions in the argument true?
2. Are the premises relevant to the conclusion?
3. Does the conclusion follow from the premises?

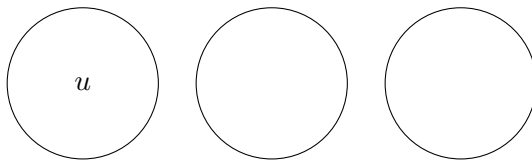
In the previous section, we clarified what it means for a conclusion to follow from the premises with the idea of validity. In addition, we noted that some of an argument's propositions may be false but the argument may still be valid. As such, this means that our answers to the first and third questions above may differ. That is, an argument may be "bad" in that it has false premises but be "good" in that it is valid. But what about an argument that is "good" in both ways, one where the premises of the argument are true and the argument is valid? That is, what if the argument is considered "good" on both ways of evaluating the argument?

An argument that is both deductively valid and has true premises is known as a "sound" argument. A sound argument is an argument where not only is it impossible for the premises to be true and the conclusion false, but it is also the case where the premises (and therefore the conclusion) are, in fact, true.

Definition 1.2.3: Sound

An argument is sound if and only if all of the premises are (in fact) true and it is deductively valid.

Suppose Tek and Liz are looking at a batch of circles. Let's refer to one *Sound argument* of these circles as u :



Now suppose Tek puts forward the following argument:

- P1: No circles are squares.
- P2: u is a circle
- C: Therefore, this circle u is not a square.

Tek has expressed a sound argument in that P2 is true because we said that u is the case, P1 is true in that it is true that no circles are squares, and the argument is valid since it is impossible for P1 and P2 to be

true and the conclusion false. Given the truth of P1 and P2, and the argument's validity, the argument is *sound*.

Definition 1.2.4: Unsound

An argument is unsound if and only if the argument is either (1) deductively invalid or (2) deductively valid yet has at least one false premise, or both.

Unsound but valid argument

What is sometimes perplexing is that an argument can be deductively valid yet unsound. Consider the following example:

- P1: Either Jennifer Lopez or Mario Lopez is the president of the United State of America (USA)
- P2: Mario Lopez is not the president of USA.
- C: Therefore, Jennifer Lopez is the president of USA.

Notice that premise (P1) is false. It is not the case that either Jennifer Lopez or Mario Lopez is the president of the USA. Nevertheless, if the premises (P1) and (P2) were true, would the conclusion also be true? The answer is a resounding Yes! It is impossible for (P1) and (P2) to be true while (C) is false. That is, it is necessarily the case that if (P1) and (P2) are true, then (C) is true. Thus, the argument is deductively valid yet unsound.

Exercise 1.6

1. What does it mean to say that an argument is sound?
2. Can an argument be sound if it is not valid?
3. Can an argument be sound if it has at least one false premise?
4. If the premises of an argument are true and the argument is deductively valid, does this mean that the conclusion is also true?

1.3 TESTING FOR DEDUCTIVE VALIDITY

In the previous section, we defined what it *means* for a conclusion to follow from a set of premises. One way we clarified this idea was in terms of truth-preservation and this idea was clarified in terms of the idea of deductive validity. However, note that it is one thing to be able to define a property of something and another to be able to *determine* whether something has that property. For example, a new pawn shop owner may know that "gold" is defined as a yellow, malleable, and precious metal,

having the chemical symbol Au, and atomic number 79, but the pawn shop owner may not know whether a particular item is gold. Similarly, we can define what it means for a number to be prime (a number is prime if and only if it is divisible by only itself and 1), but fail to know how to determine whether a number is prime (especially larger numbers).

With this in mind, while we have a definition of deductive validity, we may not be able to determine for any given argument, whether or not it is valid. What we want then is not merely to know what it *means* for an argument to be valid, but also to know how to *determine* whether an argument is valid. In this section, we will consider two different informal methods for determining or "testing" whether an argument is deductively valid. The first method is known as the **logical intuition test** and the second method is known as the **logical imagination test**.

1.3.1 *The logical intuition test*

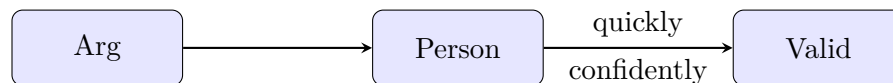
Perhaps the most common method used to determine whether an argument is deductively valid involves an appeal to a power of logical intuition. The power of logical intuition (also referred to as "logical perception", "logical sense", or even crudely as "one's gut feeling about an argument") is a supposed power found in human beings that allows them to directly evaluate arguments as being good or bad. For our purposes, we will call the evaluation of an argument through an appeal to one's logical intuition, the "logical intuition test for validity" (or "intuition test" for short).

How does this test work? To test whether an argument is deductively valid, simply use your power of intuition to intellectually "see" whether the argument is valid or invalid. To use your power of intuition, read the argument and then check whether you get a feeling that the argument is valid or invalid. If you get the feeling that the argument is valid, then it is valid. If you get the feeling that it is invalid, then it is invalid. The method does not require you to engage in conscious, deliberate thought nor does it require you to engage in a step-by-step method. Instead, it assumes that you have a power that allows you to *directly know* whether an argument is valid or invalid. Another way of explaining this test is to consider the fact that people often refer to having "gut feelings". You hear expressions like "I'm going with my gut on this one" or "I'm going to trust my gut". The logical intuition test works by simply asking you to trust your gut when evaluating an argument.

Definition 1.3.1: Logical Intuition Test

The intuition test works as follows: a subject looks at an argument and then intuits (immediately judges) the argument to be valid or invalid.

There are several arguments that we have an intuitive power that directly knows whether an argument is good or bad. Let's consider a relatively simple argument for its existence. Suppose Tek firmly believes that Liz should be the next President. Now suppose Tek overhears someone put forward an argument that concludes that Liz should be the next President. In general, those in a similar situation to Tek will have typically have two responses. First, they will quickly tell you whether the argument is good or bad. Second, the person will also be confident that their evaluation of the argument is correct.



What could possibly explain the fact that people are capable of quickly deciding the quality of an argument with such confidence? That is, what best explains the *speed* and *confidence* with which individuals can judge the quality of an argument? It isn't the fact that they reason to these judgments since reasoning often takes time and is sometimes accompanied by doubt concerning whether one has reasoned correctly. A better explanation then is that people have an intuitive power that allows them to immediately and directly know whether an argument is valid or invalid by simply looking at the argument. The *speed* with which they judged the quality of the argument suggests that they evaluated the argument without any conscious thought or self-controlled reasoning, while the *confidence* in their assessment suggests that they judged using a method that is immune to error (a kind of direct knowing). Since the existence of an intuitive power would best explain these facts, it follows that we must have a power that allows us to immediately and directly know that an argument is valid or invalid. Let's put our argument in standard form:

- P1: People are able to quickly and confidently judge whether an argument is valid or invalid.
- P2: The best explanation for the speed and confidence with which people judge the quality of an argument is that they have a power of logical intuition.
- C: Therefore, people have a power of logical intuition.

In sum, the argument above contends that if positing the existence of something best explains some experience or fact in the world, there is justification for believing that the thing exists. In the case of the intuition test, the fact that people are able to quickly and confidently judge the quality of an argument is best explained by the existence of a power of logical intuition. Thus, there is justification for believing that such a power exists.

1.3.2 Problems with the intuition test

There are several problems with the logical intuition test. First, the intuition method produces verifiably incorrect results. Suppose an argument is presented about some political topic. Tek may say that Tek's power of intuition says this argument is valid while Liz may say that Liz's power of intuition says the argument is invalid. Since an argument can only be valid or invalid (not neither and not both), one of these individuals is incorrect. The use of the logical intuition test is thus problematic since the disagreement between Tek and Liz is evidence that the test has, at least for one of them, produced an incorrect result. Let's consider two illustrations. First, consider the following question:

Problem 1: Incorrect Results.

A bat and a ball cost \$1.10. The bat costs one dollar more than the ball. How much does the ball cost?

As has been reported by Kahneman[5], many individuals answer this question by saying that the ball costs 10 cents. Presumably, the reasoning of these individuals is as follows:

- P1: The bat and ball costs \$1.10.
- P2: The bat costs one dollar more than the ball.
- P3: $\$1.10 - \$1.00 = \$0.10$
- C: Therefore, the ball costs \$0.10

But this argument is invalid. It is possible for the premises of the argument to be true and the conclusion false. In fact, the conclusion is false since ball costs \$0.05. Note that if the ball costs \$0.10 and the bat is one dollar more than the ball, then the bat is \$1.10. This cannot be the case since it would mean that the bat and the ball cost \$1.20 rather than \$1.10. In sum, if there were a power of intuition, we would not expect so many people to judge the above argument valid when it is invalid.

Let's consider a second illustration. Consider the following argument (one that I have given to college students for over ten years):

- P1: Some basketball players are millionaires.

- P2: Some millionaires drive fancy cars.
- C: Therefore, some basketball players drive fancy cars.

After learning the concept of validity, I ask students to review the above argument and write down whether the above argument is deductively valid (impossible for P1 and P2 to both be true and C to be false). Roughly 85% of students judge the above argument to be valid, 10% think it is invalid, and the remaining 5% are unsure. Since the argument is either valid or invalid (not both and not neither), the fact that students disagree about the validity of the argument is evidence that the logical intuition test has produced incorrect results.

Problem 2: Limited scope. Second, the intuition test appears to only be able to apply to a limited number of arguments. That is, it has limited applicability or scope. While individuals may use the logical intuition test when presented with simple arguments on familiar topics (e.g., politics, family life, sports), their power of intuition appears not to work on arguments involving complex, technical subjects that they have limited prior knowledge. In other words, our logical intuition seems to activate when individuals are presented with arguments on familiar topics but not activate when considering unsolved problems in mathematics (e.g. whether $P = NP$ in computer science). The fact that the logical intuition test only applies to arguments discussing familiar topics is problematic since it has limited use.

Problem 3: Large arguments. Third, the intuition test appears to only be able to apply to arguments that are relatively small in size. That is, it has limited applicability or scope. While individuals may use the logical intuition test when presented with simple arguments, their power of intuition appears not to work on arguments that have many premises and many intermediate conclusions. In other words, our logical intuition seems to activate when individuals are presented with arguments that are short but not activate when considering arguments that are long. The fact that the logical intuition test only applies to arguments of a certain (small) size is problematic since it means the test only has limited applicability.

Problem 4: No power of intuition. Fourth, in the previous section, we considered an argument in favor of the power of intuition. This argument noted that positing the existence of a power of intuition best explains the *speed* and *confidence* with which people judge whether a conclusion follows from its premises. However, this argument would not be effective if there were a better explanation that does not posit a power of intuition. And, there is a better explanation. Consider that many of our beliefs are due to education, culture, repeated experience, and evolution. If individuals are repeatedly educated

that "democracy is the best form of government" or "capitalism is the best economic system", then many (but not all) people are likely to take arguments that support these beliefs to be good arguments, while arguments that conflict with these beliefs are taken to be bad arguments. Similarly, if you have repeated experience of burning your hand on a hot stove, many individuals are likely to judge arguments that support the conclusion that "touching a hot stove is painful" to be good, and those that run contrary to this conclusion to be bad.

What we have done here however is explain the speed and confidence with which people judge arguments to be the result of education, repeated experience, habituation, and evolution rather than a special, internal logical power of intuition. It is instead the result of comparing the argument (or what the argument aims to establish) with our preexisting stock of beliefs, recognizing that the argument and our stock of beliefs conflict, and thus rejecting the argument (or its conclusion) outright. And so, while a supporter of a power of logical intuition says that such a power allows humans to quickly and confidently judge arguments as valid or invalid, if we can account for these results without positing the existence of such a power, there is reason to think that no such power exists.

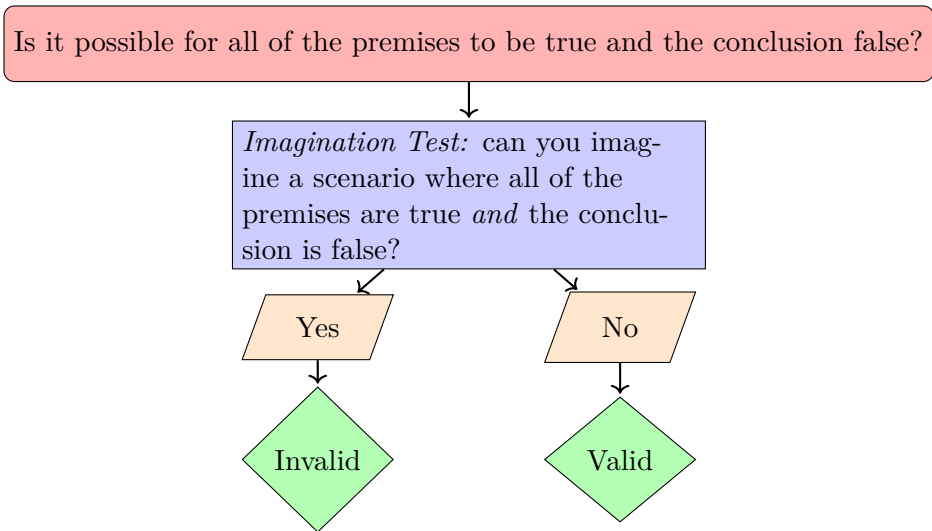
This concludes our discussion of the logical intuition test. We now turn to the second test: the imagination test.

1.3.3 The imagination test

In the previous section, several shortcomings were identified concerning the logical intuition test. In this section, let's consider a second way of determining whether an argument is valid or invalid. This test is known as the "human-imagination test" or the imagination test for short. Here is how the imagination test works. Take an argument and try to imagine a scenario where the premises are true and the conclusion is false. If you can, then the argument is invalid. If you cannot, then the argument is valid.

Key to the imagination test is the capacity to imagine or picture or think of various hypothetical or possible scenarios. For example, imagine an argument with two premises $P1$ and $P2$ and a conclusion C . To use the imagination test, we try to create a scenario where $P1$ and $P2$ are both true and C is false. If we can imagine such a scenario, then the argument is invalid. If we cannot imagine such a scenario, then the argument is valid.

The logic behind the imagination test and the concept of validity involves



connecting human imagination to logical possibility. If we are *able* to *imagine* the premises being true and the conclusion being false, then it is *possible* for the premises to be true and the conclusion false. This is because anything that we can imagine is possible (even if it is not actual). And, if it is *possible* for the premises to be true and the conclusion false, then the argument is invalid. In contrast, if we are unable to *imagine* the premises being true and the conclusion being false, then it is *not possible* for the premises to be true and the conclusion false. The logic here is that anything that we cannot imagine must not be possible (that is, it must be impossible). And if it is *not possible* (impossible) for the premises to be true and the conclusion false, then the argument is valid.

Now that we have a sense of how the test works in the abstract and the central rationale underlying the test, let's illustrate this test with some examples. First, let's consider an argument where someone reasons that John plays basketball because John is tall.

- P1: John is tall.
- C: Therefore, John plays basketball.

In applying the imagination test, we try to imagine a scenario where P1 is true and C is false. This is simple enough since we can imagine that John is tall but does not play sports: a scenario that makes P1 true and C false. Since we can imagine the premise being true and the conclusion false, the above argument is invalid. Let's consider another example.

- P1: Liz is a great philosopher.

- P2: Liz is a great politician with plenty of experience.
- C: Therefore, Liz will be the next president of the United States of America.

In applying the imagination test, we try to imagine a scenario where P1 and P2 are both true and C is false. We can imagine such a scenario since we can imagine that Liz is a great philosopher and a great politician with plenty of experience but does not become the next president of the United States of America. For example, suppose she decides not to run for president and focuses on her philosophical career. Or, perhaps Liz does run for president but does not win since the voters find another candidate more persuasive. Since we can create a scenario where the premises are true and the conclusion is false, the above argument is invalid.

Let's consider a third example.

- P1: Either Jennifer Lopez is a great tennis player or Mario Lopez is a great tennis player.
- P2: Mario Lopez is not a great tennis player.
- C: Therefore, Jennifer Lopez is a great tennis player.

If an argument is valid or invalid, we need to know whether it is possible for all of the premises to be true and the conclusion false. We determine this using the imagination test by trying to imagine a scenario where all of the premises are true and the conclusion false. In the above example, we can imagine various scenarios where P1 and P2 would be true but there is no such scenario where the conclusion is also false. That is, if P1 is true, then either Jennifer Lopez or Mario Lopez (or both) are great tennis players. If we take P2 also to be true, then Mario Lopez is not a great tennis player. But, if Mario Lopez is not a great tennis player, then the only scenario where P1 is true is the one where Jennifer Lopez is great tennis player. But, imagining the conclusion of the above argument to be false would require us to imagine that Jennifer Lopez is *not* a great tennis player. In short, imagining the premises to be true and the conclusion to be false requires us to imagine that Jennifer Lopez is both a great tennis player and not a great tennis player. But, this is impossible. Since we cannot imagine a scenario where all of the premises are true and the conclusion false, the above argument is valid.

Exercise 1.7

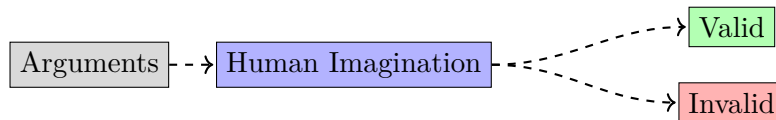
Determine whether the following arguments are valid or invalid.

1. All pigs fly. Babe is a pig. Therefore, Babe flies.

2. Some people are happy. Other people are not happy. Therefore, everyone is happy.
3. God exists or he doesn't. Well, I see no good reason for thinking God exists. Therefore, God doesn't exist.
4. Some gamblers are profitable. Some gamblers lose money. Therefore, every gambler is either profitable or lose money.
5. Some smokers are not happy. All happy people are tennis players. Therefore, some smokers are tennis players.

1.3.4 Problems with the Imagination Test

In the previous section, the imagination test was described and illustrated with several examples. In short, this test works by asking human beings to consider an argument, try to imagine the premises true and the conclusion false, and if they can imagine such a scenario, then the argument is invalid. If they cannot imagine such a scenario, then the argument is valid. In this section, several problems with this method are considered.



In applying the test to these examples, the test seems to produce the right results. We might draw the conclusion that this is all logicians and people in general need to determine whether a conclusion follows from a set of premises. More cautiously, we might say that our results suggest that whenever humans are given short, uncontroversial, and clearly-expressed arguments, the “imagination test” will correctly determine whether the argument is valid or invalid.

Unfortunately, the imagination test is highly problematic for several reasons. The fundamental problem is that the test equates what users of the test can imagine with what is logically possible. This assumption is problematic since there may be things that are possible that some individuals either cannot or will not imagine. This creates a scenario where individuals just an argument to be valid (since they believe there is no possible scenario where the premises are true and the conclusion is false) when in fact the argument is invalid since there is a possible scenario where the premises are true and the conclusion is false. In what follows, several problems with the imagination test are considered.

1.3.4.1 *Incorrect Results.*

One problem logical intuition test was that it produced incorrect results. That is, the test produced results where individuals disagreed about whether an argument was valid or invalid. Since arguments are either valid or invalid (not both and not neither), the logical intuition test was clearly not working for at least some contingent of people. Similarly, the imagination test produces incorrect results. Recall an argument that we considered earlier:

- P1: Some basketball players are millionaires.
- P2: Some millionaires drive fancy cars.
- C: Therefore, some basketball players drive fancy cars.

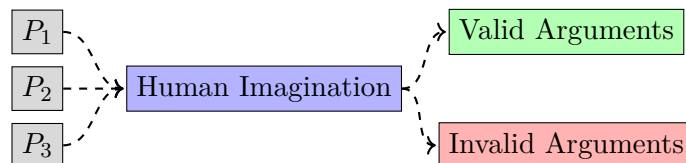
After teaching students the concept of validity, the logical intuition test, and the logical imagination test, I ask students to use the logical imagination test to determine whether the above argument is valid or invalid. Roughly 50% of students take the above argument to be valid, while the other half takes it to be invalid. The fact that one of these groups is wrong suggests either (1) students do not understand the concept of validity, or (2) students do not understand how they are supposed to use the logical imagination test, or (3) the test itself is defective insofar as even a correct understanding of the test does not produce correct results. There is good reason to believe that (3) is the case since students correctly apply the test in several other cases (like those considered above).

1.3.4.2 *Too Many Premises*

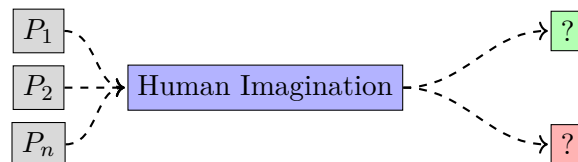
As its name implies, the imagination test requires the use of the imagination: the ability to construct various scenarios. However, consider that our power of imagination is not perfect. For example, imagine a horse. When you imagine a horse, you likely have a vague image in your mind. You can picture the horse's head and eyes and hair, but you do not picture every minute detail of the horse. For each hair on the horse, you likely don't have a picture in your mind that contains information about the hue of each color. If the horse is perspiring, it is unlikely that you are picturing each bead of sweat. The point then is that your power of imagination is limited by how much information you can imagine at a given time. What does this imply for arguments?

Let's take an argument that involves three simple propositions. Let's refer to them as P1, P2, and C. If we were to apply the imagination test to this argument, presumably the test would work correctly since the scenario we

are required to imagine is relatively simple.



On the other hand, consider a large argument with not merely one or two propositions but a large number (P_n) of propositions (where n is greater than 100). In such an argument, individuals might contend that they cannot imagine a scenario where the premises are true and the conclusion is false. But from the fact that we cannot imagine such a scenario does not mean that the argument is invalid. Similar to the case involving the horse where every detail of the horse cannot be imagined, we may also be unable to imagine all of the premises being true because there are too many premises to imagine.



In short, the imagination test for validity depends upon the limited powers of human beings to imagine scenarios. While such a test might be useful for simple arguments, this method fails when it asks its users to imagine large, detailed arguments.

1.3.4.3 Bias

When a person uses the imagination test, they are asked to imagine a scenario where the premises are true and the conclusion is false. However, human beings are not perfect users of the test. Even for a relatively simple arguments, people sometimes contend that they cannot imagine a scenario where the premises are true and the conclusion is false when such a scenario exists. One explanation is that the person has a bias in favor of the conclusion. This bias interferes with their capacity to imagine a scenario where the premises are true and the conclusion is false.

A *bias* is a tendency toward or against something or someone. Many biases have good outcomes. For example, a bias against smoking is good for your health. Other biases are negative. For example, being biased against people of a certain ethnicity is morally wrong. A cognitive bias is a specific type of bias; it is a systematic error of thinking that is due

to how the mind manages and processes information. In plainer terms, a cognitive bias is a mistaken tendency toward (or against) thinking certain things.

There are many different cognitive biases. To illustrate just one of these, *Forer effect* consider the cognitive bias known as the “Forer effect”. The Forer effect is a cognitive bias where people tend to incorrectly overestimate the accuracy of a personality description when they are told that (1) the description was tailored specifically to them and (2) when it is given by a person who is believed to be an authority. For example, suppose I (the author) tell you (the reader) that the following description is a personality sketch of you:

You can be extroverted in situations where you feel comfortable (e.g., around close friends or family members), but are also introverted (even shy) when you are in stressful situations. You are naturally curious by things that interest you, but can be bored by overly dry topics or those that don’t seem to have much practical value. You have a good sense of humor and this sense of humor shines when you are around those with whom you are close. You value these close relationships, the comfort and stability these relationships bring you, and (while you do not always show it) you genuinely appreciate the fact that you can be your true self around these people.

If you took the above description to be a personality sketch of you, then you have fallen victim to the Forer effect. In a classic experiment, 39 students were given a personality test that was said to evaluate their personality. After completing a test, each student was then given a “personality sketch” (similar to the sketch above) that they were told was tailored specifically to them. The students were asked to rate the personality sketch in terms of how accurately it described them. In general, students rated the accuracy of the test as 4.05 on a scale of 0 (poor) to 5 (perfect). It was later revealed to the students that all the personal sketches were identical.

Once students discovered that the personality descriptions were identical, they realized the test was not as accurate as they thought. In short, students who overestimated the accuracy of the personality sketch were victims of the Forer effect.

Thus far, we have defined "bias" and "cognitive bias" and then illustrated the idea of cognitive bias with an example (the Forer effect). However, *Cognitive bias and the imagination test.* what do cognitive biases have to do with the imagination test? Since the

imagination test requires imagining scenarios where the premises are true and the conclusion is false, it requires that human beings are free from cognitive biases which would lead them to ignore certain scenarios. In particular, the successful use of the imagination test assumes that people do not ignore scenarios that would show the argument to be invalid. For example, suppose Tek reads an argument that is in perfect accord with his political beliefs. He may say to himself, “ah! I have applied the imagination test and I simply cannot imagine a scenario where the premises of this argument are true and the conclusion is false. The argument is therefore valid.” Tek’s evaluation is good if there is no such scenario, but perhaps he simply ignored it. What if such a scenario exists but Tek is prone to a cognitive bias that keeps him from thinking about such a scenario? If such a bias existed, then it would be a serious problem for the imagination since it would mean that there is a tendency for people to use the imagination test incorrectly. The question then is the following:

Is there a bias that interferes with our capacity to imagine certain scenarios?

While there does not appear to be direct empirical evidence evaluating the relation of bias to the imagination test, there is empirical evidence that indicates certain biases interfere with our capacity to consider possibilities that would confirm that an argument is invalid. Consider the following definition of “confirmation bias”:

Definition 1.3.2: Confirmation Bias

Confirmation bias is the tendency to look for information that confirms our beliefs and avoid information that might disconfirm them.

Confirmation bias

Confirmation bias exists in human beings and evidence for its existence is found in many different experiments. An example of confirmation bias is found in the beliefs you have about yourself, viz., beliefs you have about what type of person you are. For example, Swann and Read[8, p. 352] have noted that we tend to remember and seek out information that *confirms* the beliefs we have about our self and avoid situations and information that would undermine these beliefs. For example, if we believe ourselves to be a dominant person, we tend to recall conversations where we spoke in an authoritative manner and ignore occasions where we acted timidly. In addition, it is claimed that we will also seek out social interactions where we can act in a dominant manner to further confirm our self-conception. In short, confirmation bias clouds our ability to accurately assess our own personality.

Confirmation bias is not only present in those deeply held beliefs we have about ourselves but also when we are simply given some new hypothesis to try out. To illustrate, suppose you are asked to determine if a person is an extrovert by asking a random person a series of yes / no questions. In this case, you do not know whether the person is, in fact, an extrovert. Trying to discover whether they are an extrovert is simply a working hypothesis. Studies show that when people are asked to test this hypothesis, people will ask questions in a way that aim to *confirm* rather than disconfirm that the person is an extrovert. In short, confirmation bias is present even when we are simply given a working hypothesis (test to see if this person is an extrovert) to try out.

Earlier, we saw that confirmation bias refers to the tendency to look for information that confirms our beliefs and avoid information that might disconfirm them. Confirmation bias thus has the potential to cause problems for users of the imagination test since it implies that if an individual thinks an argument is “good”, then they tend to avoid considering scenarios that would show that the argument “bad”. And, one way an argument may be bad is if it is invalid. In other words, confirmation bias may interfere with a person’s imagination test by causing them to avoid thinking of scenarios that would show the argument to be invalid.

One possible rebuttal is the following. Yes, people are subject to confirmation bias, but the imagination test corrects this bias. For consider that the imagination test actively tells us to engage in a type of deliberation about arguments by trying to think of scenarios where the premises are true and the conclusion is false. Perhaps then, the mere fact that we deliberate about arguments when using the imagination test protects us from confirmation bias. So the question then is whether confirmation bias is mitigated (lessened) by reflection or deliberation about a topic or argument. Surprisingly, research has shown that increased reflection or deliberation about a topic *promotes* rather than mitigates confirmation bias. One study[2] separated college-age students into two groups:

- sleep-restricted
- well-rested

The study additionally sorted the groups based on their political preference (liberal or conservative) and then asked them to select six arguments from either conservative or liberal sources on the topic of gun-control as well as rate arguments “for” and “against” gun-control. The subjects were also:

- asked to rate how much they deliberated on each argument and

- took a short cognitive reflection test to measure the degree to which they deliberated on the arguments.

This study found that:

deliberation promotes a stronger confirmation bias. Participants who had thought more about gun control were found to have a more precisely estimated confirmation bias regarding selective information exposure and also a stronger magnitude of effect regarding perceived argument strength.[2]

In other words, subjects who thought more about gun control were more likely to select sources that supported their beliefs and ignore sources that challenged them. What this shows then is that the further deliberation or reflection about a topic does not correct confirmation bias.

In sum, if human beings have confirmation bias and deliberation via the imagination test does not remove this bias, then there is reason to think that confirmation bias undermines our ability to successfully use the imagination test for determining if an argument is deductively valid. More precisely, we can make two claims about the relation of confirmation bias to the imagination test. First, if confirmation bias is extended to the imagination test, then biased individuals may tend to incorrectly evaluate arguments as being valid since they will ignore possibilities that undermine the argument's validity. Second, there is no reason that the imagination test will correct this problem given that confirmation bias is shown to be made worse by extended deliberation. In other words, biased individuals will take arguments that support their view as valid, ignoring disconfirming possibilities, and the act of further imagining possibilities will not correct this error.

Exercise 1.8

1. What are the problems associated with the logical intuition test for deductive validity?
2. What are some problems associated with the human imagination test for validity? Can any of these problems be fixed or corrected?
3. Create a scenario where a person uses the imagination test but nevertheless judges an invalid argument to be deductively valid. Explain why the person made this mistake.

Part II

Propositional Logic

In the previous chapter, some key concepts of logic were defined, e.g., argument, validity, and soundness. In addition, two different methods for testing whether an argument is valid were introduced and shown to be problematic: the intuition test and the imagination test. One way to potentially remedy these problems is to construct a formal (symbolic) language — "the language of propositional logic" (**PL**) — and then use this language to develop alternative methods for testing whether an argument is valid. The focus of this chapter is on the construction of the formal language, while the focus of subsequent chapters will be on how to use this language to test whether an argument is valid.

Our plan for introducing this language is as follows. First, we will introduce the symbols of the language. Second, we will introduce the syntax (grammar) of the language. Third, we will introduce the semantics of the language. Fourth and finally, we will introduce a method for translating English sentences into the language.

In addition, a semantics for *PL* is given. The semantics for *PL* is not, strictly speaking, part of the formal language. It is, instead, an *interpretation* of the formal language. Roughly put, it is a specification of what the different symbols and formulas of the language mean.

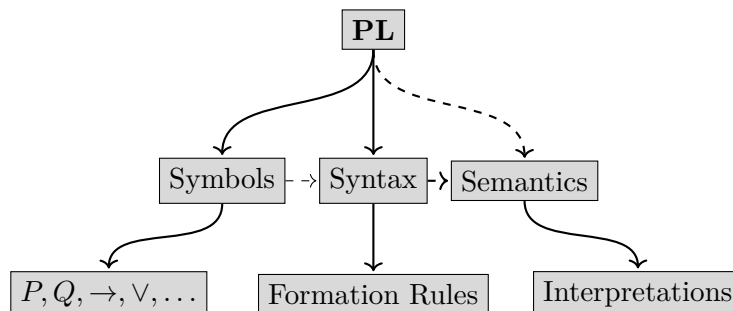


Figure 2.1: A formal language consists of a set of symbols and a syntax. The semantics of a formal language concerns the interpretation of the symbols and/or well-formed formulas of that language.

2.1 PL SYMBOLS

Let's begin our introduction to the language of propositional logic (PL) by stating the symbols (or characters) that compose the language. The symbols of *PL* are divided into the following three main types:

1. an infinite number of "propositional letters": uppercase Roman (unbolded) letters with or without subscripted integers, e.g. $A_1, A_2, A_3, B, C, \dots, Z$.
2. five truth-functional operators: $\vee, \rightarrow, \leftrightarrow, \neg, \wedge$
3. a left and right parenthesis: "(" and ")"

Let's briefly discuss each one of these symbols. First, PL contains a set of propositional letters. These are uppercase Roman letters. For example, A is a propositional letter, B is a propositional letter, and Z is a propositional letter. In contrast p is not a propositional letter because it is a lowercase letter. Similarly, ϕ is not a propositional letter since it is not one of the 26 Roman letters. In order to have an infinite number of propositional letters, we subscript the propositional letters with integers. For example, A_1 is a propositional letter, A_2 is a propositional letter, A_3 is a propositional letter, and so on. In contrast, A_x is not a propositional letter since it is not subscripted with an integer.

The second type of PL symbol is a truth-functional operator. There are five truth-functional operators: $\neg, \wedge, \vee, \rightarrow, \leftrightarrow$. If you are new to logic, then these symbols are likely unfamiliar and so it is unclear what to call them. Here are the five operators along with their names:

In some logic texts, a tilde is used (the tilde is the top half of an "n", for "not")

1. $\neg$, "not" or "negation"
2. $\wedge$, "wedge" or "and"
3. $\vee$, "vee" or "or"
4. $\rightarrow$, "rightarrow"
5. $\leftrightarrow$, "doublearrow" or "if and only if" or "iff"

The third main type of symbol is a left and right parenthesis: (and). These symbols are used to indicate the scope of truth-functional operators (more on this later). What is important to know now is that they are symbols of *PL*.

Let's add some text here to see if that is the problem.

Exercise 2.9

For each symbol below, indicate its type, viz., propositional letter, operator, etc.

1. P
2. Q
3. $\wedge$
4. $\neg$

5. (
6. $\leftrightarrow$

2.2 PL SYNTAX

In the previous section, the symbols of PL were presented. Now that we have the symbols, we want to know the right and wrong ways that these symbols can be combined. The rules that specify the right and wrong ways that the PL symbols can be combined is its syntax (or grammar). The syntax of PL consists of a set of rules known as "formation rules". When the symbols of PL are combined together according to the formation rules, the resulting combination of symbols is called a well-formed formula (abbreviated as "wff", pronounced "woof").

Definition 2.2.1: Well-formed Formula in PL

A well-formed formula (or wff) ϕ is any formula that is capable of being generated by some combination of the seven formation rules in [Definition 2.2.2](#).

Key to the definition of a wff are the formation rules. These formation rules define what is and is not a wff in PL .

Definition 2.2.2: Formation rules of PL

Let ϕ and ψ are variables for well-formed formulas in PL :

1. Every propositional letter of PL (e.g. A, B, C) is a wff.
2. If ϕ is a wff, then $\neg(\phi)$ is a wff.
3. If ϕ and ψ are wffs, then $(\phi \wedge \psi)$ is a wff.
4. If ϕ and ψ are wffs, then $(\phi \vee \psi)$ is a wff.
5. If ϕ and ψ are wffs, then $(\phi \rightarrow \psi)$ is a wff.
6. If ϕ and ψ are wffs, then $(\phi \leftrightarrow \psi)$ is a wff.
7. Nothing else is a wff except what can be formed by repeated applications of 1–6.

Let's explain each of these seven formation rules. First, rule (1) states that every propositional letter of PL is a wff. This means that by each propositional letter, all on its own, is a wff. For example, P is a wff, Q is a wff, R is a wff, and so on. Rule (2) states that if ϕ is a wff, then so is $\neg(\phi)$. In this rule, ϕ stands for any combination of PL symbols that has the status of being a wff. So, for example, since P is a wff by rule (1), then so is $\neg(P)$. Similarly, since Q is a wff by rule (2), then so is $\neg(Q)$.

The syntax defines which combinations of PL characters are legal combinations.

More compactly: if ϕ_1 and ψ are wffs, then so are $\neg(\phi)$, $(\phi \wedge \psi)$, $(\phi \vee \psi)$, $(\phi \rightarrow \psi)$, $(\phi \leftrightarrow \psi)$

Rule (2) is not saying that only singly negated propositional letters are wffs. Rather, it is saying that for any combination of symbols ϕ , if ϕ is a wff, then taking that combination of symbols, putting parentheses around it, and the negation to the left of the entire combination of symbols is a wff.

Since P is a wff, so is $\neg(P)$. And since $\neg(P)$ is a wff, so is $\neg(\neg(P))$. Since $\neg(\neg(P))$ is a wff, ..., well, you see where this is going.

Notice that rules (3)-(6) all begin the same way: "if ϕ and ψ are wffs, then ...". These four rules allow for combining (or "connecting") two wffs together with the connectives: $\wedge, \vee, \rightarrow, \leftrightarrow$. Rule (3) states that if ϕ and ψ are wffs, then the $\wedge$ can be placed between both of these wffs and parentheses put around the entire combination and the result is a wff. In short, if ϕ, ψ are both wffs, then so is $(\phi \wedge \psi)$. So, for example, since P is a wff and $\neg(Q)$ is a wff, $(P \wedge \neg(Q))$ is a wff.

Rules (4)-(6) follow the same format as rule (3), they serve to connect two wffs ϕ and ψ but with different operators. Rule (4) contends that if ϕ and ψ are wffs, then so is $(\phi \vee \psi)$. Notice that ϕ and ψ are connected using $\vee$ instead of $\wedge$. Similarly, in rule (5) ϕ and ψ are connected using $\rightarrow$; thus $(\phi \rightarrow \psi)$ is a wff if ϕ and ψ are wffs. Finally, in rule (6) ϕ and ψ are connected using $\leftrightarrow$. That is, $(\phi \leftrightarrow \psi)$ is a wff if ϕ and ψ are wffs.

Here are some *PL* wffs: $P, (A \rightarrow \neg(B)), \neg((P \vee \neg(Q))), \neg(\neg((P \wedge \neg(Q))))$

Here are combinations of *PL* symbols that are not wffs: $P\neg, PQ, \vee\neg(Q), \neg\neg P \wedge \neg Q \vee S$

The final rule is rule (7). This rule states that the only combinations of symbols that are wffs are those that can be constructed using rules (1)-(6). In other words, if a combination of symbols cannot be constructed using rules (1)-(6), then it is not a wff. For example, $P\neg$ is not a wff since it cannot be constructed using rules (1)-(6): there is no rule that allows you to place a negation to the right of a wff. Similarly, PQ is not a wff since it cannot be constructed using rules (1)-(6), as there is no rule that allows for placing two propositional letters immediately next to each other.

Now that we have defined a well-formed formula and explained the basics behind formation rules, let's consider some examples of how formation rules can be used to construct wffs. Let's start by trying to show that $(\neg(P) \rightarrow R)$ is a wff:

1. Every propositional letter is a wff, so P and R are wffs (rule 1).
2. If P is a wff, then $\neg(P)$ is a wff. (line 1 + rule 2)
3. If $\neg(P)$ and R are wffs, then $(\neg(P) \rightarrow R)$ is a wff. (lines 1,2 + rule 5)
4. $(\neg(P) \rightarrow R)$ is a wff. (lines 1-3)

Notice how the wff $(\neg(P) \rightarrow R)$ was constructed using the formation rules. The formula was constructed or "built up" from the propositional letters

(line 1) and repeated use of the rules that allow for adding operators (see line 2) or connecting wffs with connectives (see line 3).

Let's consider one more example. Here we will show that $(P \rightarrow (P \vee Q))$ is a wff.

1. Every propositional letter is a wff, so P and Q are wffs (rule 1).
2. Since P and Q are wffs, then $(P \vee Q)$ is a wff (line 1 + rule 4).
3. Since P and $(P \vee Q)$ are wffs, then $(P \rightarrow (P \vee Q))$ is a wff (lines 1,2 + rule 5).
4. $(P \rightarrow (P \vee Q))$ is a wff (lines 1-3).

Again, the wff is constructed first by starting with the propositional letters (line 1) and then adding operators (or connectives) to formulas that are already established as wffs (see lines 2 and 3).

Exercise 2.10

Use the formation rules to show that the following formulas are wffs.

1. $\neg(A)$
2. $\neg(\neg(B))$
3. $(A \wedge \neg(B))$
4. $(A \rightarrow \neg(B))$
5. $((A \vee B) \wedge (B \rightarrow C))$

2.2.1 Variables

In presenting the language of propositional logic, it is helpful to use variables for propositional letters or wffs. Generally, the use of the Greek letters ϕ ("phi") and ψ ("psi") are used. If we let ϕ stand for any propositional wff, then ϕ is a placeholder for any wff that can be constructed using the propositional logic formation rules. In other words, ϕ can be P , Q , $(P \wedge Q)$, $(\neg(P) \rightarrow Q)$ and so on.

Sometimes it is convenient to use variables and operators together. For example, if we let ϕ stand for any wff and $\neg$ stand for the negation operator, then $\neg(\phi)$ is a placeholder for any wff that can be constructed by placing a negation operator to the left of a wff. In other words, $\neg(\phi)$ can be $\neg(P)$, $\neg(Q)$, $\neg((P \wedge Q))$, $\neg((\neg(P) \rightarrow Q))$ and so on. Here is another example. Suppose we let ϕ and ψ serve as variables for any wff in **PL**. We might then want to talk about any wff that can be constructed by placing a $\wedge$ between ϕ and ψ . That is, we don't want to talk simply about a

specific wff $P \wedge Q$ but any wff where $\wedge$ is placed between two wffs ϕ and ψ . That is, rather than talking about the specific $P \wedge Q$, we want to talk about $(\phi \wedge \psi)$. Again, $(\phi \wedge \psi)$ can be any wff that has this structure, including but not limited to any of the following: $(P \wedge Q)$, $(P \wedge R)$, $(Q \wedge R)$, $(\neg(P) \wedge Q)$, $(\neg(P) \wedge \neg(Q))$, $(\neg(P) \wedge (Q \rightarrow R))$.

2.2.2 Three types of wffs

It is worthwhile to classify there different general types of well-formed formulas:

1. atomic wffs
2. complex wffs
3. literal wffs

Definition 2.2.3: atomic wff

A **PL**-well-formed formula an atomic wff ϕ in *PL* iff it consists of only a single propositional letter (with or without positive integer subscripts).

In short, an atomic wff is a wff that consists only of a single propositional letter and nothing else. For example, P is an atomic wff, so is Q , and so is Z_1 . In contrast, $(P \wedge Q)$ is not an atomic wff since it consists of two propositional letters and an operator (connective). Similarly, $\neg(P)$ is not an atomic wff since it consists of a propositional letter and a negation operator.

Definition 2.2.4: complex wff

A **PL**-well-formed formula is a complex wff ϕ in *PL* iff ϕ is a wff with at least one propositional letter and at least one truth-functional operator.

A complex wff is sometimes called a "compound wff" or "compound sentence", a "molecular wff" or a "composite sentence".

In short, a complex wff is a wff that consists of at least one propositional letter and at least one of the five truth-functional operators: $\neg, \wedge, \vee, \rightarrow, \leftrightarrow$. For example, $(P \wedge Q)$ is a complex wff since it consists of two propositional letters and a truth-functional operator. Similarly, $\neg(P)$ is a complex wff since it consists of a propositional letter and a truth-functional operator. In contrast, P is not a complex wff since it consists of only a propositional letter and no truth-functional operators. In addition, the negation operator $\neg$ is not a wff since (1) it is not a wff and (2) it does not contain at least one propositional letter.

Definition 2.2.5: literal wff

A **PL**-well-formed formula ϕ is a literal wff in *PL* iff it consists of an atomic wff ψ or a singly negated atomic wff $\neg(\psi)$.

A literal wff is a wff that consists of either an atomic wff or a negated atomic wff. For example, P is a literal wff since it is an atomic wff. Similarly, $\neg(P)$ is a literal wff since it is a negated atomic wff. In contrast, $(P \wedge Q)$ is not a literal wff since it is neither an atomic wff nor a negated atomic wff. Similarly, $\neg(\neg(P))$ is not a literal wff since it is also neither an atomic wff nor a singly negated atomic wff.

As a final note, it is important to note that while a literal wff may be atomic (when it is a single letter) or complex (when it is a singly negated atomic wff). For instance, $\neg(P)$ is a literal wff that is complex but not atomic. In addition, P is a literal wff that is also an atomic wff.

A literal wff has one foot in each camp: some are atomic, some are complex.

Exercise 2.11

Identify whether the following wff is an atomic, complex, and/or literal wff

1. P
2. Q
3. $(P \wedge Q)$
4. $\neg(P)$
5. $\neg(Q)$
6. $(\neg(P) \wedge \neg(Q))$

2.2.3 Parts and subformulas

Increasingly complex wffs are constructed from the *PL* formation rules. It will be convenient to develop some terminology for talking about some of the wffs created in the process of creating a wff ϕ .

Definition 2.2.6: Proper part of a wff

Let ϕ and ψ be any *PL*-wffs. A wff ϕ is a proper part of another wff ψ if and only if ϕ is a wff that is constructed by the formation rules in the process of constructing ψ but not ψ itself.

To illustrate, P and Q are proper parts of $(P \rightarrow Q)$ since both wffs are constructed using the formation rules in constructing $(P \rightarrow Q)$.

A proper part of a wff ψ are all the subformulas constructed in the process of constructing ψ *excluding* ψ . A subformula (or part) of ψ are all the proper parts of ψ and ψ .

Definition 2.2.7: Subformula (part)

A subformula (or part) ϕ of ψ is any wff occurring as a proper part of ψ , including ψ itself.

So, while P and Q are the proper parts of $(P \rightarrow Q)$, the subformulas of $(P \rightarrow Q)$ are the following: $P, Q, (P \rightarrow Q)$.

Let's look at a few additional examples. Consider the wff $\neg(P)$. The subformulas of $\neg(P)$ are all of the wffs created by the formation rules in the process of showing that $\neg(P)$ is a wff, including $\neg(P)$ itself. This means that P and $\neg(P)$ are subformulas of $\neg(P)$. Some additional examples:

- $P, Q, (P \vee Q)$ are subformulas of $(P \vee Q)$
- $P, \neg(P), Q, (\neg(P) \vee Q)$ are subformulas of $(\neg(P) \vee Q)$
- $P, Q, (P \rightarrow Q), \neg((P \rightarrow Q))$ are subformulas of $\neg((P \rightarrow Q))$

If you are ever unsure of whether a wff is a subformula of another wff, you can construct the wff using the formation rules. For example, suppose we were unsure whether P is a subformula of $\neg((P \rightarrow Q))$. To determine this, we would construct $\neg((P \rightarrow Q))$ using the formation rules:

1. P and Q are wffs (rule 1)
2. If P and Q are wffs, then $(P \rightarrow Q)$ is a wff (line 1 + rule 5)
3. If $(P \rightarrow Q)$ is a wff, then $\neg((P \rightarrow Q))$ is a wff (line 2 + rule 2).
4. Therefore, $\neg((P \rightarrow Q))$ is a wff (lines 1-3).

Since P is used in the construction of $\neg((P \rightarrow Q))$, P is a subformula of $\neg((P \rightarrow Q))$.

Exercise 2.12

Determine the proper parts and subformulas of the following wffs:

1. $(P \rightarrow Q)$
2. $(P \wedge \neg(Q))$
3. $\neg((P \wedge Q))$
4. $(P \vee (P \wedge Q))$
5. $\neg((\neg(P) \leftrightarrow \neg(Q)))$

2.2.4 Occurrences

Some wffs have multiple instances (or occurrences) of the same operator. For example, $(\neg(\neg(P) \rightarrow Q))$ contains two occurrences of the $\neg$ operator. For convenience, we refer to each of the different instantiations of an operator as an "occurrence" of that operator.

Definition 2.2.8: occurrence

An occurrence of an operator is an instance (a physical embodiment in space and time) of an operator type.

To further illustrate, consider how many occurrences there are of $\wedge$ in the following wff: $((P \wedge Q) \wedge (R \wedge S))$. While there is only one $\wedge$ operator, there are three occurrences of $\wedge$ in that formula.

Exercise 2.13

1. How many occurrences of $\neg$ are in the following wff: $\neg(P)$
2. How many occurrences of $\neg$ are in the following wff: $\neg((P \wedge \neg(Q)))$
3. How many occurrences of $\wedge$ are in the following wff: $((P \wedge Q) \vee R)$
4. How many occurrences of $\wedge$ are in the following wff: $((P \wedge Q) \wedge R)$
5. How many occurrences of $\wedge$ are in the following wff: $((P \rightarrow Q) \wedge R)$

2.2.5 Scope of an operator

With our understanding of the notion of a subformula and the occurrence of an operator, we can now define the notion of the "scope" of an operator. The truth-functional operators of PL are said to have "scope".

Definition 2.2.9: scope of PL operator

The scope of an occurrence of an operator in a wff ϕ is the smallest subformula of ϕ that contains that occurrence of that operator.

Informally, the scope of an operator is the part of the wff that the operator is "attached" to. It is "how much" of the wff the operator "operates" on.

Let's consider a few examples. Take the single occurrence of $\neg$ in $\neg(P)$. Next, let's identify the subformulas of $\neg(P)$. These are the following:

1. P

2. $\neg(P)$

To determine the scope of $\neg$, we need to identify the "smallest subformula" of $\neg(P)$ that contains $\neg$. Since only $\neg(P)$ contains $\neg$, the smallest subformula that contains $\neg$ is $\neg(P)$. And so, the scope of $\neg$ in $\neg(P)$ is $\neg(P)$. That is, the scope is the entire wff.

The idea of scope is not a complex concept, but it isn't the easiest to grasp. This is because scope requires you to know (1) how to construct a wff using the formation rules and (2) what a subformula is. If you are struggling with the notion of scope, then you might consider reviewing the formation rules in [Definition 2.2.2](#) and the discussion of subformulas in [subsection 2.2.3](#).

Next, consider the wff $(P \vee \neg(Q))$. For this wff, we can ask about the scope of two different operators: the $\vee$ and the $\neg$. Let's consider the scope of each. To do this, let's identify all of the subformulas of this wff:

1. P
2. Q
3. $\neg(Q)$
4. $(P \vee \neg(Q))$

To identify the scope of $\vee$, identify the smallest subformula that contains $\vee$. This is the entire wff. Next, consider the smallest subformula that contains $\neg$. Notice that two wffs contain $\neg$: $\neg(Q)$ and $(P \vee \neg(Q))$. However, since $\neg(Q)$ is smaller than $(P \vee \neg(Q))$, the scope of $\neg$ is $\neg(Q)$.

Another way to identify the scope of an occurrence of an operator is to construct the wff using the formation rules. The scope of an operator corresponds to the wff that is constructed when the operator is first introduced into the construction. Let's illustrate this idea with an example. Consider $(\neg(P) \rightarrow Q)$. If we were to construct this wff using the formation rules, the wffs we would construct the wffs as follows:

1. P
2. Q
3. Since P is a wff, then $\neg(P)$ is a wff.
4. Since $\neg(P)$ and Q are wffs, $(\neg(P) \rightarrow Q)$ is a wff.

The scope of $\neg$ corresponds to the wff constructed when the $\neg$ is first introduced in the construction. Notice that this occurs at step (3) in the above construction. Hence the scope of $\neg$ is $\neg(P)$. Similarly, since $\rightarrow$ is first introduced at step (4), the scope of $\rightarrow$ corresponds to the wff constructed at step (4).

Finally, consider $\neg((P \rightarrow Q))$. If we were to construct this wff using the formation rules, the wffs we would create are $P, Q, (P \rightarrow Q), \neg((P \rightarrow Q))$. In this case, the smallest subformula containing $\rightarrow$ is $(P \rightarrow Q)$ and so the scope of $\rightarrow$ is $(P \rightarrow Q)$. However, the smallest subformula containing $\neg$ is $\neg((P \rightarrow Q))$, so the scope of $\neg$ is the entire wff $\neg((P \rightarrow Q))$.

Exercise 2.14

Identify the scope of each occurrence of the operators in the following wffs:

1. $(P \wedge Q)$
2. $(P \rightarrow Q)$
3. $\neg(\neg(P))$
4. $(P \wedge (Q \leftrightarrow R))$
5. $\neg((P \wedge \neg(Q)))$

2.2.6 Main Operator

Every complex wff will have one and only one main operator. The main operator of a complex wff is the operator the greatest or most scope.

Definition 2.2.10: main operator

The main operator of a *PL* wff ϕ is the truth-functional operator whose scope is ϕ

Let's consider a few examples. Let's start with $\neg(P)$. In this example, the scope of $\neg$ is $\neg(P)$. Since $\neg$ is the operator whose scope is $\neg(P)$, it is the main operator. Since every complex wff has exactly one main operator and $\neg(P)$ only has one operator, $\neg$ is the main operator by default. This is also true of wffs like $(P \wedge Q)$, $(P \vee Q)$, $(P \rightarrow Q)$, and $(P \leftrightarrow Q)$. The main operator of each of these wffs is the only operator in the wff.

The main operator of a wff is the operator that is "doing the most work" or has the "most scope".

But let's consider a wff where there are at least two operators. Take the wff $(\neg(P) \vee Q)$. In this wff, there are two operators: $\neg$ and $\vee$. To determine the main operator of this wff, we want to identify the operator whose scope is the entire wff $(\neg(P) \vee Q)$. Let's consider the scope of each of the operators. The scope of $\neg$ is $\neg(P)$ and the scope of $\vee$ is $(\neg(P) \vee Q)$. Since the scope of $\vee$ is the entire wff, it is the main operator.

The formation rules can be used to determine the main operator of a wff. To see this clearly, first notice that other than Rule 1, all of the formation rules are associated with a truth-functional operator, e.g. Rule 6 with ' $\leftrightarrow$ '. Second, the main operator of a wff is the truth-functional operator associated with the last formation rule applied to create the wff. For example, consider the use of the formation rules to show that $\neg(P \rightarrow \neg(R))$ is a wff:

1. Every propositional letter is a wff, so P and R are wffs (rule 1).

2. If R is a wff, then $\neg(R)$ is a wff (line 1 + rule 2).
3. If P and $\neg(R)$ are wffs, then $(P \rightarrow \neg(R))$ is a wff (lines 1,2 + rule 5).
4. If $(P \rightarrow \neg(R))$ is a wff, then $\neg((P \rightarrow \neg(R)))$ is a wff (line 3 + rule 2).

Notice that the last formation rule applied is Rule 2 (associated with ' $\neg$ '). The ' $\neg$ ' applied to ' $P \rightarrow \neg(R)$ ' is the main operator.

Exercise 2.15

Identify the main operator of the following wffs:

1. $\neg(P)$
2. $\neg(\neg((P \wedge \neg(R))))$
3. $(P \vee (Q \wedge R))$
4. $M \wedge \neg(Q)$
5. $(\neg(\neg(M)) \vee R)$
6. $\neg((P \vee Q))$
7. $(P \vee Q) \vee (R \vee S)$
8. $((P \wedge Q) \vee (A \rightarrow S))$
9. $(P \rightarrow (Q \rightarrow R))$
10. $((P \rightarrow Q) \rightarrow R)$

2.2.7 Literal Negation

The literal negation of a wff ϕ is created either by applying or removing one instance of the formation rule for negation to ϕ . That is, two wffs ϕ and ψ are literal wffs of each other if and only if either one can be formed from the other by the negation formation rule (rule 2).

So, for example, suppose we have the wff P . The literal negation of P is $\neg(P)$ since $\neg(P)$ would be the result of applying the formation rule for negation to P . Similarly, suppose we have the wff $\neg(P)$. The literal negation of $\neg(P)$ is either (1) P since P would be the result of removing the negation from $\neg(P)$ or (2) $\neg(\neg(P))$ as this would be the result of applying the formation rule for negation to $\neg(P)$.

The literal negation is not simply adding or removing a negation to a wff. Say you have $(P \wedge \neg(Q))$. The literal negation of this wff is not $(P \wedge Q)$; rather, it is $\neg((P \wedge \neg(Q)))$

Let's look at a few more examples. Consider $(P \rightarrow R)$. The literal negation of this wff is $\neg((P \rightarrow R))$ since this would be the result of applying the formation rule for negation to $(P \rightarrow R)$. Now consider the wff $\neg(\neg(P) \wedge R)$. The literal negation of this wff is either (1) $(\neg(P) \wedge R)$ since this would be the result of removing the negation from $\neg(\neg(P) \wedge R)$

or (2) $\neg(\neg((\neg(P) \wedge R)))$ since this would be the result of applying the formation rule for negation to $\neg((\neg(P) \wedge R))$.

Exercise 2.16

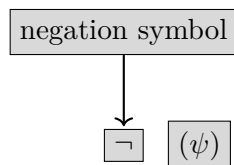
Determine the literal negation of the following formulas:

1. A
2. $\neg(A)$
3. $\neg(\neg(A))$
4. $(A \rightarrow B)$
5. $(\neg(A) \rightarrow \neg(B))$
6. $\neg(\neg((P) \rightarrow \neg(R)))$

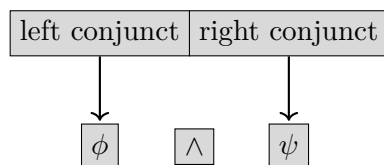
2.2.8 Types of wffs: Some terminology

With the notion of the main operator and the literal negation of a wff developed, let's develop some terminology for talking about different types of wffs. Earlier, we distinguished between atomic wffs and complex wffs, the latter being wffs that are composed of at least one propositional letter and at least one truth-functional operator. Next, let's further distinguish between different types of complex wffs. Given that every complex has exactly one main operator, we can distinguish between different types of complex wffs based on the main operator.

First, there are wffs whose main operator is the negation. These wffs are called "negated wffs" or "negations". Negations are thus complex wffs of the form: $\neg(\phi)$.

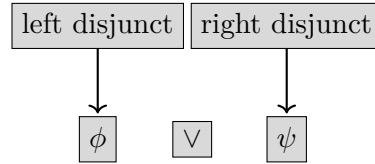


Second, there are wffs whose main operator is the $\wedge$. These wffs are called "conjunctions". Conjunctions then are complex wffs of the form $(\phi \wedge \psi)$. The subformulas on the left and right sides of the $\wedge$ are called "conjuncts".



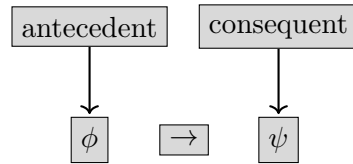
Disjunctions: $(\phi \vee \psi)$

Third, there are complex wffs whose main operator is the $\vee$. These wffs are called disjunctions. Disjunctions then are complex wffs of the form: $(\phi \vee \psi)$. The subformulas on the left and right sides of the $\vee$ are called "disjuncts".



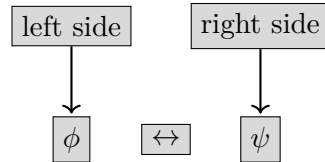
Conditionals: $(\phi \rightarrow \psi)$

Fourth, there are complex wffs whose main operator is the $\rightarrow$. These wffs are called conditionals. Conditionals are complex wffs of the form: $(\phi \rightarrow \psi)$. The subformula to the left of the $\rightarrow$ is called the "antecedent" of the conditional and the subformula to the right of the $\rightarrow$ is called the "consequent" of the conditional.



Biconditionals: $(\phi \leftrightarrow \psi)$

Fifth and finally, there are complex wffs whose main operator is the $\leftrightarrow$. These wffs are called biconditionals. Biconditionals then are complex wffs of the form: $(\phi \leftrightarrow \psi)$. Typically, the subformula to the left of the double arrow is called the "leftside of the biconditional" and the subformula to the right of the double arrow is called the "rightside of the biconditional".



With terms for each of the five types of complex wffs, we can now classify wffs based on their main operator. For example, $\neg(P)$ is a negation, $(P \wedge Q)$ is a conjunction, $(P \vee Q)$ is a disjunction, $(P \rightarrow Q)$ is a conditional, and $(P \leftrightarrow Q)$ is a biconditional. In addition, $\neg(P)$ is a negation.

Negated wffs of each type

In addition to the five main types of complex wffs, by taking the literal negation of each of these basic types, we can obtain the negated version of each of the five types of complex wffs. For example, $(P \wedge Q)$ is a conjunction. The literal negation of this wff is $\neg((P \wedge Q))$, a negated conjunction. Similarly, $\neg(P)$ is a negation. The literal negation of this wff is $\neg(\neg(P))$, a negated negation (or double negation).

Exercise 2.17

Identify the type of wff for each of the following wffs:

1. A
2. $\neg(A)$
3. $(A \wedge B)$
4. $(A \vee B)$
5. $(A \rightarrow B)$
6. $(A \leftrightarrow B)$
7. $\neg(\neg(A))$
8. $\neg(A \wedge B)$
9. $\neg((A \vee B))$
10. $((A \wedge B) \vee C)$

2.2.9 Two Conventions Concerning Parentheses

In order to make formulas more readable, three conventions are invoked. First, if a wff is fully surrounded by parentheses, then these outermost parentheses may be omitted. Consider the following examples:

1. $(P \wedge Q)$
2. $(P \vee Q)$
3. $(P \rightarrow Q)$
4. $(P \leftrightarrow Q)$

In each of the above wffs, the wff is surrounded by a pair of parentheses. The first simplification convention allows for omitting these parentheses and so each of the above wffs may be rewritten as follows:

1. $P \wedge Q$
2. $P \vee Q$
3. $P \rightarrow Q$
4. $P \leftrightarrow Q$

The second convention for simplification involves the scope of the negation operator. Our convention is the following:

In the absence of parentheses, let $\neg$ apply to the smallest subformula to its immediate right.

Here is another way of expressing this convention. Let's introduce a way of reading the scope of the negation operator. We will say that in the absence of parentheses, the negation operator should be read as having *narrow scope*. By "narrow scope" what is meant is that the negation operator

should be read as "applying" to the smallest subformula to its immediate right. That is, in the absence of parentheses, the scope of the negation operator should include (1) itself and (2) the smallest subformula to its immediate right. Let's illustrate this idea by considering the following wffs:

1. $\neg P$
2. $\neg P \wedge Q$
3. $\neg(P \wedge Q)$
4. $\neg\neg P$

Notice that in the above wffs, parentheses are missing. For example, (1) is typically written as $\neg(P)$ rather than $\neg P$. However, with our convention for reading the scope of negation as being narrow in hand, consider that we read $\neg P$ in the same way as $\neg(P)$. That is, the scope of the negation in $\neg P$ is the entire wff. Informally, the negation "applies" to the P just like it does in $\neg(P)$.

Next, consider (2). At first glance, this wff may appear ambiguous since it may be read in two different ways:

- 2a $(\neg(P) \wedge Q)$
- 2b $\neg((P \wedge Q))$

In (2a) the scope of the negation is $\neg(P)$ but in (2b), the scope of negation is the entire wff. However, consider how we ought to read (2) in light of how to read the scope of negation when parentheses are absent: its scope should extend to the smallest subformula to its immediate right. Thus, (2) and (2a) express the same wff.

Now consider (3). Notice in this wff, parentheses surround the conjunction $P \wedge Q$. These parentheses are present to indicate that the scope of negation extends to the entire conjunction. That is, (3) and (2b) express the same wff.

Let's put this point in a different way. If we want the scope of negation to only apply to the smallest subformula to its immediate right, then we can omit the use of parentheses. That is, if we want the negation to only apply to "P" in $P \wedge Q$, then we write $\neg P \wedge Q$. However, if we want the negation to apply to more than the smallest subformula to its immediate right, then we can use parentheses. That is, if we want the negation to apply to the entire conjunction $P \wedge Q$, then we write $\neg(P \wedge Q)$.

Finally, let's consider wff (4). This wff is $\neg\neg P$. This is a simplification of $\neg(\neg(P))$. How should we read the scope of the different negation op-

erators? Fortunately for us, there is only one way. Consider the leftmost negation. In the absence of parentheses, this negation should be read as applying to the smallest subformula to its immediate right. The smallest subformula is $\neg P$. What about the rightmost negation? Again, in the absence of parentheses, this negation should be read as applying to the smallest subformula to its immediate right. This is P . Thus, in the absence of parentheses, our convention makes it clear the scope of these negations and so the use of parentheses is not necessary.

1. $\neg(P)$
2. $\neg(P) \wedge \neg(Q)$
3. $\neg(\neg(P))$

Let's conclude by considering a third convention. Some wffs are difficult to read because of the sheer number of parentheses. In those cases, it can sometimes be helpful to use square brackets $[]$ and curly braces $\{ \}$ instead. Let's consider a few examples:

1. $(P \rightarrow Q) \vee ((S \vee R) \vee T)$
2. $\neg((P \wedge Q) \wedge ((L \wedge R) \wedge Z))$

Since there are so many parentheses, some people would prefer to insert brackets or curly braces to track where a pair of parentheses begins and ends. With that in mind, the above wffs may be rewritten as follows:

1. $(P \rightarrow Q) \vee ([S \vee R] \vee T)$
2. $\neg\{(P \wedge Q) \wedge ([L \wedge R] \wedge Z)\}$

Finally, why would we ever want to simplify a wff? There are two reasons. The first is readability. When a wff has many parentheses, it can sometimes be more difficult to read than its simplified counterpart.

3. $\neg((\neg(P) \rightarrow \neg(Q)))$
4. $\neg(\neg P \rightarrow \neg Q)$

For many people, wff (2) is easier to read than wff (1). The second reason is efficiency. In contrast to wff (2) with its two parentheses, wff (1) contains a whopping eight parentheses. It is thus less time-consuming to write wff (2) than wff (1).

Exercise 2.18

Make use of the above three conventions to simplify the following wffs:

1. $\neg(Q)$
2. $(P \rightarrow Q)$
3. $\neg(\neg(P))$
4. $\neg(\neg((P \rightarrow Q)))$

2.3 PL SEMANTICS

PL is a set of symbols and a set of syntax (formation) rules for putting those symbols together. As such, the symbols and the wffs that are the result of applying the formation rules are *meaningless*. PL gets its meaning by being "interpreted". This section explains what an interpretation is and formulates a set of rules that allow for the assignment of truth values to wffs given an interpretation. However, before the notion of an interpretation is introduced, we need to introduce the notion of a function.

2.3.1 Functions

What is a function? A function is a specific kind of relationship between two groups of things. The first group of things, known as the "domain", is composed of objects known as the function's "inputs" or "arguments". The second group of things, known as the "range", is composed of objects known as the function's "outputs" or "values". The relation between the inputs and the outputs is such that each input is related to one and only one output.

Definition 2.3.1: function

A function is a relation between two sets of things (the inputs and the outputs) such that each input is related to one and only one (exactly one) output.

Let's consider an example. Consider two sets. In the first set, we have four letters: A, B, C, D . In the second set, we have three numbers: 1, 2, 3. Now suppose we relate the two sets as follows, where A and B are in a relation to 1, C is in relation to 2, and C is in relation to 3. Let's illustrate this with the example below.

The above is an example of a relation. Notice that in the above example, each letter is related to exactly one number. No letter is unrelated and no letter is in relation to two numbers, e.g., " A " is not in relation to numbers 1 and 2.

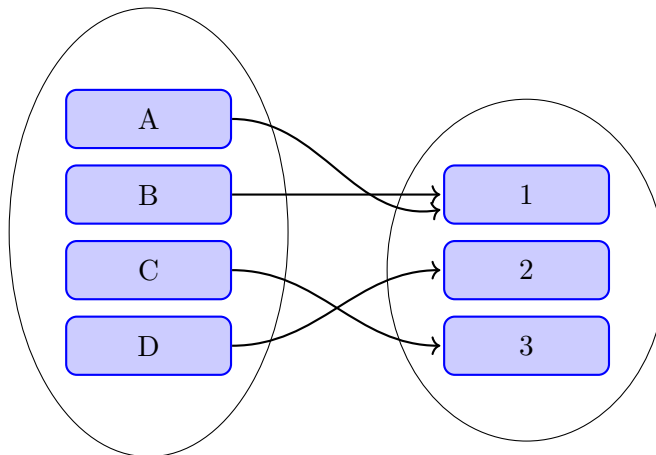


Figure 2.2: In the above example, we have each letter in the inputs (domain, arguments) related to exactly one number in the outputs (range, values).

Let's consider another example. In this example, we will consider a function (let's call it F) that relates emotional states (happy, sad) to three activities (dancing, reading, eating). In this example, notice that while each emotional state is related to exactly one output activity, the reverse is not the case: not every output activity is related to exactly one emotional state.

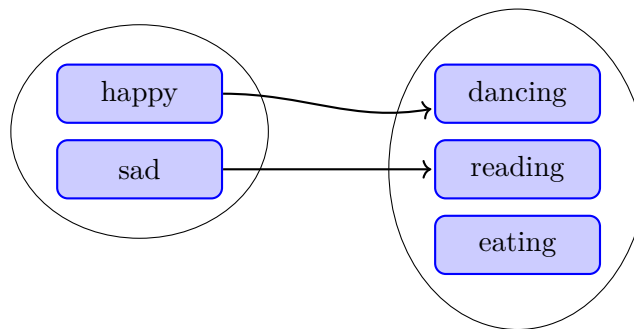


Figure 2.3: Each input emotional state (happy, sad) is related to exactly one activity (dancing, reading, eating). Notice that while each input emotional state is related to exactly one output activity, it is not the case that each activity is related to one emotional state.

Let's consider some additional functions. The first function we will consider is called the "lighter-color function". This function takes a single input color and relates it to a lighter version of that color as an output

color. In other words, it takes a single color as an argument and relates it to a color as a value.

Definition 2.3.2: Lightening Color Function

Given a color as input, generate a lighter color as output.

So, for example, the lightening-color function takes brown and generates light brown.

Input		Output
grey	$\Rightarrow$	light gray
brown	$\Rightarrow$	light brown

A function need not only involve one input. It may, for example, have two, or three, or any number of inputs. For example, consider a function that takes two colored items (red or blue) as input and produces an emotion (happy or sad) as output:

Definition 2.3.3: Color-Emotion Function

If the color-value input of both of the colored items is blue, then the output is happy. If the color-value input of either of the colored items is red, then the output is sad.

Input (color)	Input (color)		Output (emotion)
blue	blue	$\Rightarrow$	happy
blue	red	$\Rightarrow$	sad
red	blue	$\Rightarrow$	sad
red	red	$\Rightarrow$	sad

Table 2.1: Color-to-Emotion Function

Suppose that someone's emotions are determined according to the color-emotion function as represented in [Table 2.1](#) and suppose that we present them with two different pieces of clothing, a blue shirt and a blue pair of pants. The color-emotion function says that if the color-value input of both items is blue, then this individual will be happy. Alternatively, if we hand them a pair of blue pants and a red shirt, the color emotion function says that the individual will be sad.

The output of the function is determined not by anything outside (or unspecified by) the function. That is, a person's happiness or sadness (the values) are determined wholly by the input.

Note that while a function may relate a number of inputs to the same output, a function never relates a single input to more than one output. Consider a function that takes two colored and patterned rectangles as input and produces a single rectangle, which is a combination of both rectangles as output. Since there are an infinite number of color-shape combinations, an exhaustive set distinct diagrams is not possible, but the examples in Figure 2.4 and Figure 2.5 serve as useful illustrations of this function at work.

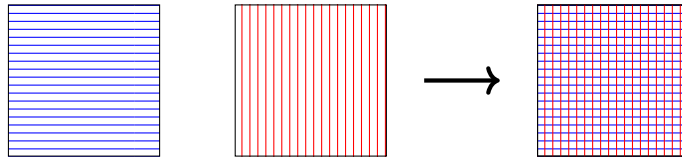


Figure 2.4: Example 1 of the Rectangle-Pattern Function

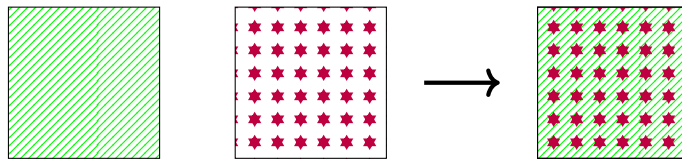


Figure 2.5: Example 2 of the Rectangle-Pattern Function

Note that the function relates two different inputs to one and only one output.

2.3.2 The interpretation function

Now that we have the notion of a function in hand, we will consider two functions that pertain to the semantics of **PL**. Recall that the formal language of *PL* consists of a set of meaningless symbols and rules for putting these symbols together in the right way. Two functions supply meaning to **PL**: one takes propositional letters as input and assigns truth values to these propositional letters, the other assigns a truth value (T or F) to a wff depending on the truth values of propositional letters and the operators that compose that wff.

Definition 2.3.4: Interpretation of **PL**

An interpretation of *PL* is a function that takes propositional letters of *PL* as input and assigns them a single truth value (T or F) as output.

61 An interpretation is sometimes said to be a "mapping" from propositional letters to truth values and that the interpretation "gives meaning" to the propositional letters.

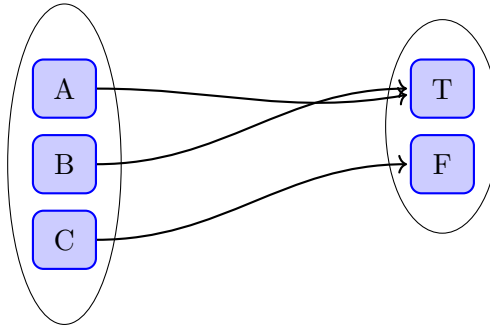


Figure 2.6: The interpretation function takes all the letters of the language of propositional logic (or those under consideration) and assigns them exactly one truth value: T or F.

In other words, an interpretation assigns propositional letters meaning by assigning them one and only one truth value (T or F).

A few points about the interpretation function. First, to symbolize the interpretation function of some proposition P , we write $\mathcal{I}(P) = T$ or $\mathcal{I}(Q) = F$, which reads " P is true under interpretation $\mathcal{I}$ " or " Q is false under interpretation $\mathcal{I}$ ".

Second, a single propositional letter can be interpreted in two different ways. Under one interpretation, it can be assigned a value of "T" while under a different interpretation it can be assigned a value of "F". In other words, if P is a propositional letter, P can be interpreted as true, i.e., $\mathcal{I}_1(P) = T$ or P can be interpreted as false, i.e., $\mathcal{I}_2(P) = F$.

Third, an interpretation is an assignment of a truth value to every propositional letter in **PL**. While **PL** has an infinite number of propositional letters, in practice interpretations are specified by assigning truth values to the propositional letters that are under analysis, e.g. $\mathcal{I}(P) = T, \mathcal{I}(Q) = T, \mathcal{I}(R) = T, \mathcal{I}(Z) = T$

Fourth, for n distinct propositional letters, there are 2^n interpretations. Thus for 1 propositional letter, $2^1 = 2$; for 2 propositional letters, $2^2 = 4$, and so on.

2.3.3 The valuation function

A second collection of functions is used to determine the truth values of all of the wffs of **PL**. This collection of functions are known as "valuations".

Recall that it is the business of the interpretation function to assign truth values to propositional letters. And, the interpretation function is a func-

tion since exactly one truth value (T, F) is assigned to each and every propositional letter. While the interpretation function determines the truth values of propositional letters, the valuation function determines the truth values of wffs.

How does the valuation function do this? The valuation function assigns truth values to wffs in a very specific way. It does this by taking the truth values of either the propositional letters or the subformulas that compose the wff as input and assigning exactly one truth value (T or F) to the wff as output. The valuation function is a function since it relates any given truth value input to exactly one truth value (T, F) output. While the output of the input and output of the valuation function are truth values, the output truth value is taken to be the value (or "meaning") of the wff itself.

