

Chapter 1

Logic and Set Theory

To criticize mathematics for its abstraction is to miss the point entirely. Abstraction is what makes mathematics work. If you concentrate too closely on too limited an application of a mathematical idea, you rob the mathematician of his most important tools: analogy, generality, and simplicity.

– *Ian Stewart*

Does God play dice? The mathematics of chaos

In mathematics, a **proof** is a demonstration that, assuming certain axioms, some statement is necessarily true. That is, a proof is a logical argument, not an empirical one. One must demonstrate that a proposition is true in all cases before it is considered a theorem of mathematics. An unproven proposition for which there is some sort of empirical evidence is known as a **conjecture**. Mathematical logic is the framework upon which rigorous proofs are built. It is the study of the principles and criteria of valid inference and demonstrations.

Logicians have analyzed set theory in great details, formulating a collection of axioms that affords a broad enough and strong enough foundation to mathematical reasoning. The standard form of axiomatic set theory is denoted ZFC and it consists of the Zermelo-Fraenkel (ZF) axioms combined with the axiom of choice (C). Each of the axioms included in this theory expresses a property of sets that is widely accepted by mathematicians. It is unfortunately true that careless use of set theory can lead to contradictions. Avoiding such contradictions was one of the original motivations for the axiomatization of set theory.

A rigorous analysis of set theory belongs to the foundations of mathematics and mathematical logic. The study of these topics is, in itself, a formidable task. For our purposes, it will suffice to approach basic logical concepts informally. That is, we adopt a naive point of view regarding set theory and assume that the meaning of a set as a collection of objects is intuitively clear. While informal logic is not itself rigorous, it provides the underpinning for rigorous proofs. The rules we follow in dealing with sets are derived from established axioms. At some point of your academic career, you may wish to study set theory and logic in greater detail. Our main purpose here is to learn how to state mathematical results clearly and how to prove them.

1.1 Statements

A proof in mathematics demonstrates the truth of certain **statement**. It is therefore natural to begin with a brief discussion of statements. A statement, or **proposition**, is the content of an assertion. It is either true or false, but cannot be both true and false at the same time. For example, the expression “There are no classes at Texas A&M University today” is a statement since it is either true or false. The expression “Do not cheat and do not tolerate those who do” is not a statement. Note that an expression being a statement does not depend on whether we personally can verify its validity. The expression “The base of the natural logarithm, denoted e , is an irrational number” is a statement that most of us cannot prove.

Statements on their own are fairly uninteresting. What brings value to logic is the fact that there are a number of ways to form new statements from old ones. In this section, we present five ways to form new statements from old ones. They correspond to the English expressions: and; or; not; if, then; if and only if. In the discussion below, P and Q represent two abstract statements.

A logical **conjunction** is an operation on two logical propositions that produces a value of true if both statements are true, and is false otherwise. The conjunction (or logical AND) of P and Q , denoted by $P \wedge Q$, is precisely defined by

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

Similarly, a logical **disjunction** is an operator on two logical propositions that is true if either statement is true or both are true, and is false otherwise. The disjunction (or logical OR) of P and Q , denoted $P \vee Q$, is defined by

P	Q	$P \vee Q$
T	T	T
T	F	T
F	T	T
F	F	F

In mathematics, a **negation** is an operator on the logical value of a proposition that sends true to false and false to true. The negation (or logical NOT) of P , denoted $\neg P$, is given by

P	$\neg P$
T	F
F	T

The next method of combining mathematical statements is slightly more subtle than the preceding ones. The **conditional connective** $P \rightarrow Q$ is a logical statement that is read “if P then Q ” and defined by the truth table

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

In this statement, P is called the **antecedent** and Q is called the **consequent**. The truth table should match your intuition when P is true. When P is false, students often think the resulting truth value should be undefined. Although the given definition may seem strange at first glance, this truth table is universally accepted by mathematicians.

To motivate this definition, one can think of $P \rightarrow Q$ as a promise that Q is true whenever P is true. When P is false, the promise is kept by default. For example, suppose your friend promises “if it is sunny tomorrow, I will ride my bike”. We will call this a true statement if they keep their promise. If it rains and they don’t ride their bike, most people would agree that they have still kept their promise. Therefore, this definition allows one to combine many statements together and detect broken promises without being distracted by uninformative statements.