The valuation function is a *truth function*. A truth function is a function from truth values (as input) to exactly one truth value (as output). The valuation function is a truth function since given any truth value input, it will always output exactly one truth value. The valuation function also *entirely determines* the truth value output of a wff by taking the truth values of the propositional letters / subformulas that compose it as input. In other words, the valuation function does not take anything other than the truth values of the propositional letters / subformulas that compose it as input. A language is said to be a truth-functional language when the truth value of a wff is *entirely determined* by the truth values of the wffs that compose it. Thus, PL is a truth-functional language.

Practically, the real question then is the following: how does the valuation function assign truth values to wffs? While we can specify the truth values of propositional letters using the interpretation function, e.g., $\mathcal{I}(P) = T, \mathcal{I}(Q) = F, \dots$, how do we determine the truth values of wffs like $P, \neg(P), (P \wedge Q)$, and so on? To answer this question, let's start by considering the subtle distinction between propositional letters and atomic wffs.

A propositional letter is a symbol of PL . For example, P, Q, R are propositional letters. These symbols get assigned a truth value by the interpretation function. An atomic wff is a wff that is defined by the formation rules. For example, P, Q, R are also atomic wffs. As wffs, it is the responsibility of the valuation function to determine their truth value. With this in mind, let's specify a function that takes the truth value of the propositional letter and uses it to assign a truth value to the atomic wff. This is relatively straightforward. We can state the following:

$v(\phi) = \mathcal{I}(R)$, where ϕ is R (let R be a variable for any propositional letter).

In other words, this says that the truth value of the wff ϕ is identical to the truth value of R when ϕ is R . If P is an atomic wff, then its truth value is determined by the interpretation of P (where P is a propositional letter). If Q is an atomic wff, then its truth value is determined by $\mathcal{I}(Q)$, and so on. More generally, if ϕ is an atomic wff, then the truth value of ϕ is the same as the truth value of its corresponding propositional letters.

What we have now is a specification of how the valuation function assigns truth values to atomic wffs using the interpretation function. But what about complex wffs? Consider the complex wff $(P \wedge Q)$. How does the valuation function assign it a truth value? What we need is a specification of a function that takes the truth values of P and Q as input, takes into account the fact that the main operator of $P \wedge Q$ is the $\wedge$, and then outputs a truth value that we take to be the truth value of $(P \wedge Q)$.

Let's consider how we can do this using truth values.

More compactly: $v(\phi \wedge \psi) = T$ iff $v(\phi) = T$ and $v(\psi) = T$, otherwise $v(\phi \wedge \psi) = F$.

1. if ϕ and ψ are both T, then the conjunction $(\phi \wedge \psi)$ is T.
2. If ϕ is T and ψ is F, then the conjunction $(\phi \wedge \psi)$ is F.
3. If ϕ is F and ψ is T, then the conjunction $(\phi \wedge \psi)$ is F.
4. If ϕ and ψ are both F, then the conjunction $(\phi \wedge \psi)$ is F.

Notice a few things about the specification of the valuation function for conjunctions. First, the valuation function v for conjunction takes the truth values of ϕ and ψ as input and then, given that ϕ and ψ are connected by $\wedge$, it specifies the truth value of the conjunction $(\phi \wedge \psi)$ itself. Second, notice that it always outputs a single truth value as output (this output value is assigned to the conjunction $(\phi \wedge \psi)$). It is never the case that, for a given input, a different output is given. Third, notice that the output of the function is determined entirely by the truth values of ϕ and ψ . It is not determined by anything other than the truth values of the wffs that compose the conjunction.

Now that we have a basic understanding of how the valuation function works for atomic wffs and conjunctions, let's consider how it works for every complex wff.

Definition 2.3.5: Valuation

For any interpretation $\mathcal{I}$, a valuation (v or $\mathcal{V}$) in PL is a function that assigns exactly one truth value (T or F) to each wff in PL in

such a way that (let ϕ and ψ be variables for any wff in **PL**, R be any propositional letter in **PL**, and "iff" be an abbreviation for "if and only if"):

1. $v(\phi) = \mathcal{I}(R)$, where ϕ is R .
2. $v(\neg(\phi)) = T$ iff $v(\phi) = F$, otherwise $v(\neg(\phi)) = F$
3. $v(\phi \wedge \psi) = T$ iff $v(\phi) = T$ and $v(\psi) = T$, otherwise $v(\phi \wedge \psi) = F$
4. $v(\phi \vee \psi) = T$ iff $v(\phi) = T$ or $v(\psi) = T$, otherwise $v(\phi \vee \psi) = F$
5. $v(\phi \rightarrow \psi) = T$ iff $v(\phi) = F$ or $v(\psi) = T$, otherwise $v(\phi \rightarrow \psi) = F$
6. $v(\phi \leftrightarrow \psi) = T$ iff $v(\phi) = T$ and $v(\psi) = T$, or $v(\phi) = F$ and $v(\psi) = F$, otherwise $v(\phi \leftrightarrow \psi) = F$

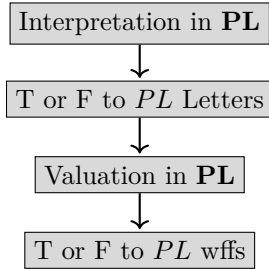
Let's briefly discuss the parts of this definition not covered earlier. First, (2) specifies how the valuation function works for negated wffs. It assigns negated wffs $\neg(\phi)$ a value of T if and only if the truth value assigned to ϕ is F. Otherwise, $\neg(\phi)$ is F. So, for example, if $v(P) = F$, then $v(\neg(P)) = T$. Alternatively, if $v(P) = T$, then $v(\neg(P)) = F$.

Next, consider (4). This specifies how the valuation function works for disjunctions. It assigns disjunctions $\phi \vee \psi$ a value of T if and only if the truth value assigned to ϕ is T or the truth value assigned to ψ is T (or both ϕ and ψ are T). Otherwise, $\phi \vee \psi$ is F. So, for example, if $v(P) = T$ and $v(Q) = F$, then $v(P \vee Q) = T$. Alternatively, if $v(P) = F$ and $v(Q) = F$, then $v(P \vee Q) = F$.

Next, consider (5). This specifies how the valuation function works for conditionals. It assigns conditionals $\phi \rightarrow \psi$ a value of T if and only if the truth value assigned to ϕ is F or the truth value assigned to ψ is T (or both ϕ is F and ψ is T). Otherwise, $\phi \rightarrow \psi$ is F. So, for example, if $v(P) = F$ and $v(Q) = T$, then $v(P \rightarrow Q) = T$. Alternatively, if $v(P) = F$ and $v(Q) = F$, then $v(P \rightarrow Q) = F$.

Finally, consider (6). This specifies how the valuation function works for biconditionals. It assigns biconditionals $\phi \leftrightarrow \psi$ a value of T if and only if the truth value assigned to ϕ is T and the truth value assigned to ψ is T, or the truth value assigned to ϕ is F and the truth value assigned to ψ is F. Otherwise, $\phi \leftrightarrow \psi$ is F. So, for example, if $v(P) = T$ and $v(Q) = T$, then $v(P \leftrightarrow Q) = T$. Alternatively, if $v(P) = F$ and $v(Q) = F$, then $v(P \leftrightarrow Q) = T$.

The semantics of *PL* provides "meaning" to the symbols and wffs of **PL**. It does this by specifying two functions: the interpretation function and



the valuation function. The interpretation function assigns truth values to propositional letters. The valuation function assigns truth values to wffs. We saw that the valuation function assigns truth values to wffs in a specific way. Namely, it assigns truth values to wffs by taking the truth values of the propositional letter (or subformulas) and the operators that compose the wff as input and then assigning exactly one truth value (T or F) to the wff as output. The fact that the valuation function determines the truth value of a wff entirely by the truth values of the propositional letters (or subformulas) and the operators that compose it is what makes *PL* a truth-functional language.

Exercise 2.19

Determine the truth value of the following wffs under the following interpretation $\mathcal{I}$ of the propositional letters: $\mathcal{I}(A) = T, \mathcal{I}(B) = F$.

1. A
2. B
3. $\neg(A)$
4. $\neg(B)$
5. $(A \wedge B)$
6. $(A \vee B)$
7. $(A \rightarrow B)$
8. $(A \leftrightarrow B)$

2.3.4 Truth Table Presentation of Valuation Function

Our presentation of the meaning of the symbols and wffs in *PL* involved specifying two functions: the interpretation and the valuation function. A second, and equivalent, method for presenting the meaning of the symbols and wffs in *PL* makes use of "truth tables". A truth table is simply a graphical display of the interpretation and valuation functions.

First, consider that any propositional letter $\mathbf{P}$ can be interpreted in two ways. On one interpretation $\mathbf{P}$ can be assigned a value of T , while on another interpretation, it can be assigned a value of F . We can display these two interpretations of $\mathbf{P}$ as follows:

	P
$\mathcal{I}_1$	T
$\mathcal{I}_2$	F

A truth table involving more than one propositional letter (e.g. P and

Q) would take into account all of the different ways that P and Q can be interpreted.

	P	Q
$\mathcal{I}_1$	T	T
$\mathcal{I}_2$	T	F
$\mathcal{I}_3$	F	T
$\mathcal{I}_4$	F	F

Since the table can be used to express every interpretation of the propositional letters, this fact can be used to also illustrate output of the valuation function under each interpretation. The simplest case involves atomic propositions. In this case, an atomic wff ϕ is true if and only if the interpretation of the propositional letter that composes that wff is true (otherwise it is false). Since an atomic wff is nothing more than a propositional letter. The valuation of ϕ is identical to the interpretation of ϕ .

	P	P
$\mathcal{I}_1$	T	T
$\mathcal{I}_2$	F	F

Next, let's show how to represent the valuation function for negated wffs. $v(\neg(\phi))$
The valuation function states that that $v(\neg(\phi)) = T$ if and only if $v(\phi) = F$, otherwise $v(\neg(\phi)) = F$. Using a table, we can represent this as follows:

P	$\neg(P)$
T	F
F	T

Next, let's consider conjunctions. In the case of conjunctions $\phi \wedge \psi$, the wff $\phi \wedge \psi$ is true if and only if ϕ is true and ψ is true. The wff is false in all other cases. To illustrate, consider the wff $P \wedge R$

P	R	$P \wedge R$
T	T	T
T	F	F
F	T	F
F	F	F

Next, disjunctions. The disjunction $\phi \vee \psi$ is true if and only if either ϕ is true or ψ is true (or both are true). The wff $\phi \vee \psi$ is false in all other cases. To illustrate, consider the wff $P \vee R$.

P	R	$P \vee R$
T	T	T
T	F	T
F	T	T
F	F	F

$v(\phi \rightarrow \psi)$

In the case of conditionals, a conditional $\phi \rightarrow \psi$ is true if and only if ϕ is false or ψ is true (or both). The wff $\phi \rightarrow \psi$ is false in all other cases. To illustrate, consider the wff $P \rightarrow R$

P	R	$P \rightarrow R$
T	T	T
T	F	F
F	T	T
F	F	T

$v(\phi \leftrightarrow \psi)$

Finally, let's consider biconditionals. A biconditional $\phi \leftrightarrow \psi$ is true if and only if ϕ and ψ have the same truth value. The wff $\phi \leftrightarrow \psi$ is false in all other cases. To illustrate, consider the wff $P \leftrightarrow R$

P	R	$P \leftrightarrow R$
T	T	T
T	F	F
F	T	F
F	F	T

Exercise 2.20

For each of the following wffs, determine if they are true or false given the interpretation provided.

1. $\neg(P)$ if $\mathcal{I}(P) = F$
2. $P \rightarrow Q$ if $\mathcal{I}(P) = F, \mathcal{I}(Q) = F$
3. $P \wedge Q$ if $\mathcal{I}(P) = T, \mathcal{I}(Q) = T$
4. $P \vee Q$ if $\mathcal{I}(P) = F, \mathcal{I}(Q) = T$
5. $P \leftrightarrow Q$ if $\mathcal{I}(P) = T, \mathcal{I}(Q) = T$

Exercise 2.21

Create your own truth-functional operator. Make a symbol and define it by creating your own valuation function.

2.4 PL TRANSLATION

One of the primary goals of logic is to provide a formal language for the evaluation of arguments. Central to this goal is the ability to translate from a natural language (English) into a formal language. We can think of a translation as a function that takes a sentence of a natural language (English) and assigns it a wff of a formal language. Translations can be good or bad. One component of a good translation is that the wff assigned to the English sentence is true if and only if the English sentence is true.

In this section, we will consider how to translate English sentences into **PL**.

2.4.1 Atomic wffs

A proposition is a sentence (or something typically expressed by a sentence) that is capable of being true or false. Propositions are the basic bearers of truth and falsity. In **PL**, atomic wffs are the basic bearers of truth and falsity. That is, P, Q, R , etc. are the basic bearers of truth and falsity in **PL**. As such, we can translate sentences that express propositions as atomic wffs in **PL**.

The wff J is true ($v(J) = T$) if and only if $\mathcal{I}(J) = T$. The wff J is false ($v(J) = F$) if and only if $\mathcal{I}(J) = F$. Whether J is true or false, thus depends upon the interpretation. Sentences that express propositions behave similarly. The sentence "John is kind" is true if and only if *John is kind*. The sentence "John is kind" is false if and only if it is not the case that *John is kind*.

So, for example, the sentence "John is kind" expresses a proposition. If *John is kind*, then the proposition is true. If *John is not kind*, then the proposition is false. Consider an atomic wff J . This wff behaves similarly. The wff J is true if and only if the interpretation of J is true. In addition, the wff J is false if and only if the interpretation of J is false.

Two final points. First, the choice of which propositional letter to use to translate an English sentence is arbitrary. In our earlier example, "John is Kind" was translated as J , but the letters A, B, C, A_1 , and so forth would serve just as well. Nevertheless, it is generally suggested that when selecting a letter to translate an English sentence, one choose either the first or the most perspicuous letter in the English sentence. For example, J was used to translate "John is kind" rather than Z_{34} because it is the only capital letter in the sentence and it is also the first letter. Similarly, suppose one wished to translate "The man wore a hat". One could trans-

late this sentence as H or M as these are the most perpicuous letters in the sentence and this selection would distinguish it from other sentences that begin with the word "The".

Second, while every proposition can be translated as an atomic wff, the general goal of translation is to represent as much of the logical structure of the proposition as possible in the language of **PL**. As we will see, this will require the use of more complex wffs than atomic wffs.

$\neg(\phi)$

2.4.2 Negated wffs

Recall from the definition of the valuation function that $\neg(\phi) = T$ if and only if $v(\phi) = F$, otherwise $\neg(\phi) = F$. Similarly, in English, a sentence of the form "not- P " is true just in the case that P is false; otherwise, P is true. For example, "It is not the case that John is kind" is true just in the case that "John is kind" is false. And, if "John is kind" is true, then "John is kind" is false. The proposal then is to translate sentences of the form "not- P " (and various equivalent sentences, e.g., "It is not the case that P ", "It is false that P ", etc.) as negated wffs $\neg(\phi)$ in **PL**.

Let's consider some additional examples. Suppose Liz sees that Tek has been behaving very badly. She remarks "It is not the case that John will go to heaven." If "John will go to heaven" is translated as H , then "It is not the case that John will go to heaven" can be translated as $\neg H$. Similarly, suppose Tek is sick. Liz is worried about him and says "It is false that Tek is young." If "Tek is young" is translated as Y , then "It is false that Tek is young" can be translated as $\neg Y$.

You can even add the negation to the end of the sentence. "That shirt looks nice, NOT!"

The cases examined thus far were of the form "not- P ". In these cases, some form of "not" is prefixed to the sentence. Another way to deny a sentence is to apply it to the main verb of the sentence. For example, suppose Tek is being cruel to others. Liz sees this and says "Tek is not kind." This sentence does not have the form "not- K " but (1) it is taken as another way of expressing "it is not the case that Tek is kind" and (2) the sentence is true just in the case that "Tek is kind" is false. As such, it can be translated as $\neg K$.

$(\phi \wedge \psi)$

2.4.3 Conjunctions

Recall from the definition of the valuation function that a conjunction $(\phi \wedge \psi)$ is true if and only if ϕ and ψ are both true. It is false just in the case that one or more of the wffs ϕ or ψ is false. In English, a proposition of the form "S and P" (where S and P are both propositions) if and only

if both S and P are true. We can thus translate propositions of the form "S and P" using $(\phi \wedge \psi)$. In short, conjunctions in PL capture truth-functional uses of 'and' in English. Therefore, if we let J stand for John is kind, H stand for John will go to heaven, and S stand for John went to the store, we can translate the "John is kind and John will go to the store" as $J \wedge S$. Similarly, we can translate "John is kind and John will go to heaven" as $J \wedge H$.

The order of the conjuncts does not matter. For example, $J \wedge H$ and $H \wedge J$ are true if and only if J and H are both true, and false otherwise. Similarly, in English, "John is kind and John will go to heaven" and "John will go to heaven and John is kind" are true and false under the same conditions. Namely, they are true just in the case that "John is kind" is true and "John will go to heaven" is true.

Some English sentences only approximate the form of "P and Q" and so it isn't immediately clear whether or not we can translate them as $P \wedge Q$. For example, consider the sentence "Tek ate a sandwich and an apple." Simply splitting this sentence at the "and" gives us (1) Tek ate a sandwich and (2) an apple. As (2) does not express a proposition, it cannot be translated into **PL**. However, this sentence, and many others, are simply abbreviated forms of a "P and Q" proposition. That is, "Tek ate a sandwich and an apple" is a shorthand way of saying "Tek ate a sandwich and Tek ate an apple." This sentence can be translated as $S \wedge A$ where S stands for "Tek ate a sandwich" and A stands for "Tek ate an apple."

Some English sentences do not have the form "P and Q" but their truth value is determined in the same way as a conjunction. For example, consider the sentence "Liz loves reading, *but* Tek loves dancing." This sentence is not of the form "P and Q" but it is true just in case Liz loves reading and Tek loves dancing. As such, this sentence can be translated as $R \wedge D$ where R stands for "Liz loves reading" and D stands for "Tek loves dancing." This is not to say that there is no difference between "and" and "but" since the latter is often used to communicate to individuals that some contrast exists between the two propositions. However, in the case of "Liz loves reading, *but* Tek loves dancing" the contrast is not relevant to the truth value of the sentence.

Another example of a sentence that can be translated into a conjunction but is not of the form "P and Q" is "P although Q". For example, suppose Tek is sick. Liz is worried about him and says "Tek is sick, although he is young." This sentence is not of the form "P and Q" but it is true if and only if both "Tek is sick" and "Tek is young" are true. It can therefore be

translated as $S \wedge Y$. Similar to "P but Q", "P although Q" is often used to communicate to individuals something more than the truth value of the sentence. For example, the use of "although" is used to communicate that Liz expects Tek to make a recovery from his sickness.

$(\phi \vee \psi)$

2.4.4 Disjunctions

A disjunction $(\phi \vee \psi)$ is true if and only if ϕ is true, or ψ is true, or both ϕ and ψ are true. It is false just in the single case that ϕ and ψ are both false. In English, certain propositions of the form "S or P" (where S and P are both propositions) are true if and only if S is true, P is true, or both S and P are true. It is thus proposed that we can translate certain propositions of the form "S or P" using $(\phi \vee \psi)$.

Let's consider some examples. Suppose Tek has been exercising regularly. Some days he runs, some days he lifts weights, and other days he runs and lifts weights. He utters to himself "I will either run or lift weights today."

1. Tek runs and lifts weights.
2. Tek runs but does not lift.
3. Tek does not run, but he does lift.
4. Tek does not run and does not lift.

It is natural to say that Tek's utterance is only false just in the case where he did not run and he did not lift. In all other cases, his utterance is true. We can translate this utterance as $R \vee L$ where R stands for "Tek runs" and L stands for "Tek lifts weights."

Here is another example. Suppose Tek is a new student at a university. He is pursuing a degree in English, but he must successfully complete two courses in mathematics. Tek is unsure which mathematics course to take and so he consults Liz, his academic advisor. Liz looks at the courses that fit his schedule and mathematical ability. Liz tells him "you can take logic or statistics." Surely, what Liz says is true if Tek takes logic. It is also true if Tek takes statistics. Finally, it is also true if Tek takes both logic and statistics. In fact, the only case where the sentence is false is if Tek can neither take logic nor statistics. We can thus translate what Liz says as $L \vee S$ where L stands for "Tek takes logic" and S stands for "Tek takes statistics."

Exclusive "or"

At first glance then, it appears that *every* sentence of the form "P or Q" can be translated into PL as $\phi \vee \psi$. This is not the case. Some examples. First, suppose Tek walks into a café and asks the cashier what he can purchase with two dollars. The cashier tells him "you can have coffee or

bottled water." In this case, the cashier is not saying that Tek can have both coffee and water for two dollars. Rather, the cashier is saying that Tek can have coffee, or he can have water, but he cannot have both coffee and water. In this case, the cashier's is using "or" in the *exclusive* sense. If we let "C" stand for "Tek can have coffee" and "W" for "Tek can have water", then we can express what the cashier is saying as follows "C or W, but not both C and W". To put this in a truth table, we can express the cashier's use of "or" as follows:

C	W	C or W, but not both C and W
T	T	F
T	F	T
F	T	T
F	F	F

Notice that the truth table for the cashier's use of "or" is different from the truth table for the inclusive use of "or". In the inclusive case, the sentence is true in all cases except where both C and W are false. In the exclusive case, the sentence is true only in the two cases where (1) C is true but not W and (2) W is true but not C. These are rows 2 and 3. In the other two cases, the sentence is false.

Given that exclusive-or expresses a different truth function than the inclusive-or, translating a sentence of the form "P or Q" as $P \vee Q$ is incorrect when exclusive-or is being used. How then should "P or Q" be translated into **PL**? There are two options. The first option is to introduce a new operator (connective) into **PL**. Exclusive-or (XOR) is captured with the $\oplus$ operator. The proposal then is that we should translate "P or Q, but not both" and "P XOR Q" sentences as $P \oplus Q$. The truth table for XOR is the following:

P	Q	$P \oplus Q$
T	T	F
T	F	T
F	T	T
F	F	F

Table 2.2: Truth Table: Exclusive Or

Notice that the truth table for XOR is the same as the truth table for the exclusive use of "or" except for row 1. Since "P XOR Q" says that P or Q is the case, but not both P or Q is the case, sentences of the form "P XOR Q" would be false when both P and Q are the case.

There is a second way of translating "P XOR Q" into **PL** that does not require the introduction of a new operator. However, we will postpone this discussion to [subsubsection 2.4.9.7](#).

In short, the use of "or" in English is ambiguous. Sometimes it is used with the exclusive sense while other times it is used in an inclusive way. When "P or Q" is uttered and "or" is used in the inclusive sense, then the sentence is false in just one case: where both P and Q are false. Such a sentence can be translated as $P \vee Q$ since $P \vee Q$ is false just in the case where $v(P) = F$ and $v(Q) = F$. In contrast, when "P or Q" is uttered and "or" is used in the exclusive sense, then the sentence is false in two cases: (1) where both P and Q are true and (2) where both P and Q are false.

In this text, we will use "or" in an inclusive way (unless we specify otherwise). That is, we will translate "P or Q" as $\phi \vee \psi$ where ϕ stands for "P" and ψ stands for "Q". When the exclusive sense is intended, we will use the phrase "P or Q, but not both" to indicate that we are using "or" in an exclusive way.

There are three remaining questions. First, there is the question of how to translate sentences that involve more than two disjuncts, e.g., $P \vee (Q \vee R)$. Alternatively put, how do we translate sentences that involve more than one occurrence of "or"? This question will be addressed in [subsubsection 2.4.9.2](#). Second, there is the question of how to translate sentences that make use of exclusive "or" into **PL**. This question we will address in [subsubsection 2.4.9.7](#). Third, there is the question of how to decide whether or not a sentence is using "or" in an inclusive or exclusive way. This is a question with no foolproof solution, but let's consider four different strategies for determining whether or not a sentence is using "or" in an inclusive or exclusive way.

First, in some cases, speakers will clearly communicate that they intend the exclusive sense of "or" by adding the "but not both" construction at the end of their sentence. When the "but not both" construction is expressed, then the speaker is explicit that mean to use the exclusive sense of "or". Second, in other cases, facts about the meaning of words will make it clear that the exclusive sense is intended. For example, suppose two sports fans are arguing whether Messi or Ronaldo is the best soccer player. One of them says "Messi or Ronaldo is the best soccer player." Both fans are likely to interpret the meaning of "best" as implying a single player who is better than all others. They do not take the meaning of the word "best" to imply that there can be more than one best soccer player. Given this fact

about their language use, it is clear that the speaker intends the exclusive sense of "or" since it is not possible for both Messi and Ronaldo to be the best soccer player. Similarly, there is a tendency to use "or" in threats or ultimatums, e.g., "you can either confess or I'm calling the police." Third, some (but not all) speakers intend to use the "or" in an exclusive way when they preface the "P or Q" construction with the word "either". For example, suppose it is Renna's birthday and Tek is offering her dessert. Compare the two different sentences Tek might utter:

1. You can have cake or ice cream.
2. You can have either cake or ice cream.

In cases (1) and (2), the "or" is ambiguous, but the speaker is more likely to intend the exclusive sense of "or" in (2).

Fourth, and finally, in some cases, contextual factors and background information about speakers will help you decide whether "or" is being used inclusively or exclusively. For example, suppose Tek's daughter Renna has been misbehaving. After dinner (one which Renna ignored most of her vegetables), Renna asks for dessert. Renna knows that Tek is not a strict parent, but he is also not a pushover. Tek looks at her, upset with her and himself, and says in a stern tone "Renna, you can have cake *or* ice cream" (raising his voice on the word "or"). Renna knows that Tek is not the type of parent to give her both cake and ice cream. She also knows that Tek is not the type of parent to give her neither cake nor ice cream. Given the context and Renna's knowledge about Tek as a person (background information), Renna interprets Tek's utterance as using "or" in an exclusive way.

Exercise 2.22

Consider the following examples of English sentences that involve the use of "or". Try to determine if they make use of the inclusive or exclusive sense of "or".

1. Jane does not know how to read or write.
2. Suppose Tek is finishing his degree in English at university. His tuition is expensive and is paid for by his parents. His parents are not happy with his grades and so they tell him "you can either get better grades or you can pay for your own tuition."
3. You are being arrested. The police officer tells you "we can do this the easy way or the hard way."

4. You are in a heated argument with your friend. Your friend utters "you can either apologize or we are no longer friends."
5. You are at an ice cream shop and looking over your options. You say to yourself "I can either have chocolate or vanilla or mint chocolate chip or strawberry or cookies and cream or rocky road or butter pecan or pistachio or mocha almond fudge or chocolate chip cookie dough."
6. You are expressing a preference for one thing over a group of others. You say "I like drawing better than painting or sculpting or photography or writing."

$(\phi \rightarrow \psi)$

2.4.5 Conditionals

Recall that $\phi \rightarrow \psi$ is true if and only if $v(\phi) = F$ or $v(\psi) = T$. This means that it is false only in one case: the case where $v(\phi) = T$ and $v(\psi) = F$. To find a suitable translation for conditionals, we need to find a class of English sentences that are false just in that one case. There are several candidates, but the most popular are sentences of the form "if P, then Q". That is, propositions of the form "if P, then Q" can be translated into *PL* as $P \rightarrow Q$, and vice versa.

"The ability to think conditional thoughts is a basic part of our mental equipment.... There's not much point in recognizing that there's a predator in your path unless you also realise that if you don't change direction pretty quickly you will be eaten." - Dorothy Edgington.

P	Q	$P \rightarrow Q$	if P then Q
T	T	T	T
T	F	F	F
F	T	T	T
F	F	T	T

Suppose Tek says "If I buy a ticket, then I will win the lotto." In making this claim, he is not saying "he has bought a ticket", nor is he claiming "he will win the lotto". Rather, Tek uses the "if P, then Q" construction to state something about the relation between these two propositions. In saying that "If I buy a ticket, then I will win the lotto", Tek is saying that the scenario where he buys a ticket and where he does *not* win the lotto is ruled out. On this interpretation, his sentence is false just in the case where he buys a ticket and he does not win the lotto. But this is the same case where $\phi \rightarrow \psi$ is false. As such, it is proposed that we translate sentences of the form $\phi \rightarrow \psi$ as "if P, then Q" and vice versa.

It is not necessary to use the "if P, then Q" construction to express a conditional. For example, suppose Tek is sick. Liz is worried about him and says "If Tek is sick, he will not be at the party." This sentence is not of the form "if P, then Q" but is rather of the form "if P, Q". This shows

that the "then" is not necessary to express the conditional. Some other

1. Q if P
2. P only if Q (see [subsection 2.4.9.8](#))
3. In order for P, it is necessary that Q
4. It is sufficient for P, in order for Q
5. P is a sufficient condition for Q
6. Q is a necessary condition for P

Finally, there are at least two problems with translating $\phi \rightarrow \psi$ as "if P, then Q". The first problem relates to when ϕ and ψ are unrelated to each other (that is, there is no relevant relation between ϕ and ψ). The second problem concerns conditionals where the antecedent is known to be false (counterfactual conditionals). We return to both of these problems in [subsection 2.4.7](#).

2.4.6 Biconditionals

$$(\phi \leftrightarrow \psi)$$

Finally, recall that biconditionals $\phi \leftrightarrow \psi$ are true if and only if $v(\phi) = v(\psi)$. In other words, a biconditional is true whenever ϕ and ψ have the same truth value. Propositions of the form "P if and only if Q" (which is abbreviated way of saying "if P, then Q; and, if Q, then P") are true whenever P and Q have the same truth value, and so we translate sentences of the form "P if and only if Q" as $\phi \leftrightarrow \psi$.

P	Q	$P \leftrightarrow Q$	P if and only if Q
T	T	T	T
T	F	F	F
F	T	F	F
F	F	T	T

For example, consider the sentence "Liz brought her umbrella if and only if it is raining." Based on the above, this translation has the form of "U if and only if R" and the proposal is to translate this type of sentence as $U \leftrightarrow R$. Now let's consider the various scenarios and see whether they match the valuation function for the biconditional:

1. Liz brought her umbrella and it is raining. The sentence is true.
2. Liz brought her umbrella and it is not raining. The sentence is false since "U if and only if R" is false when U is true and R is false. This is because "if U, then R" is false.
3. Liz did not bring her umbrella and it is raining. The sentence is false since "U if and only if R" is false when U is false and R is true. This is because "if R, then U" is false.

4. Liz did not bring her umbrella and it is not raining. The sentence is true since "U if and only if R" is true when U is false and R is false. This is because "if R, then U" is true if "R" is false *and* "if U, then R" are true if "U" is false.

If it seems strange that $v(\phi \leftrightarrow \psi) = T$ when $v(\phi) = F, v(\psi) = F$, this is likely because people find it strange that $(\phi \rightarrow \psi) = T$ when $v(\phi) = F$.

Another example. Suppose Tek wants to go to a party, but he will only go if Liz is there. Tek is also Liz's best friend and so Liz will only go to the party if Tek is there. Tek says "I will go to the party if and only if Liz goes." What he means is that "if Liz goes to the party, then he will go" and "if he goes to the party, then Liz will go." Given the proposal to translate sentences of the form "P if and only if Q" as biconditionals, Tek's sentence is translated into *PL* as $P \leftrightarrow L$, where P stands for "Tek goes to the party" and L stands for "Liz goes to the party."

Exercise 2.23

Translate the following propositions into *PL*

1. John is a zombie.
2. John is a zombie and Mary is happy.
3. If John is a zombie, then Mary is happy.
4. If John is not a zombie, then Mary is not happy.
5. Either John is a zombie or Mary is happy.
6. John is a zombie if and only if Mary is happy.
7. If John is a zombie and Mary is happy, then either John is a zombie or Mary is not happy.
8. It is not the case that John is not a zombie.

2.4.7 Problems with Conditionals

In a previous section, it was remarked that translating every "if P then Q" as $\phi \rightarrow \psi$ *feels* wrong, especially when P is false. Let's try to make this feeling more precise. In the following subsections, we consider two different problems with translating "if P then Q" sentences as *PL* conditionals. The first stems from the fact that conditionals can be true even when the antecedent and consequent have no relevant relation to each other. This is especially a concern when "if P, then Q" is true in virtue of the antecedent of the conditional is false. Let's call this problem the *Relevance Problem*. The second problem is that not every sentence of the form "if P, then Q" should be translated as $\phi \rightarrow \psi$. For some "if P, then Q" sentences, we already know the antecedent of the conditional is false, but, contrary to the truth table for $\phi \rightarrow \psi$ says, we think it is worth debating whether "if P, then Q" is true or false. Let's call this problem the *Universality*

Problem.

2.4.8 The first problem: Relevance

In this section, we consider the first problem with translating sentences of the form "if P, then Q" as $\phi \rightarrow \psi$. The problem is that if we treat the conditions under which sentences of the form "if P, then Q" as $\phi \rightarrow \psi$, then we will have to accept that sentences of the form "if P, then Q" are true even when there is no relevant relation between P and Q. To see this more clearly, consider the truth table for $\phi \rightarrow \psi$

ϕ	ψ	$\phi \rightarrow \psi$
<i>T</i>	<i>T</i>	<i>T</i>
<i>T</i>	<i>F</i>	<i>F</i>
<i>F</i>	<i>T</i>	<i>T</i>
<i>F</i>	<i>F</i>	<i>T</i>

Note that if ψ is true, then $\phi \rightarrow \psi$ is true regardless of the truth value of ϕ . Similarly, if ϕ is false in $\phi \rightarrow \psi$, then $\phi \rightarrow \psi$ is true regardless of the truth value of ψ . This means that $\phi \rightarrow \psi$ is true even when ϕ and ψ are not relevantly related to each other. In sum, since the truth value of one of the two sides of $\phi \rightarrow \psi$ is, in some cases, sufficient to determine the truth value of the wff, the truth value of the other is, in that case, irrelevant to the truth value of the entire wff. This runs contrary to how many people think about "if P then Q" sentences. They contend that the truth value of both P and Q are relevant to whether "if P then Q" is true.

Let's summarize this problem in the form of an argument:

- P1: $v(\phi \rightarrow \psi) = T$ if $v(\psi) = T$ independent of the truth value of ϕ
- P2: $v(\phi \rightarrow \psi) = T$ if $v(\phi) = F$ independent of the truth value of ψ
- IC: Since the truth value of one of the two sides of $\phi \rightarrow \psi$ is, in some cases, sufficient to determine the truth value of the wff, the truth value of the other is, in that case, irrelevant to the truth value of the entire wff.
- P3: The truth value of P and Q are both relevant to determining the truth value of "if P then Q" propositions (everyday conditionals).
- C: Therefore, "if P then Q" propositions (everyday conditionals) should not be translated as $\phi \rightarrow \psi$ (PL conditionals).

Let's illustrate this problem with some English sentences. Suppose that we do not know whether God exists or not. Tek then utters the sentence "If rain is made of spaghetti sauce, then God exists." Since rain is not made of spaghetti sauce, the antecedent (as a matter of fact) is false. If

the antecedent is false, then the conditional is true. This strikes many as strange. Even further, not only is "If rain is made of spaghetti sauce, then God exists" true, but "If rain is made of spaghetti sauce, then God does *not* exist" is also true.

Let's consider a variation on this example from [6, p. 14]. Suppose you and I make a bet about whether Tek can square the circle in under five minutes. You say he can. I say he can't. Now suppose Tek's time is up. We would then take

If Tek has squared the circle, then you win the bet.

to be true. And, we would take

If Tek has squared the circle, then I win the bet.

to be false. However, since it is impossible to square the circle, the antecedent of the conditional is false. But, if the antecedent is false, then, if we translate all "if P, then Q" sentences as $\phi \rightarrow \psi$, both of the above conditionals are true (they are both true in virtue of having false antecedents). Since this is the wrong result, we have a reason to think that not all sentences of the form "if P, then Q" should be translated as $\phi \rightarrow \psi$.

A second illustration of the lack of relevance between the antecedent and the consequent of the conditional is illustrated by conditionals whose consequents are true. As noted, when the consequent of the conditional is true, the conditional is true, regardless of the truth value of the antecedent. To illustrate, let's assume that we do not know whether the Bigfoot exists. Its truth value is unknown given its elusive nature. We do, however, know that dandelions exist on earth. Now consider the following conditional: "If Bigfoot exists, then there are dandelions on earth." Since there are dandelions on earth, the consequent is true, and so the conditional, irrespective of the truth of the antecedent, is true.

This result strikes many as counterintuitive because in order for the conditional to be true, the truth of the consequent (the existence of dandelions) should *be relevantly related* to the truth of the antecedent (the existence of Bigfoot). However, *even if the antecedent were true*, it would not imply the truth of the consequent. That is, supposing that Bigfoot did exist (which is contrary to fact), there is nothing about the truth of this proposition that would entail the existence of dandelions.

2.4.8.1 The second problem: Universality

A second problem with translating all sentences of the form "if P, then Q" as $\phi \rightarrow \psi$ is that it is unclear that *every* sentence of that form should be translated as $\phi \rightarrow \psi$. One such example are counterfactual conditionals. A counterfactual conditional is a conditional where the antecedent is known to be false. For example, compare the following two conditionals (from Ernest Adams):

1. If Oswald did not shoot Kennedy, then someone else did.
2. If Oswald hadn't shot Kennedy, then someone else would have.

Sentence (1) is true since we know that Kennedy was shot. However, (1) does not, on its own, make any claim about whether Oswald did or did not shoot Kennedy. That is, it does not convey information that the antecedent of the conditional is true or false. It leaves open whether Oswald shot Kennedy. Conditionals of this type are called *indicative conditionals* or sometimes *conditionals in the realis mood* as they concern statements about what is really the case (actuality). In contrast, in sentence (2) the antecedent is assumed to be false. It does not leave open whether Oswald shot Kennedy. It asserts that Oswald did shoot Kennedy. In other words, the antecedent assumes that Oswald did shoot Kennedy, but considers the contrary-to-fact scenario where Oswald did not shoot Kennedy. Because sentence (2) assumes a contrary-to-fact scenario, it is called a *counterfactual conditional* or sometimes *conditionals in the irrealis mood* as they concern statements about what is not the case (non-actuality).

So what if the antecedent is assumed false in the counterfactual conditional? Why would this be a problem for translating "if P, then Q" sentences as $\phi \rightarrow \psi$?

Since *PL* is a truth-functional language, the truth value of $\phi \rightarrow \psi$ is determined by the truth value of ϕ and the truth value of ψ . In the case of the conditional, the truth value of the conditional is true when ϕ, ψ are, respectively, (T,T), (F,T), and (F,F). But while $\phi \rightarrow \psi$ is truth-functional, are sentences of the form "if P, then Q" truth-functional? In asking this question, we are asking whether the truth value of "if P, then Q" is determined by the truth value of P and the truth value of Q.

2.4.8.2 A Defense of the Conditional

Part of understanding the meaning of a proposition is knowing the circumstances under which the proposition is true. With respect to wffs of the

form $\phi \rightarrow \psi$, the circumstances under which this wff is true are expressed by the valuation function. It was suggested that the circumstances under which English propositions of the form "if P, then Q" are also expressed by that same valuation function.

P	Q	$P \rightarrow Q$
T	T	T
T	F	F
F	T	T
F	F	T

P	Q	if P, then Q
T	T	T
T	F	F
F	T	T
F	F	T

However, it is sometimes thought that while rows 1 and 2 of the truth table capture the circumstances under which "if P, then Q" is true and false respectively, there is a problem with rows 3 and 4. The intuition is that while "if P, then Q" sentences are *not false* at rows 3 and 4, it feels wrong to say they are true.

Let's try to address this concern by considering several arguments that sentences of the form "if P, then Q" are true when the antecedent is false. After considering these arguments, we will turn to flesh out the feeling that sentences of the form "if P, then Q" should not be true when the antecedent is false.

Argument 1

First, consider $P \rightarrow P$ when $v(P) = F$. If we accept the translation to "if P, then P", then what we have is not a sentence saying that P is the case, but one saying that *if* P is the case, then P is the case. Surely, sentences of this form are true. But notice that P can be true or false. When P is true, we have the first row of the truth table: $T \rightarrow T$. This row is uncontroversial. But, P can also be false: $F \rightarrow F$. This is the fourth row of the truth table. This row is controversial. But, if we believe that all sentences of the form "if P, then P" are true, then we have a reason to accept that the fourth row of the truth table is true.

A similar argument can be provided for row 3. Consider $(P \wedge R) \rightarrow P$. This wff is surely true since it is a version of the $P \rightarrow P$ wff we just considered. It says "if P and R are the case, then P is the case". Let's focus on the antecedent of this conditional and suppose that $v(P \wedge R) = F$ because $v(P) = T$ and $v(R) = F$. We know that when one of the conjuncts of a conjunction is false, the conjunction is false. Since the conjunction is false, the antecedent of the conditional is false. In other words, we have a wff $F \rightarrow P$, where F is the truth value of the antecedent. Let's focus on the consequent. We stated that $v(P) = T$. Since P is true, the consequent of the conditional is true. In other words, we have a wff $F \rightarrow T$, where

F is the truth value of the antecedent and T is the truth value of the consequent. In short, we have a conditional whose truth value assignment corresponds to the controversial third row of the truth table. Consider a less formal example. Suppose "P" stands for "Paul is angry" and "R" for "Robert is sad", then $(P \wedge R) \rightarrow P$ is translated as "If Paul is angry and Robert is sad, then Paul is angry". This proposition is true for the same reason that "if Paul is angry, then Paul is angry" is true. Therefore, we have a reason to accept that the third row of the truth table is true.

Let's consider a second argument for why rows (3) and (4) should be true. Argument 2
Intuitively, sentences of the form "if P, then Q" express propositions. That is, they are either true or false. While it is uncontroversial that rows (1) and (2) are true and false respectively, it is controversial whether "if P, then Q" should be true at rows (3) and (4). However, it is more controversial to say that they are false. If Tek says that "if it rains, I will bring my raincoat" but it does not rain, we certainly do not take Tek to have said something false, regardless of whether he brought his raincoat (row 3) or did not bring his raincoat (row 4). In short, if every proposition is either true or false, and "if P, then Q" sentences express propositions, and "if P, then Q" sentences are not false when P is false, then they must be true.

A third argument goes as follows. First, it might start by saying that it is difficult to determine whether sentences of "if P, then Q" are true or false on rows (3) and (4). Nevertheless, provided there is no reason to think that they are false, we should accept that they are true. Second, it might invite us to consider that the primary goal of logic is to identify "good arguments" from "bad arguments". One feature of a good argument is that the conclusion follows from its premises. That is, which arguments are *valid*. Surely one argument that is valid has the following form: if P, then Q, P; therefore Q. Now if evaluating "if P, then Q" sentences as true when P is false led to the consequence that this argument is *invalid*, this would be a disastrous consequence that would go directly against the central aim of logic. Argument 3

But, suppose the truth value of the wffs are determined in accordance with the third line of the conditional. That is, $v(P) = F$ and $v(Q) = T$ and so $v(P \rightarrow Q) = T$. Even if this is the case, then it will never be the case that $v(P \rightarrow Q) = T$ and $v(P) = T$ and $v(Q) = F$. Thus, treating "if P, then Q" statements as true when P is false and Q is true will never yield invalid arguments. Thus, while we might not have a compelling reason to treat "if P, then Q" sentences as though they were true when P is false, we don't have a compelling reason not to treat them as true either as they

do not lead to invalid arguments.

Argument 4

This example comes from volume 1 of [4, p. 34]

A fourth argument involves a consideration of mathematical claims and what are called "generalized conditionals". Consider the claim that "if a number is larger than 5, then it is larger than 3".. This sentence can be translated into the more perspicuous universal sentence "for every number x , if $x > 5$, then $x > 3$ ". This sentence has the effect of asserting not one, but an unlimited number of conditionals. Some of these include:

1. If 17 is larger than 5, then 17 is larger than 3.
2. If 6 is larger than 5, then 6 is larger than 3.
3. If 4 is larger than 5, then 4 is larger than 3.
4. If 2 is larger than 5, then 2 is larger than 3.

Since "for every number x , if $x > 5$, then $x > 3$ " is true, then each of these conditionals is true. But notice that in the case of sentence (3), the antecedent is false, but the consequent is true. This corresponds to row (3) of the table. In the case of sentence (4), both the antecedent and consequent are false. This corresponds to row (4) of the table. Thus, we have a reason to accept that "if P, then Q" sentences are true at rows (3) and (4).

"The worst defects of the truth-functional conditional don't show up in mathematics." - Dorothy Edgington[3]

Argument 5

This is an approach articulated by [7].

The fifth and final argument we will consider for why the feeling "if P, then Q" is true at rows 3 and 4 invites us to recall that "if P, then Q" sentences do not say that "P" is true or "Q" is false. Rather, they assert a relation between P and Q, namely that it is false to say that P and not-Q are the case. In this way, sentences of the form "if P, then Q" can be understood as a type of promise. Namely, they are a promise that you will not find P and not-Q together. One way of understanding the circumstances under which conditionals are true is to think of truth conditionals as *unbroken promises* and false conditionals as *broken promises*.

To illustrate the connection between promises and conditionals, let's begin by supposing that God exists, awakens you from your slumber one evening, and tells you the following:

If you are kind, you will go to heaven.

For convenience, let's abbreviate "you are kind" as K and "you will go to heaven" as H . What is God telling you? God is saying that it will not be the case that you are kind and fail to go to heaven. This can be rephrased as a promise. God is saying "I promise that if you are kind, you will go to heaven." Now if we understand true conditionals as unbroken promises and false conditionals as broken promises, then we can consider

the circumstances under which the conditional is true and false. Let's consider each of these circumstances in turn.

First, suppose that all of your life, you are very kind. You help your neighbors, you give to the poor, and you smile at everyone you see. That is, it is true that you are kind. Now, when you die, you go to heaven. That is, it is true that you go to heaven. We said that a false conditional is like a broken promise and a true conditional is like an unbroken promise. So, has God broken his promise? No, God kept his promise. Therefore, the conditional is true under this interpretation.

K	H	if K, then H
T	T	T
T	F	
F	T	
F	F	

Next, suppose that all of your life, you are very kind. You help your neighbors, give to the poor, and smile at everyone you see. That is, it is true that you are kind. But, when you die, you do not go to heaven. Instead, you are thrust into the hottest fires in hell. That is, it is false that God sends you to heaven. Has God kept his promise? No, in uttering the conditional, God claimed that the scenario where you are both kind and not on your way to heaven would not come about. But, that scenario did come about since you were kind and you are now burning in the fires of hell. We said that a false conditional is like a broken promise and a true conditional is like an unbroken promise. Therefore, the conditional is false.

K	H	if K, then H
T	T	T
T	F	F
F	T	
F	F	

At this point, none of this should be surprising since rows 1 and 2 are uncontroversial. So, now let's consider a third scenario. Suppose that you are not kind. You make life harder for your neighbors, you steal from the poor, and you scowl at everyone you see. But, when you die, God sees what a pitiful wretch you are and decides to send you to heaven anyway. Has God, in sending you to heaven, broken God's promise? No. In uttering the conditional, God was affirming that a certain scenario would not come about, namely the scenario where you are both kind and not in heaven. That scenario did not come about since you were not

kind. We said that a false conditional is like a broken promise and a true conditional is like an unbroken promise. Therefore, if a false conditional is like a broken promise and a true conditional is like an unbroken promise, and God has not broken his promise, then the "if K, then H" sentence is

*God did not assert that
he would send you to
heaven if and only if you
were kind.*

K	H	if K, then H
T	T	T
T	F	F
F	T	T
F	F	

Finally, suppose that you are not kind. Again, you are a cruel person, diabolical to the core. And, when you die, God sees what a pitiful wretch you are and decides not to send you to heaven. Instead, God throws you into the darkest corners of hell. Has God, in sending you to hell, broken his promise? The answer is no. In uttering the conditional, God claimed that the scenario where you are both kind and not on your way to heaven would not come about. That scenario did not come about since you were not kind. We said that a false conditional is like a broken promise and a true conditional is like an unbroken promise. Therefore, if a false conditional is like a broken promise and a true conditional is like an unbroken promise, and God has not broken his promise, then the "if K, then H" sentence is true.

K	H	if K, then H
T	T	T
T	F	F
F	T	T
F	F	T

And so, if we treat true conditionals as unbroken promises and false conditionals as broken promises, then we have a reason to accept that the third and fourth rows of the truth table are true.

2.4.9 More Complex Translations

In this section, more complex translation sentences in English are translated into **PL**. We restrict this discussion to sentences of the following form:

English	Translation
P and Q and R	$(P \wedge Q) \wedge R$ or $P \wedge (Q \wedge R)$
P or Q or R	$(P \vee Q) \vee R$ or $P \vee (Q \vee R)$
P and Q or R	$(P \wedge Q) \vee R$ or $P \wedge (Q \vee R)$
neither P nor Q	$\neg P \wedge \neg Q$
not both P and Q	$\neg(P \wedge Q)$
P or Q , but not both	$(P \vee Q) \wedge \neg(P \wedge Q)$
P only if Q	$P \rightarrow Q$
P even if Q	P
P unless Q	$P \vee Q$

2.4.9.1 P and Q and R

In an earlier section, we noted that propositions of the form "P and Q" are straightforwardly translated into *PL* as $P \wedge Q$. This was because the truth of "P and Q" depends on the truth of both P and Q in such a way that "P and Q" is true if and only if "P" is true and "Q" is true (it being false in all other cases). This is the same behavior we find in conjunctions of the form $P \wedge Q$. This wff is true just in the case that $v(P) = T$ and $v(Q) = T$, and false in all other cases. Speakers however chain together sentences with the use of "and" or some equivalent. That is, there are sentences of the form "P and Q and R" or the more abbreviated "P, Q, and R".

When speakers chain together propositions using "and", these sentences are true just in the case that each proposition connected by "and" is true. That is, "P and Q and R" is true just in the case that "P" is true, "Q" is true, and "R" is true. This is the same behavior we find in conjunctions of the form $(P \wedge Q) \wedge R$. This wff is true just in the case that $v(P) = T$, $v(Q) = T$, and $v(R) = T$, and false in all other cases. We thus can translate chains of propositions connected by "and" as conjunctions of conjunctions.

A few examples. Suppose Tek utters "Liz is happy and Renna is happy and I am happy." This sentence is true just in the case that "Liz is happy" is true, "Renna is happy" is true, and "Tek is happy" is true. In **PL**, this sentence can be translated as $(L \wedge R) \wedge T$. In order for $L \wedge R$ to be true, L and R must both be true. In order for $(L \wedge R) \wedge T$ to be true, $L \wedge R$ must be true and T must be true. In sum, $(L \wedge R) \wedge T$ is true just in the case that L is true, R is true, and T is true. Thus, a good translation of "Liz is happy and Renna is happy and I am happy" is $(L \wedge R) \wedge T$.

One more example. Suppose Liz is very hungry after her morning run. She

says "I am going to have eggs and toast and orange juice for breakfast." This sentence can be interpreted as a chain of sentences all connected by the use of "and". That is, the sentence is an abbreviated version of the sentence "Liz is going to have eggs for breakfast and Liz is going to have toast for breakfast and Liz is going to have orange juice for breakfast." Based on our proposal for translating chains of sentences connected by "and", the sentence should be translated as $(E \wedge T) \wedge O$. Notice that the placement of the parentheses does not impact the accuracy of the translation. That is, $(E \wedge T) \wedge O$ and $E \wedge (T \wedge O)$ are both true just in the single case that E is true, T is true, and O is true.

2.4.9.2 P or Q or R

We have considered how to translate sentences of the form "P and Q and R". What about sentences of the form "P or Q or R"? Translation can be done in exactly the same way, replacing each instance of "or" with the $\vee$. Let's consider a few examples.

Suppose Tek plans to exercise today. He is in great shape and so is considering what activities will be a part of his workout for the day. Tek says "I will lift or bike or stretch." We interpret what he says as "Tek will lift or Tek will bike or Tek will stretch." Since we know that Tek is in great shape, it is very well likely that Tek will do all of these actions. Therefore, we take Tek to be using "or" inclusively. Tek's utterance is true provided he does at least one of the actions (if he does more, then the sentence is still true). In **PL**, we translate Tek's utterance as $(L \vee B) \vee S$ or $L \vee (B \vee S)$ as these wffs are true provided at at least one of the following is true: L is true, B is true, S is true.

2.4.9.3 P and Q or R , P or Q and R

In addition to sentences that contain chains iterations of "and" or chains of "or", there are some sentences that employ both "and" and "or". Suppose Tek wakes up early and goes to a restaurant for breakfast. As he enters the restaurant, he sees a sign that reads DAILY SPECIAL: PANCAKES OR BACON AND EGGS. Tek is quite confused since the sign can be read in several different ways. First, Tek must decide whether the "or" in the sign is to be interpreted inclusively or exclusively. If the "or" is to be interpreted inclusively, then the Daily Special is quite generous since Tek could order everything: pancakes and bacon and eggs. However, if the "or" is to be interpreted exclusively, then Tek's options are more limited. Let's suppose that "or" is to be read exclusively. Let's update the daily special

sign using "XOR" to indicate that the "or" should be read exclusively:

DAILY SPECIAL: PANCAKES XOR BACON AND EGGS

Tek's recognition that the Daily Special sign can be read in two different ways based on the ambiguous use of "or" does not end Tek's confusion. Here is a second problem. Tek does not know the scope of XOR. That is, Tek still does not know which one of the following two options is the Daily Special:

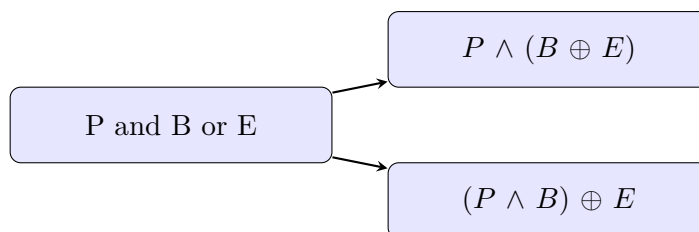
1. pancakes XOR (bacon and eggs).
2. (pancakes XOR bacon) *and* eggs.

To illustrate both of these readings of the Daily Special, consider the following scenario where reading (1) is intended. Suppose Tek asks about the Daily Special. The waiter says the following: "Our pancakes are gourmet pancakes. They are the best in the world. So, if you order these pancakes, you will only receive these pancakes, neither bacon nor eggs." However, if you don't like pancakes, you can order bacon and eggs without pancakes. In this scenario, the waiter clarifies the Daily Special to make it clear that the choice is between (a) just pancakes and (b) bacon and eggs.

Assuming the price is the same, interpretation (2) is a much better option for the consumer since it allows the consumer to order both pancakes and eggs.

In contrast, suppose reading (2) is intended. The waiter might clarify the Daily Special as follows: "The Daily Special comes with eggs, but you can choose between pancakes or bacon (you cannot have both)." In this scenario, the waiter clarifies the Daily Special to make it clear that the Daily Special comes with eggs and the choice is between (a) pancakes and (b) bacon.

Now that we have sorted through the two layers of confusion, how do we translate (1) and (2) into **PL**? First, let's translate "and" in the typical way and "or" using $\oplus$ to express exclusive "or". The result is the following: $P \oplus B \wedge E$. All that is left is to settle the scope of the operators as it is reflected in the two readings of the sentence:



On interpretation (1), the proper translation is $P \oplus (B \wedge E)$. On interpretation (2), the proper translation is $(P \oplus B) \wedge E$.

2.4.9.4 Neither P nor Q

$$\neg P \wedge \neg Q$$

Propositions of the form "neither P nor Q " are true just in the case that P is false and Q is false. In other words, "neither P nor Q " is another way of saying "not- P and not- Q ". In translating this sentence into **PL**, the goal is to select a **PL**-wff that is true just in the case that "neither P nor Q " is true. The most straightforward translation then of "neither P nor Q " is $\neg P \wedge \neg Q$ as this wff is true just in the case that P is false and Q is false. Another way of translating "neither P nor Q " is as $\neg(P \vee Q)$. Both translations are equally good because both propositions are true just in the case where P is false and Q is false.

P	Q	$\neg$	P	$\wedge$	$\neg$	Q	$\neg$	(	P	$\vee$	Q	)
T	T	F	T	F	F	T	F		T	T	T	
T	F	F	T	F	T	F	F		T	T	F	
F	T	T	F	F	F	T	F		F	T	T	
F	F	T	F	T	T	F	T		F	F	F	

To illustrate, let's consider two examples. Suppose Liz is worried about Tek. She confides in her friend that she isn't sure what is going on with him, telling her friend that "Tek is neither happy nor angry." What Liz is asserting is that it is not the case that Tek is happy and it is not the case that Tek is angry. We can translate this into PL as either $\neg H \wedge \neg A$ or $\neg(H \vee A)$. Both translations are equally good because both are true just in the case that H is false and A is false.

Let's consider a second example. Suppose Tek has been working for a company for several years now. In the previous year, he was denied a raise due to "budget issues". This year he has broached the issue of a raise or a promotion once again. Nevertheless, he is pessimistic about the outcome. In talking to a friend, Tek says the following "I will neither get a promotion nor a raise." What Tek is asserting is that *it is not the case that he will get a raise* and *it is not the case that he will get a promotion*. We can translate this into PL as either $\neg R \wedge \neg P$ or $\neg(R \vee P)$. Both translations are equally good because since both are true just in the case that R is false and P is false.

2.4.9.5 not both P and Q

Sentences of the form "not both P and Q " are denials of the sentence " P and Q ". That is, "not both P and Q " says that " P and Q " is false. If that is the case, then whenever " P and Q " is false, then "not both P and Q " is true, and vice versa. With this in mind, let's recall when " P and Q " is

true.

P	Q	P and Q
T	T	T
T	F	F
F	T	F
F	F	F

P	Q	not both P and Q
T	T	F
T	F	T
F	T	T
F	F	T

We now have a sense of when "not both P and Q" is true. How should these sentences be translated into **PL**? Some new students to logic are tempted to translate "not both P and Q" as either $\neg P \wedge \neg Q$ or as $\neg P \wedge Q$. Both of these translations are incorrect. First, as we saw in the previous section, $\neg P \wedge \neg Q$ is true when both P and Q are false. However, in the truth table for "not both P and Q", notice that this sentence is true in the row where P and Q are false. In addition, as we saw in the previous section, $\neg P \wedge \neg Q$ best captures sentences of the form "neither P nor Q". Second, $\neg P \wedge Q$ is also incorrect since this sentence is false when P is false and Q is true. But, again, if we examine row 3 of the truth table for "not both P and Q", we see that this sentence is true when P is false and Q is true. So, what is the correct translation of "not both P and Q"? The correct translation is $\neg(P \wedge Q)$. Sentences of the form "not both P and Q" are negations not merely of one of the conjuncts of the conjunction but of the entire conjunction. The negation is placed in front of the entire conjunction to express that the entire conjunction is false.

"not both P and Q" is not the same as "neither P nor Q", the former is the denial of the conjunction "P and Q", the latter is a conjunction of the denials of P and Q.

*"Not both P and Q" $\Rightarrow$
 $\neg(P \wedge Q)$*

Let's consider some examples. Suppose you are in college and you are registering for classes. Two classes you would like to take are PHIL012 - LOGIC and PHIL126 - METAPHYSICS. However, when you try to sign up for them, you receive an error that says that since these two classes are at the same time, it says "it is not the case that you can enroll in both PHIL012 and PHIL126." In this scenario, you are able to enroll in PHIL012 but not PHIL126. Alternatively, you are able to enroll in PHIL126 but not PHIL012. And, you are also able to enroll in neither course. But, it is not the case that you can enroll in both courses. The following wff can be used to translate "it is not the case that you can enroll in Logic and Metaphysics": $\neg(L \wedge M)$.

Suppose Tek's daughter Renna wants dessert. Tek tells her "it is not the case that you can have both cake and ice cream." What Tek says is false in exactly one case. This is the case where Renna can have both cake and ice cream. The case where she has both cake and ice cream is $C \wedge I$. Thus, Tek's statement is false when $C \wedge I$ is true. To express the case when his statement is true, we only need to negate the entire wff. In other words,

and so "you cannot have both cake and ice cream" is best translated as $\neg(C \wedge I)$.

2.4.9.6 *not P and Q*

Consider sentences of the form "not P and Q". At first glance, we might be tempted to take sentences of this form as a version of "not both P and Q", where the "both" is excluded to for the purpose of brevity. On this reading, sentences of this form should be translated as $\neg(P \wedge Q)$. For example, suppose Tek and Renna are planning her birthday. He might say to his daughter that "It is not the case that she can have cake and ice cream." Renna would be correct to interpret Tek as saying "It is not the case that I can have both cake and ice cream." In this case, Tek's utterance is best translated as $\neg(C \wedge I)$.

On the other hand, sentences of the form "not P and Q" might be interpreted as not denying the entire conjunction, but only the first conjunct of the conjunction. For example, suppose Tek's home has a gable vent that birds have been using to enter his attic. Tek must acquire supplies to fix the vent. There are two locations: FixItCity and HammerTown. After a quick internet search, Tek says "It is not the case that FixItCity is open, but HammerTown is open." This sentence has the form "not-F and H", but Tek means to deny that FixItCity is open and assert that HammerTown is open. In this case, Tek's utterance is best translated as $\neg F \wedge H$.

The general point then is that sentences of the form "not P and Q" are ambiguous. They can be interpreted as denying the entire conjunction or as denying the first conjunct and asserting the second conjunct. In the first case, the sentence is best translated as $\neg(P \wedge Q)$. In the second case, the sentence is best translated as $\neg P \wedge Q$. Nevertheless, there are some clues for deciding between the two translations. First, if the sentence is of the form "not P and Q" and the speaker is denying the entire conjunction, then the speaker is likely to use the word "both" in the sentence. For example, "It is not the case that you can have both cake and ice cream." Second, if the sentence is of the form "not P and Q" and the speaker is denying the first conjunct and asserting the second conjunct, then the speaker may use the word "but" or "while" rather than "and" in the sentence. For example, "It is not the case that FixItCity is open, but HammerTown is open." Third, if the sentence is of the form "not P and Q" and the speaker is denying the first conjunct and asserting the second conjunct, then the speaker may attach the negation to the main verb of the first

sentence rather than prefix the entire "and" sentence with a negation. For example, "FixItCity is not open, but HammerTown is open."

2.4.9.7 *P or Q, but not both*

In an earlier section of this chapter, two different senses of "or" were distinguished. The first sense was the inclusive sense of "or". With respect to the inclusive sense of "or" sentences of the form "P or Q" are translated as $P \vee Q$. The second sense was the exclusive sense of "or". When "or" is used exclusively, sentences of the form "P or Q" express "either P or Q, but not both P and Q".

In the section on disjunctions, we asked how to translate "P or Q" sentences when "or" was being used in the exclusive sense. In that section, we noted that there were at least two ways. The first way involved introducing a new operator: $\oplus$. So, "P or Q" ("or" used exclusively) is translated as $P \oplus Q$. However, in that section, we mentioned that it was possible to translate "P or Q" ("or" used exclusively) without introducing a new operator.

Let's consider a few examples. Suppose Tek says to Liz "you can have cake or ice cream, but not both." Let C stand for "you can have cake" and I stand for "you can have ice cream". Then Tek's statement can be translated as $C \oplus I$. Here is another example. Suppose Liz says to Tek, "either you pay me the money you owe me or I will sue you." Let P stand for "you pay me the money you owe me" and S stand for "I will sue you". Then Liz's statement can be translated as $P \oplus S$.

The second approach is to recall that when "or" is used exclusively, then "P or Q" expresses "P or Q, but not both P and Q". We can thus translate "P or Q" as $P \vee Q$, "not both P and Q" as $\neg(P \wedge Q)$, and then connect both wffs together with $\wedge$ (which is expressed by the word "but"). Thus, we can translate "P or Q, but not both P and Q" as $(P \vee Q) \wedge \neg(P \wedge Q)$.

Let's consider use the examples we used earlier involving Tek and Liz. First, suppose Tek says to Liz "you can have cake or ice cream, but not both." Tek's statement can be translated as $(C \vee I) \wedge \neg(C \wedge I)$. Next, consider Liz's utterance to Tek: "either you pay me the money you owe me or I will sue you" This sentence is translated as follows: $(P \vee S) \wedge \neg(P \wedge S)$.

Exercise 2.24

Translate the following sentences into **PL**. Use the following key:
R = "Renna is happy", H = "Renna is hungry", T = "Tek is happy",
and L = "Liz is happy"

1. Renna is happy and Tek is happy and Liz is happy.
2. Renna is happy or Tek is happy or Liz is happy.
3. Neither Renna nor Tek is happy.
4. Renna is not both happy and hungry.
5. Tek is happy or Liz is happy, but not both.

Exercise 2.25

Translate the following sentences into **PL**. Create your own translation key for each sentence.

1. Playing golf is a bourgeoisie sport and a waste of time and not a good use of land.
2. I don't like running or biking or swimming.
3. The government is neither a wise protector nor a benevolent parent.
4. Playing golf is neither beneficial for one's health nor does it promote sportsmanship.
5. It is neither just nor practical to ignore the needs of the poor.
6. We can either raise taxes or cut spending, but not both.

2.4.9.8 P only if Q

Propositions of the form "P only if Q" are typically translated as $P \rightarrow Q$ or as $\neg Q \rightarrow \neg P$. To see why this is the case, let's consider a few examples:

1. The match will light only if there is oxygen in the room.
2. The car will start only if there is gas in it.
3. I will pass this class only if I attend every class.

In case (1), what is being said is not that "if there is oxygen in the room, then the match will light". It is not saying that oxygen is sufficient (or enough) to make the match light. Rather, it is saying that "if the match lights, then there is oxygen in the room." In other words, it is saying that the match being lit guarantees that there is oxygen in the room. Given this way of thinking about (1), the correct translation of (1) is as follows: $M \rightarrow O$.

Another way of putting (1) is to say that oxygen being in the room is a (necessary) condition for the match to light. On this account, oxygen is a requirement for the match lighting so without the oxygen, the match will not light. To be more explicit, (1) says that "if it is not the case that there is oxygen in the room, then the match will not light." In thinking about (1) in this way, it seems natural to translate (1) as $\neg O \rightarrow \neg M$.

With the proposal that propositions of the form "P only if Q" are translated as either $P \rightarrow Q$ or as $\neg Q \rightarrow \neg P$, let's consider the remaining examples listed above. Proposition (2) is nearly identical to proposition (1). It says that "if the car starts, then there is gas in it." Gas being in the car is a (necessary) condition for the car to start but it being in the car does not guarantee the car will start. In other words, it may be true that the car has gas but false that the car starts since the car's starting depends upon several other conditions, e.g., the ignition works. So, just like in the case of (1), sentence (2) should be translated as the conditional $C \rightarrow G$ or as $\neg G \rightarrow \neg C$.

Finally, let's consider sentence (3). For the purpose of illustration, let's assume that the class the student is taking has a strict attendance policy: missing one class means automatic failure. When the student says "I will pass the class only if I attend every class", they are not saying that attending every class will guarantee that they receive a passing grade. That is, (3) should not be translated as "if I attend every class, then I will pass this class": $A \rightarrow P$. Rather, what is instead is that attending every class is a (necessary) condition for passing: perhaps many more things are required, e.g., getting a passing grade on every exam, doing the homework, etc. With that in mind, (3) is more naturally understood as saying "if I will pass the class, then I will attend every class."

One question that might be asked is whether there is any difference between the two translations of "P only if Q". To test this, we check to see whether $P \rightarrow Q$ and $\neg Q \rightarrow \neg P$ are true and false under every interpretation:

P	Q	$P \rightarrow Q$			$\neg Q \rightarrow \neg P$				
T	T	T	T	T	F	T	T	F	T
T	F	T	F	F	T	F	F	F	T
F	T	F	T	T	F	T	T	T	F
F	F	F	T	F	T	F	T	T	F

Notice above that whenever $P \rightarrow Q$ is true (or false), so is $\neg Q \rightarrow \neg P$. So, since the truth values of the wffs match under every interpretation, the two ways of translating "P only if Q" are equally good.

2.4.9.9 *P even if Q*

Propositions of the form "P even if Q" express the fact that *P* is the case and the truth value of *Q* has no bearing on the truth value of *P*. That is, if *Q* is true, then *P* is true, and if *Q* is false, then *P* is true. There are several ways of translating sentences of this form into **PL**.

The first, and simplest, is to translate take "P even if Q" as saying "P regardless of Q". On this apporeading the sentence is "P is the case regardless of whether Q is true or false." With that in mind, we can translate "P even if Q" sentences as simply saying "P is the case".

A second option is to take "P even if Q" as shorthand for the much longer utterance "if Q then P and if not-Q then P". On this reading, we translate both conditionals and then combine them with $\wedge$. That is, we translate "P even if Q" as $(Q \rightarrow P) \wedge (\neg Q \rightarrow P)$. Both translations are equally good since whenever *P* is true, then $(Q \rightarrow P) \wedge (\neg Q \rightarrow P)$ is true and whenever *Q* is false, then $(Q \rightarrow P) \wedge (\neg Q \rightarrow P)$ is false.

A few examples. First, suppose Tek says "Liz will go to the party even if Mary goes". Liz is saying that regardless of whether Mary goes or not, Liz will be at the party. We thus can translate this sentence as emphatically saying *L*. On the other approach, we can read the sentence as shorthand for saying the much longer "if Mary goes to the party, then Liz will go to the party *and* if Mary does not go to the party, then Liz will go to the party." On this reading, we can translate the sentence as $(M \rightarrow L) \wedge (\neg M \rightarrow L)$.

Second, suppose Renna is speculating on stock prices. She believes that stock XYZ will go up even if the economy falters. She tells her clients that "the price of stock XYZ will go up in price even if the economy falters." As mentioned before, Renna may be interpreted as saying "the stock XYZ will go up regardless of whether the economy falters or not". On this reading, her utterance can be translated as simply saying "the price of stock XYZ will go up." On the other reading, her utterance expresses the longer "if the economy falters, then the price of stock XYZ will go up *and* if the economy does not falter, then the price of stock XYZ will go up." On this reading, her utterance can be translated as $(E \rightarrow S) \wedge (\neg E \rightarrow S)$.

2.4.9.10 *P unless Q*

One of the more perplexing group of sentences to translate into *PL* are sentences of the form "P unless Q". It is commonly suggested that "P

unless Q" is another way of saying "P or Q". Since "P or Q" is ambiguous between an inclusive and exclusive reading, this suggests that "P unless Q" is also ambiguous between the two readings.

Let's consider whether this is the case with some examples. First, suppose Tek and Liz were dating, but recently ended their relationship on bad terms. Vic has invited Tek to a party, but there is the possibility that Liz will attend. Tek tells Vic the following: "I will attend the party, unless Liz is there." We would likely judge this sentence false in two cases: (1) Tek and Liz both attend the party and (2) neither Tek nor Liz attend the party. If we represent this in a table, the truth function corresponds to the same function as the exclusive or.

Tek attends	Liz attends	Tek attends unless Liz attends.
T	T	F
T	F	T
F	T	T
F	F	F

However, another way of reading Tek's utterance that "I will be at the party, unless Liz is there" is that he is saying "if Liz is not at the party, then I will be at the party." Let's consider when this sentence is true using a table (notice that the sentence is "Liz is *not* at the party"):

Tek attends	Liz attends	Liz does not attend	if Liz is not at the party, then I will be at the party
T	T	F	T
T	F	T	T
F	T	F	T
F	F	T	F

Notice that the above truth function corresponds to the same function as the inclusive or.

In sum, "P unless Q" is ambiguous between two different possible readings. It may be translated as $P \oplus Q$ where "or" is being used in the exclusive sense or $P \vee Q$ where "or" is used in the inclusive sense.

Exercise 2.26

Translate into **PL**

1. John is happy, Mary is tired, and Frank is surprised.
2. John is happy or Mary is not tired or Frank is not surprised.
3. Neither John is happy nor Mary is tired.
4. Either John is happy or Mary is tired, but not both.

5. John knows how to dance only if pigs know how to fly.
6. I will be at the party even if the party is cancelled.
7. Suppose we are Tek's friend. We want him to come to our birthday party. Unfortunately, while Tek wants to go to a party, he will not come without Liz. He tells us the following "I will go to the party only if Liz goes." Suppose Tek clarifies himself by saying "if he is at our party, then Mary is at the party." What are the two equally good ways of translating Tek's utterance "P only if L".
8. Suppose it is raining and Tek says "It rains only if I have my umbrella". What conclusion can you draw about whether Tek has his umbrella if Tek's sentence is true.

2.5 END OF CHAPTER EXERCISES

Exercise 2.27

1. List all of the symbols of **PL**
2. How many propositional letters are there in **PL**?
3. Is the following a wff in **PL**: $(\neg(P) \rightarrow \neg(Q))$?
4. Use the *PL* formation rules to prove that $(\neg(P) \vee \neg(Q))$ is a wff in **PL**.
5. What is the difference between an atomic wff and a complex wff?
6. What is the scope of the leftmost occurrence of $\neg$ in $(\neg(P) \vee \neg(Q))$?
7. What are all of the subformulas of $(\neg(P) \vee \neg(Q))$?
8. What is the main operator of a wff?
9. What is the main operator of $(\neg(P) \vee \neg(Q))$?
10. What is the literal negation of $(\neg(P) \vee \neg(Q))$?
11. How would you represent the following wff using the conventions for making wffs more readable: $((\neg(P) \vee Q) \wedge R)$
12. What is a function?
13. Can the input (item in a domain) of a function be related to more than one output (item in the range)? Why or why not?
14. What would be wrong with saying that there is a function that relates letters A, B, C (the inputs) to numbers 1, 2, 3 (the outputs) whereby A is related to 2, B is related to 1, and C is related to 2?

15. What is an interpretation in **PL**?
16. Can a single propositional letter P be assigned a truth value of both T and F under a single interpretation?
17. If a single propositional letter P receives an interpretation, is there anything wrong in saying that P is neither true nor false?
18. What is a valuation in **PL**?
19. What is the difference between a valuation and an interpretation?
20. Using the valuation function, what is the truth value for $P \rightarrow Q$ if $v(P) = T$ and $v(Q) = F$?
21. How would you translate the following English sentences in **PL**:
 - a) It is not the case that John is happy.
 - b) If John is happy, then Mary is happy.
 - c) John is happy and Mary is happy.
 - d) John is happy or Mary is smart.
 - e) John is happy if and only if Mary is happy.

In the first chapter, it was noted that there are two informal tests for determining whether or not arguments are valid or invalid: the intuition and imagination tests. With respect to the imagination test, it was said that an argument is identified as *deductively valid* if and only if it was impossible to imagine a scenario where all of the premises were true and conclusion false. If it is possible to imagine such a scenario, then the argument is invalid. It was also noted that the “imagination test” has numerous problems: the use of the imagination test will occasionally produce the wrong results, some arguments are so large that it can be difficult to imagine the argument in its totality, and biases often obstruct an individual’s openness to consider scenarios that would make the premises true and the conclusion false. What we would like then is a test that yields the same result every time, that does not rely upon the limited imaginative powers of human beings, and that works independently of human bias.

In this chapter, an alternative method for determining the validity of an argument is proposed. This alternative method is the method of truth tables. Let’s begin by defining a truth table.

Definition 3.0.1: truth table

A truth table for **PL** is a table that provides a graphical way of representing a valuation of a wff or set of interpretations.

In the previous chapter, we saw truth tables when we represented the definition of the valuation function for the various types of complex wffs. For example, the valuation function for conjunction was represented as follows:

ϕ	ψ	$\phi \wedge \psi$	$\phi \vee \psi$	$\phi \rightarrow \psi$	$\phi \leftrightarrow \psi$
<i>T</i>	<i>T</i>	<i>T</i>	<i>T</i>	<i>T</i>	<i>T</i>
<i>T</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>F</i>	<i>F</i>
<i>F</i>	<i>T</i>	<i>F</i>	<i>T</i>	<i>T</i>	<i>F</i>
<i>F</i>	<i>F</i>	<i>F</i>	<i>F</i>	<i>T</i>	<i>T</i>

Truth Table 3.1: Truth Table: Conjunction, Disjunction, Conditional, and Biconditional

A truth table can be used to determine whether sets of wffs have certain

properties. The manner in which they can do this is “mechanical”, namely they can be used, in a finite number of steps, to determine whether some set of wffs has some property (e.g. validity).

3.1 DETERMINING THE TRUTH OF A WFF

To understand how to use the method of truth tables, we first need to know how to determine the truth value of any complex wff. To do this, consider that a complex wff is composed of subformulas and some of these subformulas are propositional letters. Here then is our method:

1. Use the interpretation of propositional letters to assign truth values to the propositional letters in the wff
2. In the order in which the wff β is constructed, assign truth values to β 's subformulas $\phi_1, \phi_2, \dots$ using (1) the truth values of the subformulas $\psi_1, \psi_2, \dots$ that construct $\phi_1, \phi_2, \dots$ and (2) the valuation function that corresponds to the operator introduced in the construction of $\phi_1, \phi_2, \dots$

Both of the steps of this method are a mouthful. If you don't understand them at first, do not worry, we will break each step down and consider several examples of how the steps work.

Let's start with the wff $P \wedge Q$ where $\mathcal{I}(P) = T$ and $\mathcal{I}(Q) = F$. Step 1 tells us to assign truth values to the propositional letters in our wff $P \wedge Q$ using our interpretation. Essentially, we are using the valuation function on the propositional letters using the interpretation as input. To illustrate, let's write out our wff and write the truth values of the propositional letters under the propositional letters in the wff. That is, since $\mathcal{I}(P) = T$ and $\mathcal{I}(Q) = F$, write T under P and F under Q .

$$\frac{P \quad \wedge \quad Q}{T \quad \quad F}$$

That concludes step 1 of our method. Next, let's perform step 2. We will attack this step in two stages:

1. Construct the wff using the formation rules
2. In the order in which the wff is constructed, assign truth values to subformulas of the wff using (1) the truth values already assigned to subformulas and (2) the valuation function.

First, construct $P \wedge Q$ using the formation rules:

1. P, Q are wffs (since all propositional letters are wffs)

2. Since P and Q are wffs, $(P \wedge Q)$ is a wff (since any two wffs ϕ, ψ can be combined to create a wff $(\phi \wedge \psi)$)

Next, we will assign truth values to subformulas in the same order in which we constructed the wff. Our wff was constructed in two steps. The first step of our construction stated that the propositional letters P and Q are wffs:

1. P, Q are wffs. (Truth values already assigned!).

The second step of the construction created the conjunction $P \wedge Q$. So we will use the truth values of P and Q (which have already been determined) and the valuation function for the conjunction to assign a truth value to the subformula $P \wedge Q$.

2. Since P and Q are wffs, $(P \wedge Q)$ is a wff. (If $v(P) = T$ and $v(Q) = F$, then $v(P \wedge Q) = F$).

Let's break this step down. To determine the truth value of $P \wedge Q$ we needed two things. First, we needed the truth value of P and Q . Second, we needed the valuation function for conjunction where the leftside of the conjunction is true and the rightside of the conjunction is false. Let's look at the truth values for the operators.

ϕ	ψ	$\phi \wedge \psi$	
T	T	T	
T	F	F	$\Leftarrow$
F	T	F	
F	F	F	

We can again represent how we have gone about determining the truth value of $P \wedge Q$ using a table. Let's write the truth value assigned to $P \wedge Q$ under the operator for conjunction. Writing it under the conjunction will indicate that the subformula $P \wedge Q$ has that truth value.

P	$\wedge$	Q
T	F	F

Before considering more complex examples, let's return to our two-step process of determining the truth value of complex wffs. First, we take an interpretation and use that interpretation to determine the truth value of the propositional letters in the wff. Second, we then determine the truth values of the subformulas of that wff in the same order in which that wff was constructed. To determine the truth values of those subformulas,

we use the valuation function that corresponds to the operator in that subformula and the truth values that are assigned to the subformulas that create the subformula whose truth value we are trying to determine.

Let's consider another example. In this example, our focus will be more on using the method rather than understanding the definition of the method. To start, consider a simple example of how to determine $P \rightarrow \neg R$ under a single interpretation of P and R : $\mathcal{I}(P) = T$ and $\mathcal{I}(R) = F$. First, we write out the formula or set of formulas you want to test (see [Truth Table 3.2](#)).

$$\frac{P \rightarrow \neg R}{}$$

Truth Table 3.2: Truth Value for $P \rightarrow \neg R$

Next, we write the truth values below propositional letters. Since $\mathcal{I}(P) = T$, write T under P and since $\mathcal{I}(R) = F$, write F under R .

$$\frac{P \rightarrow \neg R}{T \qquad F}$$

Next, we want to assign truth values to the subformulas. The question then is whether we determine the truth value of $\neg R$ or $P \wedge \neg R$ first. To decide, you should again consider how the wff is constructed:

1. P, R are wffs.
2. Since R is a wff, then $\neg R$ is a wff.
3. Since $P, \neg R$ are wffs, then $P \wedge \neg R$ is a wff.

Now that we see that $\neg R$ is constructed before $P \wedge \neg R$, the truth value of $\neg R$ should be determined before $P \wedge \neg R$. Our next step then is to use the truth value assigned to R and use this truth value along with the valuation function for negation to determine the truth value of $\neg R$ (see [Truth Table 3.3](#)).

$$\frac{P \rightarrow \neg R}{T \qquad T \quad F}$$

Truth Table 3.3: Truth Value for $P \rightarrow \neg R$

The last step is to determine the truth value of the largest subformula in the wff. More frankly, which truth values should we use to determine what truth value we write under the right arrow? Should we use the truth values written under P and R or the truth values written under P and $\neg R$?

To determine this, we should again consider how the wff is constructed. When $P \rightarrow \neg R$ is constructed, it is constructed using P and $\neg R$. The truth values assigned to these subformulas then are the ones that should be run through the conditional valuation function. Since $v(P) = T$ and $v(\neg R) = T$, we should write T under the rightarrow:

$$\frac{P \quad \rightarrow \quad \neg \quad R}{T \quad T \quad T \quad F}$$

Truth Table 3.4: Truth Value for $P \rightarrow \neg R$

Exercise 3.28

Let $v(P) = T$, $v(R) = F$ and $v(S) = T$

1. $\neg P \rightarrow R$
2. $\neg(P \vee R)$
3. $\neg(\neg P \leftrightarrow \neg S)$
4. $P \wedge \neg R$
5. $(P \vee R) \leftrightarrow \neg S$
6. $P \vee \neg R$
7. $\neg(P \vee R) \leftrightarrow (\neg P \wedge \neg R)$
8. $\neg(P \rightarrow S) \rightarrow R$
9. $\neg S \vee S$
10. $(P \rightarrow R) \rightarrow \neg W$

3.2 THE TRUTH-TABLE METHOD

In the above examples, the truth value of complex wffs are determined under one interpretation of the propositional letters. The truth-table method, however, is more general than this as it allows for determining the truth value of wffs under all admissible interpretations of the propositional letters. For example, consider the following wff: $\neg P \vee \neg R$.

First, write out all of the propositional letters in the formula in a separate column and the formula or formulas you want to test to the right of it:

$$\frac{P \quad R \quad | \quad \neg \quad P \quad \vee \quad \neg \quad R}{}$$

Second, consider the different possible interpretations for the propositional letters in a wff:

Third, for each row, write the truth values under the corresponding letter in the row. In our example, P is T in row 1. Therefore, you should write

P	R	$\neg$	P	$\vee$	$\neg$	R
T	T					
T	F					
F	T					
F	F					

T under every P in row 1. In row 2, P is also T. Therefore, you should again write T under every P in row 2. In row 3, P is F. Therefore, you should write F under every P in row 3. Finally, P is F in row 4. Therefore, you should write F under every P in row 4. This same process should be performed for every propositional letter in the wff.

P	R	$\neg$	P	$\vee$	$\neg$	R
T	T		T			T
T	F		T			F
F	T		F			T
F	F		F			F

Truth Table 3.5: Truth Table for $\neg P \vee \neg R$

Fourth, for each row, determine the truth value of the subformulas in that row. Earlier in this chapter, we saw how to do this for a complex wff under one interpretation. Now we are going to do it for multiple interpretations (multiple rows). Let's consider this for one row in our example. Take row 1. At row 1, we have the truth values of P and R in the table. The next step is to determine the truth value of either $\neg P$ or $\neg R$ since these are the next two subformulas that we would construct in order to construct the wff $\neg P \wedge \neg R$. Notice that P is T at row 1. We would thus use that truth value along with the valuation function for negation to determine the truth value of $\neg P$. This would be F. We would then write F under $\neg P$. The same step should be performed for $\neg R$ at row 1. This should be done for each row in the table.

P	R	$\neg$	P	$\vee$	$\neg$	R
T	T	F	T		F	T
T	F	F	T		T	F
F	T	T	F		F	T
F	F	T	F		T	F

Truth Table 3.6: Truth Table for $\neg P \vee \neg R$

Now that we have determined the truth value for $\neg P$ and $\neg R$, use the truth values assigned to these subformulas to determine the truth value

of $\neg P \wedge \neg R$ for each row in the table. Let's consider row 1. Notice that the truth value of $\neg P$ is F and the truth value of $\neg R$ is F. Using these truth values and the valuation function for disjunction ($\vee$), we determine that the truth value of $\neg P \wedge \neg R$ at row 1 is F. Therefore, write F under the disjunction operator in row 1. This same procedure should be performed for each row in the truth table.

P	R	$\neg$	P	$\vee$	$\neg$	R
T	T	F	T	F	F	T
T	F	F	T	T	T	F
F	T	T	F	T	F	T
F	F	T	F	T	T	F

Truth Table 3.7: Truth Table for $\neg P \vee \neg R$

Our truth table is complete. **Truth Table 3.7** shows the truth value of $\neg P \vee \neg R$ under all of the different ways that P and R can be interpreted in **PL**.

Exercise 3.29

Determine the truth values of the following wffs under all of the different ways that the propositional letters in that wff can be interpreted.

1. $P \rightarrow \neg R$
2. $(P \wedge R) \rightarrow R$
3. $\neg \neg (P \leftrightarrow R) \vee Z$
4. $\neg P \vee \neg Q$
5. $\neg (P \wedge \neg Q)$
6. $(P \leftrightarrow Q) \wedge (P \rightarrow \neg Q)$
7. $P \rightarrow (\neg P \wedge \neg Q)$
8. $(P \leftrightarrow Q) \wedge (P \leftrightarrow \neg Q)$
9. $\neg \neg P \wedge (P \leftrightarrow \neg Q)$
10. $(P \rightarrow Q) \rightarrow (P \vee \neg Q)$

3.3 TRUTH-TABLE ANALYSIS

Thus far, we have shown how to use truth tables to determine the truth value of a complex wff under an interpretation. Truth tables have an additional use in that they can be employed to determine whether a certain wff or set of wffs has a certain property. This, it will be shown, is useful for it allows us a way to determine whether propositions, groups

of propositions, or arguments in English have properties that interest us, e.g. whether an argument is deductively valid.

3.3.1 Contingency, Tautology, Contradiction

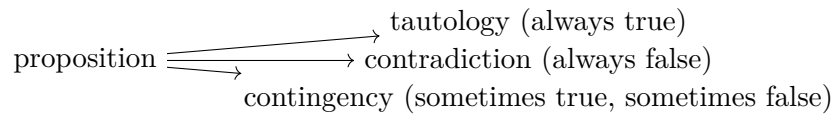
Consider whether the sentence "either there is an apple in my house or there isn't" is true or false. If there is an apple in my house, then this sentence is true. And, if there is not an apple in my house, then this "or" sentence is true. What we have then is a sentence that is *always true* regardless of whether "there is an apple in my house" is true or false.

But now notice something remarkable. Suppose we translate the "either there is an apple in my house or there isn't" as $A \vee \neg A$. And, further, suppose instead of considering the specific sentence "there is an apple in my house" (which we translated as A), we consider any sentence / wff that has the form $A \vee \neg A$. Let's represent this as $\phi \vee \neg\phi$. Notice that just as $A \vee \neg A$ is always true, we can substitute any sentence / wff for ϕ in $\phi \vee \neg\phi$ and the resulting sentence / wff will always be true. This is because if ϕ is true, then $\phi \vee \neg\phi$ is true. And, if ϕ is false, then $\phi \vee \neg\phi$ is true. Let's test this with an example where we replace "Tek is dancing" for ϕ :

- $\phi \vee \neg\phi$
- Tek is dancing or Tek is not dancing.
- Suppose "Tek is dancing" is true. It follows that "Tek is dancing or Tek is not dancing" is true.
- Suppose "Tek is dancing" is false. It follows that "Tek is dancing or Tek is not dancing" is true.
- Therefore, "Tek is dancing or Tek is not dancing" is always true (it is true regardless of whether "Tek is dancing" is true or false).

What we have found then is a sentence / wff that is always true, not in virtue of what it says, but always true in virtue of its form or structure. It is a sentence that is true regardless of the truth values assigned to the sentences that compose that sentence. Let's call these sentences "tautologies". Similarly, some sentences are *always false* in virtue of their form, e.g., "the apple both exists and does not exist". Such sentences are false regardless of the truth values assigned to the sentences that compose that sentence. Let's call these sentences "contradictions". Finally, there are sentences that are neither always true nor always false. Rather, their truth value depends upon the truth value of the sentences that compose that sentence. Let's call these sentences "contingencies".

Since propositional logic makes use of wffs rather than sentences, let's



define tautologies, contradictions, and contingencies in PL :

Definition 3.3.1: PL-Tautology

A wff ϕ is a PL-tautology (a logically valid formula) if and only if ϕ is true under every interpretation.

For example, $P \vee \neg P$ is a PL-tautology. It is a wff that is true under every interpretation of P .

Definition 3.3.2: PL-Contradiction

A wff ϕ is a PL-contradiction if and only if ϕ is false under every interpretation.

For example, $P \wedge \neg P$ is a PL-contradiction. It is a wff that is false under every interpretation of P .

Definition 3.3.3: PL-Contingency

A wff ϕ is a PL-contingency if and only if ϕ is neither always false under every valuation nor always true under every interpretation.

For example, $P \wedge Q$ is a PL-contingency. It is a wff that is neither a PL-tautology nor a PL-contradiction. The truth value of $P \wedge Q$ is true under some interpretations of P and Q and false under other interpretations of P and Q .

The next step is to determine, for any given wff ϕ , whether ϕ is a PL-tautology, PL-contradiction, or PL-contingency. To make this determination, we can use a truth table. We can simply construct a truth table for the wff ϕ , then check whether the wff is true under every interpretation (tautology), false under every interpretation (contradiction), or neither true nor false under every interpretation (contingency). Let's consider each logical property and then an example.

Definition 3.3.4: Truth-Table Test for Tautology

A truth table for a PL-tautology will have T under its main operator for every row (or in the case of no operators, under the propositional letter).

Earlier, it was claimed that $P \vee \neg P$ is a tautology. Let's use a truth table to test this claim.

P	P	$\vee$	$\neg$	P
T	T	T	F	T
F	F	T	T	F

Notice that in every row under the main operator of $P \vee \neg P$, the truth value is T. Therefore, $P \vee \neg P$ is a tautology according to the truth-table test for tautology.

Definition 3.3.5: Truth-Table Test for Contradiction

A truth table for a contradiction will have all Fs under its main operator (or in the case of no operators, under the propositional letter).

Earlier, it was claimed that $P \wedge \neg P$ is a contradiction. Let's use a truth table to test this claim.

P	P	$\wedge$	$\neg$	P
T	T	F	F	T
F	F	F	T	F

Notice that in every row under the main operator of $P \wedge \neg P$, the truth value is F. Therefore, $P \wedge \neg P$ is a contradiction according to the truth-table test for contradiction.

Definition 3.3.6: Truth-Table Test for Contingency

A truth table for a contingency will have at least one T and at least one F under its main operator (or in the case of no operators, under the propositional letter).

Earlier, it was claimed that $P \wedge Q$ is a contingency. Let's use a truth table to test this claim.

P	Q	$P \wedge Q$
T	T	T
T	F	F
F	T	F
F	F	F

If we look under the main operator for $P \wedge Q$, it is not the case that every row is T or every row is F. Instead, we have at least one T and at least one F. Therefore, $P \wedge Q$ is a contingency according to the truth-table test for contingency.

Let's consider one final illustration for a wff that we might not immediately know whether it is a contradiction, tautology, or contradiction. Consider $\neg(\neg P \rightarrow \neg Q)$. Is the following wff a contradiction, tautology, or contradiction:

P	Q	$\neg(\neg P \rightarrow \neg Q)$
T	T	F
T	F	F
F	T	T
F	F	F

The main operator of $\neg(\neg P \rightarrow \neg Q)$ is the leftmost negation. Using the truth table method, we check every row to see whether each row is T (in which case it would be a tautology), whether each row is F (in which case it would be a contradiction), or whether there is at least one T and at least one F (in which case it would be a contingency). The truth table for $\neg(\neg P \rightarrow \neg Q)$ shows that this wff is a PL-contingency since there is at least one T and at least one F under the main operator.

Exercise 3.30

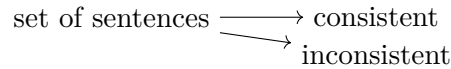
Using the truth-table method, determine whether the following wffs are tautologies, contradictions, or contingencies

1. $\neg P \rightarrow \neg P$
2. $(P \wedge \neg P) \wedge Q$
3. $P \leftrightarrow \neg R$
4. $P \rightarrow (P \vee Q)$
5. $\neg \neg P \wedge P$
6. $\neg(P \vee \neg R)$
7. $P \rightarrow (Q \rightarrow P)$

8. $R \rightarrow \neg R$
9. $(P \rightarrow Q) \wedge (\neg Q \rightarrow \neg P)$
10. $(\neg P \wedge \neg Q) \wedge P$

3.3.2 Consistency

Consider the sentences "John is tall" and "Mary is happy". Can both of these sentences be true at the same time? Sure. If John is tall is true and Mary is happy is true, then they are both true. Now consider "John is an elephant" and "John is not an elephant". Can both of these sentences be true at the same time? If "John is an elephant" is true, then "John is not an elephant" is false. Alternatively, if "John is an elephant" is false, then "John is not an elephant" is true. So, it appears that there is no way for both of these sentences to be true at the same time. Let's say that a set of sentences is *consistent* (or the sentences in the set are consistent with each other) if and only if all of the sentences in the set can be true at the same time. If a set of sentences is *not consistent*, then that set would be *inconsistent*.



Surely, there are cases where determining whether a set of sentences is consistent or inconsistent is important. A politician may make a collection of promises to their constituents. They may try to sway one swath of voters with one set of promises and another swath of voters with a different set of promises. Is the set of all of those promises consistent? Or, is the politician not capable of making good on all of their promises? Knowing that the set of promises made by the politician is inconsistent gives us grounds for doubting whether the politician is trustworthy. A scientific theory may make certain observable predictions about the world. When those predictions conflict with observation, there is an inconsistent set consisting of (1) the scientific theory, (2) its predictions what the world will look like if the scientific theory is true, and (3) our account of what the world actually looks like. Knowing that not all three can be true, at least one must be false. Perhaps the scientific theory is mistaken. Perhaps how we derived the predictions from the theory is wrong. We know that we sometimes fail to accurately describe what we observe, so perhaps our account of what we think the world actually looks like is wrong. Last, but not least, being able to determine if a set is consistent or inconsistent may

be helpful for revising our own stock of beliefs. We surely believe many things. We may have beliefs about our future, about religion and politics, about the nature of reality, or even about what movies or music are worth consuming. It is more than likely that this set of beliefs is inconsistent but we fail to recognize it. If we are interested in repairing or minimizing this inconsistency, a first step is to recognize where it occurs. Once we have done this, we can take the necessary steps of deciding how to revise our beliefs or whether some beliefs ought to be given up.

Let's return to logic. In **PL**, a non-empty set of wffs is consistent in **PL** provided all of the wffs in the set are true under at least one interpretation.

Definition 3.3.7: Consistency

A non-empty set of wffs $\{\phi_1, \phi_2, \dots, \phi_n\}$ is consistent if and only if each wff in $\{\phi_1, \phi_2, \dots, \phi_n\}$ is true under at least one interpretation.

Two points. First, we are defining consistency for non-empty sets of **PL** wffs. It would not be defined for other kinds of sets, e.g., sets of strings, sets of numbers, the empty set. Second, a set with a single wff that is either a contingency or a tautology would be consistent, while a set consisting of a contradiction would not be consistent.

More concretely, the set $\{P \wedge Q, P \vee Q\}$ is consistent since there is at least one interpretation of P and Q where each wff in that set is true: this is the interpretation where $\mathcal{I}(P) = T, \mathcal{I}(Q) = T$. One way to determine whether a set is consistent is to use a truth table. We can do this by constructing a single truth table with all of the wffs in the set and then check whether there is at least one row where all of the wffs in the set are true.

Definition 3.3.8: Truth-Table Test for Consistency

A truth table shows that a non-empty set of wffs $\{\phi_1, \phi_2, \dots, \phi_n\}$ is consistent when there is at least one row on the truth table where each wff in $\{\phi_1, \phi_2, \dots, \phi_n\}$ is true. Otherwise, the set is inconsistent.

Here is the truth table for $\{P \wedge \neg Q, P \rightarrow \neg Q\}$

P	Q	$P \wedge \neg Q$	$P \rightarrow \neg Q$
T	T	F	F
T	F	T	T
F	T	F	T
F	F	F	T

In examining the truth table above, notice that there is *at least one* row of the truth table where each of the wffs in this set are true (row 2). Since this is the case, the truth-table test for consistency determines this set to be consistent.

Let's consider another example where it might not be immediately obvious whether the set is consistent. Let's consider the set $\{P, P \vee Q, \neg(P \wedge Q)\}$. To begin, we should construct a single truth table that determines the truth value of each of these wffs.

P	Q	P	$P \vee Q$	$\neg(P \wedge \neg Q)$
T	T	T	T	F
T	F	T	T	F
F	T	F	T	T
F	F	F	F	T

To check to see whether the set of wffs is consistent, identify the main operator of each wff (or, if there are no operators, simply look at the letter). Next, go row by row and check whether all of the wffs are T at that row. Row 1 is T, T, F. Row 2 is T, T, F. Row 3 is F, T, F. Row 4 is F, F, T. If there is a row where all of the wffs are T, then the test determines the set to be consistent. Otherwise, the set is determined to be inconsistent. In this example, since there is no row where all of the wffs are T (we do not have T, T, T in at least one row), the set $\{P, P \vee Q, \neg(P \wedge Q)\}$ is inconsistent.

As a final example, let's consider whether the following set is consistent: $\{P \rightarrow Q, \neg R \vee Q, R \wedge \neg Q\}$. This example is somewhat more complicated in that (1) the set consists of three complex wffs and (2) the truth table will require 8 rows instead of the usual 4 rows.

If you go row by row, examining the truth of each wff in the table, you will see that there is no single row where each of the wffs in the set is true. The table thus shows that this set of wffs is inconsistent.

P	Q	R	$P \rightarrow Q$	$\neg R \vee Q$	$R \wedge \neg Q$
T	T	T	T	F	T
T	T	F	T	T	F
T	F	T	F	F	T
T	F	F	T	T	F
F	T	T	T	F	T
F	T	F	T	T	F
F	F	T	T	T	F
F	F	F	T	T	F

Exercise 3.31

Using a truth table, determine whether the following sets of wffs is consistent or inconsistent.

1. P, Q
2. $P \wedge Q, P$
3. $P \vee Q, P \wedge Q$
4. $P, P \wedge Q, \neg P \wedge Q$
5. $\neg P \wedge \neg R, \neg(P \vee R)$
6. $P, P \rightarrow R, R \vee P$
7. $P \vee R, \neg R, \neg P$
8. $\neg P \rightarrow R, R \rightarrow \neg P$
9. $\neg\neg P \rightarrow R, R \rightarrow Z, Z$
10. $A \leftrightarrow C, \neg C \vee \neg A$

Exercise 3.32

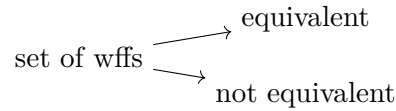
The following questions are designed to help you better understand the definition of consistency as well as make connections between consistency and other logical properties.

1. Suppose a set Γ and Γ consists of a single wff ϕ . Recall that every wff is either a PL-contingency, PL-tautology, or PL-contradiction. Under what scenario is Γ consistent?
2. What are some scenarios (other than the ones mentioned at the beginning of this section) where it might be useful to know whether a set of sentences is consistent?
3. Show that $\phi \wedge \psi$ is a contradiction if and only if $\{\phi, \psi\}$ is inconsistent.
4. Suppose a set Γ and Γ contains at least two wffs ϕ and ψ . Also

suppose that it is unknown whether ϕ is a PL-contingency, PL-tautology, or PL-contradiction, but it is known that ψ is a PL-contradiction. Is Γ consistent, inconsistent, or is it impossible to decide given that we do not know the status of ϕ .

3.3.3 *Equivalence*

Consider the sentences "John is tall and Mary is happy" and "Mary is happy and John is tall". Intuitively, these two sentences say the same thing. More precisely, they always have the same truth value: whenever one is true, the other is true; and, whenever one is false, the other is false. Let's say that a set of sentences is *equivalent* (or the sentences in the set are equivalent to each other) if and only if all of the sentences always have the same truth value.



Surely, there are cases where knowing that a set of wffs is equivalent is important. Let's consider two. First, suppose Tek and Liz are running against each other for a position on the local school board. You are a journalist assigned with covering the election. Before interviewing each candidate, you take a few moments to review statements about what each candidate believes and what they plan to do if elected. Each candidate wants to "enrich students" and "encourage critical thinking". Each candidate believes in "strong schools" and is "supportive of teachers". The candidates' platforms are, if you take them at their word, equivalent. If they are equivalent, voters might wonder why it is worth their time to vote at all. As a journalist who desires to create interest in the election, you decide to ask each candidate questions designed to tease out where their views diverge.

Suppose there are two theories T_1 and T_2 about the creation of the universe. On the surface, the two theories seem to make different claims about the world. T_1 accounts for the creation of the universe by positing an all-knowing but not all-loving being. T_2 accounts for the creation of the universe by positing an all-loving but not all-knowing being. From one perspective, the theories are not equivalent since T_1 says something that T_2 denies, and vice versa. But, from another perspective, they are

equivalent. Suppose Tek is focused on what each theory says about the observable world. Tek has a ball in his hand and plans to drop said ball. He asks a supporter of T_1 what will happen, and the supporter says it will fall to the earth. When asking a supporter of T_2 what will happen, the supporter also says it will fall to the earth. Both theorists disagree endlessly about whether the ball's falling to the earth is due to an all-loving or all-knowing being, but Tek could care less about this squabble. What is important to Tek are the practical consequences of each theory. So long as each theory makes the same observational predictions about moving bodies, each theory will be true and false under the same scenarios that are relevant to Tek. And, if that is the case, then the theories are equivalent.

Let's return to logic. In **PL** a set of wffs is equivalent provided all of the members of the set have the same truth value (T or F) for every interpretation.

Definition 3.3.9: PL-Equivalence

A set of wffs Γ is PL-equivalent if and only if there is no interpretation $\mathcal{I}$ where the valuation of the wffs in Γ fail to have the same truth value. A set of wffs Γ is not PL-equivalent if and only if there is at least one interpretation $\mathcal{I}$ where the valuation of the wffs in Γ fail to have the same truth value.

For example, the set $\{P \wedge Q, Q \wedge P\}$ is equivalent since $P \wedge Q$ and $Q \wedge P$ have the same truth value on every interpretation of P and Q . That is, for every way of going about assigning T and F to P and Q , whenever $v(P \wedge Q) = v(Q \wedge P)$. That is, whenever $P \wedge Q$ is true, $Q \wedge P$ is true, and whenever $P \wedge Q$ is false, $Q \wedge P$ is false.

One way to determine that this is the case is to use a truth table. We can do this by constructing a single truth table with all of the wffs in the set and then check whether their truth values are the same for each row.

Definition 3.3.10: Truth Table Test for Equivalence

A truth table shows that a set of wffs $\{\phi_1, \phi_2, \dots, \phi_n\}$ is PL-equivalent when the truth values of $\phi_1, \phi_2, \dots, \phi_n$ are the same for each row in the table. Otherwise, the set is not equivalent.

Here is the truth table for $\{P \wedge Q, Q \wedge P\}$:

P	Q	$P \wedge Q$	$Q \wedge P$
T	T	T T	T T
T	F	T F	F F
F	T	F F	T F
F	F	F F	F F

To determine if the set is equivalent, examine the truth value of each wff for every row of the table:

1. In row 1, $v(P \wedge Q) = v(Q \wedge P) = T$
2. In row 2, $v(P \wedge Q) = v(Q \wedge P) = F$
3. In row 3, $v(P \wedge Q) = v(Q \wedge P) = F$
4. In row 4, $v(P \wedge Q) = v(Q \wedge P) = F$

Since the truth values of the wffs are the same in each row, the truth-table test determines the set $\{P \wedge Q, Q \wedge P\}$ to be equivalent.

Let's consider another example. Suppose we wanted to know whether $\{P \rightarrow Q, \neg P \vee Q\}$ is equivalent. Rather than trying to guess, we might use the truth-table test for equivalence to determine whether the set is equivalent or non-equivalent.

P	Q	$P \rightarrow Q$	$\neg P \vee Q$
T	T	T T	F T T
T	F	T F	F T F
F	T	F T	T F T
F	F	F T	T F F

The truth-table test for $P \rightarrow Q$ and $\neg P \vee Q$ shows these two wffs to be equivalent since $v(P \rightarrow Q) = T$ whenever $v(\neg P \vee Q) = T$ and $v(P \rightarrow Q) = F$ whenever $v(\neg P \vee Q) = F$.

Finally, let's consider one more example. When testing whether a set of wffs is equivalent, it is often common to only test two wffs. For this reason, it is common to talk of two sentences being equivalent rather than a set of sentences being equivalent. However, in this example, let evaluate whether a set containing the following three wffs is equivalent: $\{P \rightarrow Q, \neg P \vee Q, \neg(P \wedge Q)\}$.

In determining whether the set is equivalent, check whether the wffs have the same truth value at each row. In this example, this set is not equivalent. Notice that in row 1, while $P \rightarrow Q$ and $\neg P \vee Q$ are T, $\neg(P \wedge Q)$ is F.

P	Q	$P \rightarrow Q$	$\neg P \vee Q$	$\neg(P \wedge Q)$
T	T	T	T	F
T	F	F	F	T
F	T	T	T	T
F	F	T	T	T

Exercise 3.33

Using a truth table, determine whether the following sets of wffs is equivalent or not equivalent.

1. P, Q
2. $P, \neg\neg P$
3. $P, P \vee P$
4. $P \wedge Q, Q \rightarrow P$
5. $\neg P \wedge \neg R, \neg(P \vee R)$
6. $P, P \rightarrow R, R \vee P$
7. $P \vee R, \neg R, \neg P$
8. $\neg P \rightarrow R, R \rightarrow \neg P$
9. $\neg\neg P \rightarrow R, R \rightarrow Z, Z$
10. $P \vee R, P \rightarrow R, P \leftrightarrow R, P \wedge R$

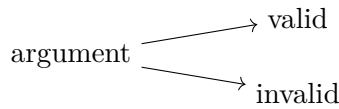
Exercise 3.34

The following questions are designed to help you better understand the definition of equivalence as well as make connections between equivalence and other logical properties.

1. Suppose a set Γ and Γ consists of a single wff ϕ . Is Γ PL-equivalent? Look closely at the definition of PL-equivalence.
2. What are some scenarios (other than the ones mentioned at the beginning of this section) where it might be useful to know whether a set of sentences is equivalent?
3. If ϕ and ψ are equivalent, does it follow that $\phi \leftrightarrow \psi$ is a tautology?
4. Suppose a set Γ and Γ contains at least two wffs ϕ and ψ . Suppose that ϕ is a contradiction and ψ is the negation of a tautology. Is Γ equivalent?
5. If $\phi \leftrightarrow \psi$, does it follow that ϕ and ψ are equivalent?
6. Create a set of wffs that is consistent but not equivalent.

3.3.4 Semantic Entailment

Consider the argument "If the money is missing, then Tek is a crook. Tek is a crook. Therefore, the money is missing." Does the conclusion of this argument follow from the premises with necessity? It does not. It is possible for the premises of this argument to be true and the conclusion false. When an argument has this property, it is said to be *invalid*. What about this argument: "If the money is missing, then Tek is a crook. The money is missing. Therefore Tek is a crook?" Does the conclusion of this argument follow from the premises with necessity? It does. It is impossible for the premises of this argument to be true and the conclusion false. When an argument has this property, it is said to be *valid*.



Surely, there are cases where it is important to know whether or not an argument is valid or invalid. Tek is a philosopher. He has all sorts of arguments for the existence of God. The people who read his books, listen to his talks, and take his classes are presented with these arguments. Suppose these individuals had no prior beliefs concerning the existence of God. When faced with these arguments for God's existence, should these individuals now accept that God does exist? Well, it seems that whether they should or they shouldn't depends upon the quality of Tek's arguments. And, the quality of Tek's arguments depend upon at least three things: (1) are the premises of his arguments true, (2) are the premises relevantly related to the conclusion, and (3) does the conclusion follow from the premises? With respect to (3), one way that a conclusion can follow from the premises is with necessity. That is, if the argument is valid, then the conclusion follows with necessity. And so, whether or not someone should have certain beliefs (e.g., a belief in God) at least partially depends upon whether or not an argument is valid.

Let's return to logic. One of the logician's tasks is to make clear what it means for the conclusion to follow from a set of premises. To some extent, we have already made progress in this effort through the introduction of the idea of *deductive validity*. Nevertheless, we can be even more precise by defining what it means for an argument to be "valid" in the language of propositional logic by introducing the notion of *semantic consequence*.

Definition 3.3.11: semantic consequence

A wff ϕ is a *semantic consequence* in **PL** of a set of wffs Γ (viz., $\Gamma \models \phi$) if and only if there is no interpretation $\mathcal{I}$ in which all of the members of Γ are true and ϕ is false. If there is an interpretation in **PL** such that all of the members of Γ are T and $v(\phi) = F$, then ϕ is not a semantic consequence of Γ , viz., $\Gamma \not\models \phi$.

The double turnstile $\models$ ("models" or "entails") expresses the fact that there is no interpretation of the wffs to the left of the turnstile such that these wffs are true and the wff to the right of the turnstile is false. Thus, when ϕ is a semantic consequence of Γ , we write $\Gamma \models \phi$. When ϕ is not a semantic consequence of Γ , we write $\Gamma \not\models \phi$.

For example, Q is a semantic consequence of $\{P \rightarrow Q, P\}$. That is, $P \rightarrow Q, P \models Q$. This is because there is no interpretation such that $v(P \rightarrow Q)$ and $v(P)$ are T and $v(Q)$ is F. In contrast, P is not a semantic consequence of $\{P \rightarrow Q, P\}$. That is, $P \rightarrow Q, P \not\models P$. This is because there is an interpretation such that $v(P \rightarrow Q)$ and $v(P)$ are T and $v(Q)$ is F.

But how do we know that this is the case? And, how would we go about checking other cases? One way to determine whether a wff ϕ is a semantic consequence of a set of wffs Γ is to use a truth table. We can do this by constructing a single truth table with all of the wffs in Γ and the wff ϕ .

Definition 3.3.12: Truth-Table Test for Semantic Consequence

A truth table determines that $\Gamma \models \phi$ if there is no row in the table where all of the members of Γ are T and ϕ is F. If there is a row where all of the members of Γ are T and ϕ is F, then $\Gamma \not\models \phi$.

Let's consider some examples. First, it was claimed that $P \rightarrow Q, P \models Q$. To test whether this is the case, let's construct a single truth table consisting of all of the wffs in the following set: $\{P \rightarrow Q, P, Q\}$.

P	Q	$P \rightarrow Q$	P	Q
T	T	T	T	T
T	F	F	T	F
F	T	T	F	T
F	F	T	F	F

Now, let's check whether there is at least one row where $P \rightarrow Q$ and P are T and Q is F.

1. In row 1, all of the wffs are T.
2. In row 2, $P \rightarrow Q$ and P are not both T.
3. In row 3, $P \rightarrow Q$ and P are not both T.
4. In row 4, $P \rightarrow Q$ and P are not both T.

In examining this table, we see that there is no row where $P \rightarrow Q$ and P are T and Q is F. Therefore, Q is a semantic consequence of $P \rightarrow Q$ and P . That is, $P \rightarrow Q, P \models Q$. Insofar as the notion of semantic consequence captures the idea of an argument being *valid* and this captures part of what it means for a conclusion to "follow from" a set of premises, we can say that Q follows from the premises $P \rightarrow Q$ and P . In other words, arguments of the form $P \rightarrow Q, P$ therefore Q are valid.

Next, it was claimed that $P \rightarrow Q, Q \not\models P$. To test whether this is the case, let's construct a single truth table consisting of all of the wffs in the following set: $\{P \rightarrow Q, Q, P\}$. Notice that we do not include the $\models$ or $\not\models$ in our truth table. Or, if it was included, it would only serve as a placeholder for the wff that is said to be the semantic consequence of another set of wffs.

P	Q	$P \rightarrow Q$	Q	P
T	T	T	T	T
T	F	F	F	T
F	T	T	T	F
F	F	T	F	F

Let's check whether there is at least one row where $P \rightarrow Q$ and Q are T and P is F. If there is such a row, $P \rightarrow Q, Q \not\models P$. If there is not such a row, then $P \rightarrow Q, Q \models P$. Notice that in row 3 $P \rightarrow Q$ and Q are both T and P is F. Therefore, P is a semantic consequence of $P \rightarrow Q, Q$. That is, $P \rightarrow Q, Q \models P$.

Let's consider one final example. Consider the claim that $P \rightarrow Q, \neg(P \vee Q) \models P$. If P is a semantic consequence of $P \rightarrow Q, \neg(P \vee Q)$, then there should be no row where $P \rightarrow Q, \neg(P \vee Q)$ are T and P is F. If P isn't a semantic consequence of $P \rightarrow Q, \neg(P \vee Q)$, then there will be at least one row where $P \rightarrow Q, \neg(P \vee Q)$ are T and P is F. Just as before, we to test $P \rightarrow Q, \neg(P \vee Q) \models P$, a single truth table containing the wffs $\{P \rightarrow Q, \neg(P \vee Q), P\}$ is constructed.

Notice that at row 4, all of the wffs in $\{P \rightarrow Q, \neg(P \vee Q)\}$ are T and P is F. Therefore, P is not a semantic consequence of $\{P \rightarrow Q, \neg(P \vee Q)\}$. In other words, $P \rightarrow Q, \neg(P \vee Q) \not\models P$.

P	Q	$P \rightarrow Q$	$\neg(P \vee Q)$	P
T	T	T	F	T
T	F	F	F	T
F	T	T	F	F
F	F	T	T	F

Exercise 3.35

Using a truth table, determine whether each wff is semantically entailed by the set of wffs in the entailment.

1. $P \models P$
2. $P \models \neg\neg P$
3. $P \wedge Q \models P$
4. $P \vee Q \models P$
5. $P \models P \vee Q$
6. $P \rightarrow Q \models P$
7. $P \rightarrow Q \models Q \rightarrow P$
8. $P \leftrightarrow Q \models Q \leftrightarrow P$
9. $P \vee R, \neg R, P \rightarrow R \models R$
10. $J \rightarrow C, \neg C \models \neg J$
11. $J \leftrightarrow C, C \models J \vee \neg\neg C$

Exercise 3.36

Translate the following arguments into **PL** and then use truth-table method to determine whether they are deductively valid or invalid.

1. John is happy or Mary is hungry. It is not the case that Mary is hungry. Therefore, John is not happy.
2. John will sell his house if and only if Mary sells her apartment. Mary will not sell her apartment. Therefore, John will not sell his house.
3. If John sells his house or Mary sells her apartment, the housing market will crash. The housing market will not crash. Therefore, John did not sell his house.

Exercise 3.37

The following questions are designed to help you better understand the definition of semantic consequence as well as help you to think

about the relation between this property and other logical properties.

1. What are some wffs such that $\phi \models \psi$ but $\psi \not\models \phi$
2. If ϕ is a tautology, does $\Gamma \models \phi$?
3. Suppose Γ is inconsistent. Does $\Gamma \models \phi$?
4. Suppose $\phi \rightarrow \psi$ is a tautology. Does $\phi \models \psi$?
5. Suppose $\phi \models \psi$. Is $\phi \rightarrow \psi$ a tautology?
6. Suppose $\{\phi, \psi\}$ is equivalent. Does $\phi \models \psi$?
7. Show that $\phi, \psi \models \chi$ if and only if $(\phi \wedge \psi) \rightarrow \chi$ is a tautology.
8. When using a truth table to test whether $\Gamma \models \phi$, a single table is constructed that contains all of the wffs in Γ and ϕ . We then evaluate the table looking to see if there is a row in the table where the wffs in Γ are T and ϕ is F. Suppose that instead, we created a table composed of all of the wffs in Γ and $\neg(\phi)$. How would we have to evaluate that table to determine whether $\Gamma \models \phi$?

3.4 LIMITATIONS OF TABLES

In this chapter, we introduced truth tables and illustrated how truth tables may be used to determine several logical properties. Of particular importance concerns the notion of semantic consequence (entailment). In this chapter, we defined the notion of semantic entailment and showed how to use a truth table to determine whether a wff is semantically entailed by a set of wffs. This is of particular importance because the notion of semantic entailment captures the idea of an argument being valid. And, the notion of an argument being valid captures the intuitive idea of a conclusion "following from" a set of premises. This, as we noted in the first chapter, is one feature of a "good argument". In this section, it is worthwhile to point out some of the strengths of truth tables over other methods of determining validity.

First, the truth-table method is a *decision procedure* or algorithm. That is, it consists of a finite sequence of rigorous instructions that, when followed, will, in a finite number of steps yield an answer of "yes" or "no" as to a question of whether a wff, set of wffs, or argument has a certain logical property.

Definition 3.4.1: decision procedure

A decision procedure is a finite sequence of rigorous instructions that, when followed, will, in a finite number of steps yield an answer of "yes" or "no" as to a question of whether a wff, set of wffs, or argument has a certain logical property.

In calling the truth-table method a decision procedure, what is being asserted is that given an argument and that it is translated into the language of propositional logic, the truth-table method provides a set of instructions that, when followed, will yield an answer of "yes" or "no" as to whether the argument is valid. This contrasts with the other two methods we considered in Chapter 1. These methods do provide a set of instructions for determining whether an argument is deductively valid, but the instructions are not guaranteed to give an answer "yes" or "no" as to whether the argument is valid. For example, an individual may use the logical intuition test and say that they don't know what their logical intuition is telling them and so are unsure whether the argument is valid or invalid. Or, in the case of the imagination test, an individual may say they think they are able to imagine a scenario where the premises are true and the conclusion is false, but are not totally sure since the scenario is difficult to imagine.

Second, the truth-table method also appears to give the right results in many cases. For example, consider the following argument:

- P1: If Jon looks fishy, then he committed the crime.
- P2: Jon looks fishy.
- C: Therefore Jon committed the crime.

This argument is intuitively valid. When it is translated into the language of propositional logic as $J \rightarrow C, J \models C$, the truth-table method confirms this intuition. That is, the table shows that there is no row where $J \rightarrow C$ and J is true and C is false.

While there are several other merits to the truth table, let's now turn to some of its problems. First, one problem with the truth-table method though is that not every English argument can be represented in a truth-functional language like **PL**. This means that while the truth-table test is an effective tool for determining the validity of *some* arguments we might express in English, it cannot express *all* English arguments. In chapter 6, we articulate a formal language capable of representing non-truth-functional arguments.

Second, from a user's standpoint, the truth-table test is relatively easy to use when dealing with arguments that involve one, two, or even three propositional letters, but they become increasingly complex the more propositional letters the argument has. For example, a truth table for $P \models P$ consists of two rows, one where $\mathcal{I}(P) = T$ and one where $\mathcal{I}(P) = F$. And, a truth table for $P \models Z$ consists of four rows. The number of rows required for a truth table of any argument is determined by 2^n where n is the number of propositional letters in the argument. Thus, $P, Q \models Z$ consists of eight rows ($2^3 = 8$), $P, Q, R \models Z$ of sixteen rows ($2^4 = 16$), and so on. Considering that there are arguments composed of dozens of sentences represented by distinct propositional letters, the truth-table test can quickly become unwieldy, requiring hundreds or thousands of rows.

Another way of thinking about this second problem, still in informal terms, involves the idea of *complexity of work*. As arguments require more and more propositional letters, the amount of work that must be performed increases: one propositional letter requires two rows, two propositional letters requires four rows, and so on. But how much more work is added with each additional propositional letter? If my boss gives me three more units of work each day, then my work increases linearly: 3, 6, 9, 12, 15, and so on. But, if my boss doubles my work each day, then it increases exponentially: 3, 6, 12, 24, 48, and so on. The method of truth tables is similar to the boss who doubles your work each day.

There are a few ways to address this problem for truth tables. One way is to use a computer to generate the truth table. From a practical point of view, truth tables are a time-consuming way of determining whether an argument is valid or invalid since it requires a considerable amount of time and energy to write out thousands of rows. However, with a computer, the amount of time and energy required to generate a truth table is greatly reduced. Several online programs exist that take a set of wffs as input and quickly generate a truth table as output. However, even with a computer, the amount of work required to generate a truth table increases exponentially. If a simpler, less time-consuming method exists, then it would be preferable to use that method.

A second solution aims to reduce the amount of work required. This method is sometimes called the "short truth-table method". This method only requires completing only one line of the truth table. The method works by (1) assigning the conclusion a value of F and the premises a value of T, and then (2) assigning truth values to the subformulas of the wffs

until a coherent interpretation is discovered. If a coherent interpretation is discovered, then the argument is invalid. If a coherent interpretation is not discovered, then the argument is valid. Let's consider an example.

Suppose we have the claim that $\neg Q \wedge R, P \rightarrow Q \models M \vee Q$. Since there are four propositional letters, the truth table would consist of $2^4 = 16$ rows. The short method, however, only requires one row. To apply the short method, first, assign the wffs in $\{\neg Q \wedge R, P \rightarrow Q\}$ a value of T and $M \vee Q$ a value of F.

M	P	Q	R	$\neg Q \wedge R$	$P \rightarrow Q$	$M \vee Q$
				T	T	F

Next, since $M \vee Q$ is F, both M and Q must be F. Let's write F under M and F under Q every Q in the table.

M	P	Q	R	$\neg Q \wedge R$	$P \rightarrow Q$	$M \vee Q$
F		F		F T	T F	F F F

Moving to the premises, since Q is F, $\neg Q$ must be T. Let's put T under $\neg Q$. And, since we supposed that $\neg Q \wedge R$ is T, R must also be T. Let's put T under R . Moving to the second premise, if $P \rightarrow Q$ is T and Q is F, then P must be F. Let's put F under P . We now have a coherent interpretation of the wffs in $\neg Q \wedge R, P \rightarrow Q \models M \vee Q$ that makes the premises true and the conclusion false. Since there is a coherent interpretation, the argument is invalid.

M	P	Q	R	$\neg Q \wedge R$	$P \rightarrow Q$	$M \vee Q$
F	F	F	T	T F T T	F T F	F F F

However, the short truth-table method is not without problems. The primary problem is that the method can be difficult to use when subformulas can be true or false. Suppose there is an argument whose conclusion is $P \wedge (Q \wedge R)$. In order to use the short truth-table method, we would need to assign $P \wedge (Q \wedge R)$ a value of F. Since $P \wedge (Q \wedge R)$ is a conjunction, it is F because at least one of its conjuncts are F. But, which conjunct is F? Is it P , Q , R , or some combination of those subformulas? Assigning $P \wedge (Q \wedge R)$ does not force a single truth value on the subformulas that constitute that wff. And, if the subformulas are not forced to be either T or F, then using the short truth-table method becomes more a matter of trial and error than a systematic method.

In the next chapter, another method is developed to determine whether an argument is valid or invalid. This method is called the *truth-tree method*. Similar to the truth-table method, the truth-tree method is a mechanical method. It thus shares the same relative strength that truth tables

does over more informal methods. However, in contrast to truth tables, the complexity of the truth-tree method is not a function of the number of propositional letters in the argument. Instead, the complexity of the truth-tree method is a function of the number (and type) of wffs in the argument. Thus, the truth-tree method is a more efficient method for determining the validity of arguments with many propositional letters.

Exercise 3.38

A set of wffs Γ **semantically entails** a wff ϕ (i.e., $\Gamma \models \phi$) if and only if there is no interpretation such that the members of Γ are true and ϕ is false. That is, $\Gamma \models \phi$ says that ϕ is true when the members of Γ are true. A wff ϕ is said to be a **tautology** if and only if there is no interpretation where ϕ is false. That is, ϕ is a tautology if and only if ϕ is true under every interpretation.

Is it possible to express every semantic entailment ($\Gamma \models \phi$) as a tautology? Provide an argument for your answer.

Exercise 3.39

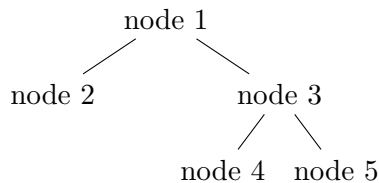
1. What is a truth table?
2. What is a decision procedure?
3. What does it mean for a wff **Q** to be a logical (semantic) consequence of a set of wffs Γ ?
4. What does it mean for an argument to be valid and invalid in **PL**?
5. What is the difference between a semantic entailment and a valid argument?
6. Under what condition does the truth-table method show an argument to be valid? Under what condition does it show an argument to be invalid?
7. What does it mean to say that two wffs in **PL** are equivalent?
8. Under what condition does the truth-table method show a pair of wffs to be equivalent?
9. What does it mean to say that a set of wffs in **PL** are consistent?
10. Under what condition does the truth-table method show a set of wffs to be consistent?
11. What does it mean to say that a wff in **PL** is a tautology, a contradiction, and a contingency?
12. Under what conditions does the truth-table method show that a

- wff in **PL** is a tautology, a contradiction, and a contingency?
13. What are two problems with using truth tables to check whether an argument expressed in English is valid or invalid?
 14. Determine the truth value of $\neg(P \rightarrow \neg Q)$ using the following interpretation: $\mathcal{I}(P) = T, \mathcal{I}(Q) = F$
 15. Using the truth-table method, determine the truth value of $\neg((P \vee Q) \leftrightarrow Q)$ for each interpretation of P and Q
 16. Using the truth-table method, determine whether $Q \rightarrow P$ is a semantic consequence of $P \vee \neg Q$. That is determine whether $P \vee \neg Q \models Q \rightarrow P$.
 17. Using the truth-table method, determine whether $\neg(P \vee Q)$ and $\neg P \wedge \neg Q$ are equivalent. If they are not, for what interpretation of P and Q are they not equivalent.
 18. Using the truth-table method, determine whether $P \rightarrow Q, Q \rightarrow P, P \leftrightarrow Q$ are consistent. If they are not, for what interpretation of P and Q are they not consistent.
 19. Using the truth-table method, determine whether $\neg P \vee \neg\neg P$ is a tautology, contradiction, or contingency.
 20. Translate the following propositions from this argument into **PL**, then use the truth-table method to determine whether the argument is valid or invalid: If John is tall, then Frank is not tall. Frank is not tall. Therefore, John is tall.
 21. Translate the following propositions from this argument into **PL**, then use the truth-table method to determine whether the argument is valid or invalid: If God exists, then there is not any unhappiness in the world. There is unhappiness in the world. Therefore, God does not exist.
 22. Translate the following propositions into **PL**, then use the truth-table method to determine whether the propositions are consistent: If taxes go up, there will be riots. Either taxes did not go up, or there are riots.
 23. Translate the following propositions into **PL**, then use the truth-table method to determine whether the propositions are logically (semantically) equivalent: (a) Neither ice cream nor cake is on the menu. (b) Ice cream and cake are not both on the menu.
 24. Translate the following proposition into **PL**, then use the truth-table method to determine whether the proposition is a tautology, contradiction, or contingency: God's existence implies if bananas are yellow, then god exists.

In chapter 2, we introduced the language of **PL**. This language consists of a set of symbols, a syntax, and a semantics. In chapter 3, we introduced the truth-table method for determining whether a wff, set of wffs, or argument has a particular logical property. In this chapter, we introduce the truth-tree method for determining whether a wff, set of wffs, or argument has a particular logical property. One immediate question concerning the truth-tree method is why it is needed. First, the general goal of logic is to separate good arguments from bad arguments. Truth trees help in achieving this goal in that they offer another method for testing arguments for validity (semantic entailment). Second, the truth-tree method also offers a solution to the problem of the truth-table method's complexity and the truth-table method's inability to handle more complex logical languages (we will discuss one such language in chapter 6).

4.1 INTRODUCTION TO TREES

What is a truth tree? A tree is a collection of nodes and branches. A node is whatever is at a particular location in a tree and a branch is all of the nodes of a tree starting from bottom nodes of the tree and moving upward through the tree to the top node. Consider the following example:



In the above tree, there are five nodes. Node 1 at the top is the root node. Nodes 2, 4, and 5 are bottom nodes. We can also call them "leaf nodes". The tree also contains three branches. Branches are all of the nodes starting at the bottom nodes and moving upward through the tree to the top node.

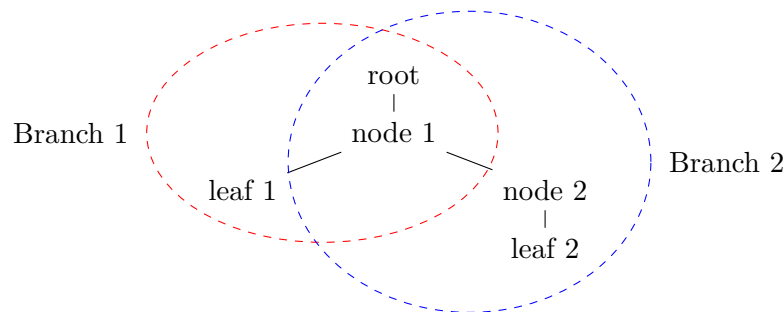
Definition 4.1.1: branch

A branch in a tree consists of all of the nodes starting from a leaf node and moving upward to the root node.

So, for example, one branch of the tree contains nodes 2 and 1. A second branch of the tree contains nodes 4, 3, and 1. Finally, the third branch of

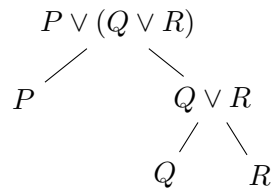
the tree contains nodes 5, 3, and 1.

Let's consider another abstract example. In the tree below, the tree has a root node at the top of the tree, a node directly below the root (node 1), then the tree splits into two different branches. On the left branch, the tree immediately terminates with a bottom node (leaf node). On the right branch, there is a node, followed by a bottom node (leaf node).



There are two branches in the above tree. The leftmost branch consists of leaf 1, node 1, and the root node. The rightmost branch consists of leaf 2, node 2, node 1, and the root node.

Now that we have considered trees from a somewhat abstract perspective, let's consider truth trees in propositional logic. In a propositional logic truth tree (or just a tree), wffs are at the various nodes of the tree and so a branch of a truth tree is a collection of wffs. Let's reconstruct the tree above using wffs. Suppose we have the following wff: $P \vee (Q \vee R)$. A tree for this wff would look as follows:



In the above tree, $P \vee (Q \vee R)$ is the root node. It is the node found at the top of the tree. The bottom (leaf) nodes are P , Q , and R . In addition, there are three branches in this tree. These branches are identified by starting at the leaves and moving upward through the tree to its root. The leftmost branch contains P and $P \vee (Q \vee R)$. The middle branch contains Q , $Q \vee R$, and $P \vee (Q \vee R)$. The third, and rightmost, branch contains R , $Q \vee R$, and $P \vee (Q \vee R)$.

4.2 BRANCHES AND DECOMPOSITION

In this section, we develop some vocabulary for categorizing different types of branches and we introduce the concept of decomposition. By the end of this section, you will be able to determine whether a branch is open or closed and have a basic idea of how to create a truth tree yourself.

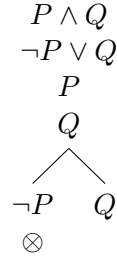
4.2.1 Open and Closed Branches

Branches are either open or closed. A closed branch is a branch that contains a wff and its literal negation.

Definition 4.2.1: Closed Branch

A branch is closed when it contains a wff ϕ and its literal negation $\neg(\phi)$. A closed branch is represented by writing $\otimes$ under the bottom node in the closed branch.

To illustrate, consider the following tree:

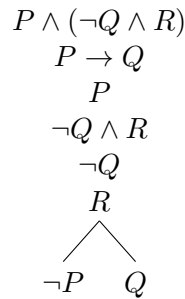


In the above tree, the left branch consists of the wffs: $\{P \wedge Q, \neg P \vee Q, P, Q, \neg P\}$. This branch contains the wff P and its literal negation $\neg P$. As such, this branch is closed. The right branch consists of the wffs: $\{P \wedge Q, \neg P \vee Q, P, Q, Q\}$. This branch does not contain a wff and its literal negation. As such, this branch is not closed. A branch that is not closed is an open branch.

Definition 4.2.2: Open Branch

An open branch is a branch that is not closed. That is, a branch that does not contain a wff ϕ and its literal negation $\neg(\phi)$.

Consider the following tree:



In this tree, we can determine whether a branch is open or closed by starting at the bottom node of the tree, moving upward through the branch, and looking for a wff ϕ on one line and its literal negation $\neg(\phi)$ on another. Starting with the left branch, notice that this branch contains the wff $\neg P$ and P . This branch is thus closed. Next, consider the right branch. This branch contains the wffs Q and $\neg Q$. This branch is also closed.

4.2.2 Decomposition

With our basic understanding of a truth tree and the distinction between an open and closed branch, let's consider how trees are constructed. In order to understand how trees are constructed, we must understand the concept of decomposition. We can express the conditions under a complex sentence is true by stating which simpler sentences are true if the complex sentence is true. For example, suppose Tek says to Liz that "Today he will go for a bike ride and tomorrow he will go for a run." We can represent what sentence or sentences are true if Tek's sentence is true by writing "Today he will go for a bike ride" and "Tomorrow he will go for a run" under Tek's sentence. The same is true for **PL**. The decomposition of a wff ϕ is the expression of what wff or wffs are true if ϕ is true. For example, the decomposition of $P \wedge Q$ is P and Q . The decomposition of $R \vee S$ is R or S .

Definition 4.2.3: decomposition

A **decomposition** of a wff ϕ is a representation of what wff or wffs are true if ϕ is true.

The decomposition of a wff may be represented using a truth tree. Let's consider the wff $P \wedge Q$. The wff $P \wedge Q$ is true in one case. This is the case where P is true *and* Q is true. To represent that this wff is true provided both of these simpler wffs are true, we will write (or "stack") P and Q directly under $P \wedge Q$.

$$\begin{array}{c}
 P \wedge Q \\
 P \\
 Q
 \end{array}$$

By writing P and Q directly under $P \wedge Q$, we are expressing the decomposition that if $P \wedge Q$ is true, then P is true *and* Q is true. Next, let's consider the disjunction $P \vee Q$. The wff $P \vee Q$ is true not in just *one* case but in the following *three* cases:

1. P is true and Q is false
2. P is false and Q is true
3. P is true and Q is true

To represent that $P \vee Q$ is true in these three cases, we will use lines to branch from $P \vee Q$ and write P in the left branch and Q in the right branch.

$$\begin{array}{c}
 P \vee Q \\
 \swarrow \quad \searrow \\
 P \quad Q
 \end{array}$$

The use of branching in the tree above represents the decomposition of $P \vee Q$ into the three cases in which $P \vee Q$ is true. This wff is true provided P is true (we move down the left branch), Q is true (we move down the right branch), or both P and Q is true (we move down both branches).

Let's express things more generally. Whenever there is a wff ϕ that is true in *one* case, we will make use of a stacking convention. The *stacking convention* represents the decomposition of a wff ϕ by writing a certain wff or wffs that are true if ϕ is true directly under the bottom node of every open branch that descends from ϕ . Whenever there is a wff ϕ that is true in *three* cases, we will make use of a branching convention. The *branching convention* represents the decomposition of a wff ϕ by writing a certain wff or wffs that are true by branching under the bottom node of every open branch that descends from ϕ .

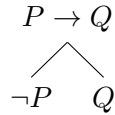
Let's illustrate both of these conventions with a few examples. First, let's consider the conditional $P \rightarrow Q$. To determine whether we ought to employ the stacking or branching convention, let's reexamine the truth table for the conditional.

P	Q	P	→	Q
T	T	T	T	T
T	F	T	F	F
F	T	F	T	T
F	F	F	T	F

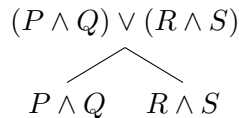
Notice that $P \rightarrow Q$ is true in any of the following three cases:

1. P is false
2. Q is true
3. P is false and Q is true

Since the conditional is true in these three cases, we can use the branching rule to represent the decomposition of the conditional. In the left branch, we write $\neg(P)$ since the conditional is true whenever P is false. In the right branch, we write Q since the conditional is true whenever Q is true.



Let's consider a slightly more complicated example in order to grasp a better understanding of the two conventions. Consider the wff $(P \wedge Q) \vee (R \wedge S)$. In examining this wff, notice that it is a disjunction. As such, it will be true if either (1) the left disjunct is true, (2) the right disjunct is true, or (3) both disjuncts are true. We thus can employ the branching rule to decompose this wff.

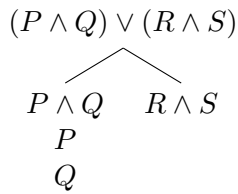


In the tree above, the wff $(P \wedge Q) \vee (R \wedge S)$ is decomposed into two smaller wffs: (1) $(P \wedge Q)$ which is located in the left branch and (2) $(R \wedge S)$ which is located in the right branch. Notice that it is possible to decompose both of these wffs further. For example, we can represent the fact that $P \wedge Q$ is true if both P and Q are true. But where do we write P and Q ? Under the left branch, under the right branch, under both branches, or in a new branch off to the side?

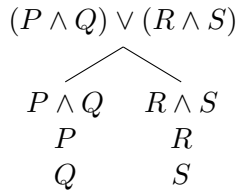
When we decompose the root node, we create two branches. In creating two branches, we are representing the fact that the root node is true if either the left branch is true, the right branch is true, or both branches are true. In our example, the conditions under which the wff in the left

branch is true are different from the conditions under which the wff in the right branch is true. That is, $P \wedge Q$ is true if both P and Q are true, while $R \wedge S$ is true if both R and S are true. As such, we will decompose $P \wedge Q$ under the left branch and $R \wedge S$ under the right branch. More generally, whenever a wff is decomposed, it should be decomposed under every leaf node in every open branch that descends from the wff.

Let's begin with the left branch. The wff $P \wedge Q$ is a conjunction and, as we have seen, it will be true if both conjuncts are true. We thus can employ the stacking convention to decompose this wff. In decomposing $P \wedge Q$ we will decompose it only under $P \wedge Q$ (the only leaf node in the left branch). In other words, we will not decompose $P \wedge Q$ under $R \wedge S$.



Next, we will decompose $R \wedge S$ in the right branch in a similar fashion.



The above tree represents the decomposition of the wff $(P \wedge Q) \vee (R \wedge S)$. We can see this decomposition by reading through the branches of the tree. That is, the wff $(P \wedge Q) \vee (R \wedge S)$ is true if the following is the case:

1. Left Branch: $P \wedge Q$ is true, and this wff is true if and only if both P and Q are true, OR
2. Right Branch: $R \wedge S$ is true, and this wff is true if and only if both R and S are true, OR
3. Both the left and right branches are true.

Thus far, we have explained the concept of a decomposition of a wff and how this decomposition can be expressed in the form of a truth tree. Let's consider one final example to gain a mastery over the stacking and branching conventions. Consider the wff $(P \vee Q) \wedge (R \vee S)$. Again, since this is a conjunction, it will be true if both conjuncts are true. We thus can employ the stacking rule to decompose this wff.

$$\begin{array}{c}
 (P \vee Q) \wedge (R \vee S) \\
 P \vee Q \\
 R \vee S
 \end{array}$$

Now that the root node has been decomposed into $P \vee Q$ and $R \vee S$, let's decompose both of these wffs. Let's focus our attention on $P \vee Q$. As we have seen, since this is a disjunction, it will be true if either (1) the left disjunct is true, (2) the right disjunct is true, or (3) both disjuncts are true. We thus can employ the branching rule to decompose this wff. The branching rule states that the wff is to be decomposed under the bottom node of every open branch that descends from the wff. Since the bottom node of our tree is $R \vee S$, we will decompose $P \vee Q$ directly under $R \vee S$.

$$\begin{array}{c}
 (P \vee Q) \wedge (R \vee S) \\
 P \vee Q \\
 R \vee S \\
 \swarrow \quad \searrow \\
 P \quad Q
 \end{array}$$

Next, let's decompose $R \vee S$. As this is a disjunction, we will branch R and S . As the branching convention states, we will decompose $R \vee S$ under the bottom node of *every open branch* that descends from $R \vee S$. Notice that there are two open branches that descend from $R \vee S$. The first is the left branch that contains P and the second is the right branch that contains Q . Since there are two open branches descending from $R \vee S$, we will decompose $R \vee S$ under the bottom node of both of these branches.

$$\begin{array}{c}
 (P \vee Q) \wedge (R \vee S) \\
 P \vee Q \\
 R \vee S \\
 \swarrow \quad \searrow \\
 P \quad Q \\
 \swarrow \searrow \quad \swarrow \searrow \\
 R \quad S \quad R \quad S
 \end{array}$$

The above tree represents the decomposition of the wff $(P \vee Q) \wedge (R \vee S)$. As with the previous example, we can see this decomposition by reading through the branches of the tree. Starting from the top and reading through the tree downward, we can see that the wff $(P \vee Q) \wedge (R \vee S)$ is true if the following is the case: $P \vee Q$ and $R \vee S$ are true, and these wffs are true if and only if (1) P is true and R or S is true, or (2) Q is true and R or S is true.

4.2.3 Decomposition Rules

In the previous section, we consider how we might represent the decomposition of a wff using two conventions: the stacking convention and the branching convention. In this section, we will consider how we might represent the decomposition of various wffs using rules for decomposition. The set of rules for decomposition are known as the "**PL** decomposition rules."

In the previous section, we noted that when representing the decomposition of the conjunction $P \wedge Q$, we employ the stacking convention. This convention has us write P and Q directly under the bottom node of every branch that descends from $P \wedge Q$. But this is true not merely for the conjunction $P \wedge Q$ but for *any* conjunction: any wff that has the form $\phi \wedge \psi$. With this in mind, let's formulate a rule for decomposing conjunctions in a truth tree. We will call this rule the *conjunction decomposition rule* and abbreviate it as $\wedge D$. In expressing this rule, we will use ϕ and ψ as variables for any wff and we will write $\wedge D$ to the right of each wff that results from decomposing the conjunction $\phi \wedge \psi$. The conjunction decomposition rule is as follows:

$$\begin{array}{ll} \phi \wedge \psi & \\ \phi & \wedge D \\ \psi & \wedge D \end{array}$$

To illustrate the use of conjunction decomposition rule, let's consider the following tree:

$$\begin{array}{ll} \neg A \wedge \neg B & \\ \neg A & \wedge D \\ \neg B & \wedge D \end{array}$$

In the tree above, the $\wedge D$ is used to represent the decomposition of the conjunction $\neg A \wedge \neg B$ into the two simpler wffs $\neg A$ and $\neg B$. The rightmost column indicates that the rule applied to $\neg A \wedge \neg B$ is the conjunction elimination rule.

A second wff that we considered in an earlier section is $P \vee Q$. We noted that when representing the decomposition of the disjunction $P \vee Q$, we employ the branching convention. This convention has us branch from the bottom node of every open branch, writing the left disjunct P on the left branch and the right disjunct Q on the right branch. But this way of branching applies not merely to the disjunction $P \vee Q$ but to *any* disjunction. Similar then to the rule for conjunction, let's formulate the following decomposition rule for disjunctions. Just as we called the

decomposition rule for conjunctions "conjunction decomposition", we will call the decomposition rule for disjunctions "disjunction decomposition" (abbreviating it as $\vee D$).

$$\begin{array}{ccc} \phi \vee \psi & & \\ \swarrow & & \searrow \\ \phi & \psi & \vee D \end{array}$$

Just as we did with conjunction decomposition, let's provide an illustration of the disjunction decomposition rule on a novel case. Consider the wff $\neg A \vee \neg B$. The disjunction decomposition rule applied to this wff is as follows:

$$\begin{array}{ccc} \neg A \vee \neg B & & \\ \swarrow & & \searrow \\ \neg A & \neg B & \vee D \end{array}$$

In the tree above, the disjunction decomposition rule is applied to $\neg A \vee \neg B$ to illustrate the decomposition of this wff into the two simpler wffs $\neg A$ and $\neg B$. The rightmost column indicates that the rule applied to $\neg A \vee \neg B$ is the disjunction elimination rule.

How many decomposition rules are there? As there are nine different non-literal wffs in **PL**. There are the wffs defined in terms of their operators:

1. conjunctions: $\phi \wedge \psi$
2. disjunctions: $\phi \vee \psi$
3. conditionals: $\phi \rightarrow \psi$
4. biconditionals: $\phi \leftrightarrow \psi$

and then there are the negated versions of each one of the above wff types, along with wffs that are doubly-negated:

5. negated conjunctions: $\neg(\phi \wedge \psi)$
6. negated disjunctions: $\neg(\phi \vee \psi)$
7. negated conditionals: $\neg(\phi \rightarrow \psi)$
8. negated biconditionals: $\neg(\phi \leftrightarrow \psi)$
9. double-negations: $\neg(\neg\phi)$

For each of these nine, non-literal wff types, there is a corresponding decomposition rule. First, there are the rules for conjunction decomposition and negated conjunction decomposition:

$$\begin{array}{ccc} \phi \wedge \psi & & \neg(\phi \wedge \psi) \\ \swarrow & & \searrow \\ \phi & \wedge D & \neg(\phi) \quad \neg(\psi) \\ \psi & \wedge D & \neg \wedge D \end{array}$$

Notice that conjunction decomposition is a stacking rule, negated conjunction decomposition is a branching. Second, there are the rules for disjunction decomposition and negated disjunction decomposition:

$$\begin{array}{ccc}
 \phi \vee \psi & & \neg(\phi \vee \psi) \\
 \swarrow \quad \searrow & & \neg(\phi) \quad \neg \vee D \\
 \phi \quad \psi & \vee D & \neg(\psi) \quad \neg \vee D
 \end{array}$$

Again, it is worth noting that while the disjunction decomposition rule branches, the negated disjunction decomposition is a stacking rule. Third, there are the rules for conditional decomposition and negated conditional decomposition:

$$\begin{array}{ccc}
 \phi \rightarrow \psi & & \neg(\phi \rightarrow \psi) \\
 \swarrow \quad \searrow & & \phi \quad \neg \rightarrow D \\
 \neg(\phi) \quad \psi & \rightarrow D & \neg(\psi) \quad \neg \rightarrow D
 \end{array}$$

In the case of biconditionals and negated biconditionals both the stacking and branching conventions are used. This is because a biconditional $\phi \leftrightarrow \psi$ is true provided $v(\phi) = v(\psi)$ and false otherwise. Therefore, the biconditional is true not in a single scenario (so we won't use the stacking convention) and not in three scenarios (so we won't use the branching convention), but in two scenarios. To represent this, we branch a stack of ϕ and ψ under the left branch (as $\phi \leftrightarrow \psi$ is true provided both ϕ and ψ are true) and we branch a stack of $\neg(\phi)$ and $\neg(\psi)$ under the right branch (as ψ is true provided both ϕ and ψ are false). The rules for biconditional decomposition and negated biconditional decomposition are as follows:

$$\begin{array}{ccc}
 \phi \leftrightarrow \psi & & \neg(\phi \leftrightarrow \psi) \\
 \swarrow \quad \searrow & & \swarrow \quad \searrow \\
 \begin{array}{ccc}
 \phi & \neg(\phi) & \leftrightarrow D \\
 \psi & \neg(\psi) & \leftrightarrow D
 \end{array} & & \begin{array}{ccc}
 \phi & \neg(\phi) & \neg \leftrightarrow D \\
 \neg(\psi) & \psi & \neg \leftrightarrow D
 \end{array}
 \end{array}$$

Finally, there is a decomposition rule for double-negations. The rule for double-negation decomposition is as follows:

$$\begin{array}{ccc}
 \neg \neg \phi & & \\
 \phi & \neg \neg D &
 \end{array}$$

Let's illustrate the use of some of these rules. However, before we do this, let's introduce two additional features to our construction of truth trees. First, we will add a column on the leftside of the tree that numbers each node of the tree. Second, when citing the decomposition rule we used to

obtain a new node in the tree, we will not only cite the abbreviation of the decomposition rule we used, but also the number of the wff to which that decomposition rule is applied. What this means is that our truth trees moving forward will always consist of three columns:

1. Leftmost column: for numbering the wffs,
2. Middle column: for nodes and branches,
3. Rightmost column: for justification of wffs

Let's consider some examples. Let's begin with the wff $(R \vee P) \wedge (C \wedge D)$. As this wff is a conjunction, I will start by writing down the wff, numbering it in the left column, and writing "P" in the right column.

1. $(R \vee P) \wedge (C \wedge D)$ P

Next, as this wff is a conjunction, let's use $\wedge D$ on it. This is a stacking rule and so we write the left conjunct on line 2 and the right conjunct on line 3. We then write "1 $\wedge D$ " in the right column to indicate that we used the conjunction decomposition rule on line 1.