Logicians draw a firm distinction between the **conditional connective** and the **implication relation**. They use the phrase “if P then Q ” for the conditional connective and the phrase “ P implies Q ” for the implication relation. They explain the difference between these two forms by saying that the conditional is the contemplated relation, while the implication is the asserted relation. We will discuss this distinction in the Section 1.2, where we formally study relations between statements. The importance and soundness of the conditional form $P \rightarrow Q$ will become clearer then.

The logical **biconditional** is an operator connecting two logical propositions that is true if the statements are both true or both false, and it is false otherwise. The biconditional from P to Q , denoted $P \leftrightarrow Q$, is precisely defined by

P	Q	$P \leftrightarrow Q$
T	T	T
T	F	F
F	T	F
F	F	T

We read $P \leftrightarrow Q$ as “ P if and only if Q .” The phrase “if and only if” is often abbreviated as “iff”.

Using the five basic operations defined above, it is possible to form more complicated compound statements. We sometimes need parentheses to avoid ambiguity in writing compound statements. We use the convention that $\neg$ takes precedence over the other four operations, but none of these operations takes precedence over the others. For example, let P , Q and R be three propositions. We wish to make a truth table for the following statement,

$$(P \rightarrow R) \wedge (Q \vee \neg R). \quad (1.1)$$

We can form the true table for this statement, using simple steps, as follows

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

We conclude this section with a brief mention of two important concepts. A **tautology** is a statement that is true in every valuation of its propositional variables, independent of the truth values assigned to these variables. The proverbial tautology is $P \vee \neg P$,

P	$P \vee \neg P$
T	T
F	T
	1
	3
	2

For instance, the statement “The Aggies won their last football game or the Aggies did not win their last football game” is true regardless of whether the Aggies actually defeated their latest opponent.

The negation of a tautology is a **contradiction**, a statement that is necessarily false regardless of the truth values of its propositional variables. The statement $P \wedge \neg P$ is a contradiction, and its truth table is

P	$P \wedge \neg P$
T	F
F	F
	1
	3
	2

Of course, most statements we encounter are neither tautologies nor contradictions. For example, (1.1) is not necessarily either true or false. Its truth value depends on the values of P , Q and R . Try to see whether the statement

$$((P \wedge Q) \rightarrow R) \rightarrow (P \rightarrow (Q \rightarrow R))$$

is a tautology, a contradiction, or neither.

1.2 Relations between Statements

Strictly speaking, relations between statements are not formal statements themselves. They are *meta-statements* about some propositions. We study two types of relations between statements, *implication* and *equivalence*. An example of an implication meta-statement is the observation that “if the statement ‘Robert graduated from Texas A&M University’ is true, then it implies that the statement ‘Robert is an Aggie’ is also true.” Another example of a meta-statement is “the statement ‘Fred is an Aggie and Fred is honest’ being true is equivalent to the statement ‘Fred is honest and Fred is an Aggie’ being true.” These two examples illustrate how meta-statements describe the relationship between statements. It is also instructive to note that implications and equivalences are the meta-statement analogs of conditionals and biconditionals.

Consider two compound statements P and Q that depend on other logical statements (e.g., $P = (R \rightarrow S) \wedge (S \rightarrow T)$ and $Q = R \rightarrow T$). A **logical implication** from P to Q , read as “ P implies Q ”, asserts that Q must be true whenever P is true (i.e., for all possible truth values of the dependent statements R, S, T). Necessity is the key aspect of this sentence; the fact that P and Q both happen to be true cannot be coincidental. To state that P implies Q , denoted by $P \Rightarrow Q$, one needs the conditional $P \rightarrow Q$ to be true under all possible circumstances.

Meta-statements, such as “ P implies Q ”, can be defined formally only when P and Q are both logical functions of other propositions. For example, consider $P = R \wedge (R \rightarrow S)$ and $Q = S$. Then, the truth of the statement $P \rightarrow Q$ depends only on the truth of external propositions R and S .