1. $(R \vee P) \wedge (C \wedge D)$ P
2. $R \vee P$ 1 $\wedge D$
3. $C \wedge D$ 1 $\wedge D$

Let's go ahead and decompose the disjunction $R \vee P$. This is a branching rule and so we will write the left disjunct on line 4 and the right disjunct on line 4. We then write "2 $\vee D$ " in the right column to indicate that we used the disjunction decomposition rule on line 2.

1. $(R \vee P) \wedge (C \wedge D)$ P
2. $R \vee P$ 1 $\wedge D$
3. $C \wedge D$ 1 $\wedge D$
4. $\begin{array}{c} \diagup \quad \diagdown \\ R \quad P \end{array}$ 2 $\vee D$

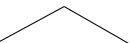
This leaves $C \wedge D$ at line 3. At first glance, it might not be clear how to decompose this wff. Should you decompose it under the left branch, the right branch, or both? The answer is that you should decompose it under both branches. This is because whenever a wff is decomposed, it is always decomposed under the bottom node of every open branch that descends from the wff being decomposed. Since the left and right branch both descend from $C \wedge D$, this wff is decomposed under both branches.

1.	$(R \vee P) \wedge (C \wedge D)$	P
2.	$R \vee P$	$1 \wedge D$
3.	$C \wedge D$	$1 \wedge D$
		
4.	$R \quad P$	$2 \vee D$
5.	$\boxed{C} \quad \boxed{C}$	$3 \wedge D$
6.	$\boxed{D} \quad \boxed{D}$	$3 \wedge D$

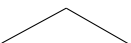
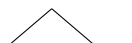
Next, let's consider the decomposition of $(P \wedge Q) \rightarrow \neg\neg S$. To set up the tree, write down the wff, number it in the left column, and write "P" in the right column.

1.	$(P \wedge Q) \rightarrow \neg\neg S$	P
----	---------------------------------------	---

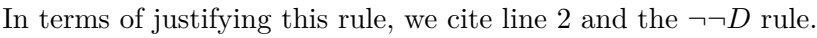
Next, since this wff is a conditional, we will decompose it using the conditional decomposition rule. To do this, we will write the decomposition rule to the right of the wff and then write the wffs that result from the decomposition of the conditional under the bottom node of every open branch that descends from the conditional. In this case, there is only one open branch that descends from the conditional. As such, we will write the decomposition of the conditional under the bottom node of this branch.

1.	$(P \wedge Q) \rightarrow \neg\neg S$	P
		
2.	$\neg(P \wedge Q) \quad \neg\neg S$	$\rightarrow D$

Line 2 in the left branch is the wff $\neg(P \wedge Q)$. This is a negated conjunction and so to decompose it, we will use the negated conjunction decomposition rule ($\neg \wedge D$). In contrast to the conjunction decomposition rule, which is a stacking rule, $\neg \wedge D$ is a branching rule.

1.	$(P \wedge Q) \rightarrow \neg\neg S$	P
		
2.	$\neg(P \wedge Q) \quad \neg\neg S$	$1 \rightarrow D$
		
3.	$\neg P \quad \neg Q$	$2 \neg \wedge D$

Finally, line 2 in the right branch is the wff $\neg\neg S$. This is a double-negation and so to decompose it, we will use the double-negation decomposition rule ($\neg\neg D$). This is a stacking rule.



Using the decomposition rules, decompose the following wffs:

- #### 4.2.4 Fully Decomposed Branches

Definition 4.2.4: Fully decomposed branch

Let's illustrate a fully decomposed branch. Consider the following tree:

- At this point, the tree consists of one branch and this branch is not fully decomposed since it contains a non-literal wff that has not had a decomposition rule applied to it. In other words, it contains a wff that has not yet been decomposed. Let's decompose the wff by using $\rightarrow D$. In addition, to indicate that it has been decomposed, let's place a checkmark next to the wff after we have applied the decomposition rule to it.

1. $(P \rightarrow S) \rightarrow Q \checkmark$ P
2. $\neg(P \rightarrow S) \quad Q$ $1 \rightarrow D$

Now the tree contains two branches. In the left branch, while we have decomposed the wff at line 1, we have not decomposed $\neg(P \rightarrow S)$. This is a non-literal that can be decomposed and so this branch is not fully decomposed. The right branch contains the wff Q . This is a literal wff and cannot have a decomposition rule applied to it. Since line 1 is decomposed and Q cannot be decomposed, the right branch is a fully decomposed branch.

Let's complete the decomposition of the left branch by applying the $\neg \rightarrow D$ rule to $\neg(P \rightarrow S)$. In addition, just like before, let's place a checkmark next to the wff to indicate that it has been decomposed.

1. $(P \rightarrow S) \rightarrow Q \checkmark$ P
2. $\neg(P \rightarrow S) \checkmark \quad Q$ $1 \rightarrow D$
3. P $2 \neg \rightarrow D$
4. $\neg S$ $2 \neg \rightarrow D$

The left branch now contains two non-literal wffs that have had decomposition rules applied to them (lines 1 and 2) and two literal wffs. Since this is a branch where every non-literal wff has had a decomposition rule applied to it, this branch is fully decomposed.

Exercise 4.41

Create a truth tree and decompose the following wffs until all of the branches in the tree are fully decomposed. In cases where there are multiple wffs, simply stack each wff, number it, and justify it with "P".

1. $A, A \wedge \neg B$
2. $A \wedge B, \neg(A \vee B)$
3. $\neg A \wedge B, P \vee \neg Q$
4. $P \rightarrow Q, P \leftrightarrow Q$
5. $\neg(A \rightarrow B), \neg(A \leftrightarrow B)$
6. $\neg \neg A, (A \vee B) \rightarrow C$
7. $\neg B \rightarrow C, (A \wedge B) \rightarrow (C \vee B)$
8. $\neg A \vee \neg B, \neg(A \wedge B) \wedge C, \neg A \leftrightarrow B$

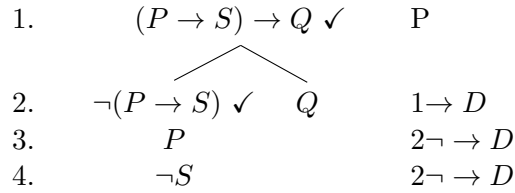
4.2.5 Completed Open Branches

Earlier we defined a closed branch as a branch that contains a wff ϕ and its literal negation $\neg(\phi)$. In contrast, an open branch was defined as branch that has not been closed. We can now define a completed open branch. A completed open branch is a branch that is open and fully decomposed.

Definition 4.2.5: Completed Open Branch (COB)

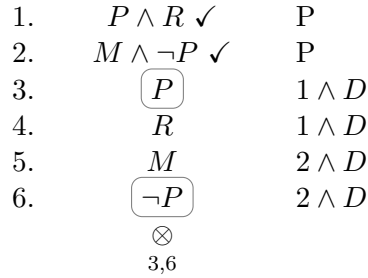
A completed open branch (COB) is a fully decomposed branch that is not closed. That is, it is a fully-decomposed branch that does not contain a wff P and its literal negation $\neg(P)$.

To illustrate, let's consider a tree we considered in the prior section:



This tree contains two branches. The left branch is fully decomposed and open: (1) each wff that can be decomposed has been decomposed and (2) there is not a wff ϕ and its literal negation $\neg(\phi)$ in the branch. As such, it is a completed open branch. The right branch is also fully decomposed and open. As such, it is also a completed open branch.

Let's consider another example. Consider the following tree:



In this tree, each wff that can be decomposed has been decomposed. The single branch of the tree is therefore a fully decomposed branch. However, the branch is not open (it is closed) as there is a wff ϕ and its literal negation $\neg\phi$ in the branch. Namely, P and its literal $\neg P$ are in the branch. Therefore, the branch is a closed branch rather than a completed open branch.

4.2.6 Completed Open Trees and Closed Trees

Now that we have defined and illustrated the difference between a closed branch and a completed open branch, let's define a closed tree and a completed open tree. Once a tree is fully decomposed, the tree is either a completed open tree or a closed tree (not both and not neither). A completed open tree is a tree containing at least one completed open branch.

Definition 4.2.6: Completed Open Tree

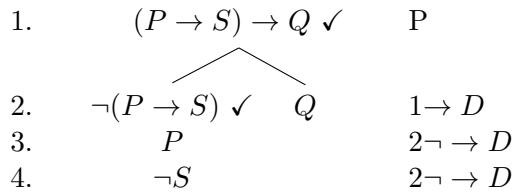
A tree is a completed open tree if and only if it has at least one completed open branch. That is, a tree is a completed open tree if and only if it contains at least one fully decomposed branch that is not closed. A completed open tree is a tree where there is at least one branch that has an **O** under it.

A closed tree is a tree containing only closed branches.

Definition 4.2.7: Closed Tree

A tree is closed when all of the tree's branches are closed. A closed tree will have an $\oplus$ under every branch.

Let's consider some examples of completed open trees and closed trees. First, let's consider a tree we considered earlier:



In our discussion of the completed open branch, we noted that this tree contains two completed open branches: both branches are fully decomposed and neither contains a wff ϕ and its literal negation $\neg(\phi)$. Since this tree contains *at least one completed open branch*, it is a completed open tree. In contrast, consider the following tree:

1.	$P \wedge R \checkmark$	P
2.	$M \wedge \neg P \checkmark$	P
3.	$\boxed{P}$	$1 \wedge D$
4.	R	$1 \wedge D$
5.	M	$2 \wedge D$
6.	$\boxed{\neg P}$	$2 \wedge D$
	$\otimes$	
	3,6	

When discussing this tree earlier, we noted that the single branch of this tree is closed. Since all of this tree's branches are closed, it is a closed tree. Let's consider a third example. Consider the following tree:

1.	$(\neg P \rightarrow Q) \wedge \neg(P \vee Q) \checkmark$	P
2.	$\neg P \rightarrow Q \checkmark$	$1 \wedge D$
3.	$\neg(P \vee Q) \checkmark$	$1 \wedge D$
4.	$\neg P$	$3 \neg \vee D$
5.	$\neg Q$	$3 \neg \vee D$
	$\swarrow \quad \searrow$	
6.	$\neg \neg P \quad Q$	$2 \rightarrow D$
	$\otimes \quad \otimes$	
	4,6 5,6	

Notice that this tree contains two branches. The left branch is closed as it contains a wff $\neg \neg P$ and its literal negation $\neg P$. The right branch is also closed as it contains a wff Q and its literal negation $\neg Q$. Since all of the branches in this tree are closed, this tree is a closed tree.

Exercise 4.42

Use the truth tree method to determine whether the following set of wffs produces a completed open tree or a closed tree.

1. $A \wedge (\neg A \wedge B)$
2. $A \wedge \neg(A \vee B)$
3. $A \wedge B, \neg(A \rightarrow B)$
4. $A \vee B, \neg A \wedge \neg B$
5. $\neg(A \rightarrow B), \neg(A \leftrightarrow B)$
6. $A \rightarrow \neg(\neg A \wedge B), \neg \neg B \vee C$

4.3 TRUTH TREES: DECOMPOSITION STRATEGIES

Before examining how truth trees can be used to determine various logical properties, it is worth discussing some strategies for decomposing trees.

By "strategy" I mean a rule of thumb for producing smaller trees. One caveat. It is not necessary to employ strategies for decomposing **PL** trees. You can randomly pick wffs to decompose and, in the end, you will have a tree that will be a completed open tree or a closed tree. You also can use said tree to determine the relevant logical properties that trees are used to test for. However, employing strategies can make the decomposition of trees more efficient (i.e., simpler trees are produced). Here are three strategic rules for decomposing trees:

1. Use only the decomposition rules you need to determine if the tree is a completed open tree or a closed tree.
2. Use rules that close branches.
3. Use stacking rules before branching rules.

Let us illustrate with some examples where these rules are applied. First, if the goal in decomposing a truth tree is to determine whether the tree is a completed open tree or a closed tree, then it is not necessary to decompose every wff in the tree. This is the case for two reasons. First, if a branch is closed, then all branches that descend from that branch are closed. Further decomposition of the branch would not change the fact that the branch is closed. Let's illustrate this with an example. Consider the following decomposition of the wff $((A \vee B) \wedge (R \vee \neg S)) \wedge (P \wedge \neg P)$.

- | | | |
|----|---|--------------|
| 1. | $((A \vee B) \wedge (R \vee \neg S)) \wedge (P \wedge \neg P) \checkmark$ | P |
| 2. | $(A \vee B) \wedge (R \vee \neg S)$ | $1 \wedge D$ |
| 3. | $P \wedge \neg P$ | $1 \wedge D$ |

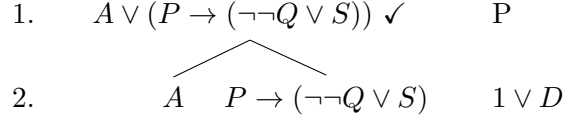
Suppose we decompose line 3 involving $P \wedge \neg P$.

- | | | |
|----|---|--------------|
| 1. | $((A \vee B) \wedge (R \vee \neg S)) \wedge (P \wedge \neg P) \checkmark$ | P |
| 2. | $(A \vee B) \wedge (R \vee \neg S) \checkmark$ | $1 \wedge D$ |
| 3. | $P \wedge \neg P$ | $1 \wedge D$ |
| 4. | P | $3 \wedge D$ |
| 5. | $\neg P$ | $3 \wedge D$ |

At this point in the tree, we might try to decompose the remaining wffs by decomposing line 2 and then decomposing the disjunctions that result from that decomposition. But, this further decomposition is unnecessary since all branches that descend from line 5 will also be closed as they will contain the wffs P and $\neg P$ in them. As such, this tree is a closed tree.

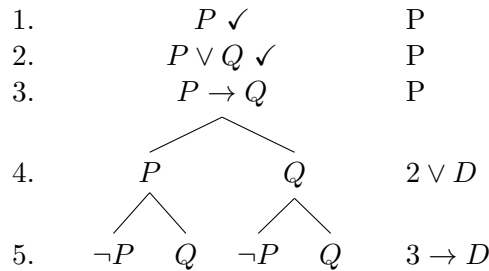
Second, recall that a completed open tree is defined as a tree with at least one completed open branch. What this means then is that as soon as the tree has a single completed open branch, we have enough information to

determine that the tree is a completed open tree. For this reason, it is not necessary to decompose every wff in the tree to determine if a tree is a completed open tree. To illustrate, consider the following tree:

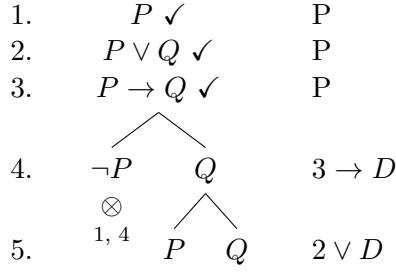


At this point in the tree, there is enough information to determine whether the tree is a completed open tree or a closed tree. Notice that the left branch of the tree is fully decomposed and there is not a wff ϕ and its literal negation $\neg\phi$ in the branch. As such, the left branch is a completed open branch. Since the tree contains at least one completed open branch, it is a completed open tree. Even though the right branch is not fully decomposed, it is thus not necessary to decompose the right branch of the tree to determine if the tree is a completed open tree or a closed tree since the left branch is sufficient for making this determination.

The second strategic rule is try to use rules that close branches rather than rules that do not close branches. The rationale behind this rule is when a wff is decomposed, it must be decomposed under every open branch that descends from that wff. Therefore, whenever it is possible to use a rule that closes a branch, there are fewer branches in which a wff must be decomposed. To illustrate, we can consider the previous example in this section, where we choose to decompose the conjunction $P \wedge \neg P$ rather than the conjunction $(A \vee B) \wedge (R \vee \neg S)$. In making this decision, we chose to decompose a wff that closes a branch rather than a wff that keeps a branch open. However, let's consider another example for further illustration. Consider the following tree:

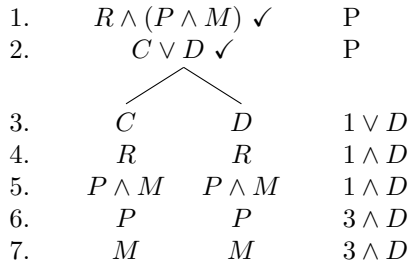


Notice that with this tree, $P \vee Q$ at line 2 is decomposed first, and then $P \rightarrow Q$ at line 3 is decomposed. Notice that in decomposing the disjunction before the conditional, both branches are kept open. In contrast, consider the following tree where $P \rightarrow Q$ is decomposed before $P \vee Q$:

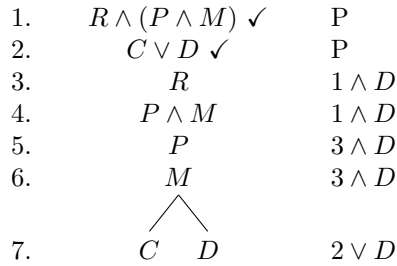


Notice that since the conditional is decomposed first, this results in closing the left branch. Since the left branch is closed, we are not required to decompose the disjunction under this branch. Instead, it only must be decomposed under the right branch. As a result, we obtain a smaller tree.

The third and final strategic rule is that more efficient trees are produced by using stacking rules before branching rules. Or, more simply "stack before you branch". The rationale behind this rule is that since branching rules create branches and when wffs are decomposed, they are decomposed under every open branch that descends from that wff, by creating more branches, one creates more open branches in which that wffs must be decomposed. On the other hand, stacking rules do not increase the number of branches in a tree. To illustrate, compare the trees in [Tree 4.1](#) and [Tree 4.2](#). In [Tree 4.1](#), a branching rule is employed before a stacking rule. In contrast, in [Tree 4.2](#), a stacking rule is employed before a branching rule. While the number of lines in each tree is the same, notice that the "stack first tree" ([Tree 4.2](#)) has fewer nodes than the "branch first tree" ([Tree 4.1](#)).



Tree 4.1: Branch first tree



Tree 4.2: Stack first tree

In short, while the use of the strategic rules is not necessary, tidier trees can sometimes be produced by employing the strategic rules for decomposition. To create these more economical trees, one can try to use no more rules than needed to determine if the tree is a completed open tree

or a closed tree, use rules that close branches rather than those that keep them open, and use stacking rules before branching rules.

Exercise 4.43

Use a truth tree method and various strategies for simplifying truth trees to determine if the tree is a completed open tree or a closed tree.

1. $P \vee Q, A \wedge \neg A, B \rightarrow C$
2. $P \rightarrow \neg P, A \vee \neg(B \wedge C)$
3. $A \wedge \neg B, B, C \rightarrow Q$
4. $A \vee B, \neg A \wedge \neg B$
5. $P \vee (Q \vee R), P \wedge \neg P$
6. $(P \vee P) \wedge (R \wedge \neg R), S \leftrightarrow T, P \vee M$
7. $(P \wedge Q) \wedge R, M \vee (Q \vee \neg R)$

4.4 TRUTH TREES: ANALYSIS

Truth trees can be used to determine various logical properties about wffs, sets of wffs, and arguments. Using truth trees to do this requires that you (1) set up the tree in a specific way to test for a specific property (you can't just stack the wffs in every instance), (2) decompose the tree using decomposition rules to a point where it is a closed or completed open tree, and (3) know what logical property is associated with the fact that the tree is closed or completed and open. Similar to how we used truth tables to determine whether a wff is a tautology, contradiction, or contingent, whether a set of wffs is consistent or inconsistent, and whether an argument is valid or invalid, in this section, we will use the truth tree method to determine these same logical properties. But, first, let's consider what a completed open branch tells us about the wffs in that branch.

4.4.1 Recovering an Interpretation

What does a closed and a completed open branch tell us about the wffs in the branch? First, recall that the decomposition of a wff ϕ is a way to represent the wff or wffs that are true if ϕ to be true. For example, consider the decomposition of $P \wedge Q$ using $\wedge D$. This decomposition tells us that if $P \wedge Q$ is true, then both P and Q must be true. Second, a fully decomposed branch is a branch where every wff in the branch has been decomposed into literal wffs, i.e., atomic wffs and negated atomic

wffs. Thus, a fully decomposed branch is a way to represent the truth conditions of the wffs in the branch. For example, consider the tree below:

1. $P \wedge R \checkmark$ P
2. $\neg M \wedge P \checkmark$ P
3. P $1 \wedge D$
4. R $1 \wedge D$
5. $\neg M$ $2 \wedge D$
6. P $2 \wedge D$

There is a single branch in the tree. Since the branch is fully decomposed, the fully decomposed branch can be read as saying the following:

If the literal wffs $P, \neg M, R$ are true, then all of the wffs in the branch, including the root wffs $\{(P \wedge R, \neg M \wedge P)\}$ are true.

Third, since the above remark is true, it appears that we can obtain an interpretation that would make all of the wffs in the branch true. To do this, we would assign T to the propositional letters in atomic wffs and F to the letter in negated atomic wffs. For example, in the above example, we would assign T to P, R , and F to M . Notice that we could use this interpretation of the propositional letters and a truth table to show that the wffs in the branch are true. For example, consider the truth table below for the wffs $P \wedge R, \neg M \wedge P$.

P	R	M	$P \wedge R$	$\neg M \wedge P$
T	T	F	T	T

However, just because it is true that "if the literal wffs in the branch are true, all of the wffs in the branch are true" does not mean that we can always specify an interpretation of the propositional letters that would make all of the wffs in the branch true. To see why consider the following tree (one very similar tree to the one we just considered):

1. $P \wedge R \checkmark$ P
 2. $\neg M \wedge \neg P \checkmark$ P
 3. P $1 \wedge D$
 4. R $1 \wedge D$
 5. $\neg M$ $2 \wedge D$
 6. $\neg P$ $2 \wedge D$
- $\otimes$
 3,5

Notice that in this tree, there is also a single branch that is fully decomposed. Since the branch is fully decomposed, the fully decomposed branch

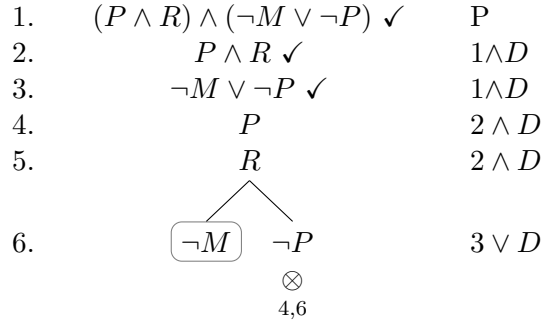
can be read as saying the following:

If the literal wffs $P, \neg M, R, \neg P$ are true, then all of the wffs in the branch, including the root wffs $\{(P \wedge R, \neg M \wedge \neg P)\}$.

However, notice that there is no *interpretation* of the propositional letters that would make all of the wffs in the branch true. This is because if $P \wedge R$ and $\neg M \wedge \neg P$ is true, then there is an interpretation where P and $\neg P$ must be true. But, this would mean that $\mathcal{I}(P) = T$ and $\mathcal{I}(P) = F$. Since this is not possible, there is no interpretation of the propositional letters that would make all of the wffs in the branch true.

What does all of this tell us? First, it tells us that if a tree contains a completed open branch (a branch that *does not* contain a wff ϕ and its negation $\neg(\phi)$), then there is an interpretation of the propositional letters that makes each of the wffs in the branch true. Second, it tells us that if the tree is closed, then there will be a wff ϕ and its literal negation $\neg(\phi)$ and there will not be an interpretation of the propositional letters that makes each of the wffs in the branch true.

Let's consider another example. The following tree decomposes $(P \wedge R) \wedge (\neg M \vee \neg P)$:



In the above tree, the left branch is a completed open branch. To construct an interpretation from this branch, list all of the literal wffs in the branch. These literal wffs are the following: $P, R, \neg M$. Next, for each unnegated literal wff, assign the corresponding propositional letter a value of T, while for each negated literal wff, assign the corresponding propositional letter a value of F.

1. Since P , $\mathcal{I}(P) = T$
2. Since R , $\mathcal{I}(R) = T$
3. Since $\neg M$, $\mathcal{I}(M) = F$

To check that this interpretation makes all of the wffs in the branch true, we can, of course, construct a truth table using this interpretation. But, let's not construct the entire table. If our process for extracting an interpretation from a completed open branch of a truth tree works, then $(P \wedge R) \wedge (\neg M \vee \neg P)$ is true in any interpretation where $\mathcal{I}(P) = T, \mathcal{I}(R) = T, \mathcal{I}(M) = F$.

M	P	R	$P \wedge R$	$\neg M \vee \neg P$
F	T	T	T	T

Notice that in the above table, when $\mathcal{I}(P) = T, \mathcal{I}(R) = T, \mathcal{I}(M) = F$, the wffs $P \wedge R, \neg M \vee \neg P$ are both true.

Exercise 4.44

Use a truth tree to determine whether the tree is a completed open tree or a closed tree. If the tree is a completed open tree, construct an interpretation (an assignment of truth values to propositional letters) that would make all of the wffs in the branch true. If the tree is a closed tree, explain why there is no interpretation that would make all of the wffs in the branch true.

1. $P \wedge \neg P$
2. $P \vee \neg P$
3. $P \wedge Q, Q \wedge R$
4. $P \wedge \neg Q, Q \wedge \neg R$
5. $\neg(P \vee Q), R \vee Q$
6. $\neg(P \rightarrow Q), S \vee T, P \rightarrow Q$
7. $P \vee \neg Q, S \leftrightarrow P, \neg S$

4.4.2 Consistency

Recall the definition of PL-consistency from chapter 3.

Definition 4.4.1: PL-Consistent

A non-empty set of wffs $\{\phi_1, \phi_2, \dots, \phi_n\}$ is PL-consistent if and only if each wff in $\{\phi_1, \phi_2, \dots, \phi_n\}$ is true under at least one interpretation.

A truth tree can be used to determine whether a set of wffs is PL-consistent. Since, a completed open branch tells us that there is an interpretation that would make all of the wffs in the branch true, if a tree has a completed open branch, then the set of wffs in the branch is

PL-consistent. What this means is that if we stack each member of the set at the top of our tree, decompose the wffs, and determine that the branch is a completed open branch, the nth collection of wffs at the top of the tree is consistent.

Definition 4.4.2: Truth-Tree Test for Consistency

A truth tree shows that a set of wffs $\{\phi_1, \phi_2, \dots, \phi_n\}$ is PL-consistent if there is at least one completed open branch for a tree whose root wffs are: $\phi_1, \phi_2, \dots, \phi_n$. If such a tree is closed, then $\{\phi_1, \phi_2, \dots, \phi_n\}$ is inconsistent.

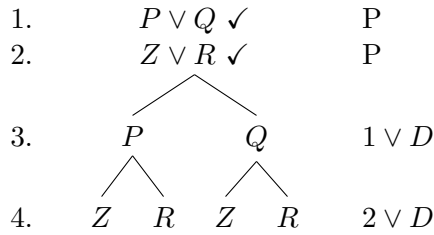
In other words, to test whether a set of wffs $\{A, B, C\}$ is consistent, start by creating a tree where A , B , and C , are put in the initial stack (root wffs). Next, decompose the tree. If the tree contains at least one completed open branch, then the tree test determines the set $\{A, B, C\}$ to be consistent. If the tree does not contain at least one completed open branch (the tree is closed), then the tree test determines the set $\{A, B, C\}$ to be inconsistent.

Let's illustrate the use of the truth-tree test for consistency. Suppose wish to test whether the set $\{P \vee Q, Z \vee R\}$ is consistent or inconsistent. In prior tree decompositions, we have examined the decomposition of single wffs. But, in this case, we have a set of wffs containing more than one element. To set up this test, we start by writing each wff in the set on its own line. In other words, we stack the wffs in the set, numbering each wff, and justifying each with "P" as it is a provided wff.

1. $P \vee Q$ P
2. $Z \vee R$ P

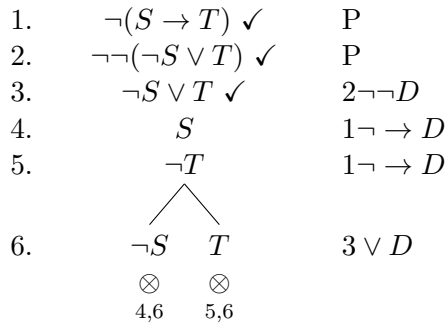
The reason for setting up the tree in this way is because what we want to test whether there is an interpretation that would make both $P \vee Q$ and $Z \vee R$ true. As such, we employ the stacking convention to indicate that we are testing whether there is an interpretation that would make both $P \vee Q$ and $Z \vee R$ true.

With the tree setup, the next step is to decompose the wffs in the tree. Since both wffs are disjunctions, the $\vee D$ rule is employed twice, first on line 1 and then on line 2. Since both branches are open after the first use of $\vee D$, the second use of $\vee D$ decomposes the wff at line 2 under both branches. The tree is found to have a total of four open branches, but a single completed open branch is sufficient for demonstrating that the set of wffs in the stack is consistent.



Next, we analyze or “read” the tree by examining whether the tree has at least one completed open branch. Since it does, the truth-tree test reveals that $\{P \vee Q, Z \vee R\}$ is *consistent*. Per the definition of consistency, there is at least one interpretation of P, Q , and Z that would make $P \vee Q$ and $Z \vee R$ true. Using the leftmost branch, and the method of obtaining that interpretation from a completed open branch, one such interpretation is the following: $\mathcal{I}(P) = T, \mathcal{I}(Q) = T, \mathcal{I}(Z) = T, \mathcal{I}(R) = T$.

If a tree does not have at least one completed open branch, the tree method determines the set of wffs being tested to be *inconsistent*. In other words, if the tree is a closed tree, the set of wffs is inconsistent. For example, consider the truth tree below:



Starting from the base of the tree, notice that there is a wff P and its literal negation $\neg(P)$ in both branches. In the left branch, there is $\neg S$ and S while in the right branch there is T and $\neg T$. As such, there is no way coherent interpretation of the propositional letters would make $\{\neg(S \rightarrow T), \neg\neg(\neg S \vee T)\}$ true. Thus, the tree method determines that $\{\neg(S \rightarrow T), \neg\neg(\neg S \vee T)\}$ is inconsistent.

Exercise 4.45

Using the truth-tree method, determine whether the following sets of wffs are consistent or inconsistent.

1. $P, \neg P$

2. $\neg P \wedge Q, Q \wedge \neg R$
3. $P \rightarrow Q, P, \neg Q$
4. $P \vee \neg Q, Q \vee \neg P$
5. $\neg\neg(P \vee Q), P \vee Q$
6. $P \rightarrow Q, Q \rightarrow \neg P$
7. P
8. $P \vee \neg Q, S \leftrightarrow P, \neg S$
9. $((P \wedge Q) \wedge S), (S \wedge \neg Q) \vee \neg S$
10. $\neg(P \rightarrow Q), \neg(R \vee Q) \vee S$

4.4.3 Tautology, Contradiction, Contingency

In chapter 3, we learned that a wff ϕ is a *PL-tautology* if and only if ϕ is true under every interpretation, a *PL-contradiction* if and only if ϕ is false under every interpretation, and a *PL-contingency* if and only if ϕ is neither always false under every interpretation nor always true under every interpretation. Truth trees can be used to determine whether a wff is a PL-tautology, PL-contradiction, or PL-contingency.

Definition 4.4.3: Truth-Tree Test for Tautology

A truth tree shows that ϕ is a PL-tautology if a tree of the stack of $\neg(\phi)$ determines a closed tree.

Definition 4.4.4: Truth-Tree Test for Contradiction

A truth tree shows that ϕ is a PL-contradiction if a tree of the stack of ϕ determines a closed tree.

Definition 4.4.5: Truth-Tree Test for Contingency

A truth tree shows that ϕ is a contingency if a tree of ϕ does not determine a closed tree and a tree of $\neg(\phi)$ does not determine a closed tree.

Let's consider some examples. Suppose we wish to test whether $P \vee \neg P$ is a contradiction, tautology, or contingency. There are two tests: the test for contradiction and the test for tautology. The test to see whether $P \vee \neg P$ involves setting up the tree by writing $P \vee \neg P$ at line 1, decomposing it, and then checking to see if the tree is a closed tree or a completed open tree. If the tree is closed, then there is no interpretation such that

$P \vee \neg P$ is true, and so $P \vee \neg P$ is a contradiction (as it is false under every interpretation). If the tree is a completed open tree, then there is at least one interpretation such that $P \vee \neg P$ is true, then it is the wff is not a contradiction. If it is not a contradiction, then $P \vee \neg P$ is either a contingency (true under at least one interpretation and false under at least one interpretation) or tautology (true under every interpretation).

1. $P \vee \neg P \checkmark$ P
- $\swarrow \quad \searrow$
 2. $P \quad \neg P$ 1 $\vee D$

Notice that the contradiction test for $P \vee \neg P$ reveals that there is at least one completed open branch and so this wff is not a contradiction.

To test to see whether $P \vee \neg P$ is a tautology, the tree is set up by creating a tree where the root node is the literal negation of $P \vee \neg P$. In this case, the root node would be $\neg(P \vee \neg P)$. From there, the wff is decomposed. If the tree is a closed tree, then there is no interpretation where $v(\neg(P \vee \neg P)) = T$. But, if that is the case, then $v(P \vee \neg P) = T$ for every interpretation. Hence $P \vee \neg P$ would be a tautology. However, if the tree is not a closed tree, then there is at least one interpretation where $v(\neg(P \vee \neg P)) = T$. But, if that is the case, then $v(P \vee \neg P) = F$ for at least one interpretation. Hence, $P \vee \neg P$ is not a tautology. And, if we show that $P \vee \neg P$ is neither a contradiction nor a tautology, then it is a contingency. Let's test $P \vee \neg P$ to see if it is a tautology.

1. $\neg(P \vee \neg P) \checkmark$ P
2. $\neg P$ 1 $\neg \vee D$
3. $\neg \neg P$ 1 $\neg \vee D$
- $\otimes$
 2,3

The tree for $\neg(P \vee \neg P)$ is closed. Again, this tells us that there is no interpretation such that $\neg(P \vee \neg P)$ is true. And, if that is the case, then there is no interpretation that would make $P \vee \neg P$ false. And, if there is no interpretation that would make $P \vee \neg P$ false, then $P \vee \neg P$ is true under every interpretation. Hence, $P \vee \neg P$ is a tautology.

Let's consider another example. Suppose we wish to test whether $P \wedge \neg P$ is a tautology, contradiction, or contingency. As mentioned before, the setup is important. We can set up the tree to test for tautology or contradiction. If we wish to test whether a wff is a contradiction, we stack the wff itself. In our example, we would write $P \wedge \neg P$ at line 1. If we wish to test whether a wff is a tautology, we stack the literal negation of the wff. In

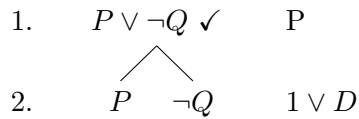
our example, we would write $\neg P \wedge \neg P$ at line 1. In our example, let's test $P \wedge \neg P$ to see if it is a contradiction. If it is not a contradiction, then it is a tautology or contingency, and we can then use the test for tautology to determine which one of the two it is.

1.	$P \wedge \neg P$	$\checkmark$	P
2.	P	$\checkmark$	$1 \wedge D$
3.	$\neg P$		$1 \wedge D$
	$\otimes$		
	2,3		

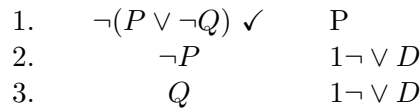
Notice that our tree produces a closed tree. Since $P \wedge \neg P$ produces a closed tree, the tree test determines that $P \wedge \neg P$ is a contradiction ($P \wedge \neg P$ is false under every interpretation of P). The reason this works is more straightforward than why the tree method works for tautology. Since the tree is closed, there is no interpretation that would make $P \wedge \neg P$ true. And, if there is no interpretation that would make $P \wedge \neg P$ is true, then $P \wedge \neg P$ is, by definition, a contradiction.

Next, let's consider an example where we test whether a wff is a contingency. Suppose we wish to test whether $P \vee \neg Q$ is a tautology, contradiction, or contingency. As mentioned before, the setup is important. We can set up the tree to test for tautology or contradiction. If we wish to test whether a wff is a contradiction, we stack the wff itself. In our example, we would write $P \vee \neg Q$ at line 1. If we wish to test whether a wff is a tautology, we stack the literal negation of the wff. In our example, we would write $\neg(P \vee \neg Q)$ at line 1. Note that it is a somewhat common mistake to write either $\neg P \vee Q$ or $\neg P \vee \neg Q$ at line 1. This is not the literal negation of $P \vee \neg Q$ and so would not be the correct setup for testing whether a wff is a tautology. In our example, let's start by testing $P \vee Q$ to see if it is a contradiction. If it is a contradiction, then no more trees are needed. If it is not a contradiction, then it is either a tautology or contingency. If it is a tautology, then we are done. However, it is neither a contradiction nor a tautology, then, by process of elimination, it is a contingency.

First, let's test whether $P \vee \neg Q$ is a contradiction. To do this, we stack $P \vee \neg Q$ at line 1. Next, we decompose the tree using the $\vee D$ rule. Using this rule produces a completed open tree (Tree 4.3). Since $P \vee \neg Q$ is not a contradiction, we can use the test for tautology to determine whether $P \vee \neg Q$ is a tautology or contingency. To do this, we stack $\neg(P \vee \neg Q)$ at line 1. Next, we decompose the tree using the $\neg \vee D$ rule. Using this rule also produces a completed open tree (Tree 4.4). Since $\neg(P \vee \neg Q)$ is not a tautology, then, by process of elimination, $P \vee \neg Q$ is a contingency.



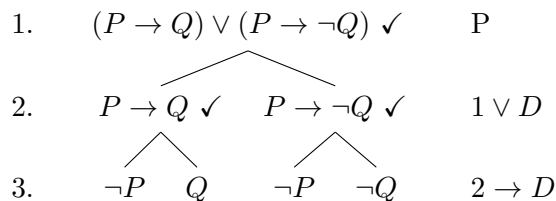
Tree 4.3: Contradiction test



Tree 4.4: Tautology test

The examples that we have considered thus far have been relatively straightforward. Perhaps you were able to determine whether the wff was a tautology, contradiction, or contingency simply by looking at the wff and without the help of a truth tree. However, this will not always be the case. For some large wffs or wffs with complex structure, it may not be immediately obvious whether the wff is a tautology, contradiction, or contingency. Consider the wff $(P \rightarrow Q) \vee (P \rightarrow \neg Q)$. Translated into pseudo-English, this wff can be read as saying "If P, then Q or if P, then not-Q". If "P" stands for "it is snowing" and "Q" stands for "the ground is wet", then this sentence says something to the effect of "If it is snowing, then the ground is wet or if it is snowing, then the ground is not wet". Is this sentence a tautology, contradiction, or contingency? For some, it is not immediately obvious since it is unclear whether a proposition either implies a sentence or its negation. However, we can use the truth-tree method to determine whether this wff is a tautology, contradiction, or contingency.

Let's begin by testing the wff to see if it is a contradiction. Start by writing the wff at line 1 of the tree. Next, decompose the tree using the $\vee D$ rule, followed by the $\rightarrow D$ rule. The tree is as follows:



The left branch of the above tree is a completed open branch. Since this completed open branch tells us that there is an interpretation such that $(P \rightarrow Q) \vee (P \rightarrow \neg Q)$ is true, it follows that this wff is not a contradiction. The wff then is either a tautology or contingency.

Since the wff is either a tautology or a contingency, let's test the wff to see if it is a tautology. To do this, we stack $\neg((P \rightarrow Q) \vee (P \rightarrow \neg Q))$ at line 1. Next, we decompose the tree using the $\neg \vee D$ rule. This rule gives us $\neg(P \rightarrow Q)$ and $\neg(P \rightarrow \neg Q)$ on lines 2 and 3, respectively. Finally, we

complete the tree by using $\neg \rightarrow D$ on both of these lines.

1.	$\neg((P \rightarrow Q) \vee (P \rightarrow \neg Q)) \checkmark$	P
2.	$\neg(P \rightarrow Q) \checkmark$	1 $\neg \vee D$
3.	$\neg(P \rightarrow \neg Q) \checkmark$	1 $\neg \vee D$
4.	$\neg P$	2 $\neg \rightarrow D$
5.	$\neg Q$	2 $\neg \rightarrow D$
6.	$\neg P$	3 $\neg \rightarrow D$
7.	$\neg \neg Q$	3 $\neg \rightarrow D$
	$\otimes$	
	5,7	

The tree above is a closed tree. What this indicates is that there is no interpretation where $\neg((P \rightarrow Q) \vee (P \rightarrow \neg Q))$ is true. It follows then that this wff is false under every interpretation. However, if this wff is false under every interpretation, then the unnegated version of this wff is true under every interpretation. And so, according to the truth-tree test for tautology, this means that $(P \rightarrow Q) \vee (P \rightarrow \neg Q)$ is a tautology.

To conclude this section, one final question concerning this test is the following: if all the branches in a tree for ϕ are open, doesn't this tell us ϕ is a tautology? For example, were we to decompose $P \rightarrow P$, the resulting tree would have two branches and both would be open. However, this is not the case with every wff. For example, consider the wff $P \rightarrow Q$. This wff is clearly not a tautology but decomposing it yields a tree with two branches, both of which are open.

Exercise 4.46

Using the truth tree method, determine whether the following wffs are PL-tautologies, PL-contradictions, or PL-contingencies.

1. $P \wedge \neg P$
2. $P \vee \neg P$
3. $P \rightarrow P$
4. $P \leftrightarrow \neg P$
5. $P \wedge (Q \rightarrow P)$
6. $P \vee (Q \rightarrow P)$
7. $P \rightarrow (Q \rightarrow \neg P)$
8. $(P \vee Q) \wedge (\neg P \vee \neg Q)$
9. $(P \rightarrow Q) \leftrightarrow \neg(P \wedge \neg Q)$

4.4.4 Equivalence

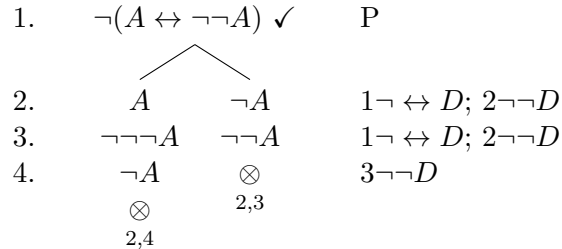
In chapter 3, we learned that a set of wffs Γ is said to be *PL-equivalent* if and only if the wffs in Γ have identical truth values under every interpretation. Concern for equivalence typically emerges when there is a question of whether two wffs ϕ and ψ are equivalent. With this in mind, our focus in this section will be on how to use the truth-tree method to determine if two wffs ϕ, ψ are PL-equivalent.

Definition 4.4.6: Truth-Tree Test for Equivalence

A truth tree shows that ϕ and ψ are PL-equivalent if a tree of $\neg(\phi \leftrightarrow \psi)$ determines a closed tree. Otherwise, the tree shows that ϕ and ψ are not PL-equivalent.

The general procedure for using the truth-tree test for equivalence is as follows: to test whether two wffs ϕ and ψ are PL-equivalent, start by stacking $\neg(\phi \leftrightarrow \psi)$ at line 1. Next, decompose the tree using the $\neg \leftrightarrow D$ rule. After that, decompose the wffs in the tree until it is determined that the tree is a completed open tree or the tree is closed. If the tree is closed, then ϕ and ψ are PL-equivalent. If the tree is a completed open tree, then ϕ and ψ are not PL-equivalent.

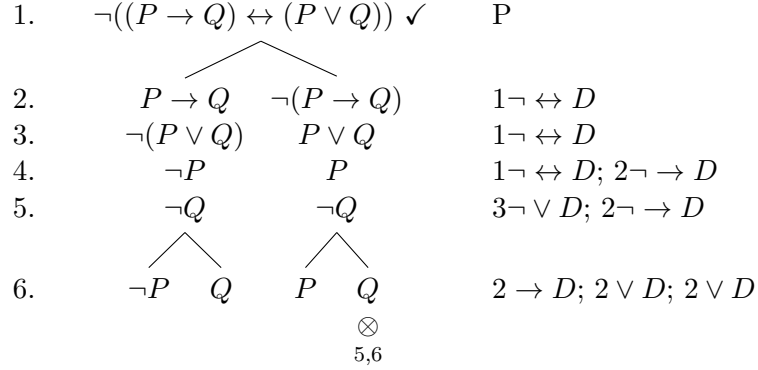
Let's illustrate the use of the truth-tree test for equivalence. We will start with the simple example of testing whether A and $\neg\neg A$ are PL-equivalent. To do this, begin by stacking $\neg(A \leftrightarrow \neg\neg A)$ at line 1. Next, decompose the tree using the $\neg \leftrightarrow D$ rule. Since the $\neg \leftrightarrow D$ rule involves both branching and stacking, the use of this rule gives us A and $\neg\neg\neg A$ on the left branch, and $\neg A$ and $\neg\neg A$ on the right branch. Next, we decompose the wffs in the tree. The right branch is closed and so there is no further work that needs to be done with it. The left branch is open and so it is necessary to decompose the triply negated $\neg\neg\neg A$ using the $\neg\neg D$ rule. After using this rule, the left branch now is closed.



In examining the tree above, we see that the tree is a closed tree. Since

the tree is a closed tree, the truth-tree method contends that A and $\neg\neg A$ are PL-equivalent. That is, that A and $\neg\neg A$ have identical truth values under every interpretation.

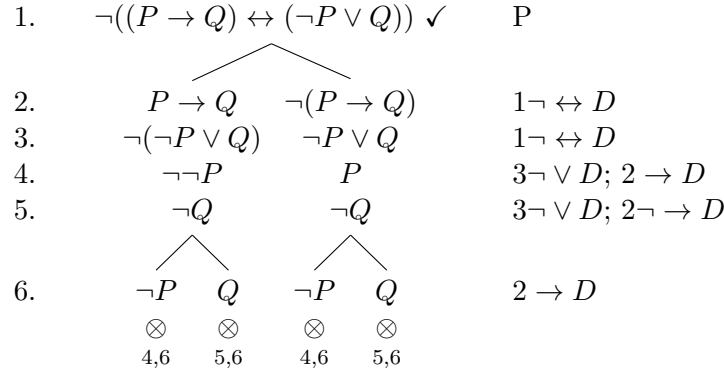
Next, let's consider a slightly more complicated example since it likely seemed obvious that A and $\neg\neg A$ are logically equivalent. Consider whether $P \rightarrow Q$ and $P \vee Q$ are PL-equivalent. To begin, stack $\neg(P \rightarrow Q \leftrightarrow P \vee Q)$ at line 1. Next, decompose the tree using the $\neg \leftrightarrow D$ rule. The rule involves stacking and branching. The left branch consists of $P \rightarrow Q$ and $\neg(P \vee Q)$, while the right branch consists of $\neg(P \rightarrow Q)$ and $P \vee Q$. Next, we decompose the remaining wffs in the tree until we can determine whether it is a completed open tree or a closed tree. On the left branch, $\neg \vee D$ is a stacking rule, let's apply this rule to $\neg(P \vee Q)$. We can finish the decomposition of the left branch by decomposing the conditional $P \rightarrow Q$. On the right branch, $\neg \vee D$ is a stacking rule, let's apply this rule to $P \vee Q$. We can finish the decomposition of the right branch by decomposing the negated conditional $\neg(P \rightarrow Q)$. The tree is a closed tree. Since the tree is a closed tree, the truth-tree method contends that $P \rightarrow Q$ and $P \vee Q$ are not PL-equivalent. That is, that $P \rightarrow Q$ and $P \vee Q$ do not have identical truth values under every interpretation.



Since the tree test shows that $P \rightarrow Q$ and $P \vee Q$ are not PL-equivalent and this tree has at least one completed open branch, it is possible to specify an interpretation where the truth values of $P \rightarrow Q$ and $P \vee Q$ are not identical. To do this, select a completed open branch. We will use the leftmost branch and then use the method of obtaining an interpretation from a completed open branch. Using this branch, on the interpretation $\mathcal{I}(P) = F$ and $\mathcal{I}(Q) = F$, the truth values of $P \rightarrow Q$ and $P \vee Q$ are not identical.

Finally, let's test whether $P \rightarrow Q$ and $\neg P \vee Q$ are PL-equivalent. It is not

immediately obvious that these two wffs are PL-equivalent since $P \rightarrow Q$ says "if P then Q" while $\neg P \vee Q$ says "not-P or Q". We can use the truth-tree method to determine whether they are PL-equivalent. To do this, we begin by stacking $\neg((P \rightarrow Q) \leftrightarrow (\neg P \vee Q))$ at line 1. Next, we decompose the tree using the $\neg \leftrightarrow D$ rule. The $\neg \leftrightarrow D$ rule is a rule that involves both stacking and branching.



Tree 4.5: Equivalence test for $P \rightarrow Q$ and $\neg P \vee Q$

Tree 4.5 is a closed tree. Since the tree is a closed tree, the truth-tree method contends that $P \rightarrow Q$ and $\neg P \vee Q$ are PL-equivalent.

Exercise 4.47

Using the truth tree method, determine whether the following pairs of wffs are PL-equivalent.

1. $P \rightarrow Q, \neg P \wedge Q$
2. $P \wedge Q, \neg P \wedge \neg Q$
3. $P \rightarrow Q, \neg P \vee \neg Q$
4. $\neg(P \vee Q), \neg P \wedge \neg Q$
5. $\neg(P \rightarrow Q), P \wedge \neg Q$

4.4.5 Semantic Entailment

In chapter 1, an argument was said to to be deductively valid if and only if it is impossible from the premises to be true and the conclusion false. In chapter 3, we tried to make this idea more precise through the notion of semantic entailment (consequence). There we said that a wff ϕ is a semantic consequence in PL of a set of wffs Γ if and only if there is no interpretation that makes all of the wffs in Γ true and ϕ false. Even further, we developed a way to test whether a set of wffs Γ entails a wff

ϕ using the truth table. However, at the end of chapter 3, we also noted several problems with this test. In this section, we define and illustrate the use of the truth-tree method for testing whether a set of wffs Γ entails a wff ϕ .

Definition 4.4.7: Truth-Tree Test for Semantic Entailment

A truth tree shows that $\Gamma \models \phi$ if a tree consisting of the members of Γ stacked with the literal negation of ϕ determines a closed tree. Otherwise, the tree shows that $\Gamma \not\models \phi$.

Let's consider a simple example where we test whether $P, P \rightarrow Q \models Q$. To test this entailment, we begin by setting up the tree. The setup involves stacking the wffs before the sign for entailment. In this case, we stack $P, P \rightarrow Q$ at lines 1 and 2, respectively. In addition, stack the literal negation of the wff after the sign for entailment. In this case, we stack $\neg(Q)$ at line 3.

- | | | |
|----|-------------------|---|
| 1. | P | P |
| 2. | $P \rightarrow Q$ | P |
| 3. | $\neg(Q)$ | P |

With the tree setup, the next step is to decompose the tree and then determine if it yields a closed tree or a completed open tree. If the tree is a closed tree, then the truth-tree method contends that $\Gamma \models \phi$. If the tree is a completed open tree, then the truth-tree method contends that $\Gamma \not\models \phi$.

- | | | |
|---------------------------|---|-------------------|
| 1. | P | P |
| 2. | $P \rightarrow Q \checkmark$ | P |
| 3. | $\neg(Q)$ | P |
| $\swarrow \quad \searrow$ | | |
| 4. | $\neg P \quad Q$ | $2 \rightarrow D$ |
| | $\otimes \quad \otimes$
1,3 3,4 | |

Since the tree for $P, P \rightarrow Q, \neg(Q)$ yields a closed tree (all closed branches), the truth-tree method contends that the $P, P \rightarrow Q \models Q$. That is, that $P, P \rightarrow Q$ semantically entails Q .

Why does the truth-tree method work for testing for semantic entailment? The truth tree method searches for an interpretation that would make all of the wffs in the stack true. When creating a tree consisting of the members of Γ stacked with the literal negation of ϕ , the truth tree method

is searching for an interpretation that would make all of the wffs in Γ true and $\neg(\phi)$ true. But this is the same thing as searching for an interpretation that would make all of the wffs in Γ true and ϕ false. Since there is no interpretation that would make all of the wffs in Γ true and ϕ true, then Γ semantically entails ϕ .

Let's consider another example. This time consider $P \rightarrow Q, Q \models P$. To test this entailment, begin by setting up the tree. The setup involves stacking the wffs before the sign for entailment. In this case, we stack $P \rightarrow Q, Q$ at lines 1 and 2, respectively. In addition, stack the literal negation of the wff after the sign for entailment. In this case, we stack $\neg(P)$ at line 3. Once the truth tree is setup, the next step is to decompose the tree and then determine if it yields a closed tree or a completed open tree.

1.	$P \rightarrow Q$	P
2.	Q	P
3.	$\neg(P)$	P
		
4.	$\neg P \quad Q$	$1 \rightarrow D$

In the case of the tree for $P \rightarrow Q, Q, \neg(P)$, there is a completed open tree (all open branches). Since the tree is a completed open tree, the truth-tree method contends that the $P \rightarrow Q, Q \not\models P$. That is, that $P \rightarrow Q, Q$ does not semantically entail P .

Exercise 4.48

Using the truth tree method, determine whether the following are cases of entailment or non-entailment.

1. $P \wedge Q \models Q$
2. $A \vee B \models B$
3. $A \rightarrow B, B \models A$
4. $(A \vee B) \rightarrow C, A \models C$
5. $(A \wedge B) \rightarrow C, A \models C$
6. $P \vee Q, P \rightarrow R, Q \rightarrow R \models R$
7. $P \rightarrow Q, Q \rightarrow R \models P \rightarrow R$
8. $P \rightarrow Q, P \wedge W \models Q \vee R$
9. $P \wedge M \models \neg(P \rightarrow Q)$

4.5 ADVANTAGES OF THE TRUTH-TREE METHOD

In chapter 1 and 3, we formulated two different accounts of what it means for a conclusion to "follow from" a set of premises. In chapter 1, we defined "following from" in terms of deductive validity. An argument was said to be "valid" if and only if it was impossible for the premises to be true and the conclusion false. In chapter 3, we defined "following from" in terms of semantic entailment. A set of wffs Γ was said to semantically entail a wff ϕ if and only if there was no interpretation that made all of the wffs in Γ true and ϕ false. In those two chapters, we also introduced various methods for determining whether an argument is valid or whether a wff is semantic consequence of a set of wffs. In chapter 1, the "imagination test" was introduced. This test checks the validity of the argument by trying to imagine a scenario where the premises are true and the conclusion is false. If we can imagine such a scenario, then the argument is *invalid*. If we can't, then the argument is said to be *valid*. The many limitations of this method calls for a better, more reliable method. In chapter 3, we developed such a method with the method of truth tables. A truth table considers the truth value of a wff or set of wffs by determining its truth value under every interpretation. This table can then be used to determine whether an argument is valid or whether a wff is a semantic consequence of a set of wffs. The truth table method has several strengths over the imagination test, but, as we saw, it is not without its own problems.

Recall that the main strength of the truth-table method over the informal methods (e.g., intuition, imagination) is that it is a *mechanical* method for checking the validity of an argument. That is, it is a method that can be carried out by a machine. The main problems with this method were as follows. First, it cannot easily handle arguments that involve many propositional letters. As the number of propositional letters increases, the number of rows required to complete the table increases exponentially. Second, we cannot represent *every* valid English argument in *PL*.

In this chapter, we introduced a new method for checking the validity of an argument. This method is called the *truth-tree method*. The truth tree method is similar to the truth-table method as both are mechanical methods for checking the validity of an argument. However, tree trees have a few advantages over tables. First, the complexity of a truth tree is not a function of the number of propositional letters. That is, where the argument $\neg(M \vee S), \neg((T \vee R) \vee Q), \neg(\neg(S \vee R) \vee Q) \models \neg A$ would require **64 rows** ($2^6 = 64$), a truth tree quickly reveals the above argument is valid. Second, the truth tree method is also more powerful in that it

can be used in other, more expressive logical languages. These logical languages are capable of representing more valid arguments in English than the language of propositional logic. This chapter has illustrated the former advantage. Chapters 6 and 7 will illustrate the latter advantage.

4.6 PL TREE RULES

$$\begin{array}{c} \phi \wedge \psi \\ \phi \quad \wedge D \\ \psi \quad \wedge D \end{array}$$

$$\begin{array}{c} \neg(\phi \wedge \psi) \\ \swarrow \quad \searrow \\ \neg(\phi) \quad \neg(\psi) \quad \neg \wedge D \end{array}$$

$$\begin{array}{c} \phi \vee \psi \\ \swarrow \quad \searrow \\ \phi \quad \psi \quad \vee D \end{array}$$

$$\begin{array}{c} \neg(\phi \vee \psi) \\ \neg(\phi) \quad \neg \vee D \\ \neg(\psi) \quad \neg \vee D \end{array}$$

$$\begin{array}{c} \phi \rightarrow \psi \\ \swarrow \quad \searrow \\ \neg(\phi) \quad \psi \quad \rightarrow D \end{array}$$

$$\begin{array}{c} \neg(\phi \rightarrow \psi) \\ \phi \quad \neg \rightarrow D \\ \neg(\psi) \quad \neg \rightarrow D \end{array}$$

$$\begin{array}{c} \phi \leftrightarrow \psi \\ \swarrow \quad \searrow \\ \phi \quad \neg(\phi) \quad \leftrightarrow D \\ \psi \quad \neg(\psi) \quad \leftrightarrow D \end{array}$$

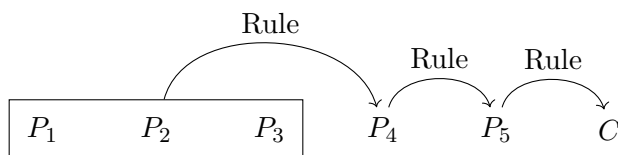
$$\begin{array}{c} \neg(\phi \leftrightarrow \psi) \\ \swarrow \quad \searrow \\ \phi \quad \neg(\phi) \quad \neg \leftrightarrow D \\ \neg(\psi) \quad \psi \quad \neg \leftrightarrow D \end{array}$$

$$\begin{array}{c} \neg\neg\phi \\ \phi \quad \neg\neg D \end{array}$$

5.1 TWO NOTIONS OF LOGICAL ENTAILMENT

In previous chapters, we clarified what it means for a proposition to "follow from" another proposition by using the notion of *deductive validity*. An argument is deductively valid if and only if it is impossible for the premises to be true and the conclusion false. However, we pointed out that it was difficult to test arguments to see if they are valid and the notion of impossibility is also difficult to grasp. For these reasons (and others), we introduced the notion of *semantic entailment*. A wff ϕ is said to "follow from" a set of wffs Γ if and only if there was no interpretation that make the members of Γ true and ϕ false. But, there is another, perhaps, more natural way of understanding what it means for a proposition to follow from a set of propositions. Namely, a conclusion can be said to follow from a set of premises provided the conclusion can be "derived" from those premises. In other words, a proposition P follows from a set of propositions Γ if and only if there is a derivation of P from Γ .

What is a derivation? Intuitively, to derive that a conclusion C from a set of propositions $P_1, P_2, \dots, P_n$ is to provide a sequence of propositions starting with $P_1, P_2, \dots, P_n$ and then reason in a step-by-step manner, where each step in the process is justified by a "rule of reason", until ultimately one can conclude the sequence with the proposition C . That is, if P_1, P_2, P_3 were our premises and C was our conclusion, then a derivation of C from P_1, P_2, P_3 might look something like this:



Let's highlight a few points about this intuitive idea of a derivation. First, a derivation is a *finite* sequence of propositions. Therefore, if in order to show that C follows from a set of premises, we needed to undertake an infinite number of steps, then we would not have a derivation. Second, a derivation of C is a finite sequence of propositions where the last proposition in the sequence is C . That is, a derivation ends with its conclusion. Third, each proposition in the sequence is either a member of $P_1, P_2, \dots, P_n$, an assumption (more on this later), or is the result of using the rules of inference (rules of reasoning or proof rules). That is, each

"Logic deals with *reasoning*. We expect it to pick out good inferences, not just good propositions; we are less interested in believing tautologies than in safely advancing from old beliefs to new ones." [10, p. 105]

proposition is one of our premises, a proposition we assume to be true for the purpose of derivation (but don't claim to be true or follows from prior propositions), or is the result of using the rules that allow us to reason from prior propositions to new propositions.

Let's conclude our discussion with a simple illustration of a derivation. Suppose we want to prove that "Frank is a criminal" from the following three premises "If John is lying and Sally is lying, then Frank is a criminal. John is lying. Sally is lying." Since a derivation must end with its conclusion and since we have premises, the basic structure of our derivation would look something like this:

1. If John is lying and Sally is lying, then Frank is a criminal. (premise)
2. John is lying. (premise)
3. Sally is lying. (premise)
4. A proposition we reasoned to using a rule
5. A proposition we reasoned to using a rule
6. A proposition we reasoned to using a rule
7. Therefore, Frank is a criminal (conclusion)

To construct our derivation, we need to supply the propositions between the premises of the above argument and its conclusion. And, we need to do so only with propositions that are permitted by our rules for writing new propositions. Suppose we have a rule that says whenever there are two propositions in a derivation, we can always reason to their conjunction. Call this the "AND-Introduction Rule" (or "And-Intro" for short).

Definition 5.1.1: AND-Introduction

From any two propositions P_1 and P_2 , it is legitimate to reason to a third proposition that is the result of placing "and" between them: " P_1 and P_2 ".

If this is a rule for writing new propositions, then we can add to our derivation in the following way.

1. If John is lying and Sally is lying, then Frank is a criminal. (premise)
2. John is lying. (premise)
3. Sally is lying. (premise)
4. Therefore, John is lying and Sally is lying (from 2, 3, using And-Intro)

At this point in the derivation, we do not have our conclusion. Let's add another rule for writing new propositions. Let's say that whenever you

have a proposition that has the form of "if P , then Q " (where P and Q are propositions) *and* you have the proposition " P " on its own line, then you can write down the proposition " Q " on a new line. If this is a rule of inference, then we can add to our derivation in the following way. Let's call this rule the "IF-THEN Rule".

Definition 5.1.2: IF-THEN Rule

From the two propositions "if P , then Q " and P , it is legitimate to reason to a third proposition Q .

Next, let's use the IF-THEN rule in our argument:

1. If John is lying and Sally is lying, then Frank is a criminal. (premise)
2. John is lying. (premise)
3. Sally is lying. (premise)
4. Therefore, John is lying and Sally is lying (from 2, 3, using And-Intro)
5. Therefore, Frank is a criminal (from 1, 4, using IF-THEN)

We now have a derivation of the conclusion from the premises. The derivation is finite (not infinite); it begins with the premises and terminates with the conclusion; and, each proposition in the derivation is either a premise, an assumption, or the result of using the rules of inference.

Definition 5.1.3: derivation

A derivation of a proposition C from a set of propositions $P_1, P_2, \dots, P_n$ is (1) a finite sequence of propositions where (2) the last proposition in the sequence is C and (3) each proposition is either a premise, an assumption, or is the result of using the rules for writing new propositions ("rules of inference" or "rules of reason").

While several questions still remain about derivation (e.g., what are the rules that allow you to write new propositions in the derivation?), we have articulated a second sense in which a conclusion can be said to follow from a set of premises. This involves, not the concept of *semantic entailment* but the concept of *derivation*. A conclusion "follows from" a set of premises if and only if there is a derivation of the conclusion from the premises. With this intuitive idea of a derivation in mind, we can now turn to the task of defining the notion of a derivation for the language of propositional logic.¹

5.2 SYNTACTIC ENTAILMENT

In the above example, we reasoned from "If John is lying and Sally is lying, then Frank is a criminal" and "John is lying and Sally is lying" to "Frank is a criminal". We reasoned from two earlier propositions to a new proposition. It was claimed that what allowed us to do this was an inference rule. But what are the inference rules? In **PL** the inference rules are the rules that permit us to write down new wffs from earlier wffs. The set of inference rules are called the deductive apparatus.

Definition 5.2.1: deductive apparatus

A deductive apparatus for **PL** is a set of rules that express whether a ϕ can be written after a set of wffs Γ in a derivation. The deductive apparatus for **PL** is hereafter abbreviated as **PD**.

In the previous section, our deductive apparatus for English consisted of two rules: AND-INTRO and the IF-THEN Rule. The deductive apparatus for **PL** are rules for writing wffs in a derivation (or proof). With that in mind, let's define a derivation (or proof) in **PL**.

Definition 5.2.2: derivation of ϕ in **PL**

A derivation in **PL** of ϕ is a finite (not infinite and not empty) string of formulas from a set Γ of **PL** wffs where (i) the last formula in the string is ϕ and (ii) each formula is either a member of Γ , an assumption, or is the result of using the deductive apparatus.

To understand the notion of a derivation, take our earlier example of an argument.

If Γ is empty and $\Gamma \vdash \phi$, then ϕ is a *theorem* and the derivation is a *proof* of ϕ .

1. If John is lying and Sally is lying, then Frank is a criminal. (premise)
2. John is lying. (premise)
3. Sally is lying. (premise)
4. Therefore, John is lying and Sally is lying (from 2, 3)
5. Therefore, Frank is a criminal (from 1, 4)

This argument is a finite sequence of propositions. The last proposition in this sequence is "Frank is a criminal". In addition, each proposition is either a premise, an assumption, or the result of using the rules of inference. For instance, the proposition "John is lying" is a premise and "John is lying and Sally is lying" is the result of using the rule of inference. As we will see, a derivation of ϕ from a set of wffs Γ is the same thing.

There is a finite sequence of wffs where the last wff is ϕ and each wff is either a member of Γ , an assumption, or the result of using the deductive apparatus.

With the notion of a deductive apparatus and a derivation in hand, we can now define the notion of a syntactic consequence (or entailment) in **PL**.

Definition 5.2.3: syntactic consequence / entailment

A formula ϕ is a syntactic consequence in **PD** of a set Γ of **PL** wffs if and only if there is a derivation in **PD** of ϕ from Γ . To express that ϕ is a syntactic consequence of Γ , we write $\Gamma \vdash \phi$.

If there is a derivation of ϕ from Γ , then ϕ is a syntactic consequence of Γ . That is, $\Gamma \vdash \phi$. In contrast, if there is no derivation of ϕ from Γ , then ϕ is not a syntactic consequence of Γ . That is, $\Gamma \not\vdash \phi$.

With these definitions in hand, we can now turn to the task of (1) learning how to set up a derivation, (2) developing a deductive apparatus for **PL**, and then (3) learning how to use the deductive apparatus to solve derivation.

5.3 SETUP

Suppose that it is claimed that $P \wedge R, P \rightarrow Z \vdash Z \vee Q$. This set of symbols says that $Z \vee Q$ is a syntactic consequence of $P \wedge R, P \rightarrow Z$. Intuitively, you can think of this as saying that $Z \vee Q$ follows from $P \wedge R, P \rightarrow Z$. And, you can think of $P \wedge R$ and $P \rightarrow Z$ as premises, while $Z \vee Q$ is the conclusion. The $\vdash$ symbol is called the “single turnstile” and is used to indicate that the proposition on the right side of the turnstile is a syntactic consequence of the propositions on the left side of the turnstile.

If $P \wedge R, P \rightarrow Z \vdash Z \vee Q$ is the case, then there is a derivation of $Z \vee Q$ from $P \wedge R, P \rightarrow Z$. How do we show that there is such a derivation? Since a derivation is nothing more than a finite sequence of wffs, technically, we would only need to provide a right sequence of wffs starting with $P \wedge R, P \rightarrow Z$ and ending in $Z \vee Q$. Such a sequence might look as follows:

$$P \wedge R, P \rightarrow Z, P, Z, Z \vee Q$$

So long as the each of the wffs in the sequence is $P \wedge R, P \rightarrow Z$, an assumption, or the result of using the deductive apparatus, and the sequence terminates in $Z \vee Q$, then we have our desired derivation. But, this is not a

very visually appealing way of providing a derivation and, if various rules of inference from our deductive apparatus allow us to take steps forward in the derivation (write new wffs), we will wonder which rule of inference is being used at various stages of the derivation. To address both of these issues, let's use a more systematic way of creating a derivation.

First, suppose we have a claim that $\Gamma \vdash \phi$. Our derivation will make use of three columns. The first column is for numbering wffs in the derivation, the second is for writing wffs in the derivation, and the third is for writing down what justifies the wffs in the middle column being there. The first column is called the "line number", the second column is called the "derivation body", and the third column is called the "justification".

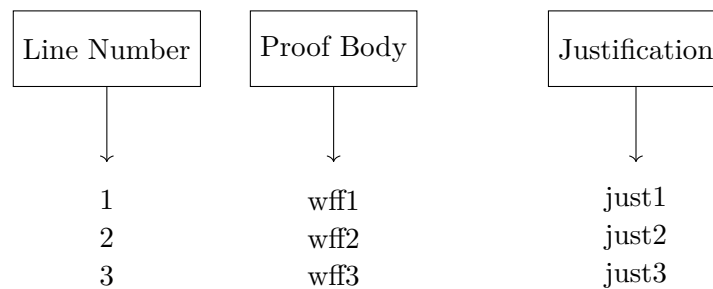


Figure 5.1: Three columns of a derivation

Let's illustrate this with our earlier example involving $P \wedge R, P \rightarrow Z \vdash Z \vee Q$. We will set up a derivation of this entailment as follows:

- | | | |
|---|-------------------|---------------|
| 1 | $P \wedge R$ | P |
| 2 | $P \rightarrow Z$ | P, $Z \vee Q$ |

Notice that there are three columns. The leftmost column numbers each row (or wff) in the derivation. The middle column $P \wedge R$ and $P \rightarrow Z$ are wffs that are members of Γ or the wffs that are said to entail Z . This is the core part of the derivation. The rightmost column is the justification column. The justification column justifies the presence of the wff at that line. Intuitively, $P \wedge R$ and $P \rightarrow Z$ can be thought of as the premises of the argument. More precisely, they are simply members of a set of wffs Γ that are claimed to entail a wff ϕ . To justify the presence of these wffs, we will write 'P' in the justification column to indicate that these wffs are "premises" or "provided".

It is sometimes customary to write the wff that the proof is supposed to derive (the conclusion) in the justification column. In our example, $Z \vee Q$ is the wff that is supposed to be derived. So, we write $Z \vee Q$ in the

justification column. This is a matter of taste, but is often helpful when trying to construct long, complex proofs that have multiple subproofs contained within them (more on this later).

Exercise 5.49

Set up the following proofs:

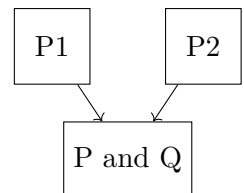
1. $\neg P, \neg P \rightarrow Z \vdash Z$
2. $P \wedge Q, Q \wedge R \vdash P \wedge R$
3. $P, Q, R \vdash (P \wedge Q) \wedge R$
4. $(\neg P \wedge \neg Q) \wedge \neg R \vdash \neg Q \wedge \neg R$
5. $\neg \neg P \vdash P$
6. $P, P \leftrightarrow Q \vdash Q$
7. $P \rightarrow Q, Q \rightarrow Z, R \wedge \neg \neg T \vdash P \rightarrow (Z \wedge R)$

5.4 THE INTELIM DERIVATION RULES

With the concepts underlying proofs along with a procedure for setting up proofs in place, the next step is to develop the deductive apparatus (**PD**). The particular type of deductive apparatus developed here is known as a system of “natural deduction” as the particular rules are akin to certain rules of inference (or reason) people use in everyday arguments. The particular rules of **PD** will be called the “derivation rules”. There are two main types of derivation rules: **introduction rules** and **elimination rules**.

An introduction rule is a rule that introduces a wff of a certain type into the proof. For example, suppose we wished to develop an rule that permits you to introduce an "and" proposition into an argument. For example, Tek and Liz may agree that whenever there are two propositions already in an argument, one is permitted to reason to a complex proposition that connects both of them with the word "and". Using this rule would allow you to reason from "snow is white" and "grass is green" to the proposition "snow is white and grass is green." But, it is important to note that this rule does not only work with the particular propositions "snow is white" and "grass is green". It is instead a rule that works with any two propositions. The rule is a template for reasoning.

Since such a rule always introduces an "and" proposition into an argument, it would be an example of an *introduction rule*. They might call it "and-introduction" to distinguish it from other types of introduction rules they might formulate.



5.4.1 Conjunction Introduction and Elimination

Before introducing our first rule, a quick note on how all of our rules will be presented. The approach of this text is to (1) define the rule, (2) provide a schematic use of the rule, and then (3) provide a minimal working example (MWE) of the rule. Once the rule has been presented, then we will discuss the rule in more detail by looking at various examples of the rule in action.

Let's introduce the first introduction rule of **PD**.

$\wedge I$

5.4.1.1 Conjunction Introduction

Definition 5.4.1: Conjunction Introduction ($\wedge I$)

$\phi, \psi \vdash \phi \wedge \psi$ or $\phi, \psi \vdash \psi \wedge \phi$. From ϕ and ψ , the wff $\phi \wedge \psi$ can be derived. Similarly, from ϕ and ψ , the wff $\psi \wedge \phi$ can be derived.

Here is a schematic use of the rule:

- | | | |
|---|--------------------|-----------------|
| 1 | ϕ | P |
| 2 | ψ | P |
| 3 | $\phi \wedge \psi$ | $\wedge I$ 1, 2 |

Conjunction introduction states that from two wffs ϕ and ψ , the conjunction of these wffs $\phi \wedge \psi$ can be derived on a new line in the proof. It is important to note that ϕ and ψ are variables for propositions and so ϕ and ψ can be any PL-wff. When conjunction introduction is used on two wffs to derive a conjunction, the line number of each wff is cited along with the abbreviation for conjunction introduction. Let's look at an MWE of conjunction introduction.

- | | | |
|---|--------------|-----------------|
| 1 | P | P |
| 2 | Q | P |
| 3 | $P \wedge Q$ | $\wedge I$ 1, 2 |

Notice that in the MWE of conjunction introduction, there are two wffs in the proof P and Q . The third line is the result of using conjunction introduction on P and Q . Since conjunction introduction requires the use of prior wffs in the proof, when line 3 is justified, we include the line numbers of the wffs used in the use of conjunction introduction. Since, as mentioned, lines 1 and 2 are used, we write "1, 2" in the justification

column next to the abbreviation for conjunction introduction, which is $\wedge I$.

Now that we have provided a definition of conjunction introduction, a schematic use of conjunction introduction, and an MWE, let's consider a few examples. To start, suppose we wished to show that $P, R, Z \vdash P \wedge Z$.

1	P	P
2	R	P
3	Z	P
4	$P \wedge Z$	$\wedge I$ 1, 3

The proof is setup as follows: the first three lines are justified as "P" as these wffs occur before the single turnstile. The fourth line is the result of using the conjunction introduction rule. The conjunction introduction rule is used to derive $P \wedge Z$ from P (line 1) and Z (line 3). Since conjunction introduction requires the use of prior wffs in the proof, it is helpful to include the lines of the wffs used in the use of conjunction introduction. Since, as mentioned, lines 1 and 3 are used, we write "1, 3" in the justification column next to the abbreviation for conjunction introduction.