The notion of implication can be rigorously defined as follows, P implies Q if the statement $P \rightarrow Q$ is a tautology. We abbreviate P implies Q by writing $P \Rightarrow Q$. It is important to understand the difference between “ $P \rightarrow Q$ ” and “ $P \Rightarrow Q$.” The former, $P \rightarrow Q$, is a compound statement that may or may not be true. On the other hand, $P \Rightarrow Q$ is a relation stating that the compound statement $P \rightarrow Q$ is true under all instances of the external propositions.

While the distinction between implication and conditional may seem extraneous, we will soon see that meta-statements become extremely useful in building valid arguments. In particular, the following implications are used extensively in constructing proofs.

Fact 1.2.1. Let P , Q , R and S be statements.

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

As an illustrative example, we show that $(P \rightarrow Q) \wedge (Q \rightarrow R)$ implies $P \rightarrow R$. To demonstrate this assertion, we need to show that

$$((P \rightarrow Q) \wedge (Q \rightarrow R)) \rightarrow (P \rightarrow R) \quad (1.2)$$

is a tautology. This is accomplished in the truth table below

P	Q	R	$((P \rightarrow Q) \wedge (Q \rightarrow R)) \rightarrow (P \rightarrow R)$
T	T	T	T
T	T	F	F
T	F	T	T
T	F	F	F
F	T	T	T
F	T	F	F
F	F	T	T
F	F	F	F

Column 11 has the truth values for statement (1.2). Since (1.2) is true under all circumstances, it is a tautology and the implication holds. Showing that the other relations are valid is left to the reader as an exercise.

Reversing the arrow in a conditional statement gives the **converse** of that statement. For example, the statement $Q \rightarrow P$ is the converse of $P \rightarrow Q$. This reversal

may not preserve the truth of the statement though and therefore logical implications are not always reversible. For instance, although $(P \rightarrow Q) \wedge (Q \rightarrow R)$ implies $P \rightarrow R$, the converse is not always true. It can easily be seen from columns 9 & 10 above that

$$(P \rightarrow R) \rightarrow ((P \rightarrow Q) \wedge (Q \rightarrow R))$$

is not a tautology. That is, $P \rightarrow R$ certainly does not imply $(P \rightarrow Q) \wedge (Q \rightarrow R)$.

A logical implication that is reversible is called a **logical equivalence**. More precisely, P is equivalent to Q if the statement $P \leftrightarrow Q$ is a tautology. We denote the sentence “ P is equivalent to Q ” by simply writing “ $P \leftrightarrow Q$.” The meta-statement $P \leftrightarrow Q$ holds if and only if $P \Rightarrow Q$ and $Q \Rightarrow P$ are both true. Being able to recognize that two statements are equivalent will become handy. It is sometime possible to demonstrate a result by finding an alternative, equivalent form of the statement that is easier to prove than the original form. A list of important equivalences appears below.

Fact 1.2.2. *Let P , Q and R be statements.*

1. $\neg(\neg P) \leftrightarrow P$.
2. $P \vee Q \leftrightarrow Q \vee P$.
3. $P \wedge Q \leftrightarrow Q \wedge P$.
4. $(P \vee Q) \vee R \leftrightarrow P \vee (Q \vee R)$.
5. $(P \wedge Q) \wedge R \leftrightarrow P \wedge (Q \wedge R)$.
6. $P \wedge (Q \vee R) \leftrightarrow (P \wedge Q) \vee (P \wedge R)$.
7. $P \vee (Q \wedge R) \leftrightarrow (P \vee Q) \wedge (P \vee R)$.
8. $P \rightarrow Q \leftrightarrow \neg P \vee Q$.
9. $P \rightarrow Q \leftrightarrow \neg Q \rightarrow \neg P$ (*Contrapositive*).
10. $P \leftrightarrow Q \leftrightarrow (P \rightarrow Q) \wedge (Q \rightarrow P)$.
11. $\neg(P \wedge Q) \leftrightarrow \neg P \vee \neg Q$ (*De Morgan's Law*).
12. $\neg(P \vee Q) \leftrightarrow \neg P \wedge \neg Q$ (*De Morgan's Law*).

Given a conditional statement of the form $P \rightarrow Q$, we call $\neg Q \rightarrow \neg P$ the **contrapositive** of the original statement. The equivalence $P \rightarrow Q \Leftrightarrow \neg Q \rightarrow \neg P$ noted above is used extensively in constructing mathematical proofs.