Let's consider another example. Suppose we wished to show that $P, R, W \vdash (P \wedge R) \wedge W$. To begin, we start by setting up the proof. To do this, we write P , R , and W in the proof body column and write "P" in the justification column.

1	P	P
2	R	P
3	W	P

Once the proof is set up, we can begin to fill in the proof body and justification columns. Note that the goal of the proof is to derive $(P \wedge R) \wedge W$ rather than $P \wedge (R \wedge W)$. This means that we will want to derive $P \wedge R$ first. This is possible with conjunction introduction given that we have P at line 1 and R at line 2.

1	P	P
2	R	P
3	W	P
4	$\boxed{P \wedge R}$	$\wedge I$ 1, 2

Next, note that $P \wedge R$ is a wff. Since it is a wff, we can use conjunction introduction on it with another wff in the proof. Our goal is to derive

the conclusion in the proof and so we can use conjunction introduction on $P \wedge R$ (line 4) and W (line 3) to derive $(P \wedge R) \wedge W$ at line 5:

1	P	P
2	R	P
3	W	P
4	$P \wedge R$	$\wedge I$ 1, 2
5	$\boxed{(P \wedge R) \wedge W}$	$\wedge I$ 4, 3

Let's consider one final example. Suppose we wished to show that $P \rightarrow Q, \neg R, S \vdash ((P \rightarrow Q) \wedge \neg R) \wedge \neg R$. To begin, we start by setting up the proof.

1	$P \rightarrow Q$	P
2	$\neg R$	P
3	S	P

Next, we can use conjunction introduction on $P \rightarrow Q$ (line 1) and $\neg R$ (line 2) to derive $(P \rightarrow Q) \wedge \neg R$ at line 4.

1	$P \rightarrow Q$	P
2	$\neg R$	P
3	S	P
4	$\boxed{(P \rightarrow Q) \wedge \neg R}$	$\wedge I$ 1, 2

Looking at our conclusion, notice that there is another $\neg R$ in the wff. To derive the conclusion, we will use $\neg R$ on line 2 again. That is, we will use line 2 and our newly formed conjunction $(P \rightarrow Q) \wedge \neg R$ (line 4).

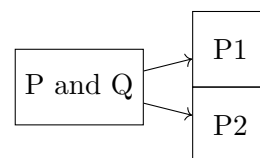
1	$P \rightarrow Q$	P
2	$\neg R$	P
3	S	P
4	$(P \rightarrow Q) \wedge \neg R$	$\wedge I$ 1, 2
5	$\boxed{((P \rightarrow Q) \wedge \neg R) \wedge \neg R}$	$\wedge I$ 4, 2

Notice that in the above example, we used the conjunction introduction on line 2 and 4. This is permitted as conjunction introduction allows us to use any two wffs in the proof.

5.4.1.2 Conjunction Elimination

Recall that the intuitive idea behind introduction derivation rules is that they introduce wffs of a certain type. With respect to English, we mentioned the hypothetical "and-introduction" rule that allows you to reason from "snow is white" and "grass is green" to "snow is white and grass is green". With respect to **PL**, we introduced the conjunction introduction rule allows you to reason from any two wffs to a conjunction of those two wffs. In addition to introduction rules, there are elimination rules. Before discussing elimination rules in **PL**, let's consider the intuitive idea behind elimination rules. Suppose Tek and Liz agree that whenever there is a proposition connected by an "and", you can reason to either side of the "and" proposition. For example, suppose there is a proposition of the form "snow is white and grass is green", then the rule would permit you to reason to "snow is white" and "grass is green". Since such a rule always reasons from an "and" proposition an argument, they decide to call this rule "and-elimination".

Let's introduce the first elimination rule of **PD**.



Definition 5.4.2: Conjunction Elimination ($\wedge E$)

$\phi \wedge \psi \vdash \phi$ or $\phi \wedge \psi \vdash \psi$. From the conjunction $\phi \wedge \psi$, the wff ϕ can be derived. Similarly, from the conjunction $\phi \wedge \psi$, the wff ψ can be derived.

Here is a schematic use of the rule:

1	$\phi \wedge \psi$	P	1	$\phi \wedge \psi$	P
2	ϕ	$\wedge E$ 1	2	ψ	$\wedge E$ 1

With this statement of conjunction elimination, let's provide a MWE of conjunction elimination. Consider $P \wedge Z \vdash Z$.

1	$P \wedge Z$	P, Z
2	P	$\wedge E$ 1
3	Z	$\wedge E$ 1

The proof is setup as follows: the first line is justified as "P" as this wff occurs before the single turnstile. The second line is the result of using the conjunction elimination rule. The conjunction elimination rule is used to derive P from $P \wedge Z$ (line 1). Since conjunction elimination requires the use of a prior wff in the proof, we include the line of the wff used in the use

of conjunction elimination. Since, as mentioned, line 1 is used, we write "1" in the justification column next to the abbreviation for conjunction elimination. Note that in addition to permitting us to derive the left wff P , we can also derive the right wff Z from $P \wedge Z$.

Let's consider another example and we will consider it in a step-by-step manner. Consider the following: $P \wedge (R \wedge W) \vdash W$. As the setup only involves listing the wffs to the left of the turnstile and numbering them, we can set up the proof as follows:

$$\begin{array}{ll} 1 & P \wedge (R \wedge W) \quad \text{P, W} \end{array}$$

Once the proof is set up, we can begin to fill in the proof body and justification columns. The temptation of students beginning logic is to directly derive W on line 2. Unfortunately, this would not be correct. Conjunction elimination allows for deriving either conjunct from a conjunction. So, if our conjunction is $\phi \wedge \psi$, using conjunction introduction on this wff allows for deriving ϕ or ψ . But what are the conjuncts of $P \wedge (R \wedge W)$? These are P and $R \wedge W$. Therefore, using conjunction elimination on this rule only permits deriving P or $R \wedge W$. Let's go ahead and derive $R \wedge W$.

$$\begin{array}{ll} 1 & P \wedge (R \wedge W) \quad \text{P} \\ 2 & \boxed{R \wedge W} \quad \wedge E \ 1 \\ 3 & W \quad \wedge E \ 2 \end{array}$$

Now that $R \wedge W$ has been derived, we can derive W from $R \wedge W$. This is because the conjunction $R \wedge W$ has two conjuncts (R and W) and conjunction elimination allows for deriving either wff.

$$\begin{array}{ll} 1 & P \wedge (R \wedge W) \quad \text{P} \\ 2 & R \wedge W \quad \wedge E \ 1 \\ 3 & \boxed{W} \quad \wedge E \ 2 \end{array}$$

Thus far, we have considered proofs that make use of conjunction introduction and conjunction elimination in isolation. Next, let's consider an example that makes use of *both* conjunction introduction and conjunction elimination. In this example, suppose we aimed to provide the proof for the following: $P \wedge Q, R \wedge S \vdash Q \wedge S$.

1	$P \wedge Q$	P
2	$R \wedge S$	P
3	Q	$\wedge E$ 1
4	S	$\wedge E$ 2
5	$Q \wedge S$	$\wedge I$ 3, 4

Once the proof was setup, we derived Q from $P \wedge Q$ using conjunction elimination and S from $R \wedge S$ using conjunction elimination. If you derived P and R using conjunction elimination, then this would be an acceptable use of the rule (and totally correct) but you would have added some extra steps to your proof. With Q and S isolated on lines 3 and 4, we used conjunction introduction to create the conjunction of these two wffs.

Exercise 5.50

Provide proofs for the following:

1. $P, Q, R, S \vdash P \wedge S$
2. $A, B, C, \neg\neg D \vdash (A \wedge B) \wedge (C \wedge \neg\neg D)$
3. $P \wedge (R \wedge M) \vdash R$
4. $P \wedge (R \wedge M) \vdash M$
5. $P \wedge R, \neg Z \wedge \neg W \vdash P \wedge \neg W$
6. $P \wedge (\neg R \wedge \neg W), L, R \wedge Z \vdash \neg R \wedge (L \wedge Z)$
7. $(M \wedge N) \wedge W, (P \rightarrow Q) \wedge Y \vdash (P \rightarrow Q) \wedge W$
8. $F \wedge (G \wedge (P \vee Q)), (L \wedge M) \vdash L \wedge (P \vee Q)$

5.4.2 Conditional Elimination and Introduction

In earlier sections, we introduced the conjunction introduction and conjunction elimination rules. These rules allow you to introduce conjunctions into the proof and derive wffs from conjunctions in the proof, respectively. In addition, we examined the idea of assumptions. In this section, the conditional introduction and conditional elimination rules. As introduction and elimination rules, conditional introduction will allow you to introduce conditionals into the proof and conditional elimination will allow you to derive wffs from conditionals in the proof.

5.4.2.1 Conditional Elimination

$\rightarrow E$

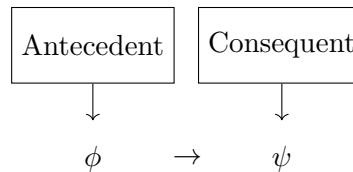
Let's consider Liz and Tek one more time. Suppose Liz and Tek want to add another rule to their set of "rules of inference". They have already agreed upon "and-introduction" and "and"-elimination, but they would

like to also reason with propositions of the form "if P, then Q". They add the following rule: whenever you have an "if P, then Q" proposition, where P and Q are propositions, *and* you have P by itself, the you can reason to Q. They decide to call this rule "if-then elimination". Here is an example:

1. If it is raining, then the ground is wet. (premise)
2. It is raining. (premise)
3. Therefore, the ground is wet. (from 1, 2, using if-then elimination)

In the above example, there is an "if-then" proposition at line 1, the proposition that is between the "if" and the "then" at line 2, and the "if-then elimination" rule allows to write the proposition to the right of the "then" at line 3.

Let's create a similar rule for our system of natural deduction. Remember that it is proposed that *some* propositions of the form "if P, then Q" (material conditionals) can be expressed as $P \rightarrow Q$. In addition, recall that the conditional has two parts: the antecedent and the consequent. The antecedent is the proposition to the left of the arrow and the consequent is the proposition to the right of the arrow. In $P \rightarrow Q$, the wff P is the antecedent and the wff Q is the consequent.



With these considerations in mind, let's add the conditional elimination rule to our system of natural deduction. Conditional elimination is one of the most famous rules of inference in logic. It has several other names, e.g., *modus ponens* or "affirming the antecedent", but we will simply call it "conditional elimination".

Definition 5.4.3: Conditional Elimination ($\rightarrow E$)

$\phi \rightarrow \psi, \phi \vdash \psi$. From $\phi \rightarrow \psi$ and ϕ , the wff ψ can be derived.

The underlying idea behind this rule is that if you have (1) a conditional and (2) the antecedent of that conditional on another line, then you can (3) derive the consequent of that conditional. Here is a schematic use of the rule:

1	$\phi \rightarrow \psi$	P
2	ϕ	P
3	ψ	$\rightarrow E$ 1, 2

Again, notice that conditional elimination involves two wffs: a conditional $\phi \rightarrow \psi$ and the antecedent ϕ of that conditional. The conditional $\phi \rightarrow \psi$ and its antecedent ϕ are used to derive the consequent ψ . In the schematic use of the rule, the line numbers of the conditional and the antecedent are cited in the justification column. Here is a minimal working example of conditional elimination:

1	$P \rightarrow Q$	P
2	P	P
3	Q	$\rightarrow E$ 1, 2

Again, notice that conditional elimination involves two wffs: a conditional $P \rightarrow Q$ and the antecedent P of that conditional. The conditional $P \rightarrow Q$ and its P are used to derive the consequent Q . In the MWE of conditional elimination, the line numbers of the conditional and the antecedent are cited in the justification column along with $\rightarrow E$, which is the abbreviation for conditional elimination.

There are two mistakes individuals make when using conditional elimination. The first mistake is to use conditional elimination on the conditional without having the antecedent of that conditional in the proof. For example, consider the following proof:

Mistake #1 using $\rightarrow E$

1	$\phi \rightarrow \psi$	P
2	ψ	$\rightarrow E$ 1

This use of conditional elimination is not correct as the antecedent ϕ of the conditional $\phi \rightarrow \psi$ is not in the proof. In other words, whenever you use conditional elimination, you must have a conditional $\phi \rightarrow \psi$ and the antecedent ϕ of that conditional in the proof. It is not likely that this is a mistake concerning how people reason as few individuals are likely to reason "if I buy a lottery ticket, then I will win the lottery. Therefore, I will win the lottery." Most individuals will assume that they need to buy a lottery ticket before they can win the lottery. Rather this type of mistake is largely due to not paying attention to the fact that both wffs need to be cited in the justification column.

The second mistake is to use conditional elimination using a conditional and the consequent of that conditional to derive the antecedent of that con-

Mistake #2 using $\rightarrow E$

ditional. In short, this mistake is to use $\phi \rightarrow \psi, \psi$ to derive ϕ . Reasoning in this way is referred to as "the fallacy of affirming the consequent". Let's consider an example of this mistake in an argument before considering how the mistake would be made in a proof. Suppose you have the following argument: "If I buy a lottery ticket, then I will win the lottery. I won the lottery. Therefore, I bought a lottery ticket." This argument is (1) not a correct use of conditional elimination as the antecedent of the conditional is not in the argument (2) not valid as it is possible to win the lottery without buying a lottery ticket (e.g., a relative may have bought you a ticket for your birthday). Let's consider this same mistake in the context of a proof:

1	$\phi \rightarrow \psi$	P
2	ψ	P
3	ϕ	$\rightarrow E$ 1, 2

Notice that in the above proof, there is the conditional at line 1, but rather than having the antecedent ϕ of the conditional on a line, there is the consequent ψ of the conditional at line 2. Conditional elimination is used incorrectly at line 3 to derive the antecedent ϕ of the conditional. This is not a correct use of conditional elimination since conditional elimination requires a conditional and the antecedent of that conditional in the proof.

Now that we have defined conditional elimination, considered a minimal working example, and considered some common mistakes, let's consider some examples where conditional elimination is used correctly. First, consider $Z \rightarrow R, Z \wedge P \vdash R$. As with prior proofs, setting up the proof involves writing $Z \rightarrow R, Z \wedge P$ at lines 1 and 2, respectively, and then trying to derive the conclusion R in the proof. Let's construct this proof in a step-by-step fashion.

1	$Z \rightarrow R$	P
2	$Z \wedge P$	P, R

Notice that we cannot use conditional elimination immediately on line 1. This is because conditional elimination requires the antecedent of the conditional. In this case, the antecedent of the conditional is Z and Z is not on its own line in the proof. However, we can use conjunction elimination on line 2 to derive Z from $Z \wedge P$.

1	$Z \rightarrow R$	P
2	$Z \wedge P$	P, R
3	$\boxed{Z}$	$\wedge E$ 2

Once we have derived Z , we now have the antecedent of the conditional at line 1 and can use conditional elimination (using lines 1 and 3).

1	$Z \rightarrow R$	P
2	$Z \wedge P$	P, R
3	Z	$\wedge E$ 2
4	$\boxed{R}$	$\rightarrow E$ 1, 3

Let's consider another example. Suppose we wished to show that $(P \rightarrow Q) \rightarrow (S \rightarrow M), P \rightarrow Q, S \vdash M$. As with prior proofs, setting up the proof involves writing $(P \rightarrow Q) \rightarrow (S \rightarrow M), P \rightarrow Q, S$ at lines 1, 2, and 3, respectively, and then trying to derive the conclusion M in the proof.

1	$(P \rightarrow Q) \rightarrow (S \rightarrow M)$	P
2	$P \rightarrow Q$	P
3	S	P, M
4	$S \rightarrow M$	$\rightarrow E$ 1, 2
5	M	$\rightarrow E$ 4, 3

In the above example, the conditional at line 1 consists of two parts: (1) the antecedent $P \rightarrow Q$ and (2) the consequent $S \rightarrow M$. Since we have the antecedent of that conditional at line 2, we can use conditional elimination to derive the consequent $S \rightarrow M$. With $S \rightarrow M$ at line 4, notice that we have the antecedent of this conditional at line 3. We can thus use conditional elimination to derive M at line 5.

Exercise 5.51

Provide proofs for the following (each proof will involve a use of conditional elimination to complete the proof):

1. $P \rightarrow Q, P \vdash Q$
2. $A \rightarrow M, \neg A, A \vdash M$
3. $(A \wedge B) \rightarrow C, A, B \vdash C$
4. $(\neg P \wedge Q) \rightarrow (\neg S \wedge T), \neg P, Q \vdash T$
5. $R \rightarrow Z, Z \rightarrow W, R \wedge M \vdash W \wedge M$

5.4.2.2 Assumptions and Subproofs

Before introducing conditional introduction, we need to introduce the idea of assumptions and subproofs. An assumption (abbreviated as 'A') is a proposition that is supposed or taken as a hypothesis. It is not a propo-

sition that is asserted to be the case or said to directly follow from other propositions. Let's consider an example. Suppose Liz and Tek are newly married and have been living in a friend's basement to save money. After several months, they think they have enough to buy their first house. They are not a rich couple and so they are worried about this purchase. There are many propositions they take to be true about this upcoming purchase. They both believe that "if they buy a house, they will have to pay taxes" and that "if they buy a house, they will do repairs on the home (as the house is not in good condition)". Let's number each of these premises and mark each with a "P" to indicate that they are premises (or propositions that Liz and Tek take to be true):

- 1 If we buy a house, we'll pay taxes. P
- 2 If we buy a house, we'll do repairs P

Now Liz and Tek are not sure if they want to buy a house. So, they say the following: "Assume we did buy a house." To make it clear that they are not saying they bought a house or that buying a house follows from (1) and (2), let's indent from the premises, write out their assumption, and justify this proposition in their line of reasoning with an "A" (for assumption).

- 1 If we buy a house, we'll pay taxes. P
- 2 If we buy a house, we'll do repairs P
- 3 | Suppose we buy a house Assumption
- 4 | .
- 5 | .

With this assumption in place (with their dream of buying in front of them), Liz and Tek might do a little reasoning. For example, Liz might say that *assuming* they buy a house (line 3) and assuming (1), it would follow they would have to pay taxes. So, Liz might develop their proof as follows:

- 1 If we buy a house, we'll pay taxes. P
- 2 If we buy a house, we'll pay for the repairs P
- 3 | Suppose we buy a house Assumption
- 4 | We'll pay taxes $\rightarrow E, 1-3$

Tek might add, "ah, yes, still under this assumption that we buy a house, not only will we have to pay taxes but we'll also do our own repairs."

1	If we buy a house, we'll pay taxes.	P
2	If we buy a house, we'll do repairs.	P
3	Suppose we buy a house	Assumption
4	We'll pay taxes	$\rightarrow E, 1-3$
5	We'll do repairs	$\rightarrow E, 2-3$

"Yes," Liz says, "so we'll not only pay taxes but we'll also do repairs."

1	If we buy a house, we'll pay taxes.	P
2	If we buy a house, we'll do repairs.	P
3	Suppose we buy a house	Assumption
4	We'll pay taxes	$\rightarrow E, 1-3$
5	We'll do repairs	$\rightarrow E, 2-3$
6	We'll pay taxes and do repairs	$\wedge E, 4,5$

A few points to make about our example. First, when Liz and Tek made an assumption, then indented from the point they were in the proof and drew a vertical line to indicate that the reasoning they were about to do was *under* the assumption. Second, they justified their assumption with "A" for assumption. In doing this, they effectively started a new proof within their proof. This is called a *subproof*. It is a proof that is nested within another proof.

1	Main Line	
2	Assumption	Start of Subproof
3	.	
4	.	
5	.	

In making the assumption and then reasoning under the assumption, Liz and Tek are not saying that she and Tek will have to pay taxes, nor are they saying that they have to do repairs. Rather, they are saying that given the existing premises, on the assumption they buy a house, it would follow that they we'll pay taxes and do repairs.

Let's pivot from Liz and Tek to **PL** proofs. In **PL**, you can make assumptions and derive wffs from those assumptions. Let's consider an example where a proof begins with *S* and then makes the assumption *B*.

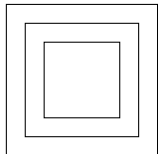
1	S	P
2	B	A
3	$S \wedge B$	$\wedge I$ 1, 2
4	.	
5	.	
6	.	

Similar to the Liz and Tek example, once B is assumed, it is possible to derive wffs from that assumption. In this case, the wff $S \wedge B$ is derived from the assumption B using line 1. This is done by using conjunction introduction on S and B .

How many assumptions can you make? You can make as many assumptions as you want. Let's consider an example where a proof begins with Q and then makes the assumption S and then makes the assumption W .

1	Q	P
2	S	A
3	W	A
4	.	
5	.	
6	.	

You can think of assumptions as containers. You can have containers within containers within containers. The outer container is the main line of the proof. The inner containers are the subproofs.



Some proofs will contain a single subproof (or subproofs within subproofs), but other proofs will contain multiple distinct subproofs. Let's consider our earlier example involving Tek and Liz. Suppose Tek and Liz are still considering whether or not to buy a house. They have considered what follows under the assumption that they do buy the house, but now they wish to consider what follows under the assumption that they do not buy the house. Their reasoning might look like the following:

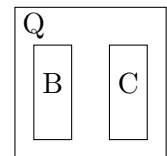
1	If we buy a house, we will have to pay taxes	prem
2	If we buy a house, we will have to pay for repairs	prem
3	Suppose we buy a house	Assumption
4	We will have to pay taxes	if-then-elim, 1,3
5	We will have to pay for repairs	if-then-elim, 2,3
6	Suppose we do not buy a house	Assumption
7	.	
8	.	

In this scenario, Tek and Liz are not reasoning "let's assume we buy a house and under this assumption, let's now assume we do not buy a house". That would be a rather confusing conversation! Rather, they are reasoning "let's assume we buy a house" and then drawing out the consequences from this assumption. Then, they are separately considering what follows under the assumption that they do not buy a house. In other words, they are making two distinct assumptions.

Let's illustrate how this would look in a **PL** proof. Suppose we have the following proof:

1	Q	P
2	B	A
3	$Q \wedge B$	$\wedge I$ 1, 2
4	.	
5	.	
6	C	A
7	$Q \wedge C$	$\wedge I$ 1, 2
8	.	
9	.	

Thinking of nests: B is nested within Q and C is nested within Q , but B and C are not nested within each other.



In this example, the proof begins with Q as a premise. Next, the assumption B is made. From this assumption, the wff $Q \wedge B$ is derived using conjunction introduction. Next, the assumption C is made *but not under the assumption B* . Rather C is a separate assumption.

Two final points about assumptions and subproofs before we return to the discussion of conditional introduction. First, one common mistake is to derive a wff from subproof that only follows in that part of the subproof. Let's consider an extreme illustration.

1	If I buy a lottery ticket, I will win the lottery.	prem
2	Assume I will buy a lottery ticket.	Assumption
3	I will win the lottery.	if-then-elim, 1,2

Suppose Tek puts forward the above argument. At line 1, Tek does not assert that he will win the lottery, only that he has exceptional luck or power over the lottery system. He asserts that *if* he buys a ticket, then he will win the lottery. From this claim, it does not directly follow that he will win the lottery. At line 2, he makes the assumption that he will buy a ticket. In this case, Tek is not asserting he will buy a ticket. Instead, he

is considering what would follow if he bought a ticket. Insofar as Tek is making an assumption, he begins to reason in a subproof. At line 3, Tek uses conditional elimination to derive the conclusion that he will win the lottery. This is not a correct use of conditional elimination. The reason is that while the conclusion "I will win the lottery" only follows under the assumption "I will buy a lottery ticket". In writing "I will win the lottery" at line 3, Tek is asserting that he will win the lottery follows from the premise "if I buy a lottery ticket, I will win the lottery". This is not the case. What went wrong is that Tek derived a wff from a subproof that only follows within that part of the proof. That is, it only follows within the subproof that begins with the assumption "I will buy a lottery ticket".

To ensure we don't create this mistake again, let's introduce some terminology. When making an assumption, we indent and draw vertical lines to indicate the *scope* of the assumption. The scope of the assumption is the part of the proof that depends upon the assumption. When an assumption is made, any further reasoning involving the assumption is done within the scope of the assumption. So, if we wanted to derive a wff using the assumption "I will buy a lottery ticket" in the example above, this reasoning would be done within the scope of the assumption. Our earlier example can be rewritten as follows:

1	If I buy a lottery ticket, I will win the lottery.	prem
2	Assume I will buy a lottery ticket.	Assumption
3	I will win the lottery.	if-then-elim, 1,2

Notice that in our new example, "I will win the lottery" at line 3 is written within the scope of the assumption. This is to signal that the reasoning at line 3 depends upon the assumption at line 2.

Let's consider another example. In this example, we'll start the proof with two premises, create an assumption, and then use our premises, the assumption, and the "if-then-elim" rule to derive some new propositions.

1	If we buy a house, we will have to pay taxes	prem
2	If we buy a house, we will have to pay for repairs	prem
3	Suppose we buy a house	Assumption
4	We will have to pay taxes	if-then-elim, 1,3
5	We will have to pay for repairs	if-then-elim, 2,3

Notice that in the above example, the proposition "we will have to pay for repairs" is derived within the scope of the assumption that begins at

line 3.

Second, you can use propositions that are outside of the subproof to derive wffs in the subproof. When Tek and Liz are considering whether to buy a house, they have two premises: "if we buy a house, we will have to pay taxes" and "if we buy a house, we will have to pay for repairs". When Liz makes the assumption "suppose we buy a house", this assumption is made in the context of these two premises already present in the proof. Liz and Tek are not considering what follows from the assumption "suppose we buy a house" in isolation from the premises. Rather, they are considering what follows from the assumption "suppose we buy a house" *and* the premises. In other words, they are considering what follows from the assumption "suppose we buy a house" *given* the premises.

Since that is the case, Liz and Tek can use the premises to derive wffs in the subproof. For example, they can use the premise "if we buy a house, we will have to pay taxes" to derive the wff "we will have to pay taxes" in the subproof.

Let's illustrate this with a **PL** proof. Consider the following proof:

1	Q	P
2	$\begin{array}{ l} B \end{array}$	A
3	$Q \wedge B$	$\wedge I$ 1, 2, OK!
4	B	R 3, NO!

Notice that in this example, the proof starts with Q and then B is assumed. In assuming B , the proof is effectively saying "given Q , assume B ". Given that the assumption B depends upon the premise Q , the proof can use Q to derive $Q \wedge B$ in the subproof.

5.4.2.3 Conditional Introduction

We have considered conditional elimination and engaged in a short discussion of assumptions and subproofs, let's now consider conditional introduction. Recall that introduction rules allow you to introduce wffs of a specific type into the proof. As such, conditional introduction will allow you to introduce conditionals into the proof. In other words, conditional introduction allows you to reason from a wff ϕ to a conditional $\phi \rightarrow \psi$.

Before considering conditional introduction, let's consider one question you may have had about assumptions and subproofs. What good are assumptions and subproofs if you are "stuck" within them? Alternatively,

is there anything that we can derive (or reason to) by using subproofs? The answer is yes. One way we can use subproofs is to derive conditionals. In other words, we can use subproofs to derive wffs of the form $\phi \rightarrow \psi$. Consider our earlier example where Liz and Tek are considering whether or not to buy a house.

1	If we buy a house, we will have to pay taxes	prem
2	If we buy a house, we will have to pay for repairs	prem
3	Suppose we buy a house	Assumption
4	We will have to pay taxes	if-then-elim, 1,3
5	We will have to pay for repairs	if-then-elim, 2,3
6	We will have to pay taxes and pay for repairs	and-intro, 4,5

Notice that Liz and Tek have reasoned to a conjunction within the subproof. They have reason that under the assumption they buy a house, then they will have to pay taxes and pay for repairs. They might express this reasoning as "if we buy a house, then we will have to pay taxes and pay for repairs". In other words, they have reasoned to an if-P, then-Q proposition. Let's call the rule that they used to reason to this conditional "if-then introduction".

1	If we buy a house, we will have to pay taxes	prem
2	If we buy a house, we will have to pay for repairs	prem
3	Suppose we buy a house	Assumption
4	We will have to pay taxes	if-then elim, 1,3
5	We will have to pay for repairs	if-then-elim, 2,3
6	We will have to pay taxes and pay for repairs	and-intro, 4,5
7	If we buy a house, then we will have to pay taxes and for repairs	if-then intro, 3-6

In justifying their reasoning, Liz and Tek would cite the lines of the subproof that began with the assumption "suppose we buy a house" and ended with the proposition they derived in the subproof. In other words, they would cite lines 3-6, along with their newly introduced if-then intro rule.

$\rightarrow I$

Let's turn to **PL** and consider how we might introduce conditionals into a proof.

Definition 5.4.4: Conditional Introduction ($\rightarrow I$)

If ψ is derived from the assumption ϕ , then $\phi \rightarrow \psi$ can be derived from the subproof started by ϕ .

Before providing a minimal working example, let's consider the general structure of this rule.

n	ϕ _____	A
	$\vdots$	
m	ψ	
m+1	$\phi \rightarrow \psi$	$\rightarrow I$ n-m

In the above example, a wff ϕ is assumed. This requires the use of an assumption (A) and indenting from the point in the proof when the assumption is made. In assuming ϕ , the wff ϕ is written in the proof body column and justified with "A" for assumption. Next, at some point in the proof in the same level of the subproof that begins with ϕ , a wff ψ is derived. That is, it is shown that from the assumption ϕ , the wff ψ is derived in the body of the proof. Once this is achieved, the conditional introduction rule is used. The wff $\phi \rightarrow \psi$ is derived from the subproof started by ϕ . This involves writing the conditional one level outside of subproof that began with ϕ . The wff $\phi \rightarrow \psi$ is justified by the rule $\rightarrow I$ and the lines of the subproof started by ϕ .

New students in logic often struggle with the use of conditional introduction, so it is worthwhile to consider several examples. First, consider the following entailment: $R \vdash Z \rightarrow R$.

1	R	R
2	Z _____	Z
3	$Z \wedge R$	$\wedge I$ 1, 2
4	R	$\wedge E$ 3
5	$Z \rightarrow R$	$\rightarrow I$ 2-4

Notice that the proof begins with the premise R . Next, the goal of the proof is the conditional $Z \rightarrow R$. In other words, we would like to *introduce* into the proof a conditional. Conditional introduction allows you to do this by assuming the antecedent of the conditional (the wff to the left of the rightarrow), deriving the consequent of the conditional (the wff to the right of the rightarrow), and the using conditional introduction. Since Z is the antecedent of the conditional $Z \rightarrow R$, we assume Z . Next, our goal

is to derive R . There is no direct way to derive R , but we can use $\wedge I$ on lines 1 and 2 to derive $Z \wedge R$ using conjunction introduction. Next, R is derived from $Z \wedge R$ using conjunction elimination. We now have R in the subproof at line 4. At this point, we have assumed Z , derived R in the subproof, and now can use conditional introduction to derive the conditional $Z \rightarrow R$ in the main body of the proof. In terms of citing the justification for conditional introduction, we cite the lines of the subproof started by Z (lines 2-4). In other words, we cite the lines of the entire subproof beginning with Z and ending with R .

Let's consider another example. Suppose we wished to show that $P \wedge R \vdash (Z \vee Q) \rightarrow P$. To begin, we start by setting up the proof. Once the proof is set up, we can begin to fill in the proof body and justification columns. In this proof, the key step concerns what assumption to make for the purpose of using conditional introduction. Since the goal is to derive the conditional $(Z \vee Q) \rightarrow P$, we need to assume the *antecedent* of the conditional. That is, the assumption should be $(Z \vee Q)$ rather than Z or P .

1	$P \wedge R$	P
2	$\mid Z \vee Q$	A

Once we assume $(Z \vee Q)$, the next step is to derive the consequent of the conditional. In this case, the consequent is P . Since P is a conjunct of line 1, we can use conjunction elimination to derive P from $P \wedge R$. Once we have derived the consequent of the conditional, we now have the $Z \vee Q$ as the assumption of the subproof and the P consequent of the subproof. In essence, we have shown that is on the assumption that $(Z \vee Q)$ is the case, then P is the case. At this point, we can use conditional introduction to derive the conditional $(Z \vee Q) \rightarrow P$.

3	$\mid P$	1 $\wedge E$
4	$(Z \vee Q) \rightarrow P$	2-3 $\rightarrow I$

Exercise 5.52

Provide proofs for the following (each proof will involve a use of conditional introduction to complete the proof):

1. $P \wedge R \vdash Z \rightarrow P$
2. $P \wedge R \vdash \neg\neg Z \rightarrow P$
3. $R \vdash R \rightarrow R$
4. $P \wedge R \vdash (Z \vee Q) \rightarrow (R \wedge P)$

5. $P, Q \vdash P \rightarrow Q$
6. $Q, R \wedge \neg M \vdash \neg L \rightarrow (S \rightarrow (\neg M \wedge Q))$
7. $\vdash A \rightarrow A$
8. $\vdash A \rightarrow (A \rightarrow A)$
9. $\vdash P \rightarrow (Q \rightarrow P)$

5.4.3 Reiteration

R

At this point, there are four derivation rules in our deductive apparatus: $\wedge I, \wedge E, \rightarrow E, \rightarrow I$. Each of these rules are either an introduction or elimination rule. In this section, we introduce a derivation rule that is neither an introduction nor an elimination rule. This rule is called "reiteration". Reiteration is a very simple derivation rule. It allows you to derive a wff from that wff. In other words, reiteration allows you to rewrite a wff in the proof.

Definition 5.4.5: Reiteration (R)

$\phi \vdash \phi$. From ϕ , the wff ϕ can be derived.

Here is the schematic use of reiteration:

- 1 ϕ P
- 2 ϕ R 1

Here is a MWE:

- 1 P P
- 2 P R 1

In the justification column, the abbreviation for the rule is written (this is *R*) and then the line of the wff to which the rule is applied is also written (if you apply reiteration to a wff P , then you would cite the line number of P).

You may wonder what is the point of introducing the reiteration rule into our deductive apparatus. The purpose of introducing the rule is that it is useful in that it simplifies (shortens) proofs. For example, suppose you were to construct a proof for $B \vdash B$. In looking at this entailment, it is pretty clear that B follows from B . However, the proof for B from B is not straightforward. One way you might solve it is by using conditional introduction.

1	B	P
2	$\left \frac{B}{} \right.$	A
3	$B \wedge B$	$\wedge I, 1$
4	$\left \frac{B}{} \right.$	$\wedge E, 1$
5	$B \rightarrow B$	$\rightarrow I\ 1-3$
6	B	$\rightarrow E\ 1-5$

But such a proof seems overly complex for something so simple. But suppose we were to include reiteration into our deductive apparatus. Then, we could provide the following proof:

1	B	P
2	$\boxed{B}$	R 1

Notice that the proof is much shorter and simpler. In fact, the proof is so simple that it is not clear that it is a proof.

There are other cases where reiteration allows for simpler proofs. For example, consider the following entailment: $R \rightarrow Z, R \wedge P \vdash Z \rightarrow R$ in [Proof 5.1](#) and [Proof 5.2](#). The addition of reiteration eliminates the need to use conjunction introduction and conjunction elimination to derive a wff from outside the scope of a subproof into the scope of the subproof. In short, reiteration thus shortens proofs by allowing you to rewrite wffs in the proof and also removes obvious steps in the proof with sacrificing the explicit nature of a proof.

1	R	P
2	$\left \frac{Z}{} \right.$	A
3	R	R 1
4	$Z \rightarrow R$	$\rightarrow I\ 2-3$

Proof 5.1

1	R	P
2	$\left \frac{Z}{} \right.$	A
3	$R \wedge Z$	$\wedge I\ 1, 2$
4	R	$\wedge E\ 3$
5	$Z \rightarrow R$	$\rightarrow I\ 2-4$

Proof 5.2

It is clear that reiteration shortens and simplifies proofs, but is reiteration a necessary addition to our deductive apparatus? Reiteration is not necessary. It is possible to construct a proof without using reiteration as was illustrated with both our proof of $B \vdash B$ and [Proof 5.1](#) and [Proof 5.2](#).

Exercise 5.53

Provide proofs for the following (proofs will involve some combination of the derivation rules introduced thus far):

1. $P \rightarrow Z, P \vdash P \wedge Z$
2. $(P \wedge Z) \rightarrow W, P, Z \vdash W$
3. $P \vdash R \rightarrow P$
4. $P, M \vdash (R \vee F) \rightarrow (P \wedge M)$
5. $(R \vee F) \rightarrow Z, M \wedge (R \vee F) \vdash M \wedge Z$
6. $R \rightarrow P, R \wedge L \vdash P \rightarrow (L \wedge R)$
7. $\vdash R \rightarrow (R \wedge R)$
8. $\vdash R \rightarrow [Z \rightarrow (M \rightarrow Z)]$

5.4.4 Negation Introduction and Elimination

We have discussed introduction and elimination rules for conjunctions and conditionals. Let's now consider the introduction and elimination rules for negation. These are negation introduction and elimination. Negation introduction and elimination are sometimes called "proofs by contradiction" and are a species of a more general type of argument known as *reductio ad absurdum*. Arguments by reductio ad absurdum work by assuming some proposition to be true and then showing that it is possible to reason from this assumption to a contradiction (ad absurdum), a false proposition, a proposition that is impossible or implausible, or some highly undesirable outcome, then one concludes that the assumption is rejected and the negation of the assumption is to be accepted.

Arguments of this type are extremely common in everyday conversation, political discourse, mathematics, and philosophy. For example, suppose you wanted to argue against the existence of God. One way to do this would be to assume that God exists and then consider the consequences of that assumption. If the consequences of that assumption lead to a contradiction, a false proposition, a proposition that is impossible or implausible, or some highly undesirable outcome, then the assumption that God exists is rejected and the negation of that assumption, that God does not exist, is accepted. One common version of this argument is the "problem of evil argument". In this argument, one assumes that God exists and that God is defined as an all-knowing, all-loving, and all-powerful being. From this assumption and the definition, one then considers the consequences. If God exists, then God would know about evil, would be able to prevent evil, and would want to prevent evil. That is, if God exists, then there

is no evil. But, evil does exist. So, the assumption that God exists leads to the contradiction that "evil exists and there is no evil". If arguments by reductio ad absurdum are permitted, then the conclusion to be drawn from them is that there is something wrong with the assumption: it is false that God exists.

Here is another example. Consider the following argument: "Suppose that there is a largest number. Call this number P . If P is the largest number, then there is no number greater than P . But, there is a number $P+1$ and $P+1 > P$. Therefore, P is not the largest number." In this argument, the assumption that there is a largest number is rejected and the negation of that assumption, that there is no largest number, is accepted.

With these examples in mind, let's consider the introduction and elimination rules for negation.

$\neg I$

Definition 5.4.6: Negation Introduction ($\neg I$)

From a subproof that assumes ϕ , derives ψ and $\neg(\psi)$, a wff $\neg(\phi)$ can be derived from the subproof.

$\neg E$

Definition 5.4.7: Negation Elimination ($\neg E$)

From a subproof that assumes $\neg(\phi)$, derives ψ and $\neg(\psi)$, a wff ϕ can be derived from the subproof.

The vertical dots in the subproof are placeholders for wffs derived en route to ψ and $\neg(\psi)$.

Here is the schematic use of both of these rules:

k	ϕ	A
$\vdots$	$\vdots$	
n	ψ	
n+1	$\neg(\psi)$	
n+2	$\neg(\phi)$	$\neg I$ k-(n+1)

Proof 5.3: $\neg I$

k	$\neg(\phi)$	A
$\vdots$	$\vdots$	
n	ψ	
n+1	$\neg(\psi)$	
n+2	ϕ	$\neg E$ k-(n+1)

Proof 5.4: $\neg E$

Notice that both negation introduction and elimination involve a subproof that begins with an assumption. In the case of negation introduction, the assumption is ϕ and in the case of negation elimination, the assumption is a negated wff $\neg(\phi)$. In both cases, the subproof ends with the derivation of ψ and $\neg(\psi)$, a wff and the literal negation of that wff. Once the wff and the literal negation are obtained, a wff is derived from the entire

subproof. In the case of negation introduction, the wff ϕ is assumed and $\neg(\phi)$ is derived. In the case of negation elimination, the negated wff $\neg(\phi)$ is assumed and ϕ is derived.

In other words, negation introduction allows you to add a negation to the wff you assumed provided you can show the assumption leads to a wff ϕ and its literal negation $\neg\phi$. So, if one were to assume P , derive Q and $\neg Q$, negation introduction would result in deriving $\neg(P)$ from the subproof. Alternatively, if $Q \wedge P$ were assumed, Q and $\neg Q$ derived, then negation introduction would result in deriving $\neg(Q \wedge P)$ from the subproof.

In the case of negation elimination, this rule allows you to remove the negation from a negated assumption provided you can show that the negated assumption leads to a wff ϕ and its literal negation $\neg\phi$. So, for example, if $\neg P$ is assumed, and Q and $\neg Q$ derived in the subproof, using negation elimination on the subproof would yield P . In addition, if $\neg(Q \wedge P)$ were assumed, Q and $\neg Q$ derived in the subproof, then negation elimination would yield $Q \wedge P$ from the subproof.

Let's consider some examples. First, consider the following two entailments $P \wedge \neg P \vdash S$ and $P \wedge \neg P \vdash \neg S$.

1	$P \wedge \neg P$	P	1	$P \wedge \neg P$	P
2	S	A	2	$\neg S$	A
3	P	$\wedge E$ 1	3	P	$\wedge E$ 1
4	$\neg P$	$\wedge E$ 1	4	$\neg P$	$\wedge E$ 1
5	$\neg S$	$\neg I$ 2-4	5	S	$\neg E$ 2-4

Proof 5.5: Negation Introduction Proof 5.6: Negation Elimination

Notice that both proofs are constructed in a similar manner. After the proof is set up, both proofs involve assuming a wff that is the *literal negation* (the opposite) of the conclusion. In **Proof 5.5**, the conclusion is S and so $\neg S$ is assumed. In **Proof 5.6**, the conclusion is $\neg S$ and so S is assumed. In both proofs, a wff and its literal negation is derived in the subproof (P and $\neg P$). In other words, a contradiction or a set of wffs that are explicitly inconsistent are derived in the subproof. Once the contradiction is derived, in **Proof 5.5**, the negation introduction rule derives $\neg S$, a wff that is the result of negating the assumption. the literal negation of the assumption $\neg S$ is derived. In **Proof 5.6**, the negation elimination rule derives S , a wff that is the result of removing the outermost negation from the assumption. In both cases, the contradiction is used to derive

the conclusion.

Exercise 5.54

1. $P \wedge \neg P \vdash R$
2. $P \wedge Q, Q \wedge \neg P \vdash \neg Q$
3. $L, \neg L \vdash \neg \neg \neg M$
4. $(P \vee Q) \rightarrow R, P \vee Q, \neg R \vdash \neg W$
5. $P \rightarrow Q, Q \rightarrow S, \neg S \vdash \neg Z$

5.4.5 Disjunction Introduction and Elimination

The next pair of rules we will consider are disjunction introduction and elimination. Let's start with disjunction introduction. Suppose Tek reasons as follows:

- P1: I will go to the store.
- C: Therefore, I will go to the store or the movies.

Notice that it is necessarily the case that if P1 is true, then C is true. The reason this is the case is because if "I will go to the store" is true, then the "or" sentence "I will go to the store or the movies" will be true. In other words, if $v(\phi) = T$, then $v(\phi \vee \psi) = T$.

Disjunction introduction allows you to introduce a disjunction into the proof and disjunction elimination allows you to derive wffs from a disjunction in the proof.

$\vee I$

Definition 5.4.8: Disjunction Introduction ($\vee I$)

$\phi \vdash \phi \vee \psi$ or $\phi \vdash \psi \vee \phi$

Let's consider a schematic use of the rule along with an MWE together:

1	ϕ	P		1	P	P
2	$\phi \vee \psi$	$\vee I$	1	2	$P \vee Q$	$\vee I$

The basic structure of the rule is that from a wff ϕ , a disjunction $\phi \vee \psi$ can be derived. In the justification column, the abbreviation for the rule is written (this is $\vee I$) and then the line of the wff to which the rule is applied is also written (if you apply disjunction introduction to a wff ϕ , then you would cite the line number of ϕ).

Let's consider some examples. Consider the following entailment: $P \vdash (P \vee Q) \vee \neg S$. Let's start by setting up the proof.

1 P $P, (P \vee Q) \vee \neg S$

It might be tempting to directly derive the conclusion from line 1 using $\vee I$. Unfortunately, this is not a permissible use of $\vee I$. Disjunction introduction allows us to take a wff ϕ and derive a disjunction $\phi \vee \psi$ where ϕ is one of the disjuncts of that disjunction. But notice that the conclusion is $(P \vee Q) \vee \neg S$ and the disjuncts of this wff are $P \vee Q$ and $\neg S$. However, one thing that we can do is to derive $P \vee Q$ from P . So, let's do that.

1 P P
 2 $\boxed{P \vee Q}$ $\vee I$ 1

Now that we have derived $P \vee Q$, we can use $\vee I$ again to derive the conclusion. In this case, we are using $\vee I$ on line 2 to derive the conclusion on line 3.

1 P P
 2 $P \vee Q$ $\vee I$ 1
 3 $(P \vee Q) \vee \neg S$ $\vee I$ 2

Next, let's consider the disjunction elimination rule. Disjunction elimination allows you to derive a wff from a disjunction. However, the rule requires showing that the same wff follows from both disjuncts of the disjunction. To show this, you assume each disjunct, then derive the same wff in each of the separate subproofs. Before stating the rule formally, let's consider a real-life example of disjunction elimination.

Suppose I assert the following "either I will stay home and watch a movie or I will go to a party." Let's also suppose that the proposition is true and "or" is being used inclusively. If this is the case, then one or the other sides of the "or" sentence is true (or both). Let's assume that if I stay home and watch a movie, then I will have a great time. If this is the case, then we might reason as follows:

1. I will either stay home and watch a movie or I will go to a party. (premise)
2. Assume I stay home and watch a movie. (assumption)
3. Under the assumption that I stay home and watch a movie, it follows that I will have a great time.

Consider this famous logic puzzle that involves $\vee E$. Renna is looking at Bob. Bob is looking at Marley. Renna is a teacher and Marley is not. It is unknown whether Bob is a teacher or not. Is a teacher looking at a non-teacher?

From this argument, can we conclude that I will have a great time? No. The reason we cannot conclude this is because we have not considered

the other disjunct of the "or" sentence. We have only considered the case where I stay home and watch a movie. For consider that I go to a party instead and have a terrible night. As such, we have not proven from the disjunction that I will have a great time. But suppose that we were capable of showing that if I go to a party, then I will have a great time. Let's modify our example above to reflect that this is the case:

1. I will either stay home and watch a movie or I will go to a party. (premise)
2. Assume I stay home and watch a movie. (assumption)
3. Under the assumption that I stay home and watch a movie, it follows that I will have a great time.
4. Assume I will go to the a party. (assumption)
5. Under the assumption that I go to a party, it follows that I will have a great time.

Now can we conclude that I will have a great time? Yes. The reason we can conclude this is because we have considered both sides of the "or" sentence. We have considered the case where I stay home and watch a movie and the case where I go to a party. In both cases, I will have a great time. If it follows that at least one of the sides of the "or" statement is true and that the same proposition P follows from both sides of the "or" statement, then it follows that the proposition P is true.

$\vee E$

Let's now consider the disjunction elimination rule.

Definition 5.4.9: Disjunction Elimination ($\vee E$)

From $\phi \vee \psi$ and two derivations of χ —one involving ϕ as an assumption in a subproof, the other involving ψ as an assumption in a subproof—we can derive χ from those subproofs.

Next, let's consider a schematic use of disjunction elimination in [Figure 5.2](#). First, note in order to use $\vee E$, you must have a disjunction in the proof. In the example, the disjunction $\phi \vee \psi$ is on line n . Next, notice that there are *two* subproofs. The first subproof is the left disjunct at line n . The second subproof is the right disjunct on line j . Third, notice that both

n	$\phi \vee \psi$	P
k	ψ	A
	$\vdots$	
$k+i$	χ	
j	ψ	A
	$\vdots$	
$j+1$	χ	
m	χ	$\vee E\ n,\ k-(k+i),\ j-(j+1)$

Figure 5.2: $\vee E$

subproofs derive the same wff. In the schematic example, this is χ . Finally, notice that the wff χ is derived on line m using $\vee E$. The justification is $\vee E$ and the line numbers of the disjunction (n) and the two subproofs ($n, k-(k+i), j-(j+1)$).

Next, let's consider a MWE of disjunction elimination.

1	$P \vee Q$	P
2	$P \rightarrow R$	P
3	$Q \rightarrow R$	P
4	$\begin{array}{ l} P \end{array}$	A
5	$\begin{array}{ l} R \end{array}$	$\rightarrow E$ 2,4
6	$\begin{array}{ l} Q \end{array}$	A
7	$\begin{array}{ l} R \end{array}$	$\rightarrow E$ 3,6
8	R	$\vee E$ 1, 4-5, 6-7

In the MWE, notice that the proof begins with a disjunction $P \vee Q$. At line 4, the left disjunct P is assumed and R is derived. At line 6, the right disjunct Q is assumed and R is derived. Given the disjunction and that from each disjunct R has been derived, we can now use $\vee E$ to derive R .

Exercise 5.55

1. $A \vdash A \vee \neg B$
2. $P \vdash P \vee (Q \vee R)$
3. $A, B \vdash A \vee B$
4. $A \vdash A \vee \neg \neg B$
5. $\neg A \vee B, \neg A \rightarrow S, B \rightarrow S \vdash S$
6. $(P \vee Z) \rightarrow R, Z \vdash R \vee \neg L$
7. $P \vee Q, \neg Q \vdash P$

5.4.6 Biconditional Introduction and Elimination

The next pair of rules we will consider are biconditional introduction and elimination. Biconditional introduction allows you to introduce a biconditional into the proof and biconditional elimination allows you to derive wffs from a biconditional in the proof.

$\Leftrightarrow E$

Definition 5.4.10: Biconditional Elimination ($\leftrightarrow E$)

$\phi \leftrightarrow \psi, \phi \vdash \psi$ or $\phi \leftrightarrow \psi, \psi \vdash \phi$

Next, let's consider a schematic use of biconditional elimination.

1	$\phi \leftrightarrow \psi$	P	1	$\phi \leftrightarrow \psi$	P
2	ϕ	P	2	ψ	P
3	ψ	$\leftrightarrow E$ 1, 2	3	ϕ	$\leftrightarrow E$ 1, 2

The basic structure of the rule is that from a wff $\phi \leftrightarrow \psi$, a wff ψ or ϕ can be derived. In the justification column, the abbreviation for the rule is written (this is $\leftrightarrow E$) and then the line of the wff to which the rule is applied is also written.

Next, let's consider an example. Suppose we wanted to prove the following entailment: $P \leftrightarrow (Q \leftrightarrow R), R, P \vdash Q$. Let's start by setting up the proof.

1	$P \leftrightarrow (Q \leftrightarrow R)$	P
2	R	P
3	P	P, Q

Notice that line 1 is a biconditional. As such, the biconditional has two sides: a left side and right side. If we have either side of the biconditional on a separate line in a proof, we can derive the other side. Notice that line 3 is the left side of the biconditional at line 1. As such, we can derive the right side of the biconditional on a new line.

1	$P \leftrightarrow (Q \leftrightarrow R)$	P
2	R	P
3	P	P, Q
4	$Q \leftrightarrow R$	$\leftrightarrow E$ 1, 3

Next, with $Q \leftrightarrow R$ at line 4 and the right side of that biconditional R at line 2, we can derive the left side of the biconditional Q on the next line.

1	$P \leftrightarrow (Q \leftrightarrow R)$	P
2	R	P
3	P	P, Q
4	$Q \leftrightarrow R$	$\leftrightarrow E$ 1, 3
5	Q	$\leftrightarrow E$ 4, 2

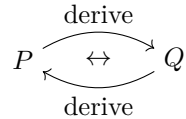
Now that we have considered biconditional elimination, let's turn to biconditional introduction.

$\leftrightarrow I$

Definition 5.4.11: Biconditional Introduction ($\leftrightarrow I$)

From two subproofs: (1) ϕ is assumed and ψ is derived; (2) ψ is assumed and ϕ is derived, a wff $\phi \leftrightarrow \psi$ can be derived.

Here is a somewhat crass way of explaining biconditional introduction. If you want to derive a biconditional, assume the left side and derive the right side. Next, assume the right side and derive the left side. Finally, derive the biconditional from the two subproofs.



Let's consider a schematic use of biconditional introduction.

n	ϕ	A
	$\vdots$	
n+i	ψ	
k	ψ	A
	$\vdots$	
k+j	ϕ	
m	$\phi \leftrightarrow \psi$	$\leftrightarrow I$ n+1, k+j

The above example illustrates this in the derivation of $\phi \leftrightarrow \psi$. First, ϕ is assumed and ψ is derived. Next, ψ is assumed and ϕ is derived. Finally, $\phi \leftrightarrow \psi$ is derived from the two subproofs.

Now that we have the basic structure of biconditional introduction, let's look at a simple MWE of biconditional introduction.

1	P	P
2	Q	P, $P \leftrightarrow Q$
3	P	A
4	Q	R 2
5	Q	A
6	P	R 1
7	$P \leftrightarrow Q$	$\leftrightarrow I$ 3-4, 5-6

In the above example, notice that the wff that we want to derive is $P \leftrightarrow Q$. In order to derive this wff, the proof involves assuming the left side of the biconditional P and deriving Q in that subproof. Then, to put it crudely,

we "do the reverse". Assume Q and then derive P . Once we have both subproofs, we can derive the biconditional $P \leftrightarrow Q$.

Exercise 5.56

1. $(P \wedge Q) \leftrightarrow R, P, Q \vdash R$
2. $(F \vee Z) \rightarrow (T \wedge P), (P \vee M) \rightarrow (R \wedge Z) \vdash P \leftrightarrow Z$
3. $(P \vee \neg M) \leftrightarrow R, P \leftrightarrow (W \vee L), L \vdash R$
4. $A \wedge \neg B \vdash A \leftrightarrow \neg B$

5.5 DERIVATION STRATEGIES

In this section, we will consider some strategies for constructing proofs. The strategies are organized around the rules that we have considered in the previous section. The strategies are not meant to be exhaustive, but rather to provide some guidance for constructing proofs. The reason for explicitly discussing strategies is fairly straightforward. While some proofs are fairly straightforward, others are rarely challenging, even to those who pick up the subject quickly. As such, it is useful to have some strategies for constructing proofs.

Strategies can be divided into two categories: (1) forward-working strategies and (2) backward-working strategies. Forward-working strategies begin with the premises and work toward the conclusion. Backward-working strategies begin with the conclusion and work toward the premises.

5.5.1 Forward-Working Strategies

Forward-working strategies begin with the premises and work toward the conclusion. It can be summarized simply in terms of the general strategy that one should *use as many elimination rules as possible*. So, for example, if your proof contains a conjunction, use $\wedge E$ (the elimination rule that corresponds to conjunctions). If your proof contains a disjunction, then use $\vee E$. If it contains conditionals, use $\rightarrow E$. If it contains biconditionals, use $\leftrightarrow E$.

Tip: If you get stuck, try to apply as many elimination rules as possible.

Let's illustrate this strategy with an example. Suppose we want to prove the following entailment: $P \wedge Q, Q \rightarrow R, R \leftrightarrow S \vdash S$. Let's start by setting up the proof.

1	$P \wedge Q$	P
2	$Q \rightarrow R$	P
3	$R \leftrightarrow S$	P
4	S	P, S

Without thinking about how to construct this proof, let's blindly try to construct this proof using the forward strategy. Line 1 is a conjunction, so let's apply the corresponding elimination rule $\wedge E$.

1	$P \wedge Q$	P
2	$Q \rightarrow R$	P
3	$R \leftrightarrow S$	P
4	S	P, S
5	$\boxed{P}$	$\wedge E$ 1
6	$\boxed{Q}$	$\wedge E$ 1

We still do not have our conclusion, so let's continue with our forward strategy by trying to use more elimination rules. Notice that line 2 is a conditional. The elimination that corresponds to conditionals is conditional elimination. Since we have Q on line 6, we can use $\rightarrow E$ to derive R .

1	$P \wedge Q$	P
2	$Q \rightarrow R$	P
3	$R \leftrightarrow S$	P
4	S	P, S
5	P	$\wedge E$ 1
6	Q	$\wedge E$ 1
7	$\boxed{R}$	$\rightarrow E$ 2,6

We still do not have our conclusion, so let's continue with our forward strategy by trying to use more elimination rules. Notice that line 3 is a biconditional. The elimination that corresponds to biconditionals is biconditional elimination. Since we have R on line 7, we can use $\leftrightarrow E$ to derive S .

1	$P \wedge Q$	P
2	$Q \rightarrow R$	P
3	$R \leftrightarrow S$	P
4	S	P, S
5	P	$\wedge E$ 1
6	Q	$\wedge E$ 1
7	R	$\rightarrow E$ 2,6
8	$\boxed{S}$	$\leftrightarrow E$ 3,7

You can think of the forward strategy as trying to "break down" the proof into its smallest parts.

The proof is complete! Notice that the proof is fairly straightforward. We simply applied the elimination rules that corresponded to the wffs in the proof.

Exercise 5.57

1. $Z \wedge (B \wedge F), (M \wedge T) \wedge (L \rightarrow P), Q \wedge (R \wedge P) \vdash \neg R \vee (S \vee T)$
2. $(S \wedge W) \wedge (T \wedge X), (P \wedge W) \wedge F, F \rightarrow R \vdash (P \wedge R) \vee (S \wedge L)$
3. $(Z \wedge Q) \wedge (F \wedge L), R \wedge P, W \wedge B \vdash (Z \vee T) \vee (M \rightarrow R)$
4. $(L \wedge F) \rightarrow S, W \wedge (F \wedge X), W \rightarrow L \vdash (S \vee R) \vee P$
5. $M \wedge (R \wedge \neg Z), S \wedge (P \wedge W), Q \vdash (S \leftrightarrow Q) \vee [M \wedge (R \wedge Z)]$
6. $[(P \wedge Q) \wedge (W \wedge L)] \wedge [R \wedge (S \wedge T)], Z \wedge [(W \wedge R) \wedge (T \wedge Z)], (F \rightarrow P) \leftrightarrow W \vdash A \vee Z$
7. $P \rightarrow (R \wedge M), (P \wedge S) \wedge Z \vdash R$
8. $P \rightarrow R, Z \rightarrow W, P \vdash R \vee W$

5.5.2 Backward-Working Strategies

In the forward-working strategy, we began with the premises and worked toward the conclusion using elimination rules. In the backward-working strategy, we begin with the conclusion and work toward the premises. In the backward-working strategy, when the conclusion contains a main operator, we can often refine this question further by asking the following question: *what introduction rule could I use to derive the conclusion?* So, for example, if the conclusion of your proof is a conjunction, consider using $\wedge I$ (the introduction rule that corresponds to conjunctions). If the conclusion is a disjunction, then consider using $\vee I$. If it is a conditional, consider $\rightarrow I$. If it is a biconditional, consider $\leftrightarrow I$.

Let's illustrate this strategy with a few examples as this strategy is more complex than the forward-working strategy. Suppose we want to construct

a proof for $Q \wedge P \vdash (S \wedge P) \rightarrow Q$. Let's start by setting up the proof.

1 $Q \wedge P$ $P, (S \wedge P) \rightarrow Q$

If we were to use the forward strategy, our first step would be to use $\wedge E$ to derive Q and P on separate lines. But, at this point, there does not appear a way to use Q and P to create a conditional. Let's let's try to use the backward strategy instead. Notice that the wff we want to derive is a conditional. As such, we should consider using the introduction rule that corresponds to conditionals ($\rightarrow I$). If we were to use $\rightarrow I$ to derive the conclusion, we would need to assume the antecedent $(S \wedge P)$ and derive the consequent Q . Let's start the proof then by assuming the antecedent $(S \wedge P)$.

It helps to take things step-by-step rather than trying to do everything at once.

1 $Q \wedge P$ $P, (S \wedge P) \rightarrow Q$
 2 $\mid S \wedge P$ A, Q

In the justification column of the subproof, "Q" is written after the assumption. This is to indicate that this is the goal of this subproof. This is the wff that we want to derive so that we can use $\rightarrow I$. Next, we need to derive Q . Notice that line 1 is $Q \wedge P$. As such, we can use $\wedge E$ to derive Q in the subproof.

If you want to derive a conditional $\phi \rightarrow \psi$, assume ϕ , derive ψ , then use $\rightarrow I$.

1 $Q \wedge P$ $P, (S \wedge P) \rightarrow Q$
 2 $\mid S \wedge P$ A, Q
 3 $\mid Q$ $\wedge E$ 1

Now since we have derived Q in the subproof, we can use $\rightarrow I$ to derive the conclusion.

1 $Q \wedge P$ $P, (S \wedge P) \rightarrow Q$
 2 $\mid S \wedge P$ A, Q
 3 $\mid Q$ $\wedge E$ 1
 4 $(S \wedge P) \rightarrow Q$ $\rightarrow I$ 2-3

The proof is complete! What it illustrates is that if you want to derive a conditional $\phi \rightarrow \psi$, consider the introduction rule that you would need to use to "introduce" into the proof that conditional. This is $\rightarrow I$. When we examine this rule, we now have clear direction about how to take the next step in the proof. When trying to derive a conditional $\phi \rightarrow \psi$, $\rightarrow I$ requires assuming ϕ (the antecedent), then deriving ψ (the consequent), then using $\rightarrow I$.

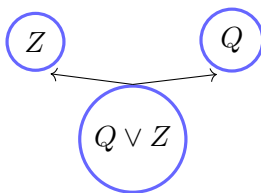
Let's consider another example. Suppose we want to prove the following entailment: $Q \vee W, P \wedge S \vdash Q \vee Z$. Let's start by setting up the proof.

- 1 $Q \vee W$ P
- 2 $P \wedge S$ P, $Q \vee S$

In this proof, our eyes might be initially attracted to line 1. This is because line 1 contains Q , the conclusion contains Q , and we might believe that we could use elimination rule that corresponds to disjunctions ($\vee E$) to derive our conclusion. Unfortunately, if you try to use line 1 and $\vee E$, you will quickly find yourself stuck.

Remember you can't derive Q from $Q \vee W$. This would be like reasoning "I have a dog or a cat. Therefore, I have a dog."

Rather than using the forward strategy, let's try to use the backward strategy. In this case, the conclusion is a disjunction. As such, we should consider using the introduction rule that corresponds to disjunctions ($\vee I$). If we were to use $\vee I$ to derive the conclusion, we would need to either have Q or Z in the proof.



- 1 $Q \vee W$ P
- 2 $P \wedge S$ P, $Q \vee S$
- ? ?
- ? Q or S ?
- ? $Q \vee S$ $\vee I$, ?

Now that we have worked backward one step in the proof, how to construct the proof becomes more apparent. We only need to derive Q or S . We cannot derive Q from line 1, but we can derive S from line 2. Let's use $\wedge E$ one line 2 to derive S and complete the rest of the proof.

- 1 $Q \vee W$ P
- 2 $P \wedge S$ P, $Q \vee S$
- 3 S $\wedge E$ 2
- 4 $Q \vee S$ $\vee I$ 3

The proof is complete! What it illustrates is that if you want to derive a disjunction $\phi \vee \psi$, consider the introduction rule that you would need to use to "introduce" into the proof that disjunction. This is $\vee I$. When we examine this rule, we now have clear direction about how to take the next step in the proof. When trying to derive a disjunction $\phi \vee \psi$, $\vee I$ requires either ϕ or ψ in the proof.

Let's consider one more example. Consider the following entailment $P \wedge (\neg Z \wedge \neg P) \vdash W$. Let's start by setting up the proof.

1 $P \wedge (\neg Z \wedge \neg P)$ P, W

Now let's consider our backward strategy. This strategy says to consider what rule we would need to use to derive the conclusion. In this case, the conclusion is W , an atomic wff. As such, we should consider what rule (or rules) would allow for deriving an atomic wff. In this case, there are several. If there was a conjunction containing W , then we could use $\wedge I$ to derive W . However, notice that there is no conjunction containing W . Alternatively, if there was a conditional containing W as a consequent (e.g., $P \rightarrow W$) and we had P , then we could use $\rightarrow I$ to derive W . However, notice that there is no conditional containing W . Finally, if there was a biconditional containing W , then we could use $\leftrightarrow I$ to derive W . However, notice that there is no biconditional containing W .

But now consider $\neg E$. This rule contends that if you assume $\neg(\phi)$, derive a wff and its literal negation $\psi, \neg(\phi)$, then you can derive ϕ . Let's try to use $\neg E$ to derive W .

1 $P \wedge (\neg Z \wedge \neg P)$ P, W
 2 $\quad \underline{\neg W}$ A, $\psi, \neg(\psi)$

In the justification column of the subproof, $\psi, \neg(\psi)$ is written after the assumption. This is to indicate that this is the goal of this subproof. Let's return to the forward strategy. Notice that line 1 is a conjunction. As such, we can use $\wedge E$ to derive P and $\neg P$ (a wff and its literal negation). Once these two wffs are in the subproof, $\neg E$ can be used to derive W .

1 $P \wedge (\neg Z \wedge \neg P)$ P, W
 2 $\quad \underline{\neg W}$ A, $\psi, \neg(\psi)$
 3 $\quad P$ $\wedge E$ 1
 4 $\quad \neg Z \wedge \neg P$ $\wedge E$ 1
 5 $\quad \neg P$ $\wedge E$ 4
 6 W $\neg E$ 2-5

In sum, there are two principal strategies for constructing proofs: (1) forward-working strategies and (2) backward-working strategies. Forward-working strategies begin with the premises and work toward the conclusion. Backward-working strategies begin with the conclusion and work toward the premises. In constructing proofs, neither strategy should be relied upon exclusively as each has its own advantages. In the forward-strategy, one "simplifies" the proof by "breaking down" the premises into smaller, more flexible parts. However, the strategy can often reach a dead-

Again, take things step-by-step. Here we are assuming "the opposite" of our conclusion in the hopes that $\neg E$ will give us our conclusion.

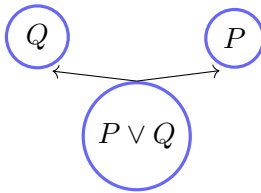
The Forward-looking strategy amounts to charging up the mountain without a clear direction, while the backward-looking strategy amounts to looking down from the mountain and determining (in broad strokes) the best paths.

end when the conclusion is complex. In the backward-strategy, one gets a broader view of the proof by considering what rule would be needed to derive the conclusion. It offers general guidelines rather than giving fine-grained directions about which rules should be used at each step.

To conclude, let's consider *part* of a proof that many students often initially struggle to construct. Consider $\neg(\neg P \wedge \neg Q) \vdash P \vee Q$. Let's start by setting up the proof.

1 $\neg(\neg P \wedge \neg Q)$ P, $P \vee Q$

At first glance, it is not clear how to solve this proof. Consider the forward-looking strategy. This strategy says to consider what elimination rules can be used. Line 1 is a negated conjunction, but there are no apparent elimination rules that can be used to "break down" the negated conjunction into simpler wffs. Now consider the backward-looking strategy. This strategy says to consider what rule can be used to derive the conclusion and try to work backward toward the premises. The conclusion is a disjunction and so trying to derive either P or Q and then using $\vee I$ seems like a good strategy.



But how do we derive P or Q ? One option is to try to assume $\neg P$, derive a wff ϕ and $\neg(\phi)$, and then use $\neg E$ to derive P .

1 $\neg(\neg P \wedge \neg Q)$ P, $P \vee Q$
 2 $\mid \neg P$ A, $\phi, \neg\phi$

But, at this point, it is not clear how to derive ϕ and $\neg\phi$. Trying this with $\neg Q$ leads to a similar dead end. So, let's try to use the backward-looking strategy again. Another approach is to try to not the negation of P or Q but the negation of the entire conclusion. That is, since our conclusion is $P \vee Q$, let's move forward in the proof by assuming $\neg(P \vee Q)$.

1 $\neg(\neg P \wedge \neg Q)$ P, $P \vee Q$
 2 $\mid \neg(P \vee Q)$ A, $\phi, \neg\phi$
 3 $\mid \vdots$

The subgoal at this point is to derive a wff ϕ and its literal negation $\neg(\phi)$. It is still not clear how to do this. So, let's try to use the backward-looking strategy again. Consider line 2 and note that if we were able to derive $P \vee Q$, we could use $P \vee Q$ with line 2 to derive a contradiction. But how do we derive $P \vee Q$? One option is to try to assume P and then use $\vee I$ to derive $P \vee Q$.

1	$\neg(\neg P \wedge \neg Q)$	P, $P \vee Q$
2	$\neg(P \vee Q)$	A, P, $\neg P$
3	P	A
4	$P \vee Q$	$\vee I$, 3

With $P \vee Q$ in the proof, let's reiterate line 2 and then use $\neg I$ to derive $\neg P$.

1	$\neg(\neg P \wedge \neg Q)$	P, $P \vee Q$
2	$\neg(P \vee Q)$	A, P, $\neg P$
3	P	A
4	$P \vee Q$	$\vee I$, 3
5	$\neg(P \vee Q)$	R 2
6	$\neg P$	$\neg I$ 3-5

This is a key step in the proof. Notice that we have derived $\neg P$ in the subproof. Try to complete the rest of your proof yourself.

Exercise 5.58

1. $P \rightarrow Q, \neg Q \vdash \neg P$
2. $R \vdash \neg(D \vee L) \rightarrow R$
3. $\neg(P \vee R) \vdash \neg P \wedge \neg R$
4. $\neg(\neg P \wedge \neg Q) \vdash P \vee Q$
5. $(\neg P \wedge L) \rightarrow \neg Q, (M \wedge T) \wedge (\neg R \wedge L), (M \wedge \neg R) \rightarrow (Z \wedge \neg P) \vdash \neg Q \vee (A \leftrightarrow B)$
6. $\neg R \vdash P \vee \neg W \rightarrow (Q \vee \neg R)$
7. $\vdash \neg(W \wedge \neg W)$
8. $P, (P \vee W) \rightarrow (R \wedge T), (T \vee \neg V) \leftrightarrow (\neg R \wedge T) \vdash S$
9. $\neg P \vee R \vdash P \rightarrow R$
10. $P \rightarrow R \vdash \neg P \vee R$

5.6 ADDITIONAL DERIVATION RULES

The set of 10 intelim rules form **PD**. At this point, we have three options. First, we might end our discussion of proofs. **PD** is "complete". This means that **PD** is capable of proving any valid argument in PL (more exactly, for any semantic entailment $\Gamma \models \phi$, there is a proof that $\Gamma \vdash \phi$). Second, we might add additional derivation rules to our proof system. We might add these rules because either (1) they make solving proofs much

easier or (2) they are "natural" and so adding them to our system would give us a proof system that better tracks how we reason. Recall that we added the reiteration rule R because it simplified (shortened) proofs. A third option is to consider alternative proof systems. For example, instead of developing a system of natural deduction, we might consider a Hilbert deductive system. In our intelim natural deduction system, we said that $\Gamma \vdash \phi$ means that there is a proof of ϕ from Γ , where the members of Γ were taken to be the "premises" of the proof. In this sense, the intelim system is a "premise-based" system. In contrast, in a Hilbert system, instead of starting with premises or assumptions, the members of Γ would consist of "logical axioms" and, quite typically, a single "logical rule of inference" (conditional elimination). In this sense, a Hilbert system is an "axiom-based" system.

In the remainder of this chapter, our focus will be on the second option. We will consider adding additional derivation rules to our intelim system. In the next section, we will consider adding additional derivation rules to our intelim system because they are "natural" and so adding them to our system would give us a proof system that better tracks how we reason.

5.6.1 Additional Derivation Rules: Naturalness

To distinguish the core derivation rules from our deductive apparatus, let's say that **PD** refers to our ten intelim rules and **PD+** consists of **PD** plus any additional rules added beyond those ten rules. At this point then, only reiteration (R) distinguishes **PD** from **PD+**. In this section, we will add three additional derivation rules to our intelim system because they are "natural" and so adding them to our system would give us a proof system that better tracks how we reason. One condition we will place on the addition of a new rule, however, is that it must be possible to derive the new rule from the existing rules.

Our approach to introducing these three derivation rules will differ slightly. We will list all three new rules together and then consider each rule in turn.

Definition 5.6.1: Disjunctive Syllogism (DS)

$\phi \vee \psi, \neg(\psi) \vdash \phi$ or $\phi \vee \psi, \neg(\phi) \vdash \psi$

Definition 5.6.2: Modus Tollens (MT)

$$\phi \rightarrow \psi, \neg(\psi) \vdash \neg(\phi)$$

Definition 5.6.3: Hypothetical Syllogism (HS)

$$\phi \rightarrow \psi, \psi \rightarrow \chi \vdash \phi \rightarrow \chi$$

Let's start with disjunctive syllogism (DS). Suppose Tek and Liz have a daughter named Pat. Tek and Liz are discussing Pat's future. Tek says "either Pat will go to college or Pat join the military." Liz says "Pat will not join the military." From this, Liz concludes that Pat will go to college. In this example, Tek utters an "or-proposition", Liz utters the negation of one of the sides of the "or-proposition", and then Liz concludes the other side of the "or-proposition" follows. This is the basic idea behind disjunctive syllogism. If there is a disjunction $\phi \vee \psi$, the negation of one of the sides of the disjunction (either $\neg(\phi)$ or $\neg(\psi)$), then DS permits the derivation of the other side of the disjunction.

Let's consider an example. Consider the entailment $P \vee (R \wedge S), \neg(R \wedge S) \vdash P$. The proof of this entailment is as follows:

- | | | |
|---|-----------------------|---------|
| 1 | $P \vee (R \wedge S)$ | P |
| 2 | $\neg(R \wedge S)$ | P |
| 3 | P | DS 1, 2 |

Notice that at line 1 there is a disjunction, at line 2 there is the negation of the right side of the disjunction, and at line 3 there is the left side of the disjunction. In citing the justification, since DS uses two lines (wffs), the line number of each line is used along with "DS" for the justification.

Recall that our criteria for adding new rules is that the reasoning the rule represents must be "natural" and it must be capable of being derived with our existing set of rules. Surely, *DS* is a natural rule. But is it capable of being derived with our existing set of rules? Yes. To illustrate, let's show that $\phi \vee \psi, \neg(\psi) \vdash \phi$ can be derived using our existing rules.

1	$\phi \vee \psi$	P
2	$\neg(\psi)$	P, ϕ
3	ϕ	A
4	ϕ	R, 3
5	ψ	A
6	$\neg\phi$	A
7	ψ	R, 5
8	$\neg\psi$	R, 6
9	ϕ	$\neg E$, 6-8
10	ϕ	$\vee E$ 1, 3-4, 5-9

Next, let's consider modus tollens (MT). Suppose Tek and Liz are returning home from work. Tek is the worrisome sort and says "if we left the stove on, then the house will have burned down." When they arrive home, Liz sees the home is intact and says "the house did not burn down." From this, Tek and Liz conclude that they did not leave the stove on. In this example, Tek utters a proposition of the form "if P, then Q". Liz utters the negation of Q, and then they both conclude that the negation of P follows. This is the basic idea behind modus tollens. If there is a conditional $\phi \rightarrow \psi$, the negation of the consequent of the conditional (i.e. $\neg(\psi)$), then MT permits the derivation of the negation of the antecedent of the conditional (i.e. $\neg(\phi)$).

Solving proofs is like solving a puzzle.

Let's consider an example. Consider the entailment $P \rightarrow Q, \neg(Q) \vdash \neg(P)$. The proof of this entailment is as follows:

1	$P \rightarrow Q$	P
2	$\neg Q$	P
3	$\neg P$	MT 1, 2

In above proof, notice that the conclusion is the literal negation of the antecedent of the conditional $P \rightarrow Q$. The conditional $P \rightarrow Q$ is the first premise and the literal negation of the consequent $\neg Q$ is the second premise. The conclusion is derived using MT and then citing the lines of both the conditional and its literal negation.

Let's consider another example of MT. First, consider the following entailment: $P \rightarrow (S \vee R), \neg(S \vee R) \vdash \neg P$.

1	$P \rightarrow (S \vee R)$	P
2	$\neg(S \vee R)$	P
3	$\neg P$	MT 1, 2

Notice that in the above example, line 1 is a conditional. This conditional has $S \vee R$ as its consequent. In order to use MT on this conditional, it is necessary to have the literal negation of $S \vee R$. This would be $\neg(S \vee R)$. With the conditional at line 1 and its literal negation at line 2, MT can be used to derive the literal negation of the antecedent at line 3.

Let's consider one more example of MT. Consider the following entailment: $(P \vee R) \rightarrow \neg S, \neg\neg S \vdash \neg(P \vee R)$.

1	$(P \vee R) \rightarrow \neg S$	P
2	$\neg\neg S$	P
3	$\neg(P \vee R)$	MT 1, 2

Notice that in the above proof, line 1 is a conditional. This conditional has $\neg S$ as its consequent. In order to use MT on this conditional, it is necessary to have the literal negation of $\neg S$. This would be $\neg\neg S$. With the conditional at line 1 and its literal negation at line 2, MT can be used to derive the literal negation of the antecedent at line 3. Note that the literal negation of the antecedent is the negation of a disjunction $\neg(P \vee R)$ rather than the negation of the disjuncts in that disjunction. That is, using MT would allow us to derive $\neg(P \vee R)$ but not $\neg P \vee \neg R$ and not $\neg P \vee R$.

Again, if we want to add MT to our system, we need to show that it can be derived from our existing rules. Consider the following derivation of $\phi \rightarrow \psi, \neg(\psi) \vdash \neg(\phi)$

1	$\phi \rightarrow \psi$	P
2	$\neg(\psi)$	P
3	ϕ	A
4	ψ	$\rightarrow E$ 1, 3
5	$\neg(\psi)$	R 2
6	$\neg\phi$	$\neg I$ 2-5

Finally, let's consider *HS*. Suppose Liz and Tek are thinking about going to a large arts festival. Tek has a habit of buying little items whenever he attends such festivals. Tek says "if I attend this festival, then I will buy several items". Liz responds, "if you buy several items, then you will have to carry them." Tek and Liz then draw the following conclusion:

"if I attend this festival, then I will have to carry several items." In this example, Tek utters an "if P, then Q" proposition. Liz utters an "if Q, then R" proposition. From this, Tek and Liz conclude that from these two "if-then" propositions, a third "if-then" proposition follows: "if P, then R". This is the basic idea behind hypothetical syllogism. If there are two conditionals $\phi \rightarrow \psi$ and $\psi \rightarrow \chi$, then HS permits the derivation of the conditional $\phi \rightarrow \chi$. Let's consider a MWE of HS. In this example, we will provide a proof for $P \rightarrow R, P \rightarrow Q \vdash Q \rightarrow R$.

- 1 $P \rightarrow R$ P
- 2 $Q \rightarrow R$ P
- 3 $P \rightarrow R$ HS 1, 2

In the above proof, notice that before using HS, there are two conditionals in the proof. The first premise is a conditional $P \rightarrow Q$ and the second premise is a conditional $Q \rightarrow R$. Notice how the consequent of $P \rightarrow Q$ is the antecedent of $Q \rightarrow R$. When this is the case, then HS allows for deriving a conditional consisting of the antecedent of the first conditional and the consequent of the second conditional. Let's consider another example of a proof involving HS. Consider the following entailment: $\neg R \rightarrow Q, S \rightarrow \neg R \vdash S \rightarrow Q$.

- 1 $\neg R \rightarrow Q$ P
- 2 $S \rightarrow \neg R$ P
- 3 $S \rightarrow Q$ HS 1, 2

In this section, we added three additional rules to **PD+**: DS, MT, and HS. We added these rules because they are "natural" and so adding them to our system would give us a proof system that better tracks how we reason. In the next section, we will add three additional rules to **PD+** because they make solving proofs much easier.

Exercise 5.59

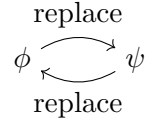
1. $(R \wedge T) \vee \neg W, S \wedge \neg \neg W \vdash R \wedge T$
2. $(P \wedge S) \rightarrow W, \neg W \wedge T \vdash \neg(P \wedge S)$
3. $(R \wedge T) \rightarrow \neg W, M \rightarrow (R \wedge T), \neg W \rightarrow (S \wedge R) \vdash M \rightarrow (S \wedge R)$
4. $P \vee \neg(R \vee S), R, L \rightarrow \neg P \vdash \neg L$
5. Challenge: Show that $P \rightarrow Q, Q \rightarrow R \vdash P \rightarrow R$ can be derived using only intelim rules and R .
6. Challenge: Create another derivation to add to **PD+** and

then use the intelim rules to show that this rule can be proven with the existing stock of derivation rules.

7. Translate the following argument into **PL**, then prove it: "Either my client is guilty or innocent. If my client is guilty, then there is evidence my client is guilty. If there is evidence that my client is guilty, then the prosecution has presented that evidence. The prosecution has not presented evidence that my client is guilty. Therefore, my client is innocent."

5.6.2 The Replacement Rules

All of the previous derivation rules have been expressed as one-direction derivation rules. That is, they allow for deriving a proposition of one form from a proposition of another form. For example, $\vee I$ permits deriving a disjunction $\phi \vee \psi$ from a disjunct ϕ , but it does not permit deriving a disjunct ϕ from a disjunction $\phi \vee \psi$. In this section, we will consider adding three additional derivation rules to **PD+** that are two-direction derivation rules. These rules are known as replacement (or equivalence) rules. Replacement rules are derivation rules that allow for interchanging certain formulas or sub-formulas in a proof.



Similar to the additional rules we added to **PD+** in the previous section, each of the replacement rules that we will add in this section are capable of being derived using **PD**.

5.6.2.1 Double Negation

The first replacement rule that we will add to **PD+** is the double negation rule. This rule allows for taking any wff ϕ and replacing it with its doubly negated form $\neg\neg\phi$ or taking any wff $\neg\neg\phi$ and replacing it with a wff that removes both negations ϕ . The double negation rule is a two-direction derivation rule. It allows for deriving $\neg\neg\phi$ from ϕ and it allows for deriving ϕ from $\neg\neg\phi$. There are two ways to express this fact. First, we might express the rule as follows: $\phi \vdash \neg\neg\phi$ and $\neg\neg\phi \vdash \phi$. Second, we might express the rule more compactly as follows: $\phi \dashv\vdash \neg\neg\phi$. The symbol $\dashv\vdash$ is known as a "turnstile" and it is used to express the fact that a rule is two-directional.

Definition 5.6.4: Double Negation (DN)

$$\neg\neg(\phi) \dashv\vdash \phi$$

Let's consider examples of DN. First, consider the following entailment: $P \rightarrow R \vdash \neg\neg(P \rightarrow R)$. The proof of this entailment is as follows:

- 1 $P \rightarrow R$ P, $\neg\neg(P \rightarrow R)$
- 2 $\neg\neg(P \rightarrow R)$ DN 1

Notice that the conclusion is the doubly negated form of the conditional $P \rightarrow R$ at line 1. To cite a use of DN, we can cite the abbreviation for the rule "DN" along with the line of the proof to which the rule is applied (line 1 in the above example). It is important to note that replacement rules can be applied not only to the entire wff ϕ but any subformula ψ of ϕ . For example, consider the use of DN in the following proof:

- 1 $P \vee (R \wedge S)$ P, $(P \vee (R \wedge S))$
- 2 $(P \vee (R \wedge \neg\neg S))$ DN 1

Notice that the use of DN in the above proof is applied to the subformula S of the formula $P \vee (R \wedge S)$ at line 1. But, be careful! DN involves replacing a wff with its doubly negated form. It does not involve adding one negation to one part of a wff and another negation to another part of a wff. For example, consider the following proof:

- 1 $P \vee (R \wedge S)$ P, $(P \vee (R \wedge S))$
- 2 $\neg P \vee (\neg R \wedge S)$, **NO!** DN 1

In the above example, notice that rather than replacing say P with $\neg\neg P$ or R with $\neg\neg R$, the subformula P is replaced with $\neg P$ and the subformula R is replaced with $\neg R$. This would not be a correct use of DN.

5.6.2.2 De Morgan's Laws

The next rule we will add to our deductive system is "De Morgan's Laws" (a pair of rules named after 19th century mathematician and logician Augustus De Morgan).

Definition 5.6.5: De Morgan's Laws (DeM)

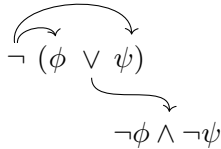
$$\neg(\phi \vee \psi) \dashv\vdash \neg(\phi) \wedge \neg(\psi)$$

$$\neg(\phi \wedge \psi) \dashv\vdash \neg(\phi) \vee \neg(\psi)$$

In the case of DeM, you can interchange a negated disjunction $\neg(P \vee Q)$ with a conjunction whose conjuncts are negated $\neg P \wedge \neg Q$ (and vice versa) and you can interchange a negated conjunction $\neg(P \wedge Q)$ with a disjunction $\neg P \vee \neg Q$ whose disjuncts are negated (and vice versa).

$\neg(\phi \vee \psi)$. The $\neg$ "leaps" over the parentheses and "lands" on the ϕ and ψ . Now the wff is $\neg(\phi) \wedge \neg(\psi)$.

DeM is like a game of leap frog. The $\neg$ "leaps" over the parentheses and "lands" on the negated propositions.



De Morgan's laws are probably not the most intuitive ("natural") rules, but they are very useful for solving proofs. Because they are so useful, let's consider several examples of DeM. First, let's consider a MWE of DeM. Rather than solving a proof, let's consider the use of DeM in a proof:

- 1 $\neg(P \vee Q)$ P
- 2 $\neg P \wedge \neg Q$ DeM 1
- 3 $\neg(P \vee Q)$ DeM 2

Notice that above proof begins with a negated disjunction. Using DeM on that negated disjunction allows for deriving a conjunction with two negated conjuncts. This is $\neg P \wedge \neg Q$ at line 2. When using DeM, we cite DeM and the line of wff to which DeM is applied. In looking at lines 2-3, recall that DeM is a rule of replacement. So, not only can we reason from a negated disjunction $\neg(P \vee Q)$ to the corresponding conjunction $\neg P \wedge \neg Q$, but we also can reason from that conjunction back to the negated disjunction $\neg(P \vee Q)$.

Since DeM is a pair of rules (one involving negated disjunctions, the other involving negated conjunctions), let's consider a MWE of the other rule.

- 1 $\neg(P \wedge Q)$ P
- 2 $\neg P \vee \neg Q$ DeM 1
- 3 $\neg(P \wedge Q)$ DeM 3

This proof illustrates a use of DeM. First, notice that the proof begins with a negated conjunction $\neg(P \wedge Q)$. Using DeM on that negated conjunction allows for deriving a disjunction with two negated disjuncts. This is $\neg P \vee \neg Q$ at line 2. In looking at lines 2-3, again recall that DeM is a rule

of replacement. So, not only can we reason from a negated conjunction $\neg(P \wedge Q)$ to the corresponding disjunction $\neg P \vee \neg Q$, but we also can reason from that disjunction back to the negated conjunction $\neg(P \wedge Q)$.

Now that we have the basic idea of how to use DeM in a proof, let's consider some more complicated examples. Consider the following entailment: $\neg(A \wedge B) \vdash B$. The proof is set up as follows:

1 $\neg(A \wedge B)$ P, B

Since line 1 is a negated disjunction, there is not a whole lot we can do with this wff. We cannot apply $\wedge E$ to it as it is not a conjunction. We might try to assume $\neg B$ and tried to derive $\phi, \neg\phi$ in the hopes of deriving B . But, once we assume $\neg B$, it is not immediately obvious what we would do next. So, let's try using DeM on line 1. This would allow us to derive a conjunction with two negated conjuncts. This is $\neg A \wedge \neg B$ at line 2.

1 $\neg(A \wedge B)$ P, B
2 $\neg A \wedge \neg B$ DeM 1

Now that we have a conjunction with two negated conjuncts, we can apply $\wedge E$ on line 2 to derive $\neg B$ at line 3. This would complete the proof. This proof highlights an important tip. When there is a negated disjunction $\neg(P \vee Q)$ or a negated conjunction $\neg(P \wedge Q)$, it is a good idea to use DeM as these rules derive conjunctions and disjunctions. Conjunctions and disjunctions can then be used to apply rules like $\wedge E$, $\vee E$, or DS .

Let's consider another example involving DeM. Consider the following entailment: $P \rightarrow (R \vee Q), \neg R \wedge \neg Q \vdash \neg P$. The proof is set up as follows:

1 $P \rightarrow (R \vee Q)$ P
2 $\neg R \wedge \neg Q$ P

As mentioned earlier, generally, it is a good idea to consider using DeM to negated conjunctions and negated disjunctions. However, it is important to remember that DeM is a two-directional rule. It can be used on the conjunction in line 2 to derive the negated disjunction $\neg(R \vee Q)$. From there, MT can be used to derive the conclusion.

1 $P \rightarrow (R \vee Q)$ P
2 $\neg R \wedge \neg Q$ P
3 $\neg(R \vee Q)$ DeM 2
4 $\neg P$ MT 1, 3

5.6.2.3 Implication

The final rule that we will add to **PD**+ is "implication" (IMP).

Definition 5.6.6: Implication (IMP)

$$\phi \rightarrow \psi \dashv\vdash \neg(\phi) \vee \psi$$

In the case of IMP, you can replace a negated conditional $P \rightarrow Q$ with a disjunction $\neg P \vee Q$, and vice versa. Similar to the illustration of DeM, let's consider an example of IMP where we are more focused on how to use the rule in a proof rather than solving a proof.

- 1 $P \rightarrow Q$ P
- 2 $\neg P \vee Q$ IMP 1
- 3 $P \rightarrow Q$ IMP 2

In the above example, the proof begins with $P \rightarrow Q$. Since this wff is a conditional, IMP can be used to derive the disjunction $\neg P \vee Q$ at line 2. In looking at lines 2-3, again recall that IMP is a rule of replacement. So, not only can we reason from a conditional $P \rightarrow Q$ to the corresponding disjunction $\neg P \vee Q$, but we also can reason from that disjunction back to the conditional $P \rightarrow Q$. As IMP is applied to only one line of the proof, we cite the line of the proof to which IMP is applied.

Next, let's consider an example of IMP where we are solving a proof. Consider the following entailment: $\neg(P \rightarrow Q) \vdash \neg Q$. The proof is set up as follows:

- 1 $\neg(P \rightarrow Q)$ P, $\neg Q$

Since IMP is a replacement rule, it can be applied to the subformula of a wff. In this case, IMP can be applied to the subformula $P \rightarrow Q$ of the wff $\neg(P \rightarrow Q)$ at line 1. This would allow for deriving the disjunction $\neg(\neg P \vee Q)$ at line 3.

- 1 $\neg(P \rightarrow Q)$ P, $\neg Q$
- 2 $\neg(\neg P \vee Q)$ 1IMP

Since line 3 is a negated disjunction, we can complete the proof by using DeM and then $\wedge E$.

1	$\neg(P \rightarrow Q)$	$P, \neg Q$
2	$\neg(\neg P \vee Q)$	IMP 1
3	$\neg\neg P \wedge \neg Q$	DeM 2
4	$\neg Q$	$\wedge E$ 3

In this section, we added three replacement rules to **PD**+: DN, DeM, and IMP. We added these rules, not because they are necessarily natural or intuitive, but because they make solving proofs much easier.

Exercise 5.60

1. $\neg\neg(\neg\neg P \rightarrow Q) \vdash P \rightarrow \neg\neg Q$
2. $\neg(\neg P \vee Q) \vdash P$
3. $S \rightarrow \neg Q \vdash \neg S \vee \neg Q$
4. $\neg\neg P \rightarrow R, P, \neg\neg R \rightarrow (W \wedge Z) \vdash \neg\neg(W \wedge \neg\neg Z)$
5. $\neg(P \vee R) \rightarrow (\neg Z \vee \neg W), \neg P \wedge \neg R \vdash \neg(Z \wedge W)$
6. $(P \rightarrow R), (\neg P \vee R) \rightarrow (Z \rightarrow \neg R) \vdash \neg Z \vee \neg R$
7. $\neg P \vee R, \neg(P \rightarrow R) \vdash S$
8. $P \rightarrow \neg(Z \vee S), \neg(P \rightarrow R) \vdash \neg Z \vee W$
9. $P, \neg(\neg P \wedge \neg R) \rightarrow \neg(S \rightarrow T) \vdash S$
10. $P \leftrightarrow (R \vee S), P \wedge \neg S, Q \rightarrow \neg R \vdash \neg Q$
11. $R \vee (M \wedge T), \neg R \wedge \neg W, L \rightarrow W \vdash \neg L$
12. $(R \vee M) \vee \neg(S \vee T), (S \vee T) \vee (Z \wedge E), \neg(R \vee M) \vdash E$
13. $\neg(P \vee R), \neg P \rightarrow \neg(M \vee S), \neg R \rightarrow \neg Q \vdash \neg M \wedge \neg Q$
14. $\neg(P \rightarrow R), P \rightarrow Z, \neg R \rightarrow M \vdash Z \wedge M$
15. $\neg(\neg P \rightarrow \neg R), Z \rightarrow P \vdash \neg Z \wedge R$
16. $\vdash \neg(P \rightarrow R) \rightarrow (S \rightarrow \neg R)$
17. $\vdash \neg(P \vee R) \rightarrow [(Z \rightarrow R) \rightarrow \neg Z]$
18. $\vdash [\neg(P \rightarrow M) \wedge \neg(T \rightarrow S)] \vee (P \vee \neg P)$
19. Recall that in an earlier section, two principal strategies for solving proofs were introduced. How would you incorporate DN, DeM, and IMP into those existing strategies?

5.6.3 Even More Derivation Rules: Theorems

Our formulation of the intelim rules was sufficient to provide a derivation of any syntactic entailment. We supplemented the intelim rules with additional derivation rules and rules of replacement. The motivation behind this supplementation was because some of these rules correspond to how people naturally reason and to simplify proofs. At this point, we might add even more rules to our deductive apparatus to make solving proofs

even easier. However, this supplementation is both not necessary and potentially comes with a cost. The cost is that the more rules we add to our deductive apparatus, the more rules it is necessary to remember. This can make solving proofs more difficult.

Rather than explicitly add more rules to our deductive apparatus, we will consider a way to create new derivation rules from the existing rules. If ϕ is derived from Γ and Γ is empty, then derivation of ϕ is a *proof* or *zero-premise deduction* of ϕ . In addition, ϕ is a *theorem*. Once it is established that ϕ is a theorem of our deductive apparatus, then ϕ may be written at any point in a proof (we only need to cite that it is a theorem).

T1. $P \rightarrow P$

1	$\frac{P}{\quad}$	A
2	$\frac{P}{\quad}$	R1
3	$P \rightarrow P$	$\rightarrow I$ 1-2

With the the proof of the above theorem, we can now write a wff of the form $\phi \rightarrow \phi$ at any point in a derivation. For example, consider the following (quite uninteresting) derivation:

1	P	P
2	$P \rightarrow P$	T1
3	P	$\rightarrow E$ 1, 2

Next, let's consider the following proof of the theorem $P \rightarrow (Q \rightarrow P)$:

T2. $P \rightarrow (Q \rightarrow P)$

1	$\frac{P}{\quad}$	A
2	$\frac{Q}{\quad}$	A
3	$\frac{P}{\quad}$	R1
4	$Q \rightarrow P$	$\rightarrow I$ 2-3
5	$P \rightarrow (Q \rightarrow P)$	$\rightarrow I$ 1-4

Now with the above theorem, we can write a wff of the form $\phi \rightarrow (\psi \rightarrow \phi)$ at any point in a derivation. In addition, we can use our theorems to derive new theorems. For example, consider the following proof of the theorem $Q \rightarrow (P \rightarrow P)$:

T3. $Q \rightarrow (P \rightarrow P)$

1	$(P \rightarrow P) \rightarrow (Q \rightarrow (P \rightarrow P))$	T2
2	$(P \rightarrow P)$	T1
3	$Q \rightarrow (P \rightarrow P)$	$\rightarrow E2,3$

In the above proof, notice that line 1 is an instance of theorem 2. Theorem 2 is any wff of the form $\phi \rightarrow (\psi \rightarrow \phi)$. In this case, ϕ is $P \rightarrow P$ and ψ is Q . Once theorem 1 is invoked, then we can use $\rightarrow E$ to derive a new theorem. Let's consider another example of a proof of a theorem. Consider the following proof of the theorem $P \rightarrow ((P \rightarrow Q) \rightarrow Q)$:

T4. $P \rightarrow ((P \rightarrow Q) \rightarrow Q)$

1	$\begin{array}{ l} P \end{array}$	A
2	$\begin{array}{ l} \begin{array}{ l} P \rightarrow Q \end{array} \end{array}$	A
3	$\begin{array}{ l} \begin{array}{ l} Q \end{array} \end{array}$	$\rightarrow 1,2$
4	$(P \rightarrow Q) \rightarrow Q$	$\rightarrow I \ 2-3$
5	$P \rightarrow ((P \rightarrow Q) \rightarrow Q)$	$\rightarrow I \ 1-4$

5.7 ADDITIONAL EXERCISES

This section provides additional exercises for practice. The exercises are organized by difficulty. The first set of exercises are easy. The second set of exercises are medium. The third set of exercises are hard. The fourth set of exercises are that do not have any premises (zero-premise proofs) and so will begin with an assumption.

5.7.1 Easy

Exercise 5.61

1. $P \wedge \neg Q, T \vee Q \vdash T$
2. $P \rightarrow Q, Q \rightarrow R, \neg R \vdash \neg P$
3. $A \rightarrow C, A \wedge D \vdash C$
4. $[(A \wedge B) \wedge C] \wedge D \vdash A$
5. $A \vdash B \rightarrow (B \wedge A)$
6. $\neg(A \vee B) \vdash \neg A \wedge \neg B$
7. $\neg(\neg A \vee B) \vdash \neg\neg A \wedge \neg B$
8. $(P \vee Q) \vee R, (T \vee W) \rightarrow \neg R, T \wedge \neg P \vdash Q$
9. $A \vee (\neg B \rightarrow D), \neg(A \vee B) \vdash D$
10. $\neg A \vee \neg B \vdash \neg(A \wedge B)$
11. $A \vee (\neg B \rightarrow D), \neg(A \vee D) \vdash B$

12. $(\neg B \wedge \neg D), \neg(B \vee D) \rightarrow S \vdash S$
13. $(B \leftrightarrow D) \vee S, \neg S, B \vdash D \vee W$
14. $\neg A \rightarrow Q, (A \vee D) \rightarrow S, \neg S \vdash Q$
15. $R \vee R \vdash R$
16. $\neg(\neg R \vee R) \vdash W$
17. $A \rightarrow B, B \rightarrow C, \neg C \vdash \neg A$
18. $A \rightarrow B, B \rightarrow C, A \vdash C$
19. $A \rightarrow B, A \vee C, \neg C \vdash B$
20. $A \rightarrow (B \rightarrow C), A, C \rightarrow D \vdash B \rightarrow D$
21. $A \rightarrow B, C \rightarrow D, B \rightarrow C \vdash A \rightarrow D$
22. $A \rightarrow (B \rightarrow C), A, \neg C \vdash \neg B \vee (W \rightarrow S)$
23. $M \vee (Q \wedge D), M \rightarrow F, (Q \wedge D) \rightarrow F \vdash F \vee S$
24. $P \vdash (\neg Q \vee P) \vee \neg W$
25. $(C \wedge D) \rightarrow (E \wedge Q), C, D \vdash Q$
26. $P \rightarrow (Q \wedge \neg S), F \rightarrow S, R \rightarrow P, R, \neg F \rightarrow \neg M \vdash \neg F \wedge \neg M$
27. $P, Q \wedge F, (P \wedge F) \rightarrow (W \vee R), \neg R, Q \rightarrow S \vdash W \wedge S$
28. $P \rightarrow Q, Q \rightarrow W, P \vdash W$
29. $P \rightarrow \neg(Q \vee \neg R), P \vdash \neg Q$
30. $\neg(\neg P \wedge \neg Q), \neg Q \vdash P$
31. $\neg(W \vee (S \vee M)) \vdash \neg W \wedge \neg M$
32. $\neg(P \vee (S \wedge M)), M \vdash \neg P \wedge \neg S$

5.7.2 Medium

Exercise 5.62

1. $(S \leftrightarrow D) \rightarrow T, P \leftrightarrow (S \wedge D), P \vdash T$
2. $\neg(P \vee Q), \neg(\neg A \vee \neg B) \vdash \neg P \wedge B$
3. $B \rightarrow \neg(S \vee T), \neg(A \vee \neg B), \neg S \rightarrow W \vdash W$
4. $P \vdash \neg P \rightarrow \neg S$
5. $\neg(A \wedge B), B, (\neg A \vee S) \rightarrow \neg(D \wedge T) \vdash \neg D \vee \neg T$
6. $P, (P \vee Q) \rightarrow W, \neg W \vdash \neg(P \vee Q)$
7. $G \rightarrow \mathcal{M} \vdash \neg \mathcal{M} \rightarrow \neg G$
8. $A \rightarrow \neg(B \rightarrow C), \neg B \vdash \neg A$
9. $(B \rightarrow C) \rightarrow \neg(D \rightarrow E), C \vdash \neg E$
10. $(S \vee W) \rightarrow M, (S \wedge T) \leftrightarrow (R \vee P), R \vdash M$
11. $L \wedge (Y \wedge \neg B), P, (P \vee \neg R) \rightarrow Z, [Z \vee (S \wedge T)] \rightarrow W \vdash W \vee \neg M$
12. $R \wedge (S \wedge T), (T \vee M) \rightarrow W, (W \vee \neg P) \rightarrow (A \wedge B) \vdash B$
13. $B \rightarrow D, \neg D \vdash \neg B \vee D$

14. $\neg P \vee (\neg Q \vee R), \neg P \rightarrow (W \wedge S), (\neg Q \vee R) \rightarrow (W \wedge S) \vdash S$
15. $\neg W \rightarrow (R \vee S), M \rightarrow \neg W, \neg S \wedge M \vdash R$
16. $A \leftrightarrow B \vdash A \rightarrow B$
17. $[P \leftrightarrow (L \vee M)] \rightarrow W, P, L \vee M \vdash W$
18. $A, B \vdash A \leftrightarrow B$
19. $A \rightarrow B, \neg A \rightarrow C \vdash \neg B \vee C$
20. $(B \rightarrow C) \rightarrow \neg(D \rightarrow E), C \vdash \neg E$
21. $(S \vee W) \rightarrow M, (S \wedge T) \leftrightarrow (R \vee P), R \vdash M$
22. $L \wedge (Y \wedge \neg B), P, (P \vee \neg R) \rightarrow Z, [Z \vee (S \wedge T)] \rightarrow W \vdash W \vee \neg M$
23. $R \wedge (S \wedge T), (T \vee M) \rightarrow W, (W \vee \neg P) \rightarrow (A \wedge B) \vdash B$
24. $B \rightarrow D, \neg D \vdash \neg B \vee D$
25. $\neg P \vee (\neg Q \vee R), \neg P \rightarrow (W \wedge S), (\neg Q \vee R) \rightarrow (W \wedge S) \vdash S$
26. $\neg W \rightarrow (R \vee S), M \rightarrow \neg W, \neg S \wedge M \vdash R$
27. $(A \vee B) \rightarrow \neg D, \neg(\neg A \vee \neg B) \rightarrow R \vdash D \rightarrow R$
28. $\neg W \rightarrow (R \vee S), M \rightarrow \neg W, \neg S \wedge M \vdash R$
29. $P \vdash \neg\neg P \vee P$
30. $(Q \vee B) \vdash (B \vee Q)$
31. $(A \wedge B) \wedge C \vdash A \wedge (B \wedge C)$
32. $\neg(A \vee B) \rightarrow D, \neg D \vdash A \vee B$
33. $\neg[(A \wedge B) \wedge C] \rightarrow R, \neg R \vdash (B \wedge A) \wedge C$
34. $R \leftrightarrow [(M \rightarrow T) \rightarrow Z], S \wedge [\neg P \wedge (\neg Q \wedge \neg S)] \vdash Z$
35. $P \wedge [S \wedge (R \wedge \neg P)], R \rightarrow [(M \rightarrow T) \rightarrow W] \vdash W$
36. $\neg\neg(\neg R \rightarrow \neg R) \vee B \vdash (R \rightarrow R) \vee \neg\neg B$

5.7.3 Difficult

Exercise 5.63

1. $\neg(A \vee [\neg(B \rightarrow R) \vee \neg(C \rightarrow R)]), \neg A \leftrightarrow (B \vee C) \vdash R$
2. $(A \rightarrow A) \rightarrow B \vdash B$
3. $A \vee B, R \vee \neg(S \vee M), A \rightarrow S, B \rightarrow M \vdash R$
4. $\neg B \wedge C \vdash \neg(B \vee \neg C) \vee (F \rightarrow M)$
5. $A \vdash \neg(\neg A \wedge \neg B)$
6. $\neg\neg(C \vee \neg\neg D), \neg D \vdash \neg(\neg C \wedge F) \vee M$
7. $A \rightarrow B, D \rightarrow E, (B \vee \neg E) \wedge (\neg A \vee \neg B) \vdash \neg A \vee \neg D$
8. $A \rightarrow \neg B, D \rightarrow \neg E, F \rightarrow \neg G, H \rightarrow \neg J, D \rightarrow G, E \rightarrow B, F \vee A \vdash \neg D \vee \neg E$
9. $(A \vee B) \rightarrow (D \vee E), [(D \vee E) \vee F] \rightarrow (G \vee H), (G \vee H) \rightarrow \neg D, E \rightarrow \neg G, B \vdash H$

10. $A \rightarrow B, A \vee (B \vee \neg D), \neg B \vdash \neg D \wedge \neg B$
11. $A \rightarrow (B \rightarrow D), \neg(D \rightarrow Y) \rightarrow \neg K, (Z \vee \neg K) \vee \neg(B \rightarrow Y) \vdash \neg Z \rightarrow \neg(A \wedge K)$
12. $\neg(A \rightarrow B), A \rightarrow D, E \rightarrow B \vdash \neg(D \rightarrow E)$
13. $\neg(P \rightarrow Q) \vdash (P \wedge \neg Q)$
14. $(P \rightarrow Q) \vdash \neg(P \wedge \neg Q)$
15. $P \vee Q \vdash Q \vee P$
16. $P \wedge Q \vdash Q \wedge P$
17. $P \vee (Q \vee R) \vdash (P \vee Q) \vee R$
18. $P \wedge (Q \wedge R) \vdash (P \wedge Q) \wedge R$
19. $P \vee (Q \wedge R) \vdash (P \vee Q) \wedge (P \vee R)$
20. $P \wedge (Q \vee R) \vdash (P \wedge Q) \vee (P \wedge R)$
21. $(P \wedge Q) \rightarrow R \vdash P \rightarrow (Q \rightarrow R)$
22. $P \rightarrow Q \vdash \neg Q \rightarrow \neg P$
23. $S \rightarrow R, P \vee S, P \rightarrow R \vdash R \vee M$
24. $P \rightarrow W, W \rightarrow S, S \rightarrow T, \neg T \vdash \neg P$
25. $A \rightarrow B \vdash (\neg A \vee B) \vee R$
26. $\neg(A \rightarrow B) \vdash (A \wedge \neg B) \vee R$
27. $F \vdash \neg F \rightarrow W$
28. $(\neg A \wedge \neg B) \rightarrow \neg C, \neg A \vdash \neg B \rightarrow \neg C$
29. $(\neg A \wedge \neg B) \rightarrow \neg D, \neg A \vdash D \rightarrow B$
30. $A \leftrightarrow (B \rightarrow C), A, (\neg B \vee C) \rightarrow D, S \rightarrow \neg(D \vee E), (\neg S \leftrightarrow A) \rightarrow L \vdash L$
31. $P \vdash [P \leftrightarrow (P \vee Q)] \leftrightarrow P$
32. $\neg(A \rightarrow R) \vee \neg(R \rightarrow A) \vdash A \vee R$
33. $[(P \rightarrow L) \vee S] \rightarrow (R \leftrightarrow S), L, R \vdash S \leftrightarrow L$
34. $Z \vdash [P \rightarrow (Z \vee \neg Z)] \leftrightarrow [Q \rightarrow (Z \vee \neg Z)]$
35. $Z \leftrightarrow F \vdash (Z \rightarrow F) \wedge (F \rightarrow Z)$
36. $P, (P \vee Q) \rightarrow (Z \wedge S), (Z \wedge P) \leftrightarrow W, (W \rightarrow W) \rightarrow (L \vee M), L \rightarrow O, M \rightarrow O, \neg(O \leftrightarrow P) \vdash U$
37. Try to solve the following only using the introduction and elimination rules (and R): $\neg E \vee D, (E \rightarrow F) \leftrightarrow G, D \rightarrow F \vdash L \rightarrow G$

5.7.4 Zero-Premise

Exercise 5.64

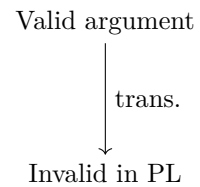
1. $\vdash P \rightarrow P$
2. $\vdash P \vee \neg P$
3. $\vdash \neg(P \wedge \neg P)$
4. $\vdash (A \wedge B) \rightarrow A$
5. $\vdash A \rightarrow \neg(B \wedge \neg B)$
6. $\vdash \neg A \rightarrow (A \rightarrow B)$
7. $\vdash A \rightarrow (A \vee \neg A)$
8. $\vdash P \rightarrow (\neg Q \rightarrow P)$
9. $\vdash [(P \rightarrow Q) \rightarrow P] \rightarrow P$
10. $\vdash \neg[(A \rightarrow \neg A) \wedge (\neg A \rightarrow A)]$
11. $\vdash A \rightarrow (A \wedge A)$
12. $\vdash [(A \rightarrow B) \wedge (A \rightarrow D)] \rightarrow [A \rightarrow (B \wedge D)]$
13. $\vdash \neg P \rightarrow \neg[(P \rightarrow Q) \rightarrow P]$
14. $\vdash [P \rightarrow (Q \rightarrow R)] \rightarrow [(P \rightarrow Q) \rightarrow (P \rightarrow R)]$
15. $\vdash (A \rightarrow B) \vee (B \rightarrow D)$
16. $\vdash A \rightarrow [A \rightarrow (A \vee A)]$
17. $\vdash \neg(P \wedge \neg P)$
18. $\vdash [(\neg A \vee B) \wedge (\neg A \vee D)] \rightarrow [\neg A \vee (B \wedge D)]$
19. $\vdash [(P \rightarrow Q) \rightarrow R] \rightarrow (\neg R \rightarrow P)$
20. $\vdash (A \rightarrow B) \rightarrow [\neg(B \wedge D) \rightarrow \neg(D \wedge A)]$
21. $\vdash (A \wedge B) \rightarrow A$
22. $\vdash \neg\neg A \rightarrow A$
23. $\vdash A \rightarrow (B \rightarrow A)$
24. $\vdash (A \vee B) \rightarrow [(\neg A \vee B) \rightarrow B]$

Part III

Predicate Logic

The language of **PL** has at least four strengths and one weakness. Let's begin by considering its strengths. First, the language of **PL** offers a more precise way of defining the informal idea of a conclusion "following from" its premises and the informal idea of a valid argument. Second, for any argument that is valid in **PL**, there is a corresponding valid argument in English. This means that some of the arguments that we wish to represent and the reasoning we do in English can be represented in the more precise language of **PL**. Third, there are ways to *identify* valid arguments in **PL**. This we saw in chapters 3 and 4 with the use of the truth-table and truth-tree methods. Fourth, not only do we have a way to identify good arguments but also *construct* them using a proof system (**PD**). That is, we have a codified set of rules that justify various derivations or moves forward in arguments.

While there are many weaknesses of **PL**, one of the main weaknesses is that there are valid arguments in English that are not identified as valid in **PL**. Let's consider this weakness in context before looking at a specific example. First, recall from chapter 1 that one of the goal's of logic is to identify good and bad arguments. In order to do this, we (1) developed a formal language that is expressive enough to capture the validity of arguments and (2) developed methods (table and tree) for identifying valid arguments in that formal language. What seems evident is that if an argument is valid in our formal language, then it is valid in English. But this does not seem to satisfy the logician's larger goal of identifying good and bad arguments since there are valid arguments in English that are not identified as valid in our formal language.



Next, let's illustrate this weakness with an example. Consider the following argument:

- P1: All humans are mortal.
- P2: Socrates is a human.
- C: Socrates is mortal.

Intuitively, this is a valid argument. If each and every human being is mortal and Socrates is a member of the larger class of human beings, then Socrates is mortal. This is a valid argument in English. However, this argument, when translated into **PL**, the argument is not valid argument in **PL**. Let H stand for "All humans are mortal", S stand for "Socrates is a

man", and M stand for "Socrates is mortal". The result of our translation is the following argument:

- P1: H
- P2: S
- C: M

Notice that the above argument is not valid in **PL**. The interpretation $\mathcal{I}(H) = T, \mathcal{I}(S) = T, \mathcal{I}(C) = F$ is such that the premises of the above argument are true and the conclusion is false. Thus, the argument is invalid in **PL**, but not valid in English. If the goal of logic is to develop criteria for identifying and constructing good and bad arguments, then some modification or extension of **PL** is needed.

In this chapter, we take a step toward this goal by building our the language of propositional logic. We extend this language by creating a new language that we will call "the language of predicate logic" (also known as "the logic of relations", or "quantificational logic", or "first-order predicate logic"). We will abbreviate this language as **QL**. One key aspect of the language of **QL** is that it is more expressive than the language of **PL** in that it allows us to capture the validity of arguments that are not valid in **PL** (like the one above). **QL** is capable of doing this because it allows us to analyze sentences at the sub-sentential level. That is, it allows us to analyze the logical relationships between *parts* of sentences or propositions.

6.1 QL SYMBOLS

First, we begin with the symbols of *QL*:

1. **names**: Lower case letters, a through v with or without numerical subscripts.
2. **n-place predicates**: Upper case letters, A through Z with or without numerical subscripts.
3. **variables**: Lower case letters, w through z with or without numerical subscripts.
4. **operators**: $\neg, \wedge, \vee, \rightarrow, \leftrightarrow$
5. **parentheses**: $(,)$
6. **quantifiers**: $\forall, \exists$

At this point, these symbols do not have meaning nor do we have any rules for how to combine said symbols together into grammatically correct expressions. We will turn to these issues in the following sections.

Exercise 6.65

Determine what kind of symbol the following symbols are:

1. a
2. c
3. g
4. z
5. x
6. $\rightarrow$
7. $\forall$
8. $\exists$

6.2 QL SYNTAX

With the symbols of **QL** specified, we now turn to the syntax of **QL**. That is, the proper way of combining these symbols.

6.2.1 QL: Formation Rules

The central concept for **QL** syntax is the idea of a *well-formed formula* (wff). In rough terms, a wff is a grammatically correct expression in **QL**. That is, it is a combination of **QL** symbols put together in the "correct" order. What defines whether something is or is not put together in the "correct" order are a set of grammatical rules known as "formation rules".

Definition 6.2.1: Well-formed Formula in QL

A well-formed formula (or wff) ϕ is any combination of **QL** symbols capable of being constructed through some combination of the formation rules in [Definition 6.2.2](#).

It is clear that the key idea for determining whether something is a wff are the formation rules. In what follows, we will first define these formation rules and then illustrate how they can be used to determine whether something is or is not a wff.

Definition 6.2.2: Formation Rules for QL

Let P be a placeholder for an n -place predicate and x be a placeholder for a variable. The formation rules for **QL** are as follows:

1. an n -place predicate P followed by n terms (names or variables) is a wff in **QL**, where $n \geq 0$.

2. If ϕ is a wff in **QL**, then $\neg(\phi)$ is a wff in **QL**.
3. If ϕ and ψ are wffs in **QL**, then $(\phi \wedge \psi)$, $(\phi \vee \psi)$, $(\phi \rightarrow \psi)$, and $(\phi \leftrightarrow \psi)$ are wffs in **QL**.
4. If ϕ is a wff containing a variable x , then $(\forall x)(\phi)$ and $(\exists x)(\phi)$ are wffs in **QL** provided that there is no other quantifier containing x in ϕ (that is, $(\exists x)$ and $(\forall x)$ are not already in ϕ).
5. Nothing else is a wff in **QL** except that which can be formed by repeated applications of the above.

With the formation rules defined, let's explain and offer illustrations for each rule. Rule 1 states that an n -place predicate followed by n terms is a wff. Thus, if P is a one-place predicate, then Pa would be a **QL**-wff. Similarly, Px would also be a **QL**-wff. Both are wffs since if we write a one-place predicate, then the creation of a wff requires one term (name or variable) written to its immediate right. If P is a one-place predicate, then both P and Pab are not **QL**-wffs. P by itself does not have a term to its immediate right and Pab has two terms (rather than one) to its immediate right.

Suppose that Q is a two-place predicate. Provided there are two terms to the right of Q , then the resulting formula is a wff. So, Qaa , Qab , Qba , Qax , Qxx , and Qzy are all wffs. However, Qa is not a wff since there is only one term to the right of Q . Similarly, $Qabc$ is not a wff since there are three terms to the right of Q .

Finally, if an n -place predicate is a 0-place predicate, then it is a propositional letter. For example, suppose S is a zero-place predicate. In this case, S is a wff. However, Sa is not a wff since there is a term to the right of S . We will treat 0-place predicates as PL propositional letters.

Rules 2 and 3 are the same as the formation rules from **PL**. They specify that if there is a wff ϕ , then $\neg(\phi)$ is a wff and if there are two wffs ϕ and ψ , then $(\phi \wedge \psi)$, $(\phi \vee \psi)$, $(\phi \rightarrow \psi)$, and $(\phi \leftrightarrow \psi)$ are wffs.

Rule 4 tends to be the most difficult to understand and different logic textbooks present this rule differently. In our presentation, we say that if there is a wff ϕ and ϕ has a variable x , then a quantifier with an x to its immediate right can be prefixed to (put to the left of) ϕ . Thus, if ϕ has the variable y in it, then formulas of the form $(\exists y)(\phi)$ and $(\forall y)(\phi)$ are wffs. More concretely, suppose F is a one-place predicate. If F is a one-place predicate, then Fx is a wff. Since the wff Fx contains a variable x , a quantifier with x to its immediate right can be prefixed to Fx . That

is, if Fx is a wff, then $(\forall x)(Fx)$ and $(\exists x)(Fx)$ are wffs.

Let's consider the use of the formation rules to show that $((\forall y)(Pyy) \wedge (\exists x)(Fx))$ is a wff.

1. Pyy is a wff (rule 1)
2. Fx is a wff (rule 1)
3. If Pyy is a wff, then $(\forall y)(Pyy)$ is a wff (line 1, rule 4)
4. If Fx is a wff, then $(\exists x)(Fx)$ is a wff (line 2, rule 4)
5. If $(\forall y)(Pyy)$ is a wff and $(\exists x)(Fx)$ is a wff, then $((\forall y)(Pyy) \wedge (\exists x)(Fx))$ is a wff (line 3, 4, rule 3)

Now that we have considered some preliminary examples, let's turn to the condition placed on the rule (4). The condition states that if ϕ is a wff containing a variable x , then $(\forall x)(\phi)$ and $(\exists x)(\phi)$ are wffs *provided* there is not already a quantifier with x to its immediate right in ϕ . For example, consider the wff $(\exists x)Lxx$. While $(\exists x)(Lxx)$ is a wff with variables, we might think we could use the fourth formation rule to construct the wff $(\forall x)((\exists x)(Lxx))$. This would not be permitted given the restriction on this rule. Since $(\exists x)(Lxx)$ contains the variable x , we cannot use the fourth formation rule to construct $(\forall x)((\exists x)(Lxx))$.

Similarly, the condition prevents us from creating the wff $(\exists x)((\exists x)(Lxx))$.

Finally, rule (5) provides closure for our list of rules. It states that there are no other rules that allow for the creation of wffs. That is, the only wffs that can be created are those that can be created by repeated applications of the above rules.

Exercise 6.66

Using the formation rules, show that the following propositions are wffs in RL, where 'Pxy' is a two-place predicate while 'Rx' and 'Zx' are one-place predicates:

1. $(Ra \wedge Paa)$
2. $(\forall x)(Pxx)$
3. $(\exists x)(Zx)$
4. $\neg((\exists y)(Pyy))$
5. $(\neg((\forall x)(Pxx)) \wedge (\exists x)(Zx))$

6.2.2 Literal wffs

In later chapters, it will be useful to refer to a wff that is either a n-place predicate followed by n-terms or the negation of such a wff. Let's define

this type of a predicate logic wff as a *literal wff*. That is, a *literal* or *literal wff* is a wff that is either a n-place predicate followed by n-terms or the negation of such a wff.

Definition 6.2.3: Literal wff

A literal wff is a wff ϕ that is either a n-place predicate followed by n-names or the negation of such a wff $\neg(\phi)$.

To illustrate, if P is a 1-place predicate, then Pa and $\neg(Pa)$ is a literal wff. Similarly, if R is a 2-place predicate, then Rax and Rab are both literal wffs; so are $\neg(Rax)$ and $\neg(Rab)$. In contrast, $(\exists x)(Fx)$ is not a literal wff as it is not simply an n-place predicate terms followed by n-terms, nor is it the negation of such a wff. Similarly, $(Pa \wedge Qa)$ is also not a literal wff.

Exercise 6.67

State which wffs are literal wffs and which are not literal wffs.

1. Fb
2. $\neg(Fb)$
3. $\neg(\neg(Fb))$
4. $(\forall x)(Fx)$
5. $Fb \wedge Fq$

6.2.3 Scope of Quantifiers and Formula Simplification

Quantifiers have scope. To determine the scope of a quantifier, we can construct the wff containing the quantifier using the formation rules. The scope of the quantifier is the wff that is constructed when the quantifier is first added to the construction. For example, consider the construction of $(\forall x)(Fx)$. First, we start with Fx . Then, we add the quantifier $(\forall x)$ to the left of Fx . The result of this construction is $(\forall x)(Fx)$. The scope of the quantifier is the wff $(\forall x)(Fx)$.

To further illustrate, consider the wff $((\exists x)(Fx) \wedge Gx)$. The construction of the wff is as follows:

1. Fx is a wff.
2. $(\exists x)(Fx)$ is a wff.
3. $((\exists x)(Fx) \wedge Gx)$ is a wff.

Notice that the existential quantifier is first introduced at line 2. As such, the scope of the existential quantifier is $(\exists x)(Fx)$. That is, the scope is

not the entire wff but only the subformula $(\exists x)(Fx)$. Let's contrast this example with $(\exists x)((Fx \wedge Gx))$. The construction of this wff is as follows:

1. Fx is a wff.
2. $(Fx \wedge Gx)$ is a wff.
3. $(\exists x)((Fx \wedge Gx))$ is a wff.

Notice that in this example, the existential quantifier is first introduced at line 3. As such, the scope of the existential quantifier is the entire wff $(\exists x)((Fx \wedge Gx))$.

While the scope of the quantifiers can always be determined by constructing the wff, generally people do not construct every wff to determine the scope of the quantifier. Instead, people make use of heuristics (shorthand methods) for identifying the scope. Let's consider three such heuristics.

First, whenever a quantifier is introduced in constructing a wff, parentheses are also introduced. The left parenthesis occurs immediately to the right of the quantifier and the right parenthesis occurs at the end of the wff. The scope of the quantifier is thus the quantifier itself and the wff between the left and right parentheses. For example, consider the wff $(\exists x)((Fx \wedge Gx))$. Notice that the existential quantifier $(\exists x)$ is at the beginning of the wff. To the immediate right of the existential quantifier is a left parenthesis and this parenthesis has its closing counterpart at the end of the wff. The scope of the quantifier $(\exists x)$ extends to the wff between the left and right parentheses. In this case, the scope of the existential quantifier is the entire wff $(\exists x)((Fx \wedge Gx))$. In contrast, the scope of the quantifier in $((\exists x)(Fx) \wedge Gx)$ is only (Fx) as the left parenthesis closes before the end of the wff.

Tip: The scope of $\forall, \exists$ is the $\forall, \exists$ and the wff in parentheses to its immediate right.

Before introducing the other heuristics, let's introduce two rules for simplifying the writing of wffs. First, we will make use of the simplification conventions from propositional logic. These include the general principle that we do not need to include the presence of parentheses that do not play a role in disambiguating the scope of an operator. So, for example, the parentheses in $(P \wedge Q)$ can be dropped and the wff can be rewritten as $P \wedge Q$. Similarly, $\neg((P \wedge Q))$ can be simplified to $\neg(P \wedge Q)$. In addition, we will make use of the convention that the scope of the operator for negation is the smallest subformula to its immediate right unless parentheses are used to extend its scope. So, for example, the scope of negation in $\neg P \wedge Q$ is $\neg P$, while the scope of negation in $\neg(P \wedge Q)$ is the entire wff. The final simplification for parentheses in propositional logic was that we can omit the presence of parentheses when negations are stacked. So, for example, $\neg(\neg(\neg(P)))$ can be simplified as $\neg\neg\neg P$.

Let's add a simplification convention for wffs containing quantifiers. Our convention for simplification will also allow for omitting certain parentheses in wffs containing quantifiers. Just as we said that the scope of the negation operator is the negation itself and the smallest subformula to its immediate right unless parentheses are used to extend the scope, we will say that *the scope of a quantifier is the quantifier itself and the smallest subformula to its immediate right unless parentheses are used to extend the scope*. Secondly, just as we can omit the presence of parentheses when negations are stacked, we will also omit parentheses when quantifiers are stacking.

Let's illustrate these simplifications with a few examples. First, since the scope of $(\exists x)(Fx)$ is $(\exists x)$ plus the smallest subformula to its immediate right, which is (Fx) , we can simplify $(\exists x)(Fx)$ as $(\exists x)Fx$, we can simplify the writing of this wff to $(\exists x)Fx$. This is because our convention is that the quantifier applies to the smallest subformula to its immediate right (unless parentheses are used to extend its scope). With this convention $(\exists x)(Fx)$ and $(\exists x)Fx$ to be read the same way, so we can simplify the writing of the wff. Second, since the scope of $(\exists x)((Fx \wedge Gx))$ is the entire wff (including the conjunction), the additional set of parentheses around the $Fx \wedge Gx$ does not play a role in disambiguating the scope of the quantifier. Therefore, the wff can be simplified to $(\exists x)(Fx \wedge Gx)$.

Next, in the case where quantifiers are stacked, the presence of certain parentheses can be omitted. For example, consider the wff $(\exists x)((\forall y)(Fx \wedge Gy))$. The scope of the existential quantifier is the entire wff, while the scope of the universal quantifier is $(\forall y)(Fx \wedge Gy)$. Let's allow the parentheses around the universal quantifier to be omitted. That is, let's allow the wff to be written as $(\exists x)(\forall y)(Fx \wedge Gy)$. In other words, $(\exists x)(\forall y)(Fx \wedge Gy)$ and $(\exists x)((\forall y)(Fx \wedge Gy))$ can be read as the same wff.

These procedures for simplification allow for introducing a second way of articulating how the scope of a quantifier is identified. The scope of a quantifier is the quantifier itself and the smallest subformula to its immediate right unless parentheses are used to extend the scope. So, for example, consider the smallest subformula to the immediate right of $(\exists x)Fx \wedge Gx$ is Fx . Therefore, the scope of $(\exists x)$ is $(\exists x)Fx$. In contrast, notice that in $(\exists x)(Fx \wedge Gx)$, there are parentheses to the right of the quantifier $(\exists x)$. This indicates that the scope of the quantifier for $(\exists x)$ extends to the point where the left parenthesis closes. Since the left parenthesis closes at the end of the wff, the scope of $(\exists x)$ is the entire wff: $(\exists x)(Fx \wedge Gx)$.

In addition, these simplification procedures also allow for introducing a

third way to identify the scope of a quantifier. Namely, a parallel can be drawn between the scope of the quantifier and the scope of the negation operator. Recall that the scope of the negation operator is the smallest subformula to its immediate right unless parentheses are used to extend the scope. For example, notice the distinction between $\neg P \wedge Q$ and $\neg(P \wedge)$. In $\neg P \wedge Q$, the negation has narrow scope. It only applies to $\neg P$. In contrast, in $\neg(P \wedge Q)$, the negation has wide scope. It applies to the entire wff. This is parallel to the contrast between $(\exists)Fx \wedge Gx$ and $(\exists x)(Fx \wedge Gx)$. In $(\exists)Fx \wedge Gx$, the quantifier has narrow scope, only extending to $(\exists x)Fx$ while in $(\exists x)(Fx \wedge Gx)$, the quantifier has wide scope, extending to the entire wff.

Let's notice a second parallel between the scope of negation and the scope of quantifiers. In wffs where negations are stacked, the further left the negation has in the stacking, the more scope it has. So, for example, in the wff $\neg\neg\neg P$, the scope of leftmost negation is $\neg\neg\neg P$ (the entire wff), the scope of the negation in the middle is $\neg\neg P$, and the rightmost negation has the least scope (only extending to $\neg P$). The same is true with respect to wffs with stacked quantifiers. For example, in the wff $(\exists x)(\forall y)(\exists z)Rxyz$, the scope of the leftmost quantifier is $(\exists x)(\forall y)(\exists z)Rxyz$ (the entire wff), the scope of the quantifier in the middle is $(\forall y)(\exists z)Rxyz$, and the rightmost quantifier has the least scope (only extending to $(\exists z)Rxyz$).

To conclude, let's consider the wff $(\exists x)(\forall y)(Px \rightarrow Ry) \wedge (\exists z)Pz$. Let's start by constructing the wff using the formation rules. The construction is as follows:

1. Px, Ry, Pz are wffs.
2. $Px \rightarrow Ry$ is a wff.
3. $(\forall y)(Px \rightarrow Ry)$ is a wff.
4. $(\exists x)(\forall y)(Px \rightarrow Ry)$ is a wff.
5. $(\exists z)Pz$ is a wff.
6. $(\exists x)(\forall y)(Px \rightarrow Ry) \wedge (\exists z)Pz$ is a wff.

The scope of each quantifier is the wff constructed when the quantifier is introduced into the construction of the wff. The scope of $(\exists x)$ is $(\exists x)(\forall y)(Px \rightarrow Ry)$. Since there is a case of stacked quantifiers, the wff is similar to $\neg\neg(P \rightarrow R) \wedge P$ where the scope of the leftmost negation is $\neg\neg(P \rightarrow R)$. In the case of our quantified wff, $(\exists x)$ has $(\forall y)$ within its scope. Notice that since the parentheses associated with $(\exists x)$ do not extend to the end of the wff, the scope of $(\exists x)$ does not include $(\exists z)Pz$. Next, the scope of $(\forall y)$ is $(\forall y)(Px \rightarrow Ry)$. The scope of $(\forall y)$ does not include $(\exists x)Px$. Again, this is similar to how the rightmost negation in

$\neg\neg(P \rightarrow R)$ is $\neg(P \rightarrow R)$ rather than $\neg\neg(P \rightarrow R)$. Finally, the scope of $(\exists z)$ is $(\exists z)Pz$.

Exercise 6.68

In the following wffs, determine the scope of each quantifier.

1. $(\forall x)Px$
2. $(\exists x)Rx$
3. $(\exists x)\neg Rx$
4. $(\forall x)Px \rightarrow Rx$
5. $(\forall x)(Px \rightarrow Rxy)$
6. $\neg(\forall x)(\exists y)Rxy$

6.2.4 The Main Operator

In the previous section, the scope of a quantifier was clarified and illustrated. In this section, we utilize this idea to define and illustrate the notion of the main operator of a wff.

As in proposition logic, the main operator of a wff can be defined in several equivalent ways. First, the main operator of a wff ϕ can be defined with respect to the construction of a wff. That is, the main operator of ϕ is the last operator introduced in the construction of ϕ . Second, it can be defined as the operator that has the widest (or most) scope. Third, it can be defined as the operator whose scope is the entire wff.

Let's consider two examples. First, consider the wff $(\exists x)Px \wedge (\exists y)Qy$. The construction of this wff is as follows:

1. Px and Qy are wffs.
2. $(\exists x)Px$ is a wff.
3. $(\exists y)Qy$ is a wff.
4. $(\exists x)Px \wedge (\exists y)Qy$ is a wff.

Notice that (1) the last operator introduced in the construction was the conjunction operator, (2) the scope of the conjunction operator has more scope than the other operators, and (3) the scope of the conjunction operator is the entire wff. As such, the main operator of $(\exists x)Px \wedge (\exists y)Qy$ is the conjunction operator.

Next, consider the wff $(\exists x)(\forall y)((Px \wedge Qy) \rightarrow (\exists z)Pz)$. The construction of this wff is as follows:

1. Px, Qy, Pz are wffs.

2. $(\exists z)Pz$ is a wff.
3. $(Px \wedge Qy)$ is a wff.
4. $((Px \wedge Qy) \rightarrow (\exists z)Pz)$ is a wff.
5. $(\forall y)((Px \wedge Qy) \rightarrow (\exists z)Pz)$ is a wff.
6. $(\exists x)(\forall y)((Px \wedge Qy) \rightarrow (\exists z)Pz)$ is a wff.

Notice that in the above example, the last operator introduced in the construction of the wff is $(\exists x)$. Second, notice that this operator has the scope of all other operators within its scope. Third, notice that the scope of this operator is the entire wff. Therefore, the main operator of $(\exists x)(\forall y)((Px \wedge Qy) \rightarrow (\exists z)Pz)$ is $(\exists x)$.

Exercise 6.69

Determine the main operator of the following wffs:

1. $(\forall x)(Px \vee Qx)$
2. $(\exists y)(Py) \wedge (\exists z)(Pz)$
3. $\neg(\exists y)(Py) \wedge \neg(\exists z)(Pz)$
4. $\neg(\exists y)(Py) \vee \neg(\exists z)(Pz)$
5. $\neg(\exists y)(Py \wedge (\forall z)Pz)$
6. $(\forall x)(Px \rightarrow (Qx \wedge Mx))$
7. $(\exists x)Px \wedge ((\exists y)Py \vee (\forall z)Qz)$
8. $\neg(\forall x)((Pb \rightarrow Qb) \leftrightarrow Px)$
9. $(\forall x)(\exists y)\neg(\forall z)(Px \rightarrow Qyz)$
10. $(\forall x)(\forall y)(\forall z)(Pxyz \wedge Rxyz) \rightarrow (\exists x)Px$

6.2.5 Free Variables and Bound Variables

Quantifiers always appear with a variable to their immediate right. For this reason, it is commonplace to refer to the quantifier along with the variable that occurs to its immediate right as the quantifier itself. When an occurrence of a quantifier that contains a variable x has within its scope other occurrences of x , those variables are said to be *bound variables*.

Definition 6.2.4: Bound Variable

A variable x is bound if and only if it is in the scope of an occurrence of a quantifier that has that variable to its immediate right.

For example, consider the following wff: $(\forall x)Fx$. The variable x in Fx is a bound variable in the scope of quantifier that has x to its immediate right: $(\forall x)$. It is common to say that the variable is "bound by" the

quantifier. And so, x in Fx is bound by the $(\forall x)$ in $(\forall x)Fx$. Another way to put this is that if we want to determine whether a variable is *bound* or *not*, we check two things:

1. Is the variable in the scope of a quantifier?
2. Does the quantifier have the variable to its immediate right?

If the answer to both of these questions is "yes", then the variable is bound. If the answer to either of these questions is "no", then the variable is not bound. If a variable is not bound, then it is a *free variable*.

Definition 6.2.5: Free Variable

A variable is *free* if it is not bound.

Let's consider a few examples to further clarify the distinction between a bound versus free variable. Consider the following wff: $(\exists x)Fx \wedge Gx$. There are two occurrences of the variable x and so there is the question about each of these variables as to whether they are bound or free. In the case of the x in Fx , this variable is bound as (1) it is in the scope of a quantifier and (2) the quantifier has x to its immediate right. However, what about x in Gx ? To determine whether this occurrence is bound or free, we again check to see if that x is in the scope of a quantifier that has x to its immediate right. In this case, notice that x in Gx is not in the scope of a quantifier. Since the variable x is not in the scope of a quantifier that has x to its immediate right, the x in Gx is not bound. And, since x in Gx is not bound, it is a free variable.

Second, consider $(\exists x)(Rx \wedge Lx)$. In this wff, there are two occurrences of the variable x : one in Rx and one in Lx . Notice that the x in Rx is in the scope of a quantifier that has x to its immediate right. Therefore, x in Rx is bound. Similarly, since the scope of the quantifier $(\exists x)$ is the entire wff $(\exists x)(Rx \wedge Lx)$, the x in Lx is also bound.

Third, consider $(\forall x)Lxy$. In this wff, there are two variables: the x and y in Lxy . We can thus consider whether each variable is bound or free. First, notice that both variables are in the scope of the quantifier $(\forall x)$. However, notice that only x in Lxy is in the scope of a quantifier that has x to its immediate right. As such, only x is bound. The y in Lxy is not bound as it is not in the scope of a quantifier that has y to its immediate right. As such, y is a free variable.

Exercise 6.70

Identify the bound and free variables in the following wffs:

1. $(\exists x)(Rx \rightarrow Ga)$
2. $(\forall y)(Mx \rightarrow Py) \vee Ga$
3. $(\exists x)Mx \vee (\forall x)Rx$
4. $Rz \wedge (\forall z)((Rz \wedge My) \rightarrow Qz)$
5. $Pa \wedge (\exists w)(Vw \wedge Lx)$
6. $\neg(\exists x)(Fx) \wedge (\exists x)(Fx)$
7. $(\exists x)(\neg Fx) \wedge (\exists x)(Fx)$
8. $(Tb \wedge Qa) \rightarrow (\forall x)(Fx \rightarrow Gy)$
9. $Tb \wedge \neg Tb$
10. $Rx \rightarrow (\forall x)Px$

6.2.6 Open and closed formulas

In this section, we will introduce the notion of an open and closed formula. When a wff contains a free variable, it is an *open formula*.

Definition 6.2.6: Open formula

An open formula is a wff containing at least one free variable.

For example, Fx , $(\exists x)Fx \wedge Fx$, and $(\forall x)Lxy$ are all open formulas as each contains a free variable. In contrast, when a wff does not contain any free variables, the wff is a *closed formula*.

Definition 6.2.7: Closed formula

A closed formula is any wff that does not contain a free variable.

For example, $(\forall x)Fx$, $(\exists x)Fx \wedge (\exists x)Fx$, and $(\forall x)(\forall y)Lxy$ are all closed wffs as there are no free variables in each wff.

Let's consider some of these examples in more detail. Compare Fx with $(\forall x)Fx$. In Fx , the x is not within the scope of a quantifier, and so the variable x is free. Since it is free, the wff is an open wff. In contrast, $(\forall x)Fx$ is a closed wff as the x in Fx is within the scope of a quantifier that has x to its immediate right ($\forall x$). Therefore, the variable is bound. Since x in Fx is the only variable, there are no other variables that are free. Therefore, $(\forall x)Fx$ is a closed wff.

Next, let's compare $(\exists x)Fx \wedge Fx$ and $(\exists x)Fx \wedge (\exists x)Fx$. In the case of

$(\exists x)Fx \wedge Fx$, there are two instances of the variable x . The first instance of x is bound as it is in the scope of $(\exists x)$. However, notice that the second instance of x is not in the scope of a quantifier. As such, it is a free variable and so $(\exists x)Fx \wedge Fx$ is an open wff. In contrast, notice that both instances of x in $(\exists x)Fx \wedge (\exists x)Fx$ are bound. The leftmost x is bound by the first quantifier while the rightmost x is bound by the second quantifier. Since all of the variables are bound, there are no free variables, and so the wff is a closed wff.

Next, let's compare $(\forall x)Lxy$ and $(\forall x)(\forall y)Lxy$. In $(\forall x)Lxy$, the variable x is bound by the quantifier $(\forall x)$. However, notice that while y is in the scope of the quantifier $(\forall x)$, it is not in the quantifier that has y to its immediate right. Therefore, y is free and so $(\forall x)Lxy$ is an open wff. In contrast, x and y are both bound in $(\forall x)(\forall y)Lxy$ since x is bound by the quantifier $(\forall x)$ and y is bound by the quantifier $(\forall y)$. Since all of the variables in this wff are bound, the wff itself is a closed wff.

Finally, let's consider a wff that sometimes gives new students a pause. Consider the wff Pb . Some individuals are inclined to say that Pb is an open wff. They reason in one of two ways.

First, they reason that since b is not within the scope of a quantifier for b , it follows that b is a free variable and so the wff Pb is an open wff. This is not correct since notice that b is not a variable. Rather, it is a name. Since a closed wff is defined as a wff that contains no free variables and Pb contains no free variables (it only contains the predicate P and b), Pb is a closed formula.

Second, they reason that a closed wff is a wff containing *only* bound variables. Since b is not bound, it follows that Pb is an open wff. This is also not correct. The source of the confusion stems from a misunderstanding of the definition of a closed wff. A closed wff is not a wff that *only* contains bound variables. Rather, it is a wff that contains *no* free variables. Again, since Pb contains no free variables, it is a closed wff.

Exercise 6.71

State whether the following wffs are open or closed. If the wff is open, identify any free variables.

1. Px
2. Lxx
3. $(\forall x)Lxx$

4. $(\forall x)Lxy$
5. $(\forall x)Px \wedge Pa$
6. Lab
7. $(\exists x)(\forall y)(\forall z)(Pxy \rightarrow Lz)$
8. $(\exists x)(\forall y)(\forall z)Pxy \rightarrow Lz$

6.3 QL SEMANTICS

In the semantics of **PL**, single propositional letters are assigned truth values (T or F) by an **interpretation function** while wffs in **PL** are assigned truth values by a **valuation function**. The semantics of **QL** is more complex in that the elementary symbols of **QL** express *parts* of propositions rather than complete propositions.

The articulation of the semantics of **QL** requires three things:

1. the domain of discourse
2. an interpretation function
3. a valuation function

(1) and (2) are known as a “model”. What is a model?

Definition 6.3.1: model

A **QL**-model ($\mathcal{M}$) is a two-part structure $M = \langle \mathcal{D}, \mathcal{I} \rangle$ where

- $\mathcal{D}$ is a non-empty set
- $\mathcal{I}$ is an interpretation function that
 1. assigns sets of n-tuples from $\mathcal{D}$ to n-place predicates
 2. assigns objects of $\mathcal{D}$ to names

Notice that a model consists of two parts. The first part is $\mathcal{D}$, which is called the “domain”. The second part is $\mathcal{I}$ which is an interpretation function. Let’s consider each part of the model in turn.

6.3.1 Domain of discourse

The first part of the model is the **domain of discourse**.¹ This part of the model is abbreviated as $\mathcal{D}$. A $\mathcal{D}$ is simply a set (or collection) of items. The items of the domain are known as its **members** or **items**. The domain can consist of anything that can be put into a set. Perhaps more intuitively, the domain of discourse is just a collection of all of the things

¹Sometimes simply called the “universe of discourse” or simply “the domain”.

we wish to talk about. You can think about it as a collection of objects in the world.

Often the domain of discourse is more restricted than *all* of the objects in the world. We can specify what items are in the $\mathcal{D}$ by writing $\mathcal{D} :$ and then either

1. individually listing each item in the domain
2. indicating some property that belongs to every item in the domain
3. indicating some method for determining every item in the domain.

Let's consider some illustrations of these different ways of specifying a domain. First, let's consider the method of listing each item in the domain (the method of enumeration). Suppose my family consists of the following three people: David, Liz, Renna. The domain can be specified by listing each item:

- $\mathcal{D} : \text{David, Liz, Renna}$

Similarly, if we wanted to talk about only the even integers from 2-10, we could list each of these items as follows:

- $\mathcal{D} : 2, 4, 6, 8, 10$

Second, let's consider the method of indicating some property that belongs to every item in the domain (the method of set-building). In the case of my family, we could indicate that each item in the domain is a member of my family. In the case of the even integers, we could indicate that each item in the domain is an even integer between 2 and 10.

- $\mathcal{D} : \{x \mid x \text{ is an immediate member of my family}\}$
- $\mathcal{D} : \{x \mid x \text{ is an even integer between 2 and 10}\}$

The method of indicating some property has at least three advantages over the method of listing each item. The first is that it is more economical to indicate a property than to list each item when there the size of the domain is large (or infinite). For example, it is easier to specify the property of being an even integer from 2-100 than to list each item. The second advantage is that the method of indicating some property allows us to specify an infinite domain. For example, we can specify the domain of all even integers as follows: $\mathcal{D} : \{x \mid x \text{ is an even integer}\}$. In short, this method is more expressive power when it comes to specifying a domain. The third is that it allows for the possibility of specifying a domain where we do not know all of the items in the domain. For example, suppose we wanted to specify a domain consisting of all the human beings on earth.

We could specify this domain as follows: $\mathcal{D} : \{x \mid x \text{ is a human being}\}$. Notice that we do not know all of the human beings on earth (some of them may be living in remote locations "off the grid"), but we can still specify the domain in terms of a property that all of the items in the domain have.

The third and final way of specifying a domain is probably the least mentioned, least used, and is possibly not distinct from the second way of specifying a domain. On this third way, the domain is specified through a recipe or set of instructions of how to determine each item in the domain. For example, suppose I have buried a chest of golden coins in the woods. I have left a map or set of instructions to the treasure in my will. I could have listed each item in the domain (each individual golden coin) or I could have indicated some property that each item in the domain has (being a golden coin). However, I have chosen to specify the items belonging to the domain in terms of a set of instructions for finding the items in the domain. In this case, the domain is specified as follows: $\mathcal{D} : \{x \mid \text{where } x \text{ is what you will discover if you enter such-and-such woods, walk 40 yards, take a left by a tree with a star carved into it, then take 40 steps, and dig 2 feet.}\}$.

One advantage over the method of specifying a common property of items (the set-builder method) is that this method is less abstract. Items from the domain can be constructed or created by following a recipe or rules for producing the items in the domain or offers a path for its users to be put in direct contact with the items from the domain. One limitation of such a method is that domains constructed in this way can often be rather limited. For suppose we wanted to talk about cakes. On the one hand, we might specify the items in this domain as follows $\mathcal{D} : \{x \mid x \text{ is a cake}\}$. On the other hand, we might specify the items in this domain as follows: $\mathcal{D} : \{x \mid x \text{ is what you will discover if you follow the following recipe ...}\}$. In the latter case, it is not clear that our instructions can produce all of the cakes in the domain since there may be some ways of making cakes that we fail to consider.

6.3.2 Interpretation of **QL**

The second part of the model is the *interpretation of **QL*** or the *interpretation*. This part of the model is abbreviated as $\mathcal{I}$ and its role (intuitively speaking) is to "give meaning" to the formal language **QL**. In other words, it "gives meaning" to the names and n -place predicates.

A little more precisely, the $\mathcal{I}$ function does two things:

1. it takes names in **QL** and assigns them to items in the domain.
2. it takes n -place predicates in **QL** and assigns them to sets of n -tuples.

That is, an interpretation of **QL** gives *meaning* to the names in **QL** by assigning each name an object in $\mathcal{D}$ and *meaning* to predicate terms by assigning each predicate term a set of items from $\mathcal{D}$.

Definition 6.3.2: Interpretation of QL

An interpretation ($\mathcal{I}$) of **QL** is a function that

1. assigns single elements (items, members) in $\mathcal{D}$ to a name in **QL**, and
2. assigns a set of n -tuples in $\mathcal{D}$ to n -place predicate terms.

Let's consider each task of the interpretation function.

First, for each name in the formal language **QL**, it refers to the corresponding object in the domain. More simply, an interpretation gives each name (a, b, c, d) a corresponding referent or item in the domain. Just as proper names in English, e.g. "George Washington", refer to people in our world, names in **QL** refer to objects in the domain.²

We abbreviate the interpretation of names as follows:

- $\mathcal{I}(a) = a$, this says the name "a" is interpreted as the object a in the domain $\mathcal{D}$
- $\mathcal{I}(\text{Mark}) = \text{Mark}$, this says the name "Mark" is interpreted as the object Mark in the domain $\mathcal{D}$
- $\mathcal{I}(b) = b$, this says the name "b" is interpreted as the object b in the domain $\mathcal{D}$

It is important to note that what is being interpreted is a *name* in terms of an item in the domain. Thus, an interpretation of the name "George Washington" is given in terms of the person *George Washington*. That is, $\mathcal{I}(\text{George Washington}) = \text{George Washington}$. Thus, $\mathcal{I}(a) = a$ says that the name "a" is interpreted in terms of an object a in the domain.

Second, an interpretation also assigns each n -place predicate in **QL** a set of objects in the domain. More simply, if we are thinking about the 1-place

²For example, if a is a name in **QL**, an interpretation function would assign it a single object in $\mathcal{D}$ as follows: $\mathcal{I}(a) = a$. What this says is the name "a" is assigned the object a in the domain of discourse. More naturally, the name "David" refers to the object *David* in the world.

predicate “is red” or R , the interpretation simply assigns that predicate all of the red objects in the domain. In other words, R or “is red” just refers to all of the red things. For example, consider the model $\mathcal{D} = a, b, c, d$. Suppose that the 1-place predicate R is interpreted as the set of objects a, b . That is, $\mathcal{I}(R) = \{a, b\}$. What this says is that the predicate “ R ” is interpreted as the set of objects a, b in the domain. More naturally, the meaning of “ x is red” or “red” is just the red things it refers to in the domain.

Part of the definition of an interpretation, however, reads as follows: the interpretation “assigns sets of n -tuples of objects of $\mathcal{D}$ to n -place predicates.” What is an n -place predicate? What is an n -tuple?