One must be careful not to allow contradictions in logical arguments because, starting from a contradiction, anything can be proven true. For example, one can verify that $P \wedge \neg P \Rightarrow Q$ is a valid logical equivalence. But, Q doesn't appear on the LHS. Thus, a contradiction in your assumptions can lead to a "correct" proof for an arbitrary statement.

Fortunately, propositional logic has an axiomatic formulation that is consistent, complete, and decidable. In this context, the term **consistent** means that the logical implications generated by the axioms do not contain a contradiction, the term **complete** means that any valid logical implication can be generated by applying the axioms, and the term **decidable** means there is a terminating method that always determines whether a postulated implication is valid or invalid.

1.2.1 Fallacious Arguments

A **fallacy** is a component of an argument that is demonstrably flawed in its logic or form, thus rendering the argument invalid. Recognizing fallacies in mathematical proofs may be difficult since arguments are often structured using convoluted patterns that obscure the logical connections between assertions. We give below examples for three types of fallacies that are often found in attempted mathematical proofs.

Affirming the Consequent: If the Indian cricket team wins a test match, then all the players will drink tea together. All the players drank tea together. Therefore the Indian cricket team won a test match.

Denying the Antecedent: If Diego Maradona drinks coffee, then he will be fidgety. Diego Maradona did not drink coffee. Therefore, he is not fidgety.

Unwarranted Assumptions: If Yao Ming gets close to the basket, then he scores a lot of points. Therefore, Yao Ming scores a lot of points.

1.2.2 Quantifiers

Consider the statements “Socrates is a person” and “Every person is mortal”. In propositional logic, there is no formal way to combine these statements to deduce that “Socrates is mortal”. In the first statement, the noun “Socrates” is called the subject and the phrase “is a person” is called the **predicate**. Likewise, in predicate logic, the statement $P(x) = “x \text{ is a person}”$ is called a predicate and x is called a **free variable** because its value is not fixed in the statement $P(x)$.

Let U be a specific collection of elements and let $P(x)$ be a statement that can be applied to any $x \in U$. In first-order predicate logic, quantifiers are applied to predicates in order to make statements about collections of elements. Later, we will see that quantifiers are of paramount importance in rigorous proofs.

The **universal quantifier** is typically denoted by $\forall$ and it is informally read “for all.” It follows that the statement “ $\forall x \in U, P(x)$ ” is true if $P(x)$ is true for all values of x in U . It can be seen as shorthand for an iterated conjunction because

$$\forall x \in U, P(x) \Leftrightarrow \bigwedge_{x \in U} P(x),$$

where $\Leftrightarrow$ indicates that these statements are equivalent for all sets U and predicates P . If $U = \emptyset$ is the empty set, then $\forall x \in U, P(x)$ is vacuously true by convention because there are no elements in U to test with $P(x)$.

Returning to the motivating example, let us also define $Q(x) = “x \text{ is mortal}”$. With these definitions, we can write the statement “Every person is mortal” as $\forall x, (P(x) \rightarrow Q(x))$. In logic, this usage implies that x ranges over the universal set. In engineering mathematics, however, the range of free variables is typically stated explicitly.

The other type of quantifier often seen in mathematical proofs is the **existential quantifier**, denoted $\exists$. The statement “ $\exists x \in U, P(x)$ ” is true if $P(x)$ is true for at least one value of x in U . It can be seen as shorthand for an iterated disjunction because

$$\exists x \in U, P(x) \Leftrightarrow \bigvee_{x \in U} P(x),$$

From these definitions, it follows naturally that $\forall x \in U, P(x) \Rightarrow \exists x \in U, P(x)$. If $U = \emptyset$ is the empty set, then $\exists x \in U, P(x)$ is false by convention because there are no elements in U .

Based on the meaning of these quantifiers, one can infer the logical implications

$$\begin{aligned}\neg(\forall x \in U, P(x)) &\Leftrightarrow \exists x \in U, \neg P(x) \\ \neg(\exists x \in U, P(x)) &\Leftrightarrow \forall x \in U, \neg P(x).\end{aligned}$$

Using the connection to conjunction and disjunction, these rules are actually equivalent to De Morgan's law for iterated conjunctions and disjunctions.