An n -place predicate is just a predicate term with n number of places where we might fill-in names of objects. Or, it is a predicate term with n number of places where it is necessary to fill-in a name so that the predicate becomes a proposition (something that can be true or false). So, for example, “ x is red” is a one-place predicate since there is one place (where the x is) where if we were to fill-in a name of an object, the expression would express a proposition.

If R is a one-place predicate term, an interpretation function assigns a set of items in $\mathcal{D}$ as follows: $\mathcal{I}(R) := \{a, b, c, d\}$. What this says is the predicate term “ R ” is assigned a group of objects a, b, c, d from the domain. More naturally, the meaning of “ x is red” or “red” is just the red things it refers to in the domain.

What is an n -tuple? Technically, an n -tuple refers to an ordered set with n elements. However, more simply, you might think of n -place predicates picking out different kinds of collections of objects. In the case of “ x is red”, when we interpret this predicate, we simply pick out all of the single red things. For example, this umbrella is red, that book is red, this pen is red, and so on. However, consider the two-place predicate “ x is taller than y ”. In this case, the predicate doesn’t just pick out tall objects. Instead, it picks out a collection of pairs of objects where the first object is taller than the second. Where each of the single objects picked out by “is red” is a 1-tuple, the pairs of objects are known as 2-tuples.

- “ x is tall” is a 1-place predicate that is interpreted as a set of 1-tuples.
- “ x is taller than y ” is a 2-place predicate that is interpreted as a set of 2-tuples.
- “ x is standing between y and z ” is a 3-place predicate that is interpreted as a set of 3-tuples.

When interpreting predicate two, or three, or n-place predicate terms, we can make clear that we are referring to a pair of objects, or triplet, or n-tuple by using angle brackets. That is, to express the interpretation of “x is taller than y”, we might write the following: $\mathcal{J}(T) := \{\langle a, b \rangle, \langle b, c \rangle, \langle a, c \rangle\}$. What this says is that the “taller than” two-place predicate is interpreted in terms of pairs of things where the first object in the pair is taller than the other.

Exercise 6.72

Create a model for the following languages:

1. A language consisting of the following names (Jon, Liz) and two-place predicate (loves)
2. A language consisting of the following names (Jon, Liz, Ryan) and the following one-place predicate (loves, likes)
3. A language consisting of the following names (a, b) and the following one-place predicate (L).
4. A language consisting of the following names (a, b, c, d), the following one-place predicate (L, R), and the following two-place predicates (S, M).

6.3.3 Valuation function

At the beginning of [section 6.3](#), it was noted that **QL**-semantics contains three key notions:

1. the domain of discourse
2. an interpretation function
3. a valuation function

With the notion of a model (domain and interpretation) clarified, a **valuation function** can be used to determine the truth value of closed wffs.

Definition 6.3.3: RL-valuation

Let $\alpha_1, \dots, \alpha_n$ be any series of names (not necessarily distinct), P be any n -place predicate, and ϕ, ψ be wffs in RL). Relative to a model (a domain $\mathcal{D}$ and an interpretation $\mathcal{J}$), a valuation function (v or V) of a wff in **QL** is a function that assigns one and only one truth value (T or F) to each closed wff in **QL** in such a way:

1. if $P\alpha_1 \dots \alpha_n$ is an atomic wff in RL, then $v_{\mathcal{M}}(P\alpha_1 \dots \alpha_n) = T$ iff $\langle \mathcal{J}(\alpha_1), \dots, \mathcal{J}(\alpha_n) \rangle \in \mathcal{J}(P)$, otherwise $v_{\mathcal{M}}(P\alpha_1 \dots \alpha_n) = F$

2. $v_{\mathcal{M}}(\neg(\phi)) = T$ iff $v_{\mathcal{M}}(\phi) = F$
3. $v_{\mathcal{M}}(\phi \wedge \psi) = T$ iff $v_{\mathcal{M}}(\phi) = T$ and $v_{\mathcal{M}}(\psi) = T$
4. $v_{\mathcal{M}}(\phi \vee \psi) = T$ iff $v_{\mathcal{M}}(\phi) = T$ or $v_{\mathcal{M}}(\psi) = T$
5. $v_{\mathcal{M}}(\phi \rightarrow \psi) = T$ iff $v_{\mathcal{M}}(\phi) = F$ or $v_{\mathcal{M}}(\psi) = T$
6. $v_{\mathcal{M}}(\phi \leftrightarrow \psi) = T$ iff $v_{\mathcal{M}}(\phi) = v_{\mathcal{M}}(\psi)$.
7. $v_{\mathcal{M}}(\exists x)\phi = T$ iff $v_{\mathcal{M}}\phi(\alpha/x) = T$ for at least one name a in **QL**.
8. $v_{\mathcal{M}}(\forall x)\phi = T$ iff $v_{\mathcal{M}}\phi(\alpha/x) = T$ for every name α in **QL**.

Let's consider each clause of the valuation function in turn. First, clause (1) states that the truth value of wffs of the form $P(\alpha_1 \dots \alpha_n)$ are determined by whether the n -tuple $\langle \mathcal{I}(\alpha_1), \dots, \mathcal{I}(\alpha_n) \rangle$ is in the set of n -tuples assigned to P by the interpretation function. To illustrate, let's consider the following model:

$$\begin{aligned}
\mathcal{D} &= \{1, 2, 3, 4\} \\
\mathcal{I}(a) &= 1, \mathcal{I}(b) = 2, \mathcal{I}(c) = 3, \mathcal{I}(d) = 4, \\
\mathcal{I}(E) &= \{2, 4\}, \mathcal{I}(O) = \{1, 3\}, \\
\mathcal{I}(G) &= \{\langle 2, 1 \rangle, \langle 3, 1 \rangle, \langle 3, 2 \rangle, \langle 4, 1 \rangle, \langle 4, 2 \rangle, \langle 4, 3 \rangle\}
\end{aligned}$$

Intuitively, we can think of this model as talking about four numbers (1, 2, 3, 4), the n -place predicate "is even" (E), the n -place predicate "is odd" (O), and the n -place predicate "is greater than" (G). As we can see, the interpretation function assigns the set of even numbers to the predicate "is even" (E), the set of odd numbers to the predicate "is odd" (O), and it assigns the set of *pairs of numbers* where the first number is greater than the second to the predicate "is greater than" (G).

With this model in place, let's consider whether the wff Ea is true or false. According to clause (1) of the valuation rule, $v(Ea) = T$ if and only if $\mathcal{I}(a) \in \mathcal{I}(E)$, viz., the value given by the interpretation of a is in the value given by the interpretation of E . Since, $\mathcal{I}(a) = 1$ and $\mathcal{I}(E) = \{2, 4\}$, we only need to check whether $1 \in \{2, 4\}$. However, since $1 \notin \{2, 4\}$, it follows that $v(Ea) = F$. In contrast, consider the wff Eb . According to clause (1) of the valuation rule, $v(Eb) = T$ if and only if $\mathcal{I}(b) \in \mathcal{I}(E)$. Since, $\mathcal{I}(b) = 2$ and $\mathcal{I}(E) = \{2, 4\}$, it follows that $v(Eb) = T$.

Let's consider one more example. Take the wff Gab . According to clause (1) of the valuation rule, $v(Gab) = T$ if and only if $\langle \mathcal{I}(a), \mathcal{I}(b) \rangle \in \mathcal{I}(G)$. That is, the pair consisting of the interpretation of "a" and the interpretation of "b" are in the interpretation of "G". Since, $\mathcal{I}(a) = 1$, $\mathcal{I}(b) = 2$, we only need to check whether the pair $\langle 1, 2 \rangle$ is in $\mathcal{I}(G)$. Notice that while

$\langle 2, 1 \rangle$ is in $\mathcal{I}(G)$, it is not the case that $\langle 1, 2 \rangle$ is in $\mathcal{I}(G)$. Another way of putting this is if we think of G as expressing "greater than" and Gab as expressing "a is greater than b", notice that the item from the domain that "a" picks out (1) is not greater than the item from the domain that "b" picks out (2). Therefore, it follows that $v(Gab) = F$.

Clauses (2)-(6) are the same as the clauses in the valuation function for **PL**. That is, the truth values for negated wffs, conjunctions, disjunctions, conditionals, and biconditionals are determined the same way as in **PL**. However, clauses (7) and (8) are new. Let's begin with clause (7), which defines how to determine the truth value of existentially quantified wffs. Clause (7) states that the truth value of a wff of the form $(\exists x)\phi$ is true if and only if the truth value of $\phi(\alpha/x)$ is true for at least one name α in **QL**. To illustrate, let's consider an example using the model we used for clause (1), but we will make two small additions:

$$\begin{aligned} \mathcal{D} &= \{1, 2, 3, 4\} \\ \mathcal{I}(a) &= 1, \mathcal{I}(b) = 2, \mathcal{I}(c) = 3, \mathcal{I}(d) = 4, \\ \mathcal{I}(E) &= \{2, 4\}, \mathcal{I}(O) = \{1, 3\}, \\ \mathcal{I}(G) &= \{\langle 2, 1 \rangle, \langle 3, 1 \rangle, \langle 3, 2 \rangle, \langle 4, 1 \rangle, \langle 4, 2 \rangle, \langle 4, 3 \rangle\}, \\ \mathcal{I}(L) &= \emptyset, \\ \mathcal{I}(N) &= \{1, 2, 3, 4\} \end{aligned}$$

Notice that we have added two new predicate terms: L and N . The interpretation of L is the empty set (symbolized by $\emptyset$). This signifies that there are no items from the domain in the interpretation of L . The second predicate term is N , which is the set of all items in the domain. Now consider the wff $(\exists x)(Lx)$. According to clause (7) of the valuation rule, $v((\exists x)Lx) = T$ if and only if $v(L(\alpha/x)) = T$ for at least one name α in **QL**. That is, there is at least one name α that, were it used to replace the variable x in Lx , would make the resulting wff true. However, notice that for our model, there is no name α that would make $L(\alpha/x)$ true. This is because the interpretation of L is the empty set. Therefore, it follows that $v((\exists x)Lx) = F$. In contrast, consider the wff $(\exists x)Ox$. According to clause (7) of the valuation rule, $v((\exists x)Ox) = T$ if and only if $v(O(\alpha/x)) = T$ for at least one name α in **QL**. That is, there is at least one name α that, were it used to replace the variable x in Ox , would make the resulting wff true. Notice that for our model, there is at least one name α that would make $O(\alpha/x)$ true: $v(Oa) = T$ and $v(Oc) = T$. Therefore, it follows that $v((\exists x)Ox) = T$.

Let's consider one more example involving clause (7). Consider $(\exists x)Gax$. According to clause (7) of the valuation rule, $v((\exists x)Gax) = T$ if and only

if $v(G(\alpha/x)) = T$ for at least one name α in **QL**. That is, there is at least one name α that, were it used to replace the variable x in Gax , would make the resulting wff true. Notice that for our model, there is no name that we could use to replace x in Gax that would make the resulting wff true. This is because the interpretation of G is the set of pairs of numbers where the first number is greater than the second. Since $\mathcal{I}(a) = 1$, there is no pair of numbers in the interpretation of G where the 1 is greater than the second number. Therefore, $(\exists x)Gax$ is false.

Let's turn our attention to clause (8). This clause deals with universally quantified wffs. Clause (8) states that the truth value of a wff of the form $(\forall x)\phi$ is true if and only if the truth value of $\phi(\alpha/x) = T$ for every name α in **QL**. To better understand this rule, let's consider a few examples. Consider our previous model (see above) and the wff $(\forall x)Nx$. Using clause (8), this wff is true provide x can be replaced by every name in Nx and the resulting wff is true. Notice that for our model, this is the case. That is, $v(Na) = T$, $v(Nb) = T$, $v(Nc) = T$, and $v(Nd) = T$. Therefore, it follows that $v((\forall x)Nx) = T$. In contrast, consider $(\forall y)Ey$. Again, using clause (8), this wff is true provide y can be replaced by every name in Ey and the resulting wff is true. However, notice that for our model, this is *not* the case. When replacing y with a , the resulting wff is false. That is, $v(Ea) = F$ since $\mathcal{I}(a) \notin \mathcal{I}(E)$. Since we have at least one replacement where Ey is false, it follows that $v((\forall y)Ey) = F$.

Let's consider one final example. Consider $(\forall x)Gax$. According to clause (8) of the valuation rule, $v((\forall x)Gbx) = T$ if and only if $v(G(b\alpha/x)) = T$ for every name α in **QL**. That is, were we to replace the variable x with any name α in our model, the resulting wff $Gb(\alpha/x)$ would be true. Let's run through each replacement to check whether each of the resulting wffs are true. In the first case, we might replace x with a (that is, $Gb(a/x)$). On this replacement, $v(Gba) = T$ since $\langle 2, 1 \rangle \in \mathcal{I}(G)$. In the second case, we might replace x with b (that is, $Gb(b/x)$). On this replacement, $v(Gbb) = F$ since $\langle 2, 2 \rangle \notin \mathcal{I}(G)$. At this point, we can stop since we have found a replacement where the resulting wff is false. Therefore, it follows that $v((\forall x)Gax) = F$.

Exercise 6.73

Let $\mathcal{D} = \{a, b, c\}$, $\mathcal{I}(a) = a$, $\mathcal{I}(b) = b$, $\mathcal{I}(c) = c$, $\mathcal{I}(Hx) = \{a, b, c\}$, $\mathcal{I}(Lxy) = \{\langle a, a \rangle, \langle b, c \rangle\}$. Determine whether the following wffs are T or F relative to this model.

1. Ha
2. $(\exists x)Hx$
3. $(\forall x)Hx$
4. $\neg(\forall x)Hx$
5. $(Ha \wedge Hb) \wedge Hc$
6. Lab
7. $\neg Laa$
8. Lba
9. $(\exists x)Lxx$
10. $(\forall x)Lxx$
11. $(\exists x)Lxx \vee (\exists x)(\exists y)Lxy$

6.3.4 Two simplifications

In the previous section, the valuation rules for **QL** were introduced. In addition, we looked at some simple examples of how to use these rules to determine the truth value of wffs. In this section, we will look at two simplifications that were made in the previous section.

The definition of our valuation rule involves two simplifications. First, the domain of the valuation function is limited to closed RL-wffs. That is, the valuation function takes as input a closed wff (a wff without any free variables) and returns a truth value. This simplification leaves the truth value of open RL-wffs (wffs with free variables, e.g. Px) undefined. That is, our system of logic has a syntax that produces formulas that are wffs but the semantic system does not state whether these wffs are true or false. There is a solution to this problem, however, and it involves the notion of a *variable assignment function*. This solution is considered in a later chapter (see ??).

The second assumption that the valuation function makes is that it assumes that there is an RL-name for every item in the domain of discourse. This assumption allows for specifying the truth value of existentially and universally quantified wffs in terms of non-quantified wffs involving names. For example, $v_{\mathcal{M}}(\forall x)\phi = T$ iff $v_{\mathcal{M}}\phi(\alpha/x)$ for every name α in RL. This simplification may be taken to be problematic for what guarantee is there that everything in the domain is named (even if we have an infinite number of names)? To state this clearly, the problem with this assumption is that even with an infinite number of names, there is no guarantee that each item in the domain is named. And if this is the case, then $v(\forall x)Px = T$ even though some unnamed item $u_1 \in \mathcal{D}$ is not in the interpretation of P .

6.3.5 Further examples of the valuation function

Let's consider a few more examples of the valuation function. Consider the following model (we used this model in an earlier example):

$$\begin{aligned}\mathcal{D} &= \{1, 2, 3, 4\} \\ \mathcal{I}(a) &= 1, \mathcal{I}(b) = 2, \mathcal{I}(c) = 3, \mathcal{I}(d) = 4, \\ \mathcal{I}(E) &= \{2, 4\}, \mathcal{I}(O) = \{1, 3\}, \\ \mathcal{I}(G) &= \{\langle 2, 1 \rangle, \langle 3, 1 \rangle, \langle 3, 2 \rangle, \langle 4, 1 \rangle, \langle 4, 2 \rangle, \langle 4, 3 \rangle\}, \\ \mathcal{I}(L) &= \emptyset, \\ \mathcal{I}(N) &= \{1, 2, 3, 4\}\end{aligned}$$

Now consider the wff $(\forall x)(Ox \rightarrow Nx)$. This wff is true provided for every name α , the wff $O(\alpha/x) \rightarrow N(\alpha/x)$ is true. Let's consider each replacement in turn.

1. Replace x with a . The result is: $Oa \rightarrow Na$. On this replacement, $v(Oa) = T$ and $v(Na) = T$ since a is in O and N , respectively. Since both sides of the conditional are true: $v(Oa \rightarrow Na) = T$.
2. Replace x with b . The result is: $Ob \rightarrow Nb$. On this replacement, $v(Ob) = F$ since $\mathcal{I}(b) \notin \mathcal{I}(O)$. Since the leftside of the conditional is false, it follows that $v(Ob \rightarrow Nb) = T$.
3. Replace x with c . The result is: $Oc \rightarrow Nc$. On this replacement, $v(Oc) = T$ since $\mathcal{I}(c) \in \mathcal{I}(O)$. In addition, $v(Nc) = T$ since $\mathcal{I}(c) \in \mathcal{I}(N)$. Since both sides of the conditional are true, it follows that $v(Oc \rightarrow Nc) = T$.
4. Replace x with d . The result is: $Od \rightarrow Nd$. On this replacement, $v(Od) = F$ since $\mathcal{I}(d) \notin \mathcal{I}(O)$. Since the leftside of the conditional is false, it follows that $v(Od \rightarrow Nd) = T$.

Notice that for every name α , the wff $O(\alpha/x) \rightarrow N(\alpha/x)$ is true. Therefore, $v((\forall x)(Ox \rightarrow Nx)) = T$.

Let's consider a similar wff. This wff is $(\forall x)(Nx \rightarrow Ox)$. We can determine whether this wff is true or false in the same way as above. Let's consider each replacement in turn.

1. Replace x with a . On this replacement, $v(Na) = T$ since $\mathcal{I}(a) \in \mathcal{I}(N)$. Similarly, $v(Oa) = T$. Since the antecedent is true and the consequent is true, it follows that $v(Na \rightarrow Oa) = T$.
2. Replace x with b . On this replacement, $v(Nb) = T$ since $\mathcal{I}(b) \in \mathcal{I}(N)$. However, $v(Ob) = F$. Since the antecedent is true and the consequent is false, it follows that $v(Nb \rightarrow Ob) = F$.

It is not necessary to consider the final two replacements since we have found a replacement where the wff $N(\alpha/x) \rightarrow O(\alpha/x)$ is false. Therefore, $v((\forall x)(Nx \rightarrow Ox)) = F$.

Let's consider two additional examples, both involving existentially quantified wffs. First, consider $(\exists x)(Ox \wedge Nx)$. This wff is true provided there is at least one name α such that the wff $O(\alpha/x) \wedge N(\alpha/x)$ is true. Let's consider each replacement of variables with names until we find a replacement where the wff is true.

1. Replace x with a . On this replacement, $v(Oa) = T$ since $\mathcal{J}(a) \in \mathcal{J}(O)$. In addition, $v(Na) = T$ since $\mathcal{J}(a) \in \mathcal{J}(N)$. Since both conjunctions are true, it follows that $v(Oa \wedge Na) = T$.

It is not necessary to consider the remaining three replacements since we have found a replacement where the wff $O(\alpha/x) \wedge N(\alpha/x)$ is true. Therefore, $v((\exists x)(Ox \wedge Nx)) = T$.

Let's consider one more example. Consider $(\exists x)(Ex \wedge Ox)$. In examining this wff, notice that both $(\exists x)Ex$ and $(\exists x)Ox$ are true. However, $(\exists x)(Ex \wedge Ox)$ is false. This is because there is no name α such that $E(\alpha/x) \wedge O(\alpha/x)$ is true. To see this, let's consider each replacement in turn.

1. Replace x with a . On this replacement, $v(Ea) = F$ since $\mathcal{J}(a) \notin \mathcal{J}(E)$. Since $v(Ea) = F$, it follows that $v(Ea \wedge Oa) = F$.
2. Replace x with b . On this replacement, $v(Eb) = T$ since $\mathcal{J}(b) \in \mathcal{J}(E)$. However, $v(Ob) = F$. Since it is not the case that both conjuncts are true, $v(Eb \wedge Ob) = F$.
3. Replace x with c . On this replacement, $v(Ec) = F$ since $\mathcal{J}(c) \notin \mathcal{J}(E)$. Since $v(Ec) = F$, it follows that $v(Ec \wedge Oc) = F$.
4. Replace x with d . On this replacement, $v(Ed) = F$ since $\mathcal{J}(d) \notin \mathcal{J}(E)$. Since $v(Ed) = F$, it follows that $v(Ed \wedge Od) = F$.

Since there is no name α such that $E(\alpha/x) \wedge O(\alpha/x)$ is true, it follows that $v((\exists x)(Ex \wedge Ox)) = F$.

Exercise 6.74

Let $\mathcal{D} = \{1, 2, 3\}$. $\mathcal{J}(a) = 1, \mathcal{J}(b) = 2, \mathcal{J}(c) = 3$.
 $\mathcal{J}(Ex) = \{2\}, \mathcal{J}(Ox) = \{1, 3\}, \mathcal{J}(Nx) = \{1, 2, 3\}, \mathcal{J}(Gxy) = \{\langle 2, 1 \rangle, \langle 3, 1 \rangle, \langle 3, 2 \rangle\}$ Determine whether the following wffs are T or F.

1. $(\forall x)(Ex \rightarrow Nx)$

2. $(\forall x)(Ex \vee Ox)$
3. $(\forall x)(Ex \rightarrow \neg Ox)$
4. $(\exists x)(Ex \wedge Nx)$
5. $(\exists x)(Ex \vee Ox)$
6. $(\exists x)\neg Ex \wedge (\exists x)\neg Ox$
7. $(\forall x)Gxx$
8. $(\exists x)Gxa$
9. $(\exists x)(\forall y)Gxy$

6.4 QL TRANSLATION

In this section, let's consider how to translate English sentences into quantified formula in **QL** and vice versa. The first step to any translation is to construct a translation key. The translation key has three parts:

1. a domain of discourse,
2. an assignment of English names to **QL** names.
3. an assignment of English predicates and relations to **QL** n-place predicates.

In short, the translation key specifies the objects that we are talking about, the names that we use to refer to these objects, and the properties and relations that these objects have. Here is an example of a translation key:

- $\mathcal{D}$: Living human beings
- a: Annie
- j: Jon
- f: Frank
- Txy : x is taller than y.
- Hx : x is happy.

Using the above translation key various English sentences can be translated into **QL** and closed wffs in **QL** can be translated into English sentences. Let's consider how this is done in four subsections. In the first subsection, we will consider how to translate **QL** wffs that do not contain quantifiers. In the second subsection, we will consider how to translate **QL** wffs that contain a single quantifier. In the third subsection, we will consider how to translate **QL** wffs that involve overlapping quantifiers. Finally, we will consider some more complex examples of translation.

6.4.1 Translating wffs without quantifiers

Let's consider several **QL** wffs that do not contain quantifiers and how to translate them into English. Consider the following wffs:

1. Ha
2. Taj
3. Tjf
4. Tfj
5. Tjj

In order to translate these wffs into **QL**, we need a translation key. Let's make use of the translation key introduced in the prior section. Next, let's take the first wff and consider what would happen if (1) write out the English predicate that corresponds to the predicate term in the wff and then (2) replace any **QL**-names with English-names

1. Ha
 - 1a. x is happy
 - 1b. Anne is happy

Notice that in the above example, we have the wff Ha . We start by writing out the English predicate that corresponds to the one-place predicate Hx . This is "x is happy." The wff that we are translating, however, is not Hx but Ha . So, let's replace the **QL**-name a with the English name "Anne". Therefore, the wff Ha translates into the English sentence "Anne is happy." Let's consider the second wff Taj .

2. Taj
 - 2a. x is taller than y
 - 2b. Annie is taller than y
 - 2c. Annie is taller than Jon

Again, notice that in the above example, we have the wff Taj . We start by writing out the English predicate that corresponds to the two-place predicate Txy . This is "x is taller than y." Next, we replace the **QL**-names a and j with the English names "Annie" and "Jon". Therefore, the wff Taj translates into the English sentence "Annie is taller than Jon."

Finally, consider the next two examples.

3. Tjf : Jon is taller than Frank:
4. Tfj : Frank is taller than Jon:

With respect to (3) and (4) notice that the order of the names in the wff matters. That is, Tjf is not the same as Tfj . If Txy is "x is taller than

y ", then Tjf is "Jon is taller than Frank" and Tfj is "Frank is taller than Jon".

Let's consider (5) Tjj . If we take the predicate "x is taller than y", and replace the first x with j (Jon) and the second x with j (Jon), then the resulting English sentence is "Jon is taller than Jon".

Finally, our approach to considering how to translate a **QL** wff into English has been mostly syntactic. That is, we have been considering the structure of the wff and how to translate the structure into English. However, another approach is to consider the truth conditions of the wff and then to find a corresponding English sentence that has the same truth conditions. This is the semantic approach to translation. For example, consider Ha . This wff is true if and only if the interpretation of a is in the interpretation of H . That is, $v(Ha) = T$ iff $\mathcal{I}(a) \in \mathcal{I}(H)$. Using our translation key, a can be translated as "Annie" and H can be translated as "is happy". An appropriate English sentence that has the same truth conditions as Ha is "Anne is happy". This is because the English sentence "Anne is happy" is true if and only if our interpretation of "Annie" (its referent) is in the set of living human beings that are happy.

Let's consider one more example using the semantic approach to translation. Recall (3) and (4) from above. In the case of (3), the wff Tjf is true if and only if the pair consisting of the interpretation of j and the interpretation of f is in the interpretation of T . That is, $v(Tjf) = T$ iff $\langle \mathcal{I}(j), \mathcal{I}(f) \rangle \in \mathcal{I}(T)$. So, an English sentence that is true under the same conditions would be one where the referent of "Jon" is taller than the referent of "Frank". What English sentence is true just in the case where Jon is taller than Frank? Quite obviously, it is the sentence "Jon is taller than Frank". So, we can translate Tjf as "Jon is taller than Frank."

In contrast to Tjf , there is the wff Tfj . Note that unlike the previous wff, $v(Tfj) = T$ iff $\langle \mathcal{I}(f), \mathcal{I}(j) \rangle \in \mathcal{I}(T)$. That is, this wff is true if the pair consisting of Frank and Jon are in the taller than interpretation. With this and our translation key in mind, we can translate Tfj as "Frank is taller than Jon".

It is worth pointing out that while the syntactic approach to translation is more straightforward, the semantic approach is more precise. As we will see later, the grammatical structure of a sentence can be misleading. In addition, when attempting to translate from one language to another, the primary is typically to capture what the sentence says (its meaning) rather than how it is said (its syntax).

Exercise 6.75

Translate the following wffs into English sentences using the translation key provided here: $\mathcal{D}$: Living human beings, a: Ann, b: Bob, c: Chris, Sx: x is sad, Hx: x is happy, Fxy: x is friends with y.

1. Sa
2. $\neg Sa$
3. $Sb \wedge Sc$
4. Fab
5. $\neg Fba$
6. $Fab \wedge \neg Fbc$
7. $(Ha \wedge Hb) \rightarrow Fab$

6.4.2 Translating wffs with a single quantifier

In this subsection, we consider how to translate **QL** wffs that contain a single quantifier (either the universal or existential quantifier) into English. Let's begin with **QL** wffs where the main operator is the universal quantifier.

6.4.2.1 Translating with the Universal Quantifier

First, let's make use of the following translation key:

- $\mathcal{D}$: human beings (living or dead)
- Hx: x is happy
- Zx: x is a zombie
- Mx: x is mortal
- Rx: x is murderer
- Wx: x is wrong

Next, let's consider the following predicate wffs:

1. $(\forall x)Hx$
2. $(\forall x)(Hx \wedge Zx)$
3. $(\forall x)(Zx \rightarrow Hx)$
4. $(\forall x)(Zx \rightarrow \neg Hx)$
5. $\neg(\forall x)(Zx \rightarrow Hx)$

Let's start by considering a syntactic approach to translating universally quantified wffs. On this approach, translate **QL**-predicate terms in the

usual way and replace the universal quantifier followed by the variable x with "For every x ". Thus, we can translate $(\forall x)Hx$ as follows:

1. $(\forall x)Hx$
2. $(\forall x)$, x is happy.
3. For every x , x is happy.

Once we have this pseudo-English sentence, we can then translate it into something more colloquial. For example, we can translate "For every x , x is happy" as "Everyone is happy" or "All people (living or dead) are happy."

The above approach is syntactic. We first create a pseudo-English sentence by reading the **QL**-wff in a particular way. Whenever we encounter a universal quantifier $(\forall x)$ we read it as "for every x " and whenever we encounter a predicate term Px we read it as " x is P ". The resulting sentence is sometimes called a "bridge translation". The next step in the process is to translate the pseudo-English sentence into a more natural English sentence. On this step, we need to make use of the translation key. We replace the **QL**-names with English names and the **QL**-predicates with English predicates.

Before considering additional examples, let's consider how to translate (1) from a semantic perspective. First, consider the truth conditions for universally quantified wffs in general. A wff $(\forall x)\phi$ is true if and only if $\phi(\alpha/x)$ is true for *every* replacement of x with a name α . Now let's consider these truth conditions with respect to wff (1). The wff $(\forall x)Hx$ is true if and only if Ha, Hb, Hc and so on are true. Using our translation key, we can say that $(\forall x)Hx$ is true if and only if Al is happy, Bob is happy, Chris is happy, and so on. Rather than specifying the truth conditions for $(\forall x)Hx$ by specifying all of its possible substitutions, we can more simply say that $(\forall x)Hx$ is true if and only if "Everyone is happy".

Let's consider some additional examples involving universally quantified wffs. Consider (2), the wff $(\forall x)(Hx \wedge Zx)$. We can translate this wff using a bridge translation as "For every x , x is happy and x is a zombie." From here, we can translate this into a more natural English sentence. Using our translation key, (2) says that "everyone is a happy zombie".

Consider (3), the wff $(\forall x)(Zx \rightarrow Hx)$. We can translate this wff using a bridge translation as "For every x , if x is a zombie, then x is happy." From here, we can translate this into a more natural English sentence. Some people, however, find it difficult to determine what a natural translation of this sentence should be. In fact, some individuals translate (3) in a

way similar to (2). Namely, they contend that the correct translation is "everyone is a happy zombie." However, this is not the correct translation. The correct translation is "every zombie is happy". There are at least two ways to see this more clearly.

The first way to see this is to consider the truth conditions for the wff $(\forall x)(Zx \rightarrow Hx)$. If $(\forall x)(Zx \rightarrow Hx)$ is true, then $Za \rightarrow Ha, Zb \rightarrow Hb, Zc \rightarrow Hc$ and so on are true. These conditionals are not saying that a or b or c or any named item in the domain is a zombie. Instead, they are saying that (1) *if* a is a zombie, then a is happy, (2) *if* b is a zombie, then b is happy, and so on.

A second approach to seeing why this wff should be translated in this way is to consider an *expanded* version of the bridge translation of (3). The bridge translation of (3) is "For every x , if x is a zombie, then x is happy." However, we might expand upon the bridge translation as follows:

Choose any object you please in the domain of discourse, if
that object is a zombie, then it will be also be happy.

With (3) expanded in this way, it becomes clearer that the correct sentence is not that everyone is a happy zombie, but instead that every zombie is happy.

Next, consider (4). This is the wff $(\forall x)(Zx \rightarrow \neg Hx)$. We can translate this wff using a bridge translation as "For every x , if x is a zombie, then x is not happy." Again, it is not entirely clear what the natural English translation of this wff should be. Some individuals incorrectly translate this wff as "everyone is an unhappy zombie." Our approach to resolving this uncertainty can be similar to the approach we took with (3): considering the truth conditions or expanding on the bridge translation. Since the wff is true if and only if $Za \rightarrow \neg Ha, Zb \rightarrow \neg Hb, Zc \rightarrow \neg Hc$ and so on are true, the wff is true provided there are no cases where a named object is a zombie and happy. We thus can translate (4) as "no zombie is happy". An expanded bridge translation of (4) would be "Choose any object you please in the domain of discourse consisting of human beings (living or dead), if that object is a zombie, then it will not be happy."

Finally, consider (5). This is the wff $\neg(\forall x)(Zx \rightarrow Hx)$. Notice that in this example, the main operator is not the universal quantifier. Instead, the main operator is negation. One way to translate this wff is to first translate the wff $(\forall x)(Zx \rightarrow Hx)$. This wff is translated as "every zombie is happy". Next, we can translate the negation in this wff by prefixing the entire sentence with "It is not the case that". Thus, (5) can be translated

as "It is not the case that every zombie is happy" or "Not every zombie is happy."

Exercise 6.76

Translate the following wffs into English sentences using the translation key provided here: $\mathcal{D}$: human beings, Fx : x is friendly, Gx : x is a ghost, Hx : x is happy, Lxy : x loves y .

1. $(\forall x)Fx$
2. $(\forall x)(Fx \rightarrow Gx)$
3. $(\forall x)(Gx \rightarrow Fx)$
4. $(\forall x)(Gx \rightarrow \neg Fx)$
5. $\neg(\forall x)(Gx \rightarrow Fx)$
6. $(\forall x)(Fx \wedge Gx)$
7. $(\forall x)(Fx \vee Gx)$

6.4.2.2 Translating with the Existential Quantifier

In the previous section, we considered how to translate **QL** wffs that contain a single quantifier where the main operator was the universal quantifier. In this section, we will consider how to translate **QL** wffs that contain a single quantifier where the main operator is the existential quantifier. We will use the translation key from the earlier section and consider the following wffs:

1. $(\exists x)Hx$
2. $(\exists x)\neg Zx$
3. $\neg(\exists x)Zx$
4. $(\exists x)(Zx \wedge Hx)$
5. $(\exists x)Zx \wedge (\exists x)Hx$

Similar to our approach to universally quantified wffs, we can approach translation involving existentially quantified wffs from a syntactic point of view. Our approach then is to translate **QL**-predicate terms in the usual way and to replace the existential quantifier followed by the variable x with "There is at least one x such that" or "For some x ". Thus, we can translate $(\exists x)Hx$ as "There is at least one x such that x is happy" or "For some x , x is happy". Once we have this pseudo-English sentence, we can then translate it into something more natural. For example, we can translate this sentence into any of the following expressions:

- 1a. At least one person is happy.

- 1b. There is at least one person that is happy.
- 1c. Someone is happy.
- 1d. There is at least one happy person.
- 1e. A happy person exists

In addition, we can also approach the translation of (1) semantically. The wff $(\exists x)Hx$ is true if and only if there is at least one wff $H(\alpha/x)$ is true. That is, $v(Ha) = T$ or $v(Hb) = T$ or $v(Hc) = T$ and so on. Using our translation key, we can say that $(\exists x)Hx$ is true if and only if Al is happy, Bob is happy, Chris is happy, and so on. Rather than specifying the truth conditions for $(\exists x)Hx$ by specifying all of its possible substitutions, we can more simply say that $(\exists x)Hx$ is true if and only if "Someone is happy".

Next, let's consider (2). Again, we can use a bridge translation, "For some x , x is not a zombie." This is saying that there exists some item in the domain of discourse that is not a zombie. More naturally, this can be expressed as "Someone is not a zombie" or "There is at least one person that is not a zombie".

In the case of (3), note that negation has wide scope. Our approach to translating this wff then is to first translate the existentially quantified wff and then to prefix it with "It is not the case that". Thus, we can translate $\neg(\exists x)Zx$ as "It is not the case that someone is a zombie." Notice that (2) and (3) say something distinct. (2) says that something exists that is not a zombie, while all (3) says is that zombies do not exist.

Finally, let's consider (4) and (5) together. The bridge translations for (4) and (5) are as follows:

- 4B. For some x , x is a zombie and x is happy.
- 5B. For some x , x is a zombie, and for some x , x is happy.

Notice that these two propositions do not say the same thing. (4) asserts that there is something that is both a zombie and happy, while (5) asserts that there is a zombie and there is someone who is happy.

Exercise 6.77

Let the domain $\mathcal{D}$ be people and Px : x is poor; Lx : x is lazy; Rx : x is rich. Translate the following into **QL** and English.

- 1. $(\forall x)Px$
- 2. $(\forall x)(Px \rightarrow Lx)$
- 3. $(\forall x)Px \wedge (\forall x)Lx$
- 4. $(\forall x)(Px \wedge Lx)$

5. $(\forall x)(Px \rightarrow \neg Lx)$
6. $(\forall x)(Px \vee Lx)$
7. $(\exists x)Px$
8. $(\exists x)Px \wedge (\exists x)Rx$
9. $(\exists x)(Px \wedge Rx)$, what is the difference between #3 and #2?
10. $(\exists x)(Px \vee Rx)$
11. $\neg(\exists x)(Px \wedge Rx)$
12. $(\exists x)\neg(Px \wedge Rx)$

6.4.3 Translating Wffs with Overlapping Quantifiers

When dealing with wffs with quantifiers whose scope overlaps, does the order of the quantifiers matter? Consider the following eight wffs (let Lxy express the two-place English expression “ x loves y ”). Let the domain consists of living human beings.

1. $(\forall x)(\forall y)Lxy$
2. $(\forall y)(\forall x)Lxy$
3. $(\exists x)(\exists y)Lxy$
4. $(\exists y)(\exists x)Lxy$
5. $(\forall x)(\exists y)Lxy$
6. $(\exists y)(\forall x)Lxy$
7. $(\forall y)(\exists x)Lxy$
8. $(\exists x)(\forall y)Lxy$

While some of these wffs entail others, only the first two pairs of wffs are equivalent. That is, $(\forall x)(\forall y)Lxy$ is equivalent to $(\forall y)(\forall x)Lxy$ and $(\exists x)(\exists y)Lxy$ is equivalent to $(\exists y)(\exists x)Lxy$.

(1) and (2) express the proposition that “everyone loves everyone”. In this scenario, every item in the domain of discourse loves every item in the domain of discourse.

(3) and (4) express the proposition that “someone loves someone”. In this scenario, at least one item in the discourse loves at least one other. In both cases, the order of the quantifiers does not impact the truth or falsity of the wff.

(5)-(8) express different propositions. Let’s characterize each in terms of a scenario. (5) is what I will call the “crush” scenario. $(\forall x)(\exists y)Lxy$ states that for every individual x , there is an individual y such that x loves y . Here is another way of putting it.

Take any person you please from the domain of discourse (let's call them x), once you have that person, then there is at least one person y that x loves.

What our sentence is committed to is everyone loves at least one person. This is why I call it the "crush" scenario since it implies that everyone has a crush on someone (some person that they love). In short, the wff says (most compactly) that *everyone loves someone*.

Two additional clarifications. First, (5) should not be translated as "everyone is loved by someone". The truth of (5) does not require that for each individual in the domain is loved. For example, suppose that there is one person that everyone loves, let's call them y . If that were the case, then it would be true that "everyone loves someone" since y would be that someone. However, it would not be true that "each person is loved by someone" since there would be many people who are unloved. Second, (5) also does not say that there is some single person who is loved by all. Consider the scenario where there are two people Tek and Liz and where all people love Tek and no one else or Liz and no one else. In such a scenario, "everyone loves someone" would be true but it would be false that "someone is loved by everyone."

(6) is what I will call the "Santa Claus scenario". $(\exists y)(\forall x)Lxy$ states that there is at least one person y such that for every x , x loves y . In other words, there exists someone that everyone loves. In short, "someone is loved by everyone." I call this scenario the "Santa Claus" scenario since he is someone that everyone loves. Here is another way of putting it.

Go through the entire domain of discourse and you will eventually find a person y . Now that person y bears a special relation to all people in the domain of discourse, including themselves. To see this relation more clearly, go through the entire domain again and select any person you please (let's call them x). You will find that x loves y .

Wff (6) is similar to (5) in that it implies that everyone loves at least one person. However, they are importantly different. In (5), every single person loves at least one person, but the loved person can differ from person to person. For example, in a scenario consisting of Jane, John, and Sally, (5) would be true if Jane loves John and John loves Jane and Sally loves herself. In contrast, (6) is true just in the case that there is at least one person loved by everyone, e.g. John loves Jane, Sally loves Jane, and Jane loves Jane.

(7) is what I will call the "Secret Admirer scenario". $(\forall y)(\exists x)Lxy$ says that "for every y , there exists an x , such that x loves y ". In short "everyone is loved by someone". This wff is true provided there is one person who loves you. Here is another way of putting it.

Take any person you please from the domain of discourse (let's call them y), once you have that person, then go through the domain again and you eventually find one person x (the Secret Admirer) that loves y .

So, if we consider a small domain consisting of Jane, John, and Sally, (7) is true in the case that, beginning with Jane, we can find at least one other person who loves Jane (e.g. John, but it could be anyone) and one person who loves John and one person who loves Sally. Note that the person doing the loving need not be the same person in each case, nor is it the case that everyone loves someone. (7) differs from (5) in that it doesn't imply that everyone loves at least one other person. (7), in contrast, is true if John loves Sally and Jane and himself (he would secretly admire several people), but neither Sally nor Jane love anyone. (7) also differs from (6) in that it doesn't imply that there is some person who is secretly admired by all.

(8) is what I will call the "Loving God scenario". We can read this wff as saying "there exists an x such that for every y , x loves y ." More naturally, it can be translated as "someone loves everyone". Here is another way of putting it.

Look through your domain and you will find a person x . Now go through the domain again and x will be in a relation to any person y you might select. Namely, that x will love any person y you select.

I call the scenario that represents (8) the "Loving God Scenario" since a Loving God would love every single person. A few clarifications. First, in contrast to (5), (8) does not imply that everyone loves at least one other person. Suppose God loves everyone but most of the people God loves do not love anyone. In that scenario (8) is true but (5) is false. In contrast to (6), (8) does not imply that there is at least one person loved by everyone. God may love everyone but there may be no person who is loved by all (including God). Finally, while (8) implies (7)—for if someone loves everyone, then everyone is loved by at least one person—(7) does not imply (8). This is because (7) can be true in a case where (8) is not, namely in the case where everyone is loved by someone, but everyone is not loved by a single person.

1. $(\forall x)(\forall y)Lxy$. Everyone loves everyone.
2. $(\forall y)(\forall x)Lxy$. Everyone loves everyone.
3. $(\exists x)(\exists y)Lxy$. Someone loves someone.
4. $(\exists y)(\exists x)Lxy$. Someone loves someone.
5. $(\forall x)(\exists y)Lxy$. Everyone loves someone.
6. $(\exists y)(\forall x)Lxy$. Someone is loved by everyone.
7. $(\forall y)(\exists x)Lxy$. Everyone is loved by someone.
8. $(\exists x)(\forall y)Lxy$. Someone loves everyone.

Exercise 6.78

Using the following symbolization key, translate the following predicate logic expressions into English: D: living humans, Hxy: x hates y, s: Sally, b: Bob, Lxy: x loves y

1. $(\forall x)Lxb$
2. $(\exists x)Hxs$
3. $(\forall x)(Lxb \rightarrow \neg Hxs)$
4. $(\exists x)(Lxb \wedge Hxs)$
5. $[(\exists x)(Lbs \wedge Hbx)] \rightarrow Lbs$
6. $(\forall x)Lxx \wedge (\exists y)Hyb$
7. $(\exists x)Lxx \wedge (\forall y)Hyy$
8. $[(\exists x)Lxb \wedge (\exists x)Lxs] \wedge [(\exists x)\neg Lxb \wedge (\exists x)\neg Lxs]$
9. $(\exists x)Lxb \wedge (\exists x)Lbx$
10. $[(\exists x)(\neg Hxs) \wedge (\exists x)(Lxs)] \wedge (\forall x)Lxb$

6.4.4 Quantificational ambiguities

One benefit of learning a logical language is its use for identifying and resolving ambiguities in natural language. In contrast to natural language, the language of predicate logic does not contain any ambiguity. Names refer to exactly one item in the domain, each n-place predicate picks out a set of n-tuples in the domain, operators have a fixed meaning, and the scope of the various operators is precise. In contrast, natural language is rife with ambiguity. In some cases, the existence of ambiguity influences our evaluation of an argument. For example, consider the following argument:

- P1. All stars are in outer space.
 P2. Socrates is a star.
 C. Therefore, Socrates is in outer space.

In the above example, the word "star" is ambiguous. It has two distinct

meanings. A "star" can refer to a celestial body that produces light and heat, or it can refer to a celebrity. Noting this ambiguity, consider three different ways in which the argument can be evaluated.

First, we take "star" to have the *same* meaning in both premises. For the sake of the example, let's suppose it means "a celestial body that produces light and heat". If this is the case, then the argument is valid. However, notice that the validity of the argument comes at a cost. If "star" means "a celestial body that produces light and heat", then P2 is false. Socrates is not a star in this sense (he is a philosophical celebrity). So, the argument is valid but unsound.

Second, we take "star" to have *different* senses in the two premises. In P1, "star" refers to a class of celestial bodies that produce light and heat. In P2, "star" refers to a celebrity. If this is the case, then the premises of the argument are true, but the argument is invalid. This way of evaluating the argument is certainly the most natural and reasonable.

Third, when evaluating the argument for validity, we take "star" to have the same meaning in both premises and so the argument is valid. However, when we evaluate the truth of the premises, we take the different instances of "star" to have different meanings (in P1 "celestial body", in P2 "celebrity"). If this is the case, then the argument has true premises. The argument then is said to be sound (valid and true premises), although this comes at the cost of an inconsistency since we are saying the two occurrences of "star" have different meanings but also the same meaning.

1. Approach 1: Argument is valid but at least one premise is false.
2. Approach 2: Argument is invalid but has true premises.
3. Approach 3: Argument is sound (valid and true premises), but one has to accept an inconsistency.

It is clear then that the existence of ambiguity in natural language can have an impact on the evaluation of arguments. Insofar as one of the above approaches is more plausible than the others, it is important to be able to identify and resolve ambiguities in natural language.

What sort of ambiguities exist in natural language that can be revealed using predicate logic? Let's consider a type of ambiguity that we will call a "quantificational ambiguity".

Definition 6.4.1

A *quantificational ambiguity* is an ambiguity that arises when a sentence is ambiguous about the presence or scope of a quantifier.

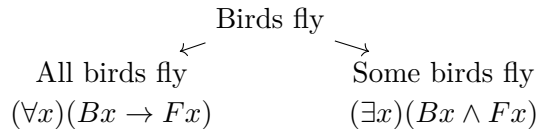
There are at least two types of quantificational ambiguities: quantificational presence ambiguities and quantificational scope ambiguities. The first type involves a sentence where the sentence suggests the presence of a quantifier, but it is ambiguous as to which quantifier is present. The second type involves a sentence where the sentence is ambiguous as to the scope of a quantifier.

Quantificational presence ambiguities

Let's begin with quantificational presence ambiguities. One simple example involves the following sentence:

Birds fly.

How should we translate this sentence into predicate logic? There appear to be two different possibilities. The first possibility is that "Birds fly" means "All birds fly" and so we would translate the sentence using the universal quantifier as follows: $(\forall x)(Bx \rightarrow Fx)$. The second possibility is that "Birds fly" means "Some birds fly" and so we would translate the sentence using the existential quantifier as follows: $(\exists x)(Bx \wedge Fx)$.



Our attempt to translate "Birds fly" into predicate logic reveals that the sentence is not explicit about which quantifier is present. Since both interpretations of the sentence are possible and these interpretations are not equivalent, an ambiguity is present. As mentioned, this type of ambiguity may impact how arguments are evaluated. Consider the following argument:

P1: Birds fly.

C: Therefore, penguins fly.

If "birds fly" expresses "all birds fly", then the above argument is valid but unsound. P1 would be false since penguins are birds, but do not fly. In contrast, if "birds fly" expresses "some birds fly", then while P1 is true, the argument is invalid. Finally, someone might contend that the argument is sound (valid and has true premises), but they would do so at the cost

of inconsistency: they would be saying that "birds fly" expresses "all birds fly" and "some birds fly".

Our previous example involving birds is straightforward and few people fail to recognize that there is something fishy about the sentence itself. In what follows, let's consider a similar, but real-life example of a quantificationally ambiguous sentence. Consider the following sentence:

The truth of the matter is that pretty much anywhere in the world men tend to think that they are much smarter than women.[qtd in 9]

Similar to the previous example, the sentence "men think that they are smarter than women" is quantificationally ambiguous. First, it is indeterminate about how many men think they are smarter than women (all, some, most, etc.). We might try to resolve this ambiguity by prefixing a different quantificational expression ("all" or "some") to each sentence.

1. *All* men think that they are smarter than women.
2. *Some* men think that they are smarter than women.

In prefixing a quantifier to "men think they are smarter than women", we have made explicit the ambiguity concerning the *presence* of a quantifier. The sentence itself does not contain any quantificational expression, but the interpretation of the sentence requires the addition of some quantificational expression. However, note that even after we have made explicit the ambiguity concerning the presence of a quantifier, there still remains *another* quantificational ambiguity. That is, not only is there an ambiguity concerning *how many men* think a certain way, but there is also ambiguity concerning *what* they think. To keep things somewhat simple, let's suppose that the author is attempting to characterize what *all* rather than *some* or *most* men think. We can then ask, *what* is it that each man thinks? Each man thinks "men are smarter than women". This sentence however has two quantificational presence ambiguities since "men are smarter than women" can mean any of the following:

1. *All* men are smarter than *all* women.
2. *All* men are smarter than *some* women.
3. *Some* men are smarter than *all* women.
4. *Some* men are smarter than *some* women.

Each of these sentences has a corresponding translation in **QL** and each of these translations is distinct.

1. $(\forall x)(Mx \rightarrow (\forall y)(Wy \rightarrow Sxy))$

2. $(\forall x)(Mx \rightarrow (\exists y)(Wy \wedge Sxy))$
3. $(\exists x)(Mx \wedge (\forall y)(Wy \rightarrow Sxy))$
4. $(\exists x)(Mx \wedge (\exists y)(Wy \wedge Sxy))$

Making the quantificational ambiguity that is present in this sentence explicit is helpful for evaluating arguments containing this sentence. Consider the following argument:

- P1: Men think that they are smarter than women.
- C: Men are chauvinistic.

Depending on how P1 and P2 are interpreted in the above argument influences how the argument is evaluated. For example, let's consider just two different interpretations of P1 and P2, noting how it impacts the truth of the premises and the argument's validity:

1. If P1 is read as "all men think that all men are smarter than all women" and C is read as "all men are chauvinistic", then the argument is valid but unsound due to the falsity of P1.
2. If P1 is "some men think they are smarter than some women" and C is read as "all men are chauvinistic", then P1 is true but the argument is invalid.

Quantificational ambiguity *scope* The first type of quantificational ambiguity we considered involved sentences that were ambiguous about the *presence* of a quantifier. The interpretation of the sentence requires the addition of some quantificational expression (e.g., "some", "all", "many") but the sentence itself does not contain any such expression. The absence of the quantifier gives rise to the ambiguity. Let's now turn to a second type of quantificational ambiguity. This second type concerns sentences that are ambiguous about the scope of a quantifier. Consider the following sentence:

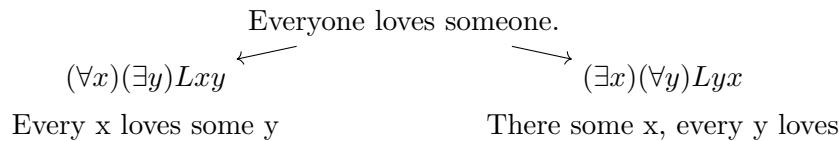
Everyone loves someone.

This sentence is ambiguous between two readings. First, this can be interpreted as saying that take any person you want, there is at least one person (not necessarily the same person) that they love. So, imagine a domain consisting of Jon, Tek, and Sal. On this reading, such a sentence would be true if the following is true:

1. Jon loves someone (e.g., Jon)
2. Tek loves someone (e.g., Jon)
3. Sal loves someone (e.g., Tek)

Notice that every single person loves at least one person from the domain. However, notice that it is not the same person (although if they loved the same person, this would also make that reading true).

On a second reading of "everyone loves someone", the reading is more restrictive. On this reading, everyone is taken to love one and the same person. That is, everyone loves the same person. Alternatively, we might express this by saying there is someone that everyone loves. Such a reading is translated as $(\exists x)(\forall y)Lyx$.



What we see then is that the sentence "everyone loves someone" is ambiguous about the scope of the quantifiers. In the first reading, the universal quantifier has wide scope. In the second reading, the existential quantifier has wide scope.

Let's consider another example of a quantificational scope ambiguity. Our previous example involved a case where it was ambiguous whether the universal or existential quantifier had wide scope. In this example, we will consider a case where it is ambiguous whether the negation or the universal quantifier has wide scope. To set up this type of ambiguity, let's consider the following sentence:

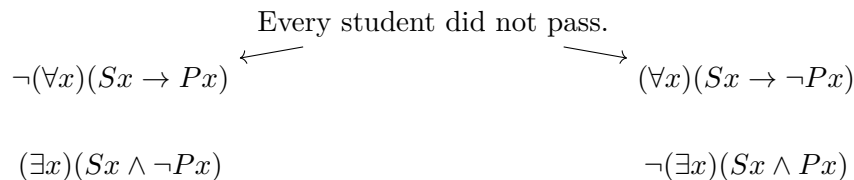
Every student passed.

The translation of this sentence into **QL** is straightforward. Let Sx express the English expression " x is a student" and let Px express the English expression " x passed the exam". The sentence is then translated as $(\forall x)(Sx \rightarrow Px)$. With this translation in mind, let's consider the following sentence:

Every student did *not* pass.

This sentence is ambiguous between two readings. First, the sentence can be read as "it is not the case that every student passed". On this reading, the sentence is translated as $\neg(\forall x)(Sx \rightarrow Px)$, the negation operator having wide scope. Equivalently, this sentence says that "there is a student who did not pass", translated as $(\exists x)(Sx \wedge \neg Px)$. Second, the sentence can be read as "no student passed". On this reading, the sentence is translated as $(\forall x)(Sx \rightarrow \neg Px)$. The negation operator having

narrow scope. Equivalently, this sentence says that "there does not exist a student that passed", translated as $\neg(\exists x)(Sx \wedge \neg Px)$.



It is sometimes remarked that both of these readings are not suggested by the sentence. That is, the sentence "every student did not pass" only refers to one of these readings. To see how the sentence is capable of suggesting both readings, consider the following two scenarios. First, suppose that a group of students are standing around comparing their scores on their logic exam. Tek begins, saying that he passed the exam with flying colors. He says he received a 100%. Liz also says that she passed the exam. She says that it actually was one of the easiest exam she has ever taken. With a sad look on her face, the professor informs the students that while she's happy they did so well, it would be a good idea not to be too vocal about their scores since "every student did not pass the exam." On this reading, the sentence "every student did not pass the exam" means "*it is not the case that every student passed the exam*": $\neg(\forall x)(Sx \rightarrow Px)$.

Second, suppose again that a group of students are standing around comparing their scores on their logic exam. Tek begins, complaining that this exam was really hard, noting he failed. Liz also says that she failed. She says that the exam was actually one of the hardest exams she has ever taken. With a sad look on her face, the professor apologizes to the students. The professor says that the exam was indeed very difficult and that "Every student did not pass the exam." On this reading, the sentence means "*no student passed the exam*".

In this section, we considered how ambiguity can arise in natural language sentences and impact how arguments are evaluated. In particular, we considered two types of quantificational ambiguities: quantificational presence ambiguities and quantificational scope ambiguities. The first type involves a sentence where the sentence suggests the presence of a quantifier, but it is ambiguous as to which quantifier is present. The second type involves a sentence where the sentence is ambiguous as to the scope of a quantifier.

Exercise 6.79

Identify the types of quantificational ambiguity present in the following sentences (involving presence or involving scope). What are the two readings of the sentence? Translate each reading into **QL**.

1. Criminals are bad.
2. All criminals are not greedy.
3. Smokers are not bad.
4. All fetuses are not persons.
5. Athletes are not heroes.

This chapter is a placeholder for a chapter.

In a previous chapter, it was noted that there are two different ways that we can understand what it means for a conclusion to follow from a set of premises. The first way is semantically. A conclusion follows from a set of premises semantically if and only if it is impossible for the premises to be true and the conclusion false. Understood semantically, we can check whether a conclusion follows from a set of premises using a table or tree. The second way is syntactic. A conclusion follows from a set of premises syntactically if and only if there is a derivation of the conclusion from the premises. Understood syntactically, we can show that a conclusion follows from a set of premises by providing a proof. While we have previously elaborated on the concept of syntactic entailment, in this chapter we redefine this notion for the language of predicate logic (**QL**). In addition, we introduce a new set of rules of derivation for *QL* that allow us to construct proofs in the language.

8.1 QL DERIVATION TERMINOLOGY

When discussing propositional logic proofs, we presented definitions for the concepts of a deductive apparatus, a derivation, and a syntactic entailment (consequence). These definitions remain the same for *QL* with one small change. The definitions only need to be modified to reflect that the deductive apparatus for *QL* involves predicate logic wffs rather than propositional logic wffs. To illustrate, consider the definition of syntactic consequence:

Definition 8.1.1: syntactic consequence of ϕ in RL

A formula ϕ is a syntactic consequence of a set Γ of *QL* wffs if and only if there is a derivation of ϕ from Γ . To express that ϕ is a syntactic consequence of Γ , we write $\Gamma \vdash_{RL} \phi$ or $\Gamma \vdash \phi$ for short.

Notice that the above definition is substantially the same as the definition of syntactic consequence for propositional logic. As mentioned, the only difference is that the deductive apparatus for *QL* involves predicate logic wffs rather than propositional logic wffs.

What does, however, need substantial elaboration are the additional rules of derivation that are needed for **QL**. In propositional logic, the deductive apparatus is *PD* which consists of a set of introduction-elimination rules.

In addition, some derived and equivalence rules were added to make the completion of proofs more natural and/or simpler. The deductive apparatus for predicate logic consists of all of the introduction-elimination rules from propositional logic as well as five new rules.

8.2 QL FOUR QUANTIFIER RULES

Let's refer to the deductive apparatus (or proofs system) for QL as RD . This proof system will consist of $PD+$ (the introduction-elimination rules for PD along with several additional rules) and all of the introduction-elimination quantifier rules articulated in this chapter.

As a reminder, recall that in an introduction-elimination system of proof, there are two types of rules: introduction rules and elimination rules. An introduction rule for an operator allows you to derive a wff that contains that operator as its main operator. For example, in propositional logic, the conjunction introduction rule ($\wedge I$) allows for deriving a wff of the form $\phi \wedge \psi$. An elimination rule for an operator allows you to derive a wff from a wff that contains that operator as its main operator. For example, in propositional logic, the conjunction elimination rule ($\wedge E$) allows for deriving a wff of the form ϕ from a wff of the form $\phi \wedge \psi$.

In what follows, four introduction-elimination rules will be articulated: universal introduction ($\forall I$), universal elimination ($\forall E$), existential introduction ($\exists I$), and existential elimination ($\exists E$). These rules will be used to construct proofs in **QL**.

8.2.1 Universal Elimination

Let's introduce the universal elimination rule by defining it, then illustrating it with a schematic example, and finally illustrating it with a minimal working example (MWE).

Definition 8.2.1: Universal Elimination ($\forall E$)

From any universally quantified wff $(\forall x)\phi(x_1, \dots, x_n)$, the wff $\phi(a_1 \dots a_n/x_1 \dots x_n)$ can be derived.
 $(\forall x)\phi(x_1 \dots x_n) \vdash \phi(a_1 \dots a_n/x_1, \dots, x_n)$

$\forall E$

Here is a schematic use of the rule:

- 1 $\forall x\phi(x_1 \dots x_n)$ P
- 2 $\phi(a_1 \dots a_n)$ $\forall E$ 1

Here is a MWE of the rule:

- 1 $(\forall x)Px$ P
- 2 Pa $\forall E$ 1

The basic idea behind $\forall E$ is that if you have a universally quantified wff $(\forall x)\phi$, you can move forward a step in the proof with a wff $\phi(a/x)$ that is the result of removing the universal quantifier and uniformly replacing each of the variables that is bound by the universal quantifier with any name of your choosing. So, for example, in the prior example, x in $(\forall x)Px$ was replaced with the name a to yield Pa . However, the name a could have been replaced with any name of your choosing, e.g., b or c . To illustrate this, consider the following example where $\forall E$ is used three times to derive Pa , Pb , and Pc .

- 1 $(\forall x)Px$ P
- 2 Pa $\forall E$ 1
- 3 Pb $\forall E$ 1
- 4 Pc $\forall E$ 1

One important feature of $\forall E$ is that when using it to replace more than one instance of a bound variable, the replacement of the variables with names must be done *uniformly*. For example, consider the following proof:

- 1 $(\forall x)Rxx$ P
- 2 Raa $\forall E$ 1
- 3 Rbb $\forall E$ 1

In the above example, line 1 contains two instances of the variable x bound by the universal quantifier. When $\forall E$ is used to derive a wff from line 1, each instance of the variable must be replaced by the same name. In line 2, the two instances of x are replaced by two instances of a . Similarly, in line 3, the two instances of x are replaced by two instances of b .

To contrast this correct use of $\forall E$ where instances of variables are replaced uniformly, let's consider an incorrect use of the rule where instances of variables are *not* replaced uniformly:

- 1 $(\forall x)Rxx$ P
- 2 Rab $\forall E$ 1, NO!

This use of $\forall E$ is incorrect since the two instances of x are replaced by two different names, a and b . That is, the replacement of variables with names is not done uniformly.

With our basic understanding of how to use $\forall E$ in place. Let's consider some examples of how to use $\forall E$ in a proof. First, let's consider an example of its use in an English argument and then that same argument in the formal language of predicate logic.

P1	Everyone is a person.	P
P2	If Alfred is a person, then Bob is a zombie.	P
IC	Therefore, Alfred is a person	from P1
C	Therefore, Bob is a zombie	from P2, IC, "if-elim"

Next, let's translate the above argument into predicate logic:

1	$(\forall x)Px$	P
2	$Pa \rightarrow Zb$	P
3	Pa	$\forall E$ 1
4	Zb	$\rightarrow E$ 2, 3

For our purposes, the key part of the above proof is line 3. In the above example, since line 1 states that *everyone* is a person, the use of $\forall E$ on line 1 justifies the new proposition (or wff) that Alfred is a person on line 3 (the IC). While we could have derived several other propositions (wffs) using $\forall E$, e.g., "Bob is a person" or "Frank is a person", "Alfred is a person" is derived since deriving this proposition (wff) allows the use of line 2 and $\rightarrow E$ to obtain the conclusion that Bob is a zombie.

Recall that one key consideration involved when using $\forall E$ is that instances of variables must be *uniformly* replaced with names. To see this more clearly, consider an English argument where $\forall E$ is used correctly and incorrectly.

- P1: Everyone loves themselves.
- C1: Therefore, Tek loves Tek. *Correct use of $\forall E$*
- C2: Therefore, Tek loves Jon. *Incorrect use of $\forall E$*

While C1 follows from P1 in the above argument, C2 does not. Just because every individual x loves themselves x , it does not follow that one person Tek loves another person Jon. Similarly, when using $\forall E$, the replacement of variables with names must be done uniformly. From $(\forall x)Lxx$, wffs like Laa, Lbb, Lcc and so forth may be derived since the variables are replaced uniformly (with the same name). The rule does not permit deriving wffs such as Lab, Lba, Lac and so other (where the names are different).

Let's consider another example involving $\forall E$. Consider the following entailment: $(\forall x)(Px \rightarrow (\forall y)(Qx \rightarrow Wy)), Pb, Qb \vdash Wt$.

1	$(\forall x)(Px \rightarrow (\forall y)(Qx \rightarrow Wy))$	P
2	Pb	P
3	Qb	P
4	$Pb \rightarrow (\forall y)(Qb \rightarrow Wy)$	$\forall E$ 1
5	$(\forall y)(Qb \rightarrow Wy)$	$\rightarrow E$ 2, 4
6	$Qb \rightarrow Wt$	$\forall E$ 5
7	Wt	$\rightarrow E$ 3, 6

In the above example, the proof begins by using $\forall E$ on line 1, uniformly replacing each x with the name b . This is possible because the wff at line 1 has, as its main operator, the universally quantifier $(\forall x)$. As the proof continues, notice that $\forall E$ is *not* used on the conditional at line 4. This is because $\forall E$ can only be used on universally quantified wffs. However, once the universally quantified wff $(\forall y)(Qb \rightarrow Wy)$ is derived at line 5, $\forall E$ can be used on it to derive line 6.

One question that is commonly asked when using $\forall E$ is the following: When using $\forall E$, I know that I can uniformly replace bound variables with any name of my choosing. However, which name should I choose? The choice of names is guided by two considerations. First, it is guided by names already occurring in the proof. Second, it is guided by names in the conclusion. Let's consider an example involving both cases.

Let's consider each one of these considerations. First, consider the entailment $(\forall x)Px, Pc \rightarrow Qd \vdash Qd$. The first step is to set up the proof:

1	$(\forall x)Px$	P
2	$Pc \rightarrow Qd$	P, Qd

Next, while we can derive Pa, Pb, Pc, Pd and so on from $(\forall x)Px$ at line 1, a more economical choice is to replace x with c since c already occurs in the proof at line 2. In addition, choosing c allows for writing Pc in the proof and this wff can be used with the conditional at line 2.

1	$(\forall x)Px$	P
2	$Pc \rightarrow Qd$	P, Qd
3	Pc	$\forall E$ 1
4	Qd	$\rightarrow E$ 3, 2

Next, let's illustrate the second consideration. Consider the entailment $(\forall x)Px \vdash Pc \wedge (Pd \wedge Pe)$. The first step is to set up the proof:

1 $(\forall x)Px$ P

Notice that the wff we want to derive contains the names c , d , and e . This helps guide what names we should use when using $\forall E$. Since the wff we wish to derive is $Pc \wedge (Pd \wedge Pe)$, we can use three separate instances of $\forall E$ to derive Pc, Pd, Pe and then use $\wedge I$ twice to derive $Pc \wedge (Pd \wedge Pe)$.

1 $(\forall x)Px$ P
2 Pc $\forall E$ 1
3 Pd $\forall E$ 1
4 Pe $\forall E$ 1
5 $Pd \wedge Pe$ $\wedge I$ 3, 4
6 $Pc \wedge (Pd \wedge Pe)$ $\wedge I$ 2, 5

In this section, we considered the use of $\forall E$ in a proof. As this is an elimination rule, this derivation rule allows for deriving wffs from universally quantified. In the next section, we consider our first introduction rule: $\exists I$.

Exercise 8.80

Focusing on $\forall E$, solve the following proofs.

1. $(\forall x)Px, (\forall z)Qz \vdash Pa \wedge Qa$
2. $(\forall x)\neg Px, (\forall x)Px \vdash \neg Pa \wedge \neg Pb$
3. $Lab \rightarrow Sa, (\forall x)(\forall y)Lxy \vdash Sa$
4. $(\forall x)(\forall y)Pxy, (\forall y)Pay \rightarrow (\forall z)Qzz \vdash Qbb$
5. $(\forall x)(\forall y)(\forall z)(Lxy \wedge Syz) \vdash Laa$

8.2.2 Existential Introduction

The next rule we will consider allows for deriving existentially quantified wffs. This rule is known as "existential introduction".

Definition 8.2.2: Existential Introduction ($\exists I$)

From $\phi(a_n)$, an existentially quantified wff $(\exists x)\phi(x_i/a_n)$ can be derived, where (1) $i \leq n$ and (2) the var x_i is not in the original $\phi(a_n)$, viz., x_i is not a variable already bound by a quantifier.

$$\phi(a_n) \vdash (\exists x)\phi(x_i/a_n)$$

The idea is that if you have a non-quantified proposition that contains a name, you can replace that name with an existentially quantified variable.

$$\begin{array}{ll} 1 & \phi(a_n) \quad \text{P} \\ 2 & (\exists x)\phi(x_i/a_n) \quad \exists I \end{array}$$

In the above example, some number n of instances of the name ' a ' is replaced with instances of the existentially quantified variable ' x '. In making this point, it is important to note that the number of instances of the name ' a ' that are replaced with instances of the existentially quantified variable ' x ' is not specified. That is, the number of instances of the name ' a ' that are replaced with instances of the existentially quantified variable ' x ' can be one or more. However, there are two important things to note. First, whatever number of instances of the name ' a ' that are replaced with instances of the existentially quantified variable ' x ', the replacement must be uniform. That is, if an instance of $\exists I$ is used to replace more than one name, then it must be the same name. Second, whatever variable is used to replace the name, that variable must not already be bound by a quantifier.

Let's consider a minimal working example (MWE) of $\exists I$.

$$\begin{array}{ll} 1 & Pa \quad \text{P} \\ 2 & (\exists x)Px \quad \exists I1 \end{array}$$

In the above example, there is a wff Pa and from this wff, the existentially quantified wff $(\exists x)Px$ is derived. This is done by replacing at least one instance of the name a with the existentially quantified variable x .

Now that we have defined $\exists I$, considered a schematic use of the rule, and looked at a MWE, let's consider some examples of how to use $\exists I$ in a proof. First, let's consider an example of its use in an English argument and then that same argument in the formal language of predicate logic. Consider the following argument:

$$\begin{array}{ll} \text{P1 Rick is a zombie.} & \text{Prem} \\ \text{C Therefore, someone is a zombie.} & \exists I, \text{ from 1} \end{array}$$

The above example involves reasoning from an expression involving a sentence containing a name "Rick". In order to derive the conclusion, the name "Rick" is replaced by a sentence that specifies a quantity of individ-

uals that have the property of being a zombie. In this case, the sentence is "someone is a zombie". In formal logic, our argument looks like this:

- 1 Zr P
- 2 $(\exists x)Zx$ $\exists I$ 1

In the above example, the wff Zr is a wff with a name r . The $\exists I$ derivation rule allows for deriving an existentially quantified wff $(\exists x)Zx$ from this wff Zr : where r has been replaced with x and an existential quantifier with x is now binding the x .

There are two key conditions placed on $\exists I$. The first is that it permits the replacement of more than one instance of a name with an existentially quantified variable provided that the replacement is uniform. What this means is that if more than one name is replaced in a single use of $\exists I$, the same name must be replaced with the same variable. For example, consider the following proof:

- 1 Lbb P
- 2 $(\exists x)Lxx$ $\exists I$ 1

Notice that the use of $\exists I$ on line 2 replaces two instances of the name b with two instances of the existentially quantified variable x . This is permitted since the same name b is being replaced with the same variable x . What is not occurring is that the name b is being replaced with two different variables, e.g., x and y , or there are two different names, e.g., a and b , being replaced with the same variable x in a single use of $\exists I$. Concerning this condition, note that the condition does not require that each name be replaced with an existentially quantified variable. Rather, it says that if more than one name is replaced, each name must be replaced in a uniform manner. For example, consider the following proof:

- 1 Lbb P
- 2 $(\exists x)Lbx$ $\exists I$ 1
- 3 $(\exists x)Lxb$ $\exists I$ 1

Notice that in the above example, the name b is replaced with the existentially quantified variable x in two different ways. In the first instance, the name b is replaced with the existentially quantified variable x in the first position of the wff Lbx . In the second instance, the name b is replaced with the existentially quantified variable x in the second position of the wff Lxb . This is permitted since the name b is replaced with the same variable x in both instances. This proof also illustrates that the condi-

tion does not require that each name be replaced with an existentially quantified variable.

To further instill understanding of this condition, consider the following proof that contains some incorrect uses of $\exists I$. As you read through the proof try to identify where $\exists I$ is used incorrectly.

1	Lab	P
2	Rbb	P
3	$(\exists x)Lax$	$\exists I$ 1
4	$(\exists x)Lxb$	$\exists I$ 1
5	$(\exists x)Lxx$	$\exists I$ 1
6	$(\exists x)Rxx$	$\exists I$ 2
7	$(\exists x)Rxy$	$\exists I$ 2

There are two errors in the above proof. The first is located at line 5. In using $\exists I$ at line 5, the replacement of names with variables is not done uniformly. Notice that there are two different names, a and b , being replaced with the same variable x . The second error is located at line 7. In using $\exists I$ at line 7, the replacement of names with variables is also not done uniformly. If you look at line 7, you will notice that both names, b , are replaced with two different variables, x and y , with y being a free variable.

Finally, the second condition placed on $\exists I$ is that the variable used to replace the name must not already be bound by a quantifier. For example, consider the following proof:

1	$(\forall x)Lxb$	P
2	$(\exists x)(\forall x)Lxx$	$\exists I$ 1. <i>Incorrect!</i>
3	$(\exists y)(\forall x)Lxy$	$\exists I$ 1. <i>Correct!</i>

Notice that line 1 contains an quantified variable x in it. Since it contains the quantified variable x , the rule does not permit replacing the name b with x . Doing so would result in the case where it is unclear which quantifier is binding the variable x in the wff $(\exists x)(\forall x)Lxx$. That is, does $(\exists x)(\forall x)Lxx$ say "some loves themselves" or "everyone loves themselves" or "someone loves everyone" or "everyone loves someone". To avoid this ambiguity, the rule requires that the variable used to replace the name must not already be bound by a quantifier. This is why the use of $\exists I$ at line 3 is permissible. In this case, the name b is replaced with the existentially

quantified variable y which is not already bound by a quantifier in the wff to which $\exists I$ is being applied.

Finally, let's conclude this section on $\exists I$ with a general guideline on when and how you should try to use $\exists I$ in a proof. First, as a general reminder, much like other introduction rules, you want to consider using $\exists I$ when you want to derive an existentially quantified wff. If the conclusion of an argument is $(\exists x)Px$, then you should consider using $\exists I$ if possible. Second, and perhaps more helpful advice, since $\exists I$ requires replacing a name with an existentially quantified variable, another goal you may have is to try to derive a wff that contains the name you want to replace. For example, if you want to derive $(\exists x)Px$, then you should try to derive a wff that contains the name a that you want to replace. Let's illustrate this second suggestion with an example. Consider the setup of the following proof

1	$(\forall x)(Px \wedge Rx)$	P
2	$Pa \rightarrow Wb$	P, $(\exists x)Wy$
.		
.		
k		
k+1	$(\exists x)Wy$	$\exists I$ k

In examining the above proof, it would be nice if there were a wff that we could immediately use $\exists I$ on to derive the conclusion $(\exists x)Wy$. However, if we were to use $\exists I$ on line 2, we would be using it on a conditional. This would result in a wff like $(\exists y)(Pa \rightarrow Wy)$. This wff is not the conclusion we want. What we would like then is Wb because if we had Wb on its own, we could use $\exists I$ on it to derive $(\exists x)Wy$. With that in mind, let's work backward one step in our proof.