One can also define predicates with multiple free variables such as $P(x, y)$ = “ x contains y ”. Once again, these statements are assumed to be true or false for every choice of x, y . There are 8 possible quantifiers for a 2-variable predicate and they can be arranged according to their natural implications:

$$\begin{array}{ccccccc}\forall x, \forall y, P(x, y) & \Rightarrow & \exists x, \forall y, P(x, y) & \Rightarrow & \forall y, \exists x, P(x, y) & \Rightarrow & \exists y, \exists x, P(x, y) \\ \Downarrow & & & & & & \Downarrow \\ \forall y, \forall x, P(x, y) & \Rightarrow & \exists y, \forall x, P(x, y) & \Rightarrow & \forall x, \exists y, P(x, y) & \Rightarrow & \exists x, \exists y, P(x, y)\end{array}$$

All of these implications follow from $\forall x \forall y = \forall y \forall x$, $\exists x \exists y = \exists y \exists x$, and the single variable inference rule $\forall x, P(x) \Rightarrow \exists x, P(x)$ except for two: $\exists x, \forall y, P(x, y) \Rightarrow \forall y, \exists x, P(x, y)$ and its symmetric pair.

To understand this last implication, consider an example where x is in a set I of images and y is in a set C of colors. Then, $\exists x, \forall y, P(x, y)$ means “there is an image that contains all the colors” (e.g., an image of a rainbow) and $\forall y, \exists x, P(x, y)$ means “for each color there is an image containing that color”. The first statement implies the second because, in the second, the rainbow image satisfies the $\exists x$ quantifier for all y . To see that the implication is not an equivalence, consider a set of pictures where each image contains exactly one color and there is one such image for each color. In this case, it is true that “for each color there is an image containing that color” but it is not true that “there is an image that contains all the colors”.

In quantified statements, such as $\exists x \in U, P(x)$, the variable x is called a **bound variable** because its value cannot be chosen freely. Similarly, in the statement $\exists y \in U, P(x, y)$, x is a free variable and y is a bound variable.

Finally, we note that first-order predicate logic has an axiomatic formulation that is consistent, complete, and semidecidable. In this context, **semidecidable** means that there is an algorithm that, if it terminates, correctly determines the truth of any postulated implication. But, it is only guaranteed to terminate for true postulates.

1.3 Strategies for Proofs

The relation between intuition and formal rigor is not a trivial matter. Intuition tells us what is important, what might be true, and what mathematical tools may be used to prove it. Rigorous proofs are used to verify that a given statement which appears intuitively true is indeed true. Ultimately, a mathematical proof is a convincing argument that starts from some premises, and logically deduces the desired conclusion. Most proofs do not mention the logical rules of inference used in the derivation. Rather, they focus on the mathematical justification of each step, leaving to the reader the task of filling the logical gaps. The mathematics is the major issue. Yet, it is essential that you understand the underlying logic behind the derivation as to not get confused while reading or writing a proof.

True statements in mathematics have different names. They can be called theorems, propositions, lemmas, corollaries and exercises. A **theorem** is a statement that can be proved on the basis of explicitly stated or previously agreed assumptions. A **proposition** is a statement not associated with any particular theorem; this term sometimes connotes a statement with a simple proof. A **lemma** is a proven proposition which is used as a stepping stone to a larger result rather than an independent statement in itself. A **corollary** is a mathematical statement which follows easily from a previously proven statement, typically a mathematical theorem. The distinction between these names and their definitions is somewhat arbitrary. Ultimately, they are all synonymous to a true statement.

A proof should be written in grammatically correct English. Complete sentences should be used, with full punctuation. In particular, every sentence should end with a period, even if the sentence ends in a displayed equation. Mathematical formulas and symbols are parts of sentences, and are treated no differently than words. One way to learn to construct proofs is to read a lot of well written proofs, to write progressively more difficult proofs, and to get detailed feedback on the proofs you write.

Direct Proof: The simplest form of proof for a statement of the form $P \rightarrow Q$ is the **direct proof**. First assume that P is true. Produce a series of steps, each one following from the previous ones, that eventually leads to conclusion Q . It warrants the name “direct proof” only to distinguish it from other, more intricate, methods

of proof.

Proof by Contrapositive: A proof by contrapositive takes advantage of the mathematical equivalence $P \rightarrow Q \Leftrightarrow \neg Q \rightarrow \neg P$. That is, a proof by contrapositive begins by assuming that Q is false (i.e., $\neg Q$ is true). It then produces a series of direct implications leading to the conclusion that P is false (i.e., $\neg P$ is true). It follows that Q cannot be false when P is true, so $P \rightarrow Q$.

Proof by Contradiction: A proof by contradiction is based on the mathematical equivalence $\neg(P \rightarrow Q) \Leftrightarrow P \wedge \neg Q$. In a proof by contradiction, one starts by assuming that both P and $\neg Q$ are true. Then, a series of direct implications are given that lead to a logical contradiction. Hence, $P \wedge \neg Q$ cannot be true and $P \rightarrow Q$.

Example 1.3.1. *We wish to show that $\sqrt{2}$ is an irrational number.*

First, suppose that $\sqrt{2}$ is a rational number. This would imply that there exist integers p and q with $q \neq 0$ such that $p/q = \sqrt{2}$. In fact, we can further assume that the fraction p/q is irreducible. That is, p and q are coprime integers (they have no common factor greater than 1). From $p/q = \sqrt{2}$, it follows that $p = \sqrt{2}q$, and so $p^2 = 2q^2$. Thus p^2 is an even number, which implies that p itself is even (only even numbers have even squares). Because p is even, there exists an integer r satisfying $p = 2r$. We then obtain the equation $(2r)^2 = 2q^2$, which is equivalent to $2r^2 = q^2$ after simplification. Because $2r^2$ is even, it follows that q^2 is even, which means that q is also even. We conclude that p and q are both even. This contradicts the fact that p/q is irreducible. Hence, the initial assumption that $\sqrt{2}$ is a rational number must be false. That is to say, $\sqrt{2}$ is irrational.

Example 1.3.2. *Consider the following statement, which is related to Example 1.3.1. “If $\sqrt{2}$ is rational, then $\sqrt{2}$ can be expressed as an irreducible fraction.” The contrapositive of this statement is “If $\sqrt{2}$ cannot be expressed as an irreducible fraction, then $\sqrt{2}$ is not rational.” Above, we proved that $\sqrt{2}$ cannot be expressed as an irreducible fraction and therefore $\sqrt{2}$ is not a rational number.*

The final proof strategy we discuss is finite induction.

Definition 1.3.3. *Let $P(n)$ be a logical statement for each $n \in \mathbb{N}$. The principle of **mathematical induction** states that $P(n)$ is true all $n \in \mathbb{N}$ if:*

1. $P(1)$ is true, and
2. $P(n) \rightarrow P(n+1)$ for all $n \in \mathbb{N}$.

From a foundational perspective, this statement is essentially equivalent to the existence and uniqueness of the natural numbers. It is taken as an axiom in the Peano axiomatic formulation of arithmetic. In contrast, the ZF axiomatic formulation of set theory defines the natural numbers as the smallest inductive set and the existence of an inductive set is taken as an axiom.

Example 1.3.4. Let $S_n = \sum_{i=1}^n i$. We wish to show that the statement $P(n) = "S_n = \frac{n^2+n}{2}"$ is true for all $n \in \mathbb{N}$. For $n = 1$, this is true because both expressions equal 1. For $P(n+1)$, we are given $P(n)$ and can write

$$S_{n+1} = S_n + (n+1) = \frac{n^2+n}{2} + n+1 = \frac{n^2+3n+2}{2} = \frac{(n+1)^2+(n+1)}{2}.$$

Thus, the result follows from mathematical induction.

More general forms of finite induction are also quite common but they can be reduced to the original form. For example, let $Q(m)$ be a predicate for $m \in \mathbb{N}$ and define $P(n) = "\forall m \in S_n, Q(m)"$ for a sequence nested finite sets $S_1 \subset S_2 \subset \dots \subseteq \mathbb{N}$. Defining $S_\infty = \bigcup_{n \in \mathbb{N}} S_n$, we see that $"\forall n \in \mathbb{N}, P(n)" \Leftrightarrow "\forall m \in S_\infty, Q(m)"$ follows from $P(1) = "\forall m \in S_1, Q(m)"$ and $"P(n) \rightarrow P(n+1)" \Leftrightarrow "\forall m \in S_n, Q(m) \rightarrow \forall m \in S_{n+1}, Q(m)"$.