1	$(\forall x)(Px \wedge Rx)$	P
2	$Pa \rightarrow Wb$	P, $(\exists x)Wy$
.		
.		
k	Wb	?
k+1	$(\exists x)Wy$	$\exists I$ k

The question then is how to derive Wb . The line that stands out is line 2, the conditional $Pa \rightarrow Wb$. Since we want to derive Wb , we will need

Pa on its own line. We can get this wff by using $\forall E$ on line 1. We can now complete the proof as follows:

1	$(\forall x)(Px \wedge Rx)$	P
2	$Pa \rightarrow Wb$	P, $(\exists x)Wy$
3	Pa	$\forall E$ 1
4	Wb	$\rightarrow E$ 2, 3
5	$(\exists x)Wy$	$\exists I$ 4

In short, once you have the idea that you would like to use $\exists I$ to derive an existentially quantified wff, it is helpful to focus your attention on trying to derive a wff that contains the name you want to replace. Let's conclude this section with one more example of this strategy for using $\exists I$. Consider the following: $(\forall x)Px, (\forall y)Zy \vdash (\exists x)(Px \wedge Zx)$.

1	$(\forall x)Px$	P
2	$(\forall y)Zy$	P

Since the wff that we want to derive is the existentially quantified wff $(\exists x)(Px \wedge Zx)$, we should try to derive a wff that contains the name a . The question then is what wff do we want to derive to ultimately use $\exists I$ on? Let's consider one incorrect approach:

1	$(\forall x)Px$	P
2	$(\forall y)Zy$	P
3	Pa	$\forall E$ 1
4	Zb	$\forall E$ 2
5	$Pa \wedge Zb$	$\wedge I$ 3, 4

In the above example, we derived $Pa \wedge Zb$. While we can use $\exists I$ on this wff, we cannot use $\exists I$ on this wff to derive $(\exists x)(Px \wedge Zx)$ since replacement of names with variables must be uniform. Since $Pa \wedge Zb$ contains two distinct names a and b , they cannot be replaced with the same variable x . This is why the tip for thinking about which wff you would like to ultimately use $\exists I$ on is helpful. Since we want to use $\exists I$ on a wff that contains two instances of the same name, our initial approach in the proof should be to derive a wff like $Pa \wedge Za$ or $Pb \wedge Zb$. Let's redo the proof with this in mind:

1	$(\forall x)Px$	P
2	$(\forall y)Zy$	P
3	Pa	$\forall E$ 1
4	Za	$\forall E$ 2
5	$Pa \wedge Za$	$\wedge I$ 3, 4
6	$(\exists x)(Px \wedge Zx)$	$\exists I$ 5

In the above example, we benefited from the use of a strategy associated with $\exists I$. Namely, if you want to derive an existentially quantified wff, direct your attention not to the wff you want to derive, but to a wff that, if you were to use $\exists I$ on it, would result in the wff you want to derive.

Exercise 8.81

Focusing on $\exists I$ and, to a lesser extent $\forall E$, solve the following proofs.

1. $Pa, Rb \vdash (\exists x)(Px \wedge Rb)$
2. $Pa, Rb \vdash (\exists x)Px \wedge (\exists x)Rx$
3. $(\exists x)Rx \rightarrow Mc, Ra \vdash (\exists x)Mx$
4. $Pa, Rb \vdash (\exists x)(\exists y)(Px \wedge Ry)$
5. $(\forall x)Lax, (\exists x)(\forall y)Lxy \rightarrow Ra \vdash (\exists x)Rx$

8.2.3 Universal Introduction

In the previous sections, we considered the use of universal elimination ($\forall E$) and existential introduction ($\exists I$). Universal elimination allows for deriving a wff from a universally quantified wff, while existential introduction allows for deriving an existentially quantified wff from a wff that contains a name. In this section and the next, we consider the introduction and elimination version of these two rules. In this section, we consider universal introduction ($\forall I$), a rule that will, provided certain strict conditions are met, allow for deriving a universally quantified wff from a wff. Let's begin by stating the rule.

Definition 8.2.3: Universal Introduction ($\forall I$)

Let α be any name. From a wff $\phi(\alpha_1 \dots \alpha_n)$, a universally quantified wff $(\forall x)\phi(x_1 \dots x_n/\alpha_1 \dots \alpha_n)$ can be derived where (1) the name α does not occur in a wff as premise or as an assumption in an open subproof, (2) the name α does not occur in the derived wff $(\forall x)\phi(x_1 \dots x_n/\alpha_1 \dots \alpha_n)$, and (3) where x is not a variable in ϕ .

$\forall I$

$$\phi(\alpha_1 \dots \alpha_n) \vdash (\forall x)\phi(x_1 \dots x_n / \alpha_1 \dots \alpha_n)$$

The above rule is rather complex so let's consider its parts. The essence of universal introduction is that it is a derivation rule that allows for deriving a universally quantified wff from a wff that contains a name. So, in essence, if there is a wff Pa , it allows for deriving a wff $(\forall x)Px$. However, this rule has three restrictions, viz., three situations in which $\forall I$ cannot be used. Let's consider these three restrictions before examining some correct uses of $\forall I$.

The first restriction on the use of $\forall I$ is that the name α does not occur in a wff as a premise or an assumption in an open subproof. In short, if the wff Pa is in premise or an assumption, we cannot use $\forall I$ to derive $(\forall x)Px$. Let's illustrate this restriction with the following English argument.

- P1: John won the lottery.
- C: Therefore, everyone won the lottery.

The above argument is clearly invalid. Just because a single individual won the lottery does not mean that everyone won the lottery. Translating this argument into predicate logic results in the following proof:

- 1 Wj P
- 2 $(\forall x)Wx$ $\forall I$ 1, *Incorrect!*

In the above example, the argument attempts to use $\forall I$ to derive $(\forall x)Wx$ from Wj . However, this is not permitted since it violates the restriction that $\forall I$ cannot be used to replace a name j that occurs in a premise.

The first restriction states that $\forall I$ cannot be used on a wff that contains a name that occurs in a premise or in an assumption in an active subproof. In other words, if you have an active subproof where you have assumed a wff that contains a name, you cannot use $\forall I$ on any wff containing that name until the subproof is closed. Let's first consider an English argument where this restriction is violated and then consider the same argument in predicate logic.

- 1 | Let's assume John won the lottery. A
- 2 | Therefore, everyone won the lottery. $\forall I$ 1, *Incorrect!*

The above example is clearly invalid. It does not follow from the *assumption* that John won the lottery that everyone won the lottery. Translating this argument into predicate logic results in the following:

1	$ Wj$	A
2	$ (\forall x)Wx$	$\forall I$ 1, <i>Incorrect!</i>

In the above example, the wff Wj is assumed at line 1. This wff contains the name j . Since the subproof is still open, $\forall I$ cannot be used on a wff that contains j .

Universal introduction is a rule that allows for deriving a universally quantified wff from a wff that contains a name provided certain strict conditions are met. As we saw, the first restriction on using $\forall I$ is that the name α does not occur in a wff as a premise or an assumption in an open subproof. Let's turn to the second restriction.

To understand this restriction, it is helpful to state when $\forall I$ is used, it is used on a wff containing a name. Again, providing certain restrictions are met, $\forall I$ can be used on a wff like Pa to derive $(\forall x)Px$. One way this rule can be understood is that the name a is being *replaced* by a universally quantified variable. In our example from Pa to $(\forall x)Px$, the name a is being replaced by the variable x in the wff $(\forall x)Px$. With this in mind, the second restriction states that when using $\forall I$ on a wff $\phi(\alpha_1 \dots \alpha_n)$, the name α cannot occur in the derived wff $(\forall x)\phi(x_1 \dots x_n/\alpha_1 \dots \alpha_n)$. To put this more simply, whenever universal introduction is used on a wff containing more than one name, replacement of a name with a variable must be comprehensive. Even more simply, if $\forall I$ were to be used on Laa (replacing a with x), both a 's would need to be replaced. That is, $(\forall x)Lxx$ could be derived, but not $(\forall x)Lax$. Let's illustrate this restriction with the following English argument.

- John loves himself.
- Therefore, everyone loves John. *Incorrect!*

The above argument is clearly invalid. Just because John loves himself does not entail that everyone loves John. Translating this argument into predicate logic results in the following:

1	Ljj
2	$(\forall x)Lxj \quad \forall I$ 1, <i>Incorrect!</i>

In the above example, the argument attempts to use $\forall I$ to derive $(\forall x)Lxj$ from Ljj . However, this is not permitted since it violates the restriction if $\forall I$ is used and the name j is to be replaced by a universally quantified variable x , it cannot be the case that j that occurs in the derived wff. In short, if $\forall I$ is used on a wff containing more than one name of the same

type, that name must be completely replaced by universally quantified variables.

We have considered two of the three restrictions on $\forall I$. The final condition states that when using $\forall I$ on a wff ϕ to derive a universally quantified wff $(\forall x)\phi$, the variable x cannot already be found in ϕ . This restriction is in place to avoid ambiguity. Consider the following use of $\forall I$ that violates this third restriction:

- 1 $(\exists x)Lxa$
- 2 $(\forall x)(\exists x)Lxx$ $\forall I$ 1, *Incorrect!*

Notice that the wff $(\exists x)Lxa$ contains a name a but also a variable x . The third restriction states that if $\forall I$ is used on this wff, it cannot be used to replace the name a with the variable x . This is because the variable x is already present in the wff $(\exists x)Lxa$. This is why the use of $\forall I$ at line 2 is incorrect.

Now that we have considered the three restrictions on $\forall I$, let's consider some correct uses of $\forall I$. Consider the following proof:

- 1 $(\forall x)Lxx$ P
- 2 Laa $\forall E$ 1
- 3 $(\forall y)Lyy$ $\forall I$ 2

Notice that a is not in a premise or in an assumption of an open subproof (restriction 1), each a is replaced by y in the derived wff so there are no a s in the resulting universally quantified wff (restriction 2), and y is not already in the wff $(\forall x)Lxx$ (restriction 3). Since this proof does not violate any of the three restrictions, it is a permissible use of $\forall I$.

Before looking at additional correct uses, let's address one possible concern. Suppose we were to translate the above proof into English. Translating this argument would result in the following:

- 1 Everyone loves themselves P
- 2 Al loves Al $\forall E$ 1
- 3 Everyone loves themselves $\forall I$ 2

One concern then is that the above argument may appear to be invalid. How can it follow from the proposition that "Al loves Al" that "Everyone loves themselves"? The reason why this argument is valid is that the name a ('Al') is arbitrarily-selected. The fact that Al loves Al is not unique to Al, but a relation that could be applied to any item in the domain. For

consider that we might have derived line 2 using *any* name b , c , d , and so on. That is, could have reasoned to "Bob loves Bob", "Chris loves Chris", "David loves David" and so on. The idea then is that if we can derived "Al loves Al" and we could replace "Al" with any name, it follows that "Everyone loves themselves". That is, since we can derive Laa , and a in Laa could be any name, it follows that $(\forall x)Lxx$ is true.

Let's consider another example of a correct use of $\forall I$.

1	$\frac{Raa}{\quad}$	A
2	$\frac{Raa}{\quad}$	R 1
3	$Raa \rightarrow Raa$	$\rightarrow I$ 1-2
4	$(\forall x)(Rxx \rightarrow Raa)$	$\forall I$ 3

In the above example, $(\forall I)$ is applied to line 3 even though a is in the assumption. Doesn't this violate our first condition that $\forall I$ cannot be used on an assumption in an open subproof? The answer is "no" since the assumption is closed at line 3. Since the name a is not a name in an assumption in an open subproof (and no other restrictions are violated), it is permissible to use $\forall I$ on $Raa \rightarrow Raa$ to derive $(\forall x)(Rxx \rightarrow Raa)$. In addition, notice that our selection of the name in the assumption is arbitrary. Since any name (b , c , d , etc.) could have been selected, we could have derived a conditional of the form $R\alpha\alpha \rightarrow R\alpha\alpha$ where α is any name. So, since we could have derived $R\alpha\alpha \rightarrow R\alpha\alpha$, it follows that $(\forall x)(Rxx \rightarrow Raa)$.

Let's consider another example of a correct use of $\forall I$. Consider the following proof:

1	$(\forall x)(Pxc \vee Qx)$	P
2	$(\forall x)\neg Qx$	P
3	$Pac \vee Qa$	$\forall E$ 1
4	$\neg Qa$	$\forall E$ 2
5	Pac	DS 3, 4
6	$(\forall x)Pxc$	$\forall I$ 5

Notice that the use $\forall I$ at line 6 does not violate any of the restrictions. First, a does not occur in any premise or open assumption in the proof. While c does occur in a premise at line 1, at no point is name c being replaced with a universally quantified wff. Second, there is not an a in $(\forall x)Pxc$. Third, the wff to which $\forall I$ is being applied does not already contain the variable x .

Exercise 8.82

Provide proofs for the following entailments.

1. $(\forall x)Px \vdash (\forall y)Py$
2. $(\forall x)(Px \wedge Rx) \vdash (\forall x)Px$
3. $(\forall x)Px \wedge (\forall y)Sy \vdash (\forall x)(Px \wedge Sx)$
4. $(\forall x)(\forall y)Lxy, (\forall x)Lxx \rightarrow (\exists x)(\exists y)Lxy \vdash (\exists x)(\exists y)Lxy$
5. $\vdash (\forall x)(Px \rightarrow Px)$

8.2.4 Existential Elimination

The final intelim rule we will consider is existential elimination ($\exists E$). This rule allows for deriving a wff from an existentially quantified wff. Let's begin by stating the rule.

$\exists E$

Definition 8.2.4: Existential Elimination ($\exists E$)

Let α be any name and x be any variable. From an existentially quantified wff $(\exists x)\phi$, a wff ψ can be derived from $(\exists x)\phi$ and a subproof that begins with $\phi(\alpha/x)$ and ends with ψ , provided (1) the name α does not occur in any premise or in an active proof (or subproof) prior to its arbitrary introduction in the assumption $\phi(\alpha/x)$ and (2) the name α does not occur in ψ .

In other words, $\exists E$ is a rule that allows you (provided certain restrictions are met) to derive a wff ψ provided you have an existentially quantified wff $(\exists x)\phi$ and a subproof that begins with the assumption $\phi(\alpha/x)$ and ends with ψ in that subproof. The rule has the following form:

n	$(\exists x)\phi x$	P
j	$\phi(\alpha/x)$	A
	$\vdots$	
j+i	ψ	
j+(i+1)	ψ	$\exists E, n, j-(j+1)$

Let's clarify this rule by highlighting several key points. First, since $\exists E$ is an elimination rule, the rule requires an existentially quantified wff to exist in the proof. Just as $\wedge E$ derives a wff ψ or ϕ from a conjunction $\phi \wedge \psi$, existential elimination involves deriving a wff ψ from an existentially quantified wff $(\exists x)\phi$. Second, the rule requires a subproof that begins with an assumption that is related to the existentially quantified wff.

That is, if the existentially quantified wff is $(\exists x)\phi$, the assumption must be of the form $\phi(\alpha/x)$. This is an assumption where each name α has replaced each variable x in the wff ϕ . Let's consider an illustration of this. Suppose we have the following existentially quantified wff: $(\exists x)(Px \wedge Qx)$. The assumption that would be required for $\exists E$ would be a wff of the form $(P\alpha \wedge Q\alpha)$ where α is any name. To illustrate further, consider the following:

1	$(\exists x)(Px \wedge Qx)$	P
2	$\mid (Pb \wedge Qb)$	A

In the above example, the existentially quantified wff is $(\exists x)(Px \wedge Qx)$. When the assumption is made at line 2, each existentially quantified variable has been uniformly replaced in the wff $(Px \wedge Qx)$ with the name b .

This feature of existential elimination takes us to our first restriction on the use of the rule. Namely, that if we wish to use $\exists E$, the name α used in the assumption $\phi(\alpha/x)$ must not occur in any premise or in an active proof (or subproof) prior to its arbitrary introduction in the assumption $\phi(\alpha/x)$. Another popular way of putting this restriction is that the name α must be *fresh* to the proof. Let's illustrate this restriction with an example.

1	Na	P
2	$Ma \rightarrow Aa$	P
3	$(\exists x)Mx$	P
4	$\mid Ma$	A, <i>Incorrect if using $\exists E$ later!</i>

Let's consider another example of an incorrect use of $\exists E$ that ignores this restriction. In this example we'll look at an example where the name α is used in an active subproof.

1	$(\exists z)Pzz$	P
2	$\mid Pbb \rightarrow Wcc$	A
3	$\mid \mid Pbb$	A
4	$\mid \mid Wcc$	$\rightarrow E$ 2, 3
5	Wcc	$\exists E$ 2, 4-5 <i>Incorrect!</i>

In the above example, notice that the name b is found in the assumption of an active subproof (line 2). According to the restriction, it is not permissible to use $\exists E$ on another subproof where b is used in the assumption.

That is, while it is permissible to assume Pbb at line 3, it is not permissible to use $\exists E$ on a subproof where b is used in the assumption.

The second restriction on $\exists E$ is that the name α used in the assumption $\phi(\alpha/x)$ must not occur in ψ . Let's illustrate this restriction with an example.

1	$(\exists x)Mx$	P
2	<u>Mb</u>	A
3	Mb	$R\ 2$
4	Mb	$\exists E\ 1, 2-3, \text{Incorrect!}$

Notice that in the above example, Mb is derived at line 3. Since b is the name used in the assumption at line 2, it cannot be used in the derived wff at line 4. To better see why this is the case, let's consider an English argument that is analogous to the above proof. Suppose you are a detective and you have come across a corpse. You have come to the conclusion that "someone is a murderer". As you want to solve the case, you begin to reason as follows:

1	Someone is a murderer.	P
2	<u>Assume Tek is a murderer.</u>	A
3	Tek is a murderer.	$R\ 2$
4	Tek is a murderer.	$\exists E\ 1, 2-3$

The reasoning performed in the above example is clearly mistaken. While it does follow from the assumption that Tek is a murderer, that he is a murderer, this subproof along with the proposition that "Someone is a murderer" does not entail that *Tek is a murderer*. Let's illustrate this same point with another English argument.

1	Al is my neighbor.	P
2	If Al is a murderer, then he should be arrested.	P
3	Someone is a murderer.	P
4	<u>Assume Al is the murderer.</u>	A
5	Therefore, Al should be arrested.	$\rightarrow E\ 2, 4$
6	Al should be arrested.	$\exists E\ 3, 4-5$

There are at least two problems with the above proof. First, it violates the first restriction that when using $\exists E$, the name used in the assumption should not be found in a premise. This violation occurs at line 4. Second,

it violates the second restriction that the name used in the assumption (at line 4) should not be found in the wff derived from the subproof. This violation occurs at line 6.

At this point, the rule for $\exists E$ has been defined and various restrictions on its use have been clarified. To conclude this section, let's consider proofs where $\exists E$ is used correctly. First, consider the following entailment: $(\exists x)Px \vdash (\exists y)Py$.

1	$(\exists x)Px$	P
2	Pa	A
3	$(\exists y)Py$	$\exists I, 2$
4	$(\exists y)Py$	$\exists E 1, 2-3$

Notice that the above proof (while trivial) does not violate any of the restrictions on $\exists E$. The name a that is assumed at line 2 is not found in any premise or in an active subproof prior to its arbitrary introduction in the assumption at line 2. In addition, the name a is not found in the derived wff at line 4.

Next, let's consider a slight variation on the prior proof. In this example, let's provide a proof for $(\exists x)Px \vdash (\exists x)(Px \vee Qx)$.

1	$(\exists x)Px$	P
2	Pa	A
3	$Pa \vee Qa$	$\vee I 2$
4	$(\exists x)(Px \vee Qx)$	$\exists I 3$
5	$(\exists x)(Px \vee Qx)$	$\exists E 1, 3-4$

Again, notice that none of the restrictions have been violated. The name a is not found in any premise or in an active subproof prior to its arbitrary introduction in the assumption at line 2. In addition, the name a is not found in the derived wff at line 5.

Finally, let's consider a proof containing wffs where there are names in the premises of the proof but where we avoid violating any of the restrictions on $\exists E$. That is, let's consider a proof for the following entailment: $Ra, (\exists x)(Px \wedge Rx) \vdash (\exists x)Px$.

1	Ra	P
2	$(\exists x)(Px \wedge Rx)$	P
3	$Pb \wedge Rb$	A
4	Pb	$\wedge E$ 3
5	$(\exists x)Px$	$\exists I$ 4
6	$(\exists x)Px$	$\exists E$ 2, 3-5

Notice that in the above proof, there is a name a in the premise at line 1. Because this name is present, it is necessary to use a name other than a in our assumption at line 3. For this reason, the name b is chosen (although we could have, in this proof, selected any name other than a , e.g., c , d , etc.). Since b this name is not found in any premise or in an active subproof prior to its arbitrary introduction in the assumption at line 3 and because the name b is not found in the derived wff at line 6, none of the restrictions on $\exists E$ have been violated.

Exercise 8.83

Provide proofs for the following entailments.

1. $(\exists x)Lxx \vdash (\exists y)Lyy$
2. $(\exists x)Lxx \vdash (\exists y)(\exists x)Lxy$
3. $(\forall x)Lax, (\exists x)Px \vdash (\exists y)Py$
4. $Pa, Pa \rightarrow (\exists x)Qx \vdash (\exists y)Qy$
5. $(\exists x)(Bx \rightarrow Mx), (\forall x)(Mx \rightarrow Dxc) \vdash (\exists x)(Bx \rightarrow Dxc)$

8.3 QUANTIFIER NEGATION

In addition to the four introduction and elimination rules for quantified propositions, we will add one final rule. The final rule to our deductive apparatus is an equivalence rule (or rule of replacement). This rule, known as "Quantifier Negation", allows us to replace negated quantified subformulas with non-negated quantified subformulas, and vice versa.

However, before we define this rule, let's consider two proofs. First, we will prove that $\neg(\forall x)Px \vdash (\exists x)\neg Px$.

1	$\neg(\forall x)Px$	P
2	$\neg(\exists x)Px$	A
3	$\neg Pa$	A
4	$(\exists x)Px$	$\exists I$ 3
5	$\neg(\exists x)Px$	R 2
6	Pa	$\neg E$ 3-5
7	$(\exists x)Px$	$\exists I$ 6
8	$\neg(\exists x)Px$	R 2
9	$(\exists x)\neg Px$	$\neg E$ 2-8

In our second proof, we will show that $(\exists x)\neg Px \vdash \neg(\forall x)Px$.

1	$(\exists x)\neg Px$	P
2	$\neg Pa$	A
3	$(\forall x)Px$	A
4	Pa	$\forall E$ 3
5	$\neg Pa$	R 2
6	$\neg(\forall x)Px$	$\neg I$ 3-5
7	$\neg(\forall x)Px$	$\exists E$ 1, 2-6

Our two proofs collectively show that $\neg(\forall x)Px$ and $(\exists x)\neg Px$ are interderivable. More generally, they indicate that any wff of the form $\neg(\forall x)\phi$ is interderivable with $(\exists x)\neg\phi$. With this in mind, we can define the rule of replacement QN as follows:

QN

Definition 8.3.1: Quantifier Negation (QN)

$$\neg(\forall x)\phi \dashv\vdash (\exists x)\neg\phi$$

$$\neg(\exists x)\phi \dashv\vdash (\forall x)\neg\phi$$

The earlier two proofs showed that QN is not strictly necessary (the proof of the second version of QN is left as an exercise). That is, any proof making use of QN can be solved without it. The second version of QN is proved in a similar manner. Our addition to our proof system is largely due to two reasons. The first reason is one of economy. Proofs involving QN are commonplace and we would like not to have to repeat these steps each time we want to prove this entailment. Instead, we can use the rule of replacement QN to derive this entailment more efficiently. Second, QN is a natural rule of replacement. That is, it is a rule that is intuitive and

easy to understand. That is, since "It is not the case that everyone is a person" is equivalent to "There is someone who is not a person", it is natural to allow for the replacement of one with the other.

Let's conclude this section by considering a few proofs that makes use of QN . Consider the following entailment: First, let's revisit an earlier proof we completed in this section, but this time, let's make use of QN . Consider the proof of the following entailment: $\neg(\forall x)Px \vdash (\exists x)\neg Px$.

- 1 $\neg(\forall x)Px$ P
- 2 $(\exists x)\neg Px$ QN 1

Notice that QN is applied to the wff $\neg(\forall x)Px$ at line 2. Since QN is an equivalence rule, we could use QN on line 2 of the above proof to derive $\neg(\forall x)Px$.

- 1 $\neg(\forall x)Px$ P
- 2 $(\exists x)\neg Px$ QN 1
- 3 $\neg(\forall x)Px$ QN 2

Since (QN) is a rule of replacement, it can be applied to subformula of wffs. For example, consider the following proof:

- 1 $\neg(\forall x)Px \rightarrow Pa$ P
- 2 $(\exists x)\neg Px \rightarrow Pa$ QN 1

Notice that while line 1 is a conditional, it is permissible to use QN to "replace" or "swap" $\neg(\forall x)Px$ with $(\exists x)\neg Px$. Let's consider another example where QN is used on a subformula. Consider the following wff:

- 1 $\neg(\forall x)\neg(\exists y)Rxy$ P

Notice that the premise, with its negations and quantifiers, is somewhat difficult to read. We can simplify this premise by applying QN and DN until we obtain a more manageable wff.

- 1 $\neg(\forall x)\neg(\exists y)Rxy$ P
- 2 $(\exists x)\neg\neg(\exists y)Rxy$ QN 1
- 3 $(\exists x)(\exists y)Rxy$ DN 2

Notice that the above proof makes use of QN and DN to simplify the premise. The first step in the proof applies QN to the wff $\neg(\forall x)\neg(\exists y)Rxy$ to derive $(\exists x)\neg\neg(\exists y)Rxy$. The second step applies DN to the wff at line 2 to derive the much simpler to read $(\exists x)(\exists y)Rxy$.

Exercise 8.84

Prove the following entailments.

1. $(\exists x)\neg Px, \neg(\forall x)Lxx \vdash \neg(\forall x)Px \wedge (\exists x)\neg Lxx$
2. $(\forall x)Pxx \rightarrow (\exists x)Rx, (\forall x)\neg Rx \vdash (\exists x)\neg Pxx$
3. $(\forall x)\neg(\exists y)Pxy \rightarrow (\forall x)Zx, (\forall x)(\forall y)\neg Pxy \vdash (\forall x)Zx$
4. $\neg(\exists x)(\exists y)Pxy, (\forall x)(\forall y)\neg Pxy \rightarrow (\forall z)Zz \vdash (\forall z)Zz$
5. $\neg(\forall x)Px \vdash (\exists x)(\neg Px \vee Zx)$
6. $\vdash (\exists x)Px \rightarrow (\forall x)(Px \rightarrow Px)$
7. $\vdash (\forall x)(\forall y)(Lxy \rightarrow (Px \rightarrow Lxy))$
8. $(\forall x)(\forall y)Lxy, (\forall x)Lxx \rightarrow (\exists x)(\exists y)Lxy \vdash (\exists x)(\exists y)Lxy$
9. $(\forall x)(Pxx \rightarrow Pxx) \rightarrow (\exists x)Mx \vdash (\exists x)(Mx \wedge Mx)$
10. $\neg(\exists x)Px \not\models (\forall x)\neg Px$ without using *QN*

Part IV

Back Matter: Solutions, Etc.

NOTES

CHAPTER 5 – PL DERIVATIONS

1. According to Copi[1, p. 30] it is cumbersome to establish the validity of an argumetn with more than two propositional leters using truth tables.He contends that a "more convenient method of establishing the validity of some arguments is to *deduce* their conclusions from their premises by a sequence of shorter, more elementary arguments already known to be valid."

BIBLIOGRAPHY

- [1] Irving M. Copi. *Symbolic Logic*. 4th. New York: Macmillan Publishing Co., Inc., 1973.
- [2] David L. Dickinson. “Deliberation Enhances the Confirmation Bias in Politics”. In: *Games* 11.4 (Nov. 2020), p. 57. ISSN: 2073-4336. DOI: [10.3390/g11040057](https://doi.org/10.3390/g11040057). (Visited on 05/28/2021).
- [3] Dorothy Edgington. “On Conditionals”. In: *Mind* 104.414 (1995), pp. 235–329.
- [4] L. T. F. Gamut. *Logic, Language, and Meaning*. Chicago: University of Chicago Press, 1991. ISBN: 978-0-226-28084-4 978-0-226-28085-1 978-0-226-28086-8 978-0-226-28088-2.
- [5] Daniel Kahneman. *Thinking, Fast and Slow*. 1st pbk. ed. New York: Farrar, Straus and Giroux, 2013. ISBN: 978-0-374-53355-7.
- [6] Edwin D. Mares. *Relevant Logic: A Philosophical Interpretation*. 1. publ. Cambridge: Cambridge University Press, 2004. ISBN: 978-0-521-03925-3 978-0-521-82923-6.
- [7] Mark T. Nelson and Philosophy Documentation Center. “Promises and Material Conditionals”. In: *Teaching Philosophy* 16.2 (1993), pp. 155–156. ISSN: 0145-5788. DOI: [10.5840/teachphil199316218](https://doi.org/10.5840/teachphil199316218). (Visited on 03/12/2019).
- [8] William B Swann and Stephen J Read. “Self-Verification Processes: How We Sustain Our Self-Conceptions”. In: *Journal of Experimental Social Psychology* 17.4 (July 1981), pp. 351–372. ISSN: 00221031. DOI: [10.1016/0022-1031\(81\)90043-3](https://doi.org/10.1016/0022-1031(81)90043-3). (Visited on 05/28/2021).
- [9] Ruchika Tulshyan and Jodi-Ann Burey. “Stop Telling Women They Have Imposter Syndrome”. In: *Harvard Business Review* (Feb. 2021).
- [10] Timothy Williamson. *Vagueness*. London and New York: Routledge, 1994.

SOLUTIONS TO EXERCISES

Solution 1.1

1. Be a yardstick of quality. (Not a Proposition)
3. How may I help you? (Not a proposition)
5. In a fixed rate par bond, the issuer issues the bond at par value.
(Proposition)

Solution 1.2

1. God does not exist. (This is debatable. If you think that the sentence can be true or false, then it expresses a proposition. If you think that sentences about God are meaningless, then it does not express a proposition. Given our definition of a proposition, it does not matter whether no one "knows" whether it is true or false.)
3. You are beautiful. (This is debatable. If you think that the sentence can be true or false, then it expresses a proposition. If you think that sentences about beauty are meaningless, then it does not express a proposition.)

Solution 1.3

1. Argument. Notice the use of the argument indicator "Therefore", which signifies the conclusion of the argument.
3. Not explicitly an argument. All three of the sentences are propositions, but it is not clear that the third proposition is the conclusion.

Solution 1.4

1. Invalid. Imagine two people: Tek and Liz. If Tek is friendly but Liz is not, then the premise of the argument is true, but the conclusion is false. This would mean it is possible for the premise to be true and the conclusion false. And so, the argument would be invalid.
3. Valid. There is no possible scenario where the premises are true and the conclusion is false. Note that an argument may be valid even though it has false premises.

Solution 1.6

1. An argument is sound if and only if (1) the argument is valid and (2) all of the premises in the argument are true.
3. No. If an argument is sound argument, then all of the premises in the argument are true. So, a sound argument cannot have at least one premise.

Solution 1.7

1. Valid.
3. Invalid. It is possible for a person to have no good reason for thinking God exists and yet God exists.

Solution 2.9

1. P is a propositional letter
3. $\wedge$ is an operator
5. (is a left parenthesis

Solution 2.10

1. A is a wff by Rule 1. Therefore, $\neg(A)$ is a wff by Rule 2.
3. A and B are wffs by Rule 1. Since B is a wff, $\neg(B)$ is a wff by Rule 2. Since A and $\neg(B)$ are wffs, $(A \wedge \neg(B))$ is a wff by Rule 3.

Solution 2.11

1. P is an atomic wff, a literal wff, but not a complex wff.
3. $(P \wedge Q)$ is a complex wff, but not an atomic wff or a literal wff.

Solution 2.12

1. The proper parts of $(P \rightarrow Q)$ are P and Q . The subformulas of $(P \rightarrow Q)$ are $(P \rightarrow Q)$, P , and Q .
3. The proper parts of $\neg((P \wedge Q))$ are P , Q , and $(P \wedge Q)$. The subformulas of $\neg((P \wedge Q))$ are all of the proper parts and $\neg((P \wedge Q))$.

Solution 2.14

1. The scope of $\wedge$ is $(P \wedge Q)$.

3. The scope of the leftmost $\neg$ is $\neg(\neg(P))$. The scope of the rightmost $\neg$ is $\neg(P)$

Solution 2.15

1. The main operator of $\neg(P)$ is $\neg$.
3. The main operator of $(P \vee (Q \wedge R))$ is $\vee$.
5. The main operator of $(\neg(\neg(M)) \vee R)$ is $\vee$

Solution 2.16

1. $\neg(A)$
3. $\neg(\neg(\neg(A)))$

Solution 2.17

1. A is an atomic wff.
3. $(A \wedge B)$ is a conjunction.
5. $(A \rightarrow B)$ is a conditional.
7. $\neg(\neg(A))$ is a negation of a negation. A double negation.
9. $\neg(A \vee B)$ is a negated disjunction.

Solution 2.18

1. Q
3. $\neg\neg P$

Solution 2.19

1. $v(A) = T$
3. $v(\neg(A)) = F$
5. $v((A \wedge B)) = F$

Solution 2.20

1. $v(\neg(P)) = T$
3. $v(P \wedge Q) = T$
5. $v(P \leftrightarrow Q) = T$

Solution 2.21

Here is a hint on how to create your own operator. First, create a truth function. Consider how many truth values it takes as input and what its truth value is as output given its input. Finally, create a symbol for your operator and define it using your truth function.

Solution 2.23

1. J
3. $J \rightarrow M$
5. $J \vee M$

Solution 2.24

1. $(R \wedge T) \wedge L$
3. $\neg R \wedge \neg T$ or $\neg(R \vee T)$
5. $(T \vee L) \wedge \neg(T \wedge L)$

Solution 2.25

1. $(P \wedge W) \wedge \neg L$
3. $\neg W \wedge \neg B$ or $\neg(W \vee B)$

Solution 2.26

1. $(H \wedge T) \wedge S$
3. $\neg H \wedge \neg T$

Solution 3.28

1. $\neg P \rightarrow R$ is F
3. $\neg(\neg P \leftrightarrow \neg S)$ is F
4. $P \wedge \neg R$ is T.
10. Hint: Consider both rows where the antecedent of the conditional is false.

Solution 3.29

1. Truth table for $P \rightarrow \neg R$

P	R	P	$\rightarrow$	$\neg$	R
T	T	T	F	F	T
T	F	T	T	T	F
F	T	F	T	F	T
F	F	F	T	T	F

3. Truth table for $\neg\neg(P \leftrightarrow R) \vee Z$

P	R	Z	$\neg$	$\neg$	$(P \leftrightarrow R)$	$\vee$	Z
T	T	T	T	F	T	T	T
T	T	F	T	F	T	T	F
T	F	T	F	T	F	F	T
T	F	F	F	T	F	F	F
F	T	T	F	T	F	T	T
F	T	F	F	T	F	T	F
F	F	T	T	F	T	F	T
F	F	F	T	F	T	F	F

Solution 3.30

1. $\neg P \rightarrow \neg P$ is a contingency. There is at least one T and at least one F under the main operator.

P	$\neg$	P	$\rightarrow$	$\neg$	P
T	F	T	T	F	T
F	T	F	T	T	F

3. $P \leftrightarrow \neg R$ is a contingency. There is at least one T and at least one F under the main operator.

P	R	P	$\leftrightarrow$	$\neg$	R
T	T	T	F	F	T
T	F	T	T	T	F
F	T	F	T	F	T
F	F	F	F	T	F

Solution 3.31

1. P, Q is consistent. There is at least one row where P, Q are both true. We can show this by simply setting up the truth table.

P	Q
T	T
T	F
F	T
F	F

3. $P \vee Q, P \wedge Q$ is consistent. There is at least one row where both wffs are true.

P	Q	$P \vee Q$	$P \wedge Q$
T	T	T	T
T	F	T	F
F	T	T	F
F	F	F	F

Solution 3.32

1. If Γ consists of a single wff ϕ , then Γ is consistent if and only if ϕ is a PL-contingency or PL-tautology. This is because if ϕ is a tautology, ϕ is true under every interpretation and so it is true under at least one interpretation. If ϕ is a PL-contingency, then ϕ is true under at least one interpretation. However, if ϕ is a PL-contradiction, then ϕ is false under every interpretation and so it is not true under at least one interpretation.
3. Let's first show that if $\phi \wedge \psi$ is a contradiction, then $\{\phi, \psi\}$ is inconsistent. If $\phi \wedge \psi$ is a contradiction, then $\phi \wedge \psi$ is false under every interpretation. This means that there is no interpretation where both ϕ and ψ are true. If there is no interpretation where ϕ and ψ are true, then $\{\phi, \psi\}$ is inconsistent. Second, let's show that if $\{\phi, \psi\}$ is inconsistent, then $\phi \wedge \psi$ is a contradiction. If $\{\phi, \psi\}$ is inconsistent, then there is no interpretation where both ϕ and ψ are true. If there is no interpretation where both ϕ and ψ are true, then $\phi \wedge \psi$ is false under every interpretation. Therefore it is a contradiction.

Solution 3.33

1. P, Q are not equivalent. We can see this simply by setting up the table. Notice that there is at least one row where P and Q do not have the same truth value.

P	Q
T	T
T	F
F	T
F	F

3. $P, P \vee P$ are equivalent. Alternatively, we can say that the set containing these wffs is equivalent: $\{P, P \vee P\}$.

P	P	P	$\vee$	P
T	T	T	T	T
F	F	F	F	F

Solution 3.34

1. If Γ consists of a single wff ϕ , then Γ is equivalent if and only if for each interpretation $\mathcal{I}$, ϕ has the same truth value as ϕ . Since for each interpretation of ϕ , ϕ has the same truth value as ϕ , it follows that a set Γ consisting of a single wff ϕ is equivalent.
3. If ϕ and ψ are equivalent, then there is no interpretation where ϕ and ψ have different truth values. A wff $\phi \leftrightarrow \psi$ is false just in the case where ϕ and ψ have different truth values. Since there is no interpretation where ϕ and ψ have different truth values, it follows that $\phi \leftrightarrow \psi$ is a tautology.

Solution 3.35

1. $P \models P$ since there is no row where P is T and P is F.

P	P
T	T
F	T

3. $P \wedge Q \models P$ since there is no row where $P \wedge Q$ is T and P is F.

P	Q	P	$\wedge$	Q	$\models$	P
T	T	T	T	T		T
T	F	T	F	F		T
F	T	F	F	T		F
F	F	F	F	F		F

Solution 3.36

1. $J \vee M, \neg M \models \neg J$ is not the case since there is a row where $J \vee M$ and $\neg M$ are T and $\neg J$ is F: Row 2.

J	M	J	$\vee$	M	$\neg$	M	$\not\models$	$\neg$	J
T	T	T	T	T	F	T		F	T
T	F	T	T	F	T	F		F	T
F	T	F	T	T	F	T		T	F
F	F	F	F	F	T	F		T	F

Solution 3.37

1. $P \wedge Q \models P$ but notice that $P \models P \wedge Q$ is not the case.
3. If Γ is inconsistent, then there is no single interpretation where all of its members are true. If that is the case, then for any wff ϕ , there is no interpretation where all of the members of Γ are true and ϕ is false. Therefore, $\Gamma \models \phi$.

Solution 4.41

1. $A, A \wedge \neg B$
 1. A P
 2. $A \wedge \neg B \checkmark$ P
 3. A $1 \wedge D$
 4. $\neg B$ $1 \wedge D$
3. $\neg A \wedge B, P \vee \neg Q$
 1. $\neg A \wedge B \checkmark$ P
 2. $P \vee \neg Q \checkmark$ P
 3. $\neg A$ $1 \wedge D$
 4. B $1 \wedge D$

$\swarrow \quad \searrow$
 $P \quad \neg Q$
5. $P \quad \neg Q$ $2 \vee D$
5. $\neg(A \rightarrow B), \neg(A \leftrightarrow B)$
 1. $\neg(A \rightarrow B) \checkmark$ P
 2. $\neg(A \leftrightarrow B) \checkmark$ P
 3. A $1 \neg \rightarrow D$
 4. $\neg B$ $1 \neg \rightarrow D$

$\swarrow \quad \searrow$
 $A \quad \neg A$
5. $A \quad \neg A$ $2 \neg \rightarrow D$
6. $\neg B \quad B$ $2 \neg \rightarrow D$

Solution 4.42

1. $A \wedge (\neg A \wedge B)$
 1. $A \wedge (\neg A \wedge B) \checkmark$ P
 2. A $1 \wedge D$
 3. $\neg A \wedge B$ $1 \wedge D$
 4. $\neg A$ $3 \wedge D$
 5. B $3 \wedge D$
- $\otimes$
- 4,2

3.	$A \wedge B, \neg(A \rightarrow B)$	
1.	$A \wedge B \checkmark$	P
2.	$\neg(A \rightarrow B) \checkmark$	P
3.	A	$1 \wedge D$
4.	B	$1 \wedge D$
5.	A	$2 \neg \rightarrow D$
6.	$\neg B$	$2 \neg \rightarrow D$
	$\otimes$	
	4,6	

Solution 4.43

1.	$P \vee Q, A \wedge \neg A, B \rightarrow C$	
1.	$P \vee Q \checkmark$	P
2.	$A \wedge \neg A \checkmark$	P
3.	$B \rightarrow C \checkmark$	P
4.	A	$2 \wedge D$
5.	$\neg A$	$2 \wedge D$
	$\otimes$	
	4,5	

3.	$A \wedge \neg B, B, C \rightarrow Q$	
1.	$A \wedge \neg B \checkmark$	P
2.	B	P
3.	$C \rightarrow Q \checkmark$	P
4.	A	$1 \wedge D$
5.	$\neg B$	$1 \wedge D$
	$\otimes$	
	2,5	

Solution 4.44

1.	$P \wedge \neg P$	
1.	$P \wedge \neg P \checkmark$	P
2.	P	$1 \wedge D$
3.	$\neg P$	$1 \wedge D$
	$\otimes$	
	2,3	

3. $P \wedge Q, Q \wedge R$. Since there is a completed open branch, an interpretation that would make each wff in the branch true is $\mathcal{I}(P) = T, \mathcal{I}(Q) = T, \mathcal{I}(R) = T$.

1. $P \wedge Q \checkmark$ P
2. $Q \wedge R \checkmark$ P
3. P $1 \wedge D$
4. Q $1 \wedge D$
5. Q $2 \wedge D$
6. R $2 \wedge D$

Solution 4.45

1. $P, \neg P$. The set is inconsistent.

1. P P
 2. $\neg P$ P
- $\otimes$
 1,2

3. $P \rightarrow Q, P, \neg Q$. The set is inconsistent.

1. $P \rightarrow Q \checkmark$ P
 2. P P
 3. $\neg Q$ P
- $\swarrow \quad \searrow$
 4. $\neg P \quad Q$ $1 \rightarrow D$
 $\otimes \quad \otimes$
 2,4 3,5

5. $\neg\neg(P \vee Q), P \vee Q$. The set is consistent.

1. $\neg\neg(P \vee Q) \checkmark$ P
 2. $P \vee Q \checkmark$ P
 3. $P \vee Q$ $1 \neg\neg D$
- $\swarrow \quad \searrow$
 4. $P \quad Q$ $2 \vee D$
 $\swarrow \quad \searrow \quad \swarrow \quad \searrow$
 5. $P \quad Q \quad P \quad Q$ $3 \vee D$

Solution 4.46

1. $P \wedge \neg P$ is a contradiction.

1. $P \wedge \neg P \checkmark$ P
 2. P $1 \wedge D$
 3. $\neg P$ $1 \wedge D$
- $\otimes$
 2,3

3. $P \rightarrow P$ is a tautology. To test this using the truth-tree method, we

test whether $\neg(P \rightarrow P)$ yields a closed or open tree. Since it yields a closed tree, there is no interpretation where $\neg(P \rightarrow P)$ is true. That is, $\neg(P \rightarrow P)$ is false under every interpretation. Since $\neg(P \rightarrow P)$ is false under every interpretation, $(P \rightarrow P)$ is true under every interpretation.

1. $\neg(P \rightarrow P) \checkmark$ P
2. P $1 \neg \rightarrow D$
3. P $1 \neg \rightarrow D$
- $\otimes$
- 2,3

5. $P \wedge (Q \rightarrow P)$ is a contingency. The first tree shows it is not a contradiction. The second shows it is not a tautology.

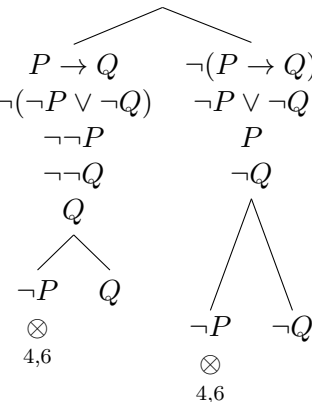
1. $P \wedge (Q \rightarrow P) \checkmark$ P
2. P $1 \wedge D$
3. $Q \rightarrow P \checkmark$ $1 \wedge D$
- $\swarrow \quad \searrow$
 4. $\neg Q \quad P$ $3 \rightarrow D$
1. $\neg(P \wedge (Q \rightarrow P)) \checkmark$ P
- $\swarrow \quad \searrow$
 2. $\neg P \quad \neg(Q \rightarrow P) \checkmark$ $1 \neg \wedge D$
3. Q $2 \rightarrow D$
4. $\neg P$ $2 \rightarrow D$

Solution 4.47

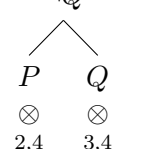
1. $P \rightarrow Q, \neg P \wedge Q$ are not equivalent.

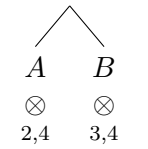
1. $\neg((P \rightarrow Q) \leftrightarrow (\neg P \wedge Q)) \checkmark$ P
- $\swarrow \quad \searrow$
 2. $P \rightarrow Q \quad \neg(P \rightarrow Q)$ $1 \neg \leftrightarrow D$
3. $\neg(\neg P \wedge Q) \quad \neg P \wedge Q$ $1 \neg \leftrightarrow D$
4. $\neg \neg P \quad P$ $3 \neg \wedge D; 2 \neg \wedge D$
5. $Q \quad Q$ $3 \neg \wedge D; 2 \neg \wedge D$
- $\swarrow \quad \searrow \quad \swarrow \quad \searrow$
 6. $P \quad Q \quad P \quad Q$ $2 \rightarrow D$
 $\otimes \quad \otimes \quad \otimes \quad \otimes$
 4,6 5,6 4,6 5,6

3. $P \rightarrow Q, \neg P \vee \neg Q$ are not equivalent.

1.	$\neg((P \rightarrow Q) \leftrightarrow (\neg P \vee \neg Q)) \checkmark$	P
		
2.	$P \rightarrow Q$	$1 \neg \leftrightarrow D$
3.	$\neg(\neg P \vee \neg Q)$	$1 \neg \leftrightarrow D$
4.	$\neg\neg P$	$3 \neg \vee D; 2 \neg \rightarrow D$
5.	$\neg\neg Q$	$3 \neg \vee D; 2 \neg \rightarrow D$
6.	Q	$3 \neg\neg D$
7.	$\neg P$	$2 \rightarrow D$
8.	$\otimes$ 4,6	$3 \vee D$

Solution 4.48

1.	$P \wedge Q \models Q$	
1.	$P \wedge Q \checkmark$	P
2.	P	P
3.	$\neg Q$	P
		
4.	P	$1 \wedge D$
	$\otimes$ 2,4	
	$\otimes$ 3,4	

3.	$A \rightarrow B, B \models A$	
1.	$A \rightarrow B \checkmark$	P
2.	B	P
3.	$\neg A$	P
		
4.	A	$1 \rightarrow D$
	$\otimes$ 2,4	
	$\otimes$ 3,4	

5. $(A \wedge B) \rightarrow C, A \not\models C$

1.	$(A \wedge B) \rightarrow C$	$\checkmark$	P
2.	A		P
3.	$\neg C$		P
	$\swarrow$	$\searrow$	
4.	$\neg(A \wedge B)$	C	$1 \rightarrow D$
	$\swarrow$	$\searrow$	$\otimes$
5.	$\neg A$	$\neg B$	$3,4$
	$\otimes$		
	$2,5$		$4 \neg \wedge D$

Solution 5.49

1. $\neg P, \neg P \rightarrow Z \vdash Z$
 - 1 $\neg P$ P
 - 2 $\neg P \rightarrow Z$ P, Z
3. $P, Q, R \vdash (P \wedge Q) \wedge R$
 - 1 P P
 - 2 Q P
 - 3 R P, $(P \wedge Q) \wedge R$

5. $\neg\neg P \vdash P$
 - 1 $\neg\neg P$ P, P

Solution 5.50

1. $P, Q, R, S \vdash P \wedge S$
 - 1 P P
 - 2 Q P
 - 3 R P
 - 4 S P
 - 5 $P \wedge S$ $\wedge I$ 1, 4
3. $P \wedge (R \wedge M) \vdash R$
 - 1 $P \wedge (R \wedge M)$ P
 - 2 $R \wedge M$ $\wedge E$ 1
 - 3 R $\wedge E$ 2

5. $P \wedge R, \neg Z \wedge \neg W \vdash P \wedge \neg W$

- 1 $P \wedge R$ P
- 2 $\neg Z \wedge \neg W$ P
- 3 P $\wedge E$ 1
- 4 $\neg W$ $\wedge E$ 2
- 5 $P \wedge \neg W$ $\wedge I$ 3, 4

Solution 5.51

1. $P \rightarrow Q, P \vdash Q$
 - 1 $P \rightarrow Q$ P
 - 2 P P
 - 3 Q $\rightarrow E$ 1, 2
3. $(A \wedge B) \rightarrow C, A, B \vdash C$
 - 1 $(A \wedge B) \rightarrow C$ P
 - 2 A P
 - 3 B P
 - 4 $A \wedge B$ $\wedge I$ 2, 3
 - 5 C $\rightarrow E$ 1, 4
5. $R \rightarrow Z, Z \rightarrow W, R \wedge M \vdash W \wedge M$
 - 1 $R \rightarrow Z$ P
 - 2 $Z \rightarrow W$ P
 - 3 $R \wedge M$ P
 - 4 R $\wedge E$ 3
 - 5 Z $\rightarrow E$ 1, 4
 - 6 W $\rightarrow E$ 2, 5
 - 7 M $\wedge E$ 3
 - 8 $W \wedge M$ $\wedge I$ 6, 7

Solution 5.52

1. $P \wedge R \vdash Z \rightarrow P$
 - 1 $P \wedge R$ P
 - 2 $\begin{array}{|l} Z \\ \hline P \end{array}$ A
 - 3 $\begin{array}{|l} Z \\ \hline P \end{array}$ $\wedge E$ 1
 - 4 $Z \rightarrow P$ $\rightarrow I$ 2-3

3. $R \vdash R \rightarrow R$

1	R	P
2	R	A
3	$R \wedge R$	$\wedge$ 1,2
4	R	$\wedge E$ 3
5	$R \rightarrow R$	$\rightarrow I$ 2-4

5. $P, Q \vdash P \rightarrow Q$

1	P	P
2	Q	Q
3	P	A
4	$P \wedge Q$	$\wedge I$ 1,2
5	Q	$\wedge E$ 3
6	$P \rightarrow Q$	$\rightarrow I$ 3-5

Solution 5.53

1. $P \rightarrow Z, P \vdash P \wedge Z$

1	$P \rightarrow Z$	P
2	P	P
3	Z	$\rightarrow E$ 1,2
4	$P \wedge Z$	$\wedge I$ 2,3

3. $P \vdash R \rightarrow P$

1	P	P
2	R	A
3	P	R 1
4	$R \rightarrow P$	$\rightarrow I$ 1-3

5. $(R \vee F) \rightarrow Z, M \wedge (R \vee F) \vdash M \wedge Z$

1	$(R \vee F) \rightarrow Z$	P
2	$M \wedge (R \vee F)$	P
3	M	$\wedge E$ 2
4	$R \vee F$	$\wedge E$ 2
5	Z	$\rightarrow E$ 1,4
6	$M \wedge Z$	$\wedge I$ 3,5

Solution 5.54

1. $P \wedge \neg P \vdash R$

1	$P \wedge \neg P$	P
2	R	A
3	P	$\wedge E$ 1
4	$\neg P$	$\wedge E$ 1
5	$\neg R$	$\neg I$ 2-4
3. $L, \neg L \vdash \neg\neg\neg M$

1	L	P
2	$\neg L$	P
3	$\neg\neg M$	A
4	L	R 1
5	$\neg L$	R 2
6	$\neg\neg\neg M$	$\neg I$ 3-5

Solution 5.55

1. $A \vdash A \vee \neg B$

1	A	P
2	$A \vee \neg B$	$\vee I$ 1
3. $A, B \vdash A \vee B$

1	A	P
2	B	P
3	$A \vee B$	$\vee I$ 1
5. $\neg A \vee B, \neg A \rightarrow S, B \rightarrow S \vdash S$

1	$\neg A \vee B$	P
2	$\neg A \rightarrow S$	P
3	$B \rightarrow S$	P
4	$\neg A$	A
5	S	$\rightarrow E$ 2, 4
6	B	A
7	S	$\rightarrow E$ 3, 6
8	S	$\vee E$ 1, 4-5, 6-7
7. $P \vee Q, \neg Q \vdash P$

- Hint 1: Start by assuming $\neg P$. In making this assumption, the goal

is to derive P using $\neg E$.

- Hint 2: After you have assumed $\neg P$, consider using $\vee E$. This will require you to make two subproofs, one for P and one for Q .

Solution 5.56

1. $(P \wedge Q) \leftrightarrow R, P, Q \vdash R$
 - 1 $(P \wedge Q) \leftrightarrow R$ P
 - 2 P P
 - 3 Q P
 - 4 $P \wedge Q$ $\wedge I$ 2,3
 - 5 R $\leftrightarrow E$ 1, 4

3. $(P \vee \neg M) \leftrightarrow R, P \leftrightarrow (W \vee L), L \vdash R$
 - 1 $(P \vee \neg M) \leftrightarrow R$ P
 - 2 $P \leftrightarrow (W \vee L)$ P
 - 3 L P
 - 4 $W \vee L$ $\vee I$ 3
 - 5 P $\leftrightarrow E$ 2, 4
 - 6 $P \vee \neg M$ $\vee I$ 5
 - 7 R $\leftrightarrow E$ 1,6

Solution 5.57

1. $Z \wedge (B \wedge F), (M \wedge T) \wedge (L \rightarrow P), Q \wedge (R \wedge P) \vdash \neg R \vee (S \vee T)$
 - 1 $Z \wedge (B \wedge F)$ P
 - 2 $(M \wedge T) \wedge (L \rightarrow P)$ P
 - 3 $Q \wedge (R \wedge P)$ P
 - 4 $M \wedge T$ $\wedge E$ 2
 - 5 T $\wedge E$ 4
 - 6 $S \vee T$ $\vee I$ 5
 - 7 $\neg R \vee (S \vee T)$ $\vee I$ 6

	1	$(Z \wedge Q) \wedge (F \wedge L)$	P
	2	$R \wedge P$	P
	3	$W \wedge B$	P
3.		$(Z \wedge Q) \wedge (F \wedge L), R \wedge P, W \wedge B \vdash (Z \vee T) \vee (M \rightarrow R)$	
	4	$Z \wedge Q$	$\wedge E$ 1
	5	Z	$\wedge E$ 4
	6	$Z \vee T$	$\vee I$ 5
	7	$(Z \vee T) \vee (M \rightarrow R)$	$\vee I$ 6

	1	$M \wedge (R \wedge \neg Z)$	P
	2	$S \wedge (P \wedge W)$	P
	3	Q	P
	4	$\begin{array}{ l} S \end{array}$	A
5.		$M \wedge (R \wedge \neg Z), S \wedge (P \wedge W), Q \vdash (S \leftrightarrow Q) \vee [M \wedge (R \wedge Z)]$	
	5	$\begin{array}{ l} Q \end{array}$	R 3
	6	$\begin{array}{ l} Q \end{array}$	A
	7	$\begin{array}{ l} S \end{array}$	$\wedge E$ 2
	8	$S \leftrightarrow Q$	$\leftrightarrow I$ 7
	9	$(S \leftrightarrow Q) \vee [M \wedge (R \wedge Z)]$	$\vee I$ 8

Solution 5.58

1.	$P \rightarrow Q, \neg Q \vdash \neg P$	
1	$P \rightarrow Q$	P
2	$\neg Q$	P
3	$\neg P$	MT 1,2

3.	$\neg(P \vee R) \vdash \neg P \wedge \neg R$	
1	$\neg(P \vee R)$	P
2	$\neg P \wedge \neg R$	DeM 1

5.	$(\neg P \wedge L) \rightarrow \neg Q, (M \wedge T) \wedge (\neg R \wedge L), (M \wedge \neg R) \rightarrow (Z \wedge \neg P) \vdash \neg Q \vee (A \leftrightarrow B)$	
----	---	--

1	$(\neg P \wedge L) \rightarrow \neg Q$	P
2	$(M \wedge T) \wedge (\neg R \wedge L)$	P
3	$(M \wedge \neg R) \rightarrow (Z \wedge \neg P)$	P
4	$M \wedge T$	$\wedge E$ 2
5	$\neg R \wedge L$	$\wedge E$ 2
6	M	$\wedge E$ 4
7	$\neg R$	$\wedge E$ 5
8	$M \wedge \neg R$	$\wedge I$ 6, 7
9	$Z \wedge \neg P$	$\rightarrow E$ 3, 8
10	$\neg P$	$\wedge E$ 9
11	L	$\wedge E$ 5
12	$\neg P \wedge L$	$\wedge I$ 11, 10
13	$\neg Q$	$\rightarrow E$ 1, 12
14	$\neg Q \vee (A \leftrightarrow B)$	$\vee I$ 13

Solution 5.59

1. $(R \wedge T) \vee \neg W, S \wedge \neg \neg W \vdash R \wedge T$

1	$(R \wedge T) \vee \neg W$	P
2	$S \wedge \neg \neg W$	P
3	$\neg \neg W$	$\wedge E$ 2
4	$R \wedge T$	DS 1, 3

3. $(R \wedge T) \rightarrow \neg W, M \rightarrow (R \wedge T), \neg W \rightarrow (S \wedge R) \vdash M \rightarrow (S \wedge R)$

1	$(R \wedge T) \rightarrow \neg W$	P
2	$M \rightarrow (R \wedge T)$	P
3	$\neg W \rightarrow (S \wedge R)$	P
4	$\begin{array}{ l} M \end{array}$	A
5	$\begin{array}{ l} R \wedge T \end{array}$	$\rightarrow E$ 2, 4
6	$\begin{array}{ l} \neg W \end{array}$	$\rightarrow E$ 1, 4
7	$\begin{array}{ l} S \wedge R \end{array}$	$\rightarrow E$ 3, 5
8	$M \rightarrow (S \wedge R)$	$\rightarrow I$ 1-6

5. Hint: Consider using $\rightarrow I$.

Solution 5.60

1. $\neg\neg(\neg\neg P \rightarrow Q) \vdash P \rightarrow \neg\neg Q$
 - 1 $\neg\neg(\neg\neg P \rightarrow Q)$ P
 - 2 $\neg\neg(P \rightarrow Q)$ DN 1
 - 3 $(P \rightarrow Q)$ DN 2
 - 4 $P \rightarrow \neg\neg Q$ DN 3

3. $S \rightarrow \neg Q \vdash \neg S \vee \neg Q$
 - 1 $S \rightarrow \neg Q$ P
 - 2 $\neg\neg S \vee \neg Q$ IMP 1

5. $\neg(P \vee R) \rightarrow (\neg Z \vee \neg W), \neg P \wedge \neg R \vdash \neg(Z \wedge W)$
 - 1 $\neg(P \vee R) \rightarrow (\neg Z \vee \neg W)$ P
 - 2 $\neg P \wedge \neg R$ P
 - 3 $\neg(P \vee R)$ DeM 2
 - 4 $\neg Z \vee \neg W$ $\rightarrow E$ 1, 3
 - 5 $\neg(Z \wedge W)$ DeM 4

Solution 5.61

1. $P \wedge \neg Q, T \vee Q \vdash T$
 - 1 $P \wedge \neg Q$ P
 - 2 $T \vee Q$ P
 - 3 $\neg Q$ $\wedge E$ 1
 - 4 T DS 3, 4

3. $A \rightarrow C, A \wedge D \vdash C$
 - 1 $A \rightarrow C$ P
 - 2 $A \wedge D$ P
 - 3 A $\wedge E$ 2
 - 4 C $\rightarrow E$ 1, 3

5. $A \vdash B \rightarrow (B \wedge A)$
 - 1 A P
 - 2 $\begin{array}{|l} B \\ \hline \end{array}$ A
 - 3 $\begin{array}{|l} B \wedge A \\ \hline \end{array}$ $\wedge I$ 1, 2
 - 4 $B \rightarrow (B \wedge A)$ $\rightarrow I$ 2-3

Solution 5.62

1. $(S \leftrightarrow D) \rightarrow T, P \leftrightarrow (S \wedge D), P \vdash T$

1	$(S \leftrightarrow D) \rightarrow T$	P
2	$P \leftrightarrow (S \wedge D)$	P
3	P	P
4	$S \wedge D$	$\leftrightarrow E$ 3, 2
5	$\begin{array}{ l} S \end{array}$	A
6	$\begin{array}{ l} D \end{array}$	$\wedge E$ 4
7	$\begin{array}{ l} D \end{array}$	A
8	$\begin{array}{ l} S \end{array}$	$\wedge E$ 4
9	$S \leftrightarrow D$	$\leftrightarrow I$ 5-8
10	T	$\rightarrow E$ 1, 9

3. $B \rightarrow \neg(S \vee T), \neg(A \vee \neg B), \neg S \rightarrow W \vdash W$

1	$B \rightarrow \neg(S \vee T)$	P
2	$\neg(A \vee \neg B)$	P
3	$\neg S \rightarrow W$	P
4	$\neg A \wedge \neg \neg B$	<i>DeM</i> 1
5	$\neg \neg B$	$\wedge E$ 4
6	B	<i>DN</i> 5
7	$\neg(S \vee T)$	$\rightarrow E$ 1, 6
8	$\neg S \wedge \neg T$	<i>DeM</i> 7
9	$\neg S$	$\wedge E$ 8
10	W	$\rightarrow E$ 3, 9

5. $\neg(A \wedge B), B, (\neg A \vee S) \rightarrow \neg(D \wedge T) \vdash \neg D \vee \neg T$

1	$\neg(A \wedge B)$	P
2	B	P
3	$(\neg A \vee S) \rightarrow \neg(D \wedge T)$	P
4	$\neg A \vee \neg B$	<i>DeM</i> 1
5	$\neg \neg B$	<i>DN</i> 2
6	$\neg A$	<i>DS</i> 4, 5
7	$\neg A \vee S$	$\vee I$ 6
8	$\neg(D \wedge T)$	$\rightarrow E$ 3, 7
9	$\neg D \vee \neg T$	<i>DeM</i> 8

Solution 5.63

1.	$\neg(A \vee [\neg(B \rightarrow R) \vee \neg(C \rightarrow R)])$	$\neg A \leftrightarrow (B \vee C) \vdash R$
1	$\neg(A \vee [\neg(B \rightarrow R) \vee \neg(C \rightarrow R)])$	P
2	$\neg A \leftrightarrow (B \vee C)$	P
3	$\neg A \wedge \neg[\neg(B \rightarrow R) \vee \neg(C \rightarrow R)]$	<i>DeM</i> 1
4	$\neg A$	$\wedge E$ 3
5	$\neg[\neg(B \rightarrow R) \vee \neg(C \rightarrow R)]$	$\wedge E$ 3
6	$\neg\neg(B \rightarrow R) \wedge \neg\neg(C \rightarrow R)$	<i>DeM</i> 5
7	$\neg\neg(B \rightarrow R)$	$\wedge E$ 6
8	$\neg\neg(C \rightarrow R)$	$\wedge E$ 6
9	$B \rightarrow R$	<i>DN</i> 7
10	$C \rightarrow R$	<i>DN</i> 8
11	$B \vee C$	$\leftrightarrow E$ 2, 4
12	B	A
13	R	$\rightarrow E$ 9, 12
14	C	A
15	R	$\rightarrow E$ 10, 12
16	R	$\vee E$ 11, 12-15

3. $A \vee B, R \vee \neg(S \vee M), A \rightarrow S, B \rightarrow M \vdash R$

1	$A \vee B$	P
2	$R \vee \neg(S \vee M)$	P
3	$A \rightarrow S$	P
4	$B \rightarrow M$	P
5	$\neg R$	A
6	$\neg(S \vee M)$	<i>DS</i> 2,5
7	$\neg S \wedge \neg M$	<i>DeM</i> 6
8	$\neg S$	$\wedge E$ 7
9	$\neg M$	$\wedge E$ 7
10	$\neg B$	<i>MT</i> 4, 9
11	A	<i>DS</i> 1, 10
12	$\neg A$	<i>MT</i> 3, 8
13	R	$\neg E$ 5-12

5. $A \vdash \neg(\neg A \wedge \neg B)$
- | | | |
|---|------------------------------|-------------|
| 1 | A | P |
| 2 | $\neg\neg A$ | $DN\ 1$ |
| 3 | $\neg\neg A \vee \neg\neg B$ | $\vee I\ 2$ |
| 4 | $\neg(\neg A \wedge \neg B)$ | $DeM\ 3$ |

Solution 5.64

1. $\vdash P \rightarrow P$
- | | | |
|---|-------------------|----------------------|
| 1 | P | A |
| 2 | P | $R\ 1$ |
| 3 | $P \rightarrow P$ | $\rightarrow I\ 1-2$ |

3. $\vdash \neg(P \wedge \neg P)$
- | | | |
|---|-------------------------|---------------|
| 1 | $P \wedge \neg P$ | A |
| 2 | P | $\wedge E\ 1$ |
| 3 | $\neg P$ | $\wedge E\ 1$ |
| 4 | $\neg(P \wedge \neg P)$ | $\neg I\ 1-3$ |

5. $\vdash A \rightarrow \neg(B \wedge \neg B)$
- | | | |
|---|---|----------------------|
| 1 | A | A |
| 2 | <div style="border-left: 1px solid black; padding-left: 10px;">$(B \wedge \neg B)$</div> | A |
| 3 | <div style="border-left: 1px solid black; padding-left: 10px;">B</div> | $\wedge E\ 2$ |
| 4 | <div style="border-left: 1px solid black; padding-left: 10px;">$\neg B$</div> | $\wedge E\ 2$ |
| 5 | $\neg(B \wedge \neg B)$ | $\neg I\ 2-4$ |
| 6 | $A \rightarrow \neg(B \wedge \neg B)$ | $\rightarrow I\ 1-5$ |

7. $\vdash A \rightarrow (A \vee \neg A)$
- | | | |
|---|--|----------------------|
| 1 | A | A |
| 2 | <div style="border-left: 1px solid black; padding-left: 10px;">$\neg(A \vee \neg A)$</div> | A |
| 3 | <div style="border-left: 1px solid black; padding-left: 10px;">$\neg A \wedge \neg\neg A$</div> | $DeM\ 2$ |
| 4 | <div style="border-left: 1px solid black; padding-left: 10px;">$\neg A$</div> | $\wedge E\ 3$ |
| 5 | <div style="border-left: 1px solid black; padding-left: 10px;">$\neg\neg A$</div> | $\wedge E\ 3$ |
| 6 | $A \vee \neg A$ | $\neg I\ 2-5$ |
| 7 | $A \rightarrow (A \vee \neg A)$ | $\rightarrow I\ 1-6$ |

Solution 6.65

1. a is a name.
3. g is a name.
5. x is a variable.

Solution 6.66

1. Ra is a wff (rule 1). Paa is a wff (rule 1). $(Ra \wedge Paa)$ is a wff (rule 3).
3. Zx is a wff (rule 1). $(\exists x)(Zx)$ is a wff (rule 4).
5. Pxx is a wff (rule 1). $(\forall x)(Pxx)$ is a wff (rule 4). $\neg((\forall x)(Pxx))$ (rule 4). Zx is a wff (rule 1). $(\exists x)(Zx)$ is a wff (rule 4). $(\neg((\forall x)(Pxx)) \wedge (\exists x)(Zx))$ is a wff (rule 3).

Solution 6.67

1. Fb is a literal wff.
3. $\neg(\neg(Fb))$ is not a literal wff.
5. $Fb \wedge Fq$ is not a literal wff.

Solution 6.68

1. The scope of $(\forall x)Px$ is $(\forall x)Px$
3. The scope of $(\exists x)\neg Rx$ is $(\exists x)\neg Rx$

Solution 6.69

1. $(\forall x)$
3. $\wedge$
5. $\neg$

Solution 6.70

1. x is bound in Rx
3. x is bound in Mx and Rx
5. w is bound in Vx but x is free in Lx

Solution 6.71

1. Open. x is free.

3. Closed
5. Closed

Solution 6.72

1. Let the domain refer to the set that includes both Jon and Liz. Let "Jon" refer to Jon and "Liz" refer to Liz. Let "loves" refer to the set of ordered pairs $\langle Jon, Liz \rangle$ and $\langle Liz, Jon \rangle$.

Solution 6.73

1. $v(Ha) = T$
3. $v((\forall x)Hx) = T$
5. $v((Ha \wedge Hb) \wedge Hc) = T$

Solution 6.74

1. $(\forall x)(Ex \rightarrow Nx)$ is T
3. $(\forall x)(Ex \rightarrow \neg Ox)$ is T.
5. $(\exists x)(Ex \vee Ox)$ is T.
7. $(\forall x)Gxx$ is F.

Solution 6.76

1. Everyone is friendly.
3. All ghosts are friendly.
5. It is not the case that all ghosts are friendly.
7. Everyone is either friendly or a ghost.

Solution 6.77

1. Everyone is poor.
3. Everyone is poor and everyone is lazy.
5. No poor people are lazy.
7. Someone is poor.

Solution 6.78

1. Everyone loves Bob.
3. No one who loves Bob hates Sally.
5. If Bob loves Sally and Bob hates someone, then Bob loves Sally.

Solution 6.79

1. Involving presence. "Criminals are bad" is ambiguous between "all criminals are bad" and "some criminals are bad". The former can be translated as $(\forall x)(Cx \rightarrow Bx)$ while the latter can be translated as $(\exists x)(Cx \wedge Bx)$.
3. Involving presence and involving scope. "Smokers are not bad" is ambiguous between "all smokers are not bad" and "some smokers are not bad". The former is additionally ambiguous between "it is not the case that all smokers are bad" and "no smokers are bad". The first sentence can be translated as $\neg(\forall x)(Sx \rightarrow Bx)$ while the second can be translated as $(\forall x)(Sx \rightarrow \neg Bx)$. On the other hand, "some smokers are not bad" can be translated as $(\exists x)(Sx \wedge \neg Bx)$.

Solution 8.80

1. $(\forall x)Px, (\forall z)Qz \vdash Pa \wedge Qa$

1	$(\forall x)Px$	P	
2	$(\forall z)Qz$	P	
3	Pa	$\forall E$ 1	3. $Lab \rightarrow Sa, (\forall x)(\forall y)Lxy \vdash Sa$
4	Qa	$\forall E$ 2	
5	$Pa \wedge Qa$	$\wedge I$ 3, 4	

1	$Lab \rightarrow Sa$	P	
2	$(\forall x)(\forall y)Lxy$	P	
3	$(\forall y)Lay$	$\forall E$ 2	
4	Lab	$\forall E$ 3	
5	Sa	$\rightarrow E$ 1, 4	

Solution 8.81

1. $Pa, Rb \vdash (\exists x)(Px \wedge Rb)$

1	Pa	P	
2	Rb	P	
3	$Pa \wedge Rb$	$\wedge I$ 1, 2	3. $(\exists x)Rx \rightarrow Mc, Ra \vdash (\exists x)Mx$
4	$(\exists x)(Px \wedge Rb)$	$\exists I$ 3	

1	$(\exists x)Rx \rightarrow Mc$	P
2	Ra	P
3	$(\exists x)Rx$	$\exists I$ 2
4	Mc	$\rightarrow E$ 1, 3
5	$(\exists x)Mx$	$\exists I$ 4

Solution 8.82

1. $(\forall x)Px \vdash (\forall y)Py$

1	$(\forall x)Px$	P
2	Pa	$\forall E$ 1
3	$(\forall y)Py$	$\forall I$ 2

3. $(\forall x)Px \wedge (\forall y)Sy \vdash (\forall x)(Px \wedge Sx)$

1	$(\forall x)Px \wedge (\forall y)Sy$	P
2	$(\forall x)Px$	$\wedge E$ 1
3	Pa	$\forall E$ 2
4	$(\forall y)Sy$	$\wedge E$ 1
5	Sa	$\forall E$ 4
6	$Pa \wedge Sa$	$\wedge I$ 3, 5
7	$(\forall x)(Px \wedge Sx)$	$\forall I$ 6

Solution 8.83

1. $(\exists x)Lxx \vdash (\exists y)Lyy$

1	$(\exists x)Lxx$	P
2	$\left \frac{Lbb}{(\exists y)Lyy} \right.$	A
3	$(\exists y)Lyy$	$\exists I$ 2
4	$(\exists y)Lyy$	$\exists E$ 1, 2-3

3. $(\forall x)Lax, (\exists x)Px \vdash (\exists y)Py$

1	$(\forall x)Lax$	P
2	$(\exists x)Px$	P
3	$\left \frac{Pd}{(\exists y)Py} \right.$	A
4	$(\exists y)Py$	$\exists I$ 3
5	$(\exists y)Py$	$\exists E$ 2, 3-4

Solution 8.84

$$1. (\exists x)\neg Px, \neg(\forall x)Lxx \vdash \neg(\forall x)Px \wedge (\exists x)\neg Lxx$$

$$\begin{array}{ll} 1 & (\exists x)\neg Px \quad \text{P} \\ 2 & \neg(\forall x)Lxx \quad \text{P} \\ 3 & \neg(\forall x)Px \quad \text{QN 1} \\ 4 & (\exists x)\neg Lxx \quad \text{QN 2} \\ 5 & \neg(\forall x)Px \wedge (\exists x)\neg Lxx \quad \wedge I \text{ 4, 5} \end{array}$$

$$3. (\forall x)\neg(\exists y)Pxy \rightarrow (\forall x)Zx, (\forall x)(\forall y)\neg Pxy \vdash (\forall x)Zx$$

$$\begin{array}{ll} 1 & (\forall x)\neg(\exists y)Pxy \rightarrow (\forall x)Zx \quad \text{P} \\ 2 & (\forall x)(\forall y)\neg Pxy \quad \text{P} \\ 3 & (\forall x)\neg(\exists y)Pxy \quad \text{QN 2} \\ 4 & (\forall x)Zx \quad \rightarrow E \text{ 1, 3} \end{array}$$