1.4 Set Theory

Set theory is generally considered to be the foundation of all modern mathematics. This means that most mathematical objects (numbers, relations, functions, etc.) are defined in terms of sets. Unfortunately for engineers, set theory is not quite as simple as it seems. It turns out that simple approaches to set theory include paradoxes (e.g., statements which are both true and false). These paradoxes can be resolved by putting set theory in a firm axiomatic framework, but that exercise is rather unproductive for engineers. Instead, we adopt what is called **naïve set theory** which rigorously defines the operations of set theory without worrying about possible contradictions. This approach is sufficient for most of mathematics and also acts as a stepping-stone to more formal treatments.

A **set** is taken to be any collection of objects, mathematical or otherwise. For example, one can think of “the set of all books published in 2007”. The objects in a set are referred to as **elements** or members of the set. The logical statement “ a is a member of the set A ” is written

$$a \in A.$$

Likewise, its logical negation “ a is not a member of the set A ” is written $a \notin A$. Therefore, exactly one of these two statements is true. In naive set theory, one assumes the existence of any set that can be described in words. Later, we will see that this can be problematic when one considers objects like the “set of all sets”.

One may present a set by listing its elements. For example, $A = \{a, e, i, o, u\}$ is the set of standard English vowels. It is important to note that the order elements are presented is irrelevant and the set $\{i, o, u, a, e\}$ is the same as A . Likewise, repeated elements have no effect and the set $\{a, e, i, o, u, e, o\}$ is the same as A . A **singleton** set is a set containing exactly one element such as $\{a\}$.

There are a number of standard sets worth mentioning: the **integers** $\mathbb{Z}$, the **real numbers** $\mathbb{R}$, and the **complex numbers** $\mathbb{C}$. It is possible to construct these sets in a rigorous manner, but instead we will assume their meaning is intuitively clear. New sets can be defined in terms of old sets using **set-builder notation**. Let $P(x)$ be a logical statement about objects x in the set X , then the “set of elements in X such that $P(x)$ is true” is denoted by

$$\{x \in X | P(x)\}.$$

For example, the set of even integers is given by

$$\{x \in \mathbb{Z} | “x \text{ is even}”\} = \{\dots, -4, -2, 0, 2, 4, \dots\}.$$

If no element $x \in X$ satisfies the condition, then the result is the **empty set** which is denoted $\emptyset$. Using set-builder notation, we can also recreate the **natural numbers** $\mathbb{N}$ and the **rational numbers** $\mathbb{Q}$ with

$$\begin{aligned}\mathbb{N} &= \{n \in \mathbb{Z} | n \geq 1\} \\ \mathbb{Q} &= \{q \in \mathbb{R} | q = a/b, a \in \mathbb{Z}, b \in \mathbb{N}\}.\end{aligned}$$

The following standard notation is used for interval subsets of the real numbers:

Open interval: $(a, b) \triangleq \{x \in \mathbb{R} \mid a < x < b\}$

Closed interval: $[a, b] \triangleq \{x \in \mathbb{R} \mid a \leq x \leq b\}$

Half-open intervals: $(a, b] \triangleq \{x \in \mathbb{R} \mid a < x \leq b\}$

$[a, b) \triangleq \{x \in \mathbb{R} \mid a \leq x < b\}$

Definition 1.4.1. For a finite set A , the **cardinality** $|A|$ equals the number of elements in A . If there is a bijective mapping between the set A and the natural numbers $\mathbb{N}$, then $|A| = \infty$ and the set is called **countably infinite**. If $|A| = \infty$ and the set is not countably infinite, then A is called **uncountably infinite**.

Example 1.4.2. The set of rational numbers is countably infinite while the set of real numbers is uncountably infinite.

Example 1.4.3 (Russell's Paradox). Let R be the set of all sets that do not contain themselves or $R = \{S \mid S \notin S\}$. Such a set is said to exist in naive set theory (though it may empty) simply because it can be described in words. The paradox arises from the fact that the definition leads to the logical contradiction $R \in R \leftrightarrow R \notin R$.

What this proves is that *naive set theory is not consistent* because it allows constructions that lead to contradictions. Axiomatic set theory eliminates this paradox by disallowing self-referential and other problematic constructions. Thus, another reasonable conclusion is that Russell's paradox shows that the set R cannot exist in any consistent theory of sets.

Another common question is whether there are sets that contains themselves. In naive set theory, the answer is yes and some examples are the “set of all sets” and the “set of all abstract ideas”. On the other hand, in the ZF axiomatic formulation of set theory, it is a theorem that no set contains itself.

There are a few standard relationships defined between any two sets A, B .

Definition 1.4.4. We say that A **equals** B (denoted $A = B$) if, for all x , $x \in A$ iff $x \in B$. This means that

$$A = B \Leftrightarrow \forall x ((x \in A) \leftrightarrow (x \in B)).$$

Definition 1.4.5. We say that A is a **subset** of B (denoted $A \subseteq B$) if, for all x , if $x \in A$ then $x \in B$. This means that

$$A \subseteq B \Leftrightarrow \forall x ((x \in A) \rightarrow (x \in B)).$$

It is a **proper subset** (denoted $A \subset B$) if $A \subseteq B$ and $A \neq B$.

There are also a number of operations between sets. Let A, B be any two sets.

Definition 1.4.6. The **union** of A and B (denoted $A \cup B$) is the set of elements in either A or B . This means that $A \cup B = \{x \in A \text{ or } x \in B\}$ is also defined by

$$x \in A \cup B \Leftrightarrow (x \in A) \vee (x \in B).$$

Definition 1.4.7. The **intersection** of A and B (denoted $A \cap B$) is the set of elements in both A and B . This means that $A \cap B = \{x \in A | x \in B\}$ is also defined by

$$x \in A \cap B \Leftrightarrow (x \in A) \wedge (x \in B).$$

Two sets are said to be **disjoint** if $A \cap B = \emptyset$.

Definition 1.4.8. The **set difference** between A and B (denoted $A - B$ or $A \setminus B$) is the set of elements in A but not in B . This means that

$$x \in A - B \Leftrightarrow (x \in A) \wedge (x \notin B).$$

If there is some implied universal set U , then the **complement** (denoted A^c) is defined by $A^c = U - A$

One can apply De Morgan's Law in set theory to verify that

$$(A \cup B)^c = A^c \cap B^c$$

$$(A \cap B)^c = A^c \cup B^c,$$

which allows us to interchange union or intersection with set difference.

We can also form the union or the intersection of arbitrarily many sets. This is defined in a straightforward way,

$$\bigcup_{\alpha \in I} S_\alpha = \{x | x \in S_\alpha \text{ for some } \alpha \in I\}$$

$$\bigcap_{\alpha \in I} S_\alpha = \{x | x \in S_\alpha \text{ for all } \alpha \in I\}.$$

It is worth noting that the definitions apply whether the index set is finite, countably infinite, or even uncountably infinite.

Another way to build sets is by grouping elements into pairs, triples, and vectors.

Definition 1.4.9. The **Cartesian Product**, denoted $A \times B$, of two sets is the set of ordered pairs $\{(a, b) | a \in A, b \in B\}$. For n -tuples taken from the same set, the notation A^n denotes the n -fold product $A \times A \times \cdots \times A$.

Example 1.4.10. If $A = \{a, b\}$, then the set of all 3-tuples from A is given by

$$A^3 = \{(a, a, a), (a, a, b), (a, b, a), (a, b, b), (b, a, a), (b, a, b), (b, b, a), (b, b, b)\}.$$

The countably infinite product of X , denoted X^ω , is the set of infinite sequences $(x_1, x_2, x_3, \dots)$ where $x_n \in X$ is arbitrary for $n \in \mathbb{N}$. If the sequences are restricted to have only a finite number of non-zero terms, then the set is usually denoted X^∞ .

One can also formalize relationships between elements of a set. A **relation** $\sim$ between elements of the set A is defined by the pairs $(x, y) \in A \times A$ for which the relation holds. Specifically, the relation is defined by the subset of ordered pairs $E \subseteq A \times A$ where the relation $a \sim b$ holds; so $x \sim y$ if and only if $(x, y) \in E$. A relation on A is said to be:

1. Reflexive if $x \sim x$ holds for all $x \in A$
2. Symmetric if $x \sim y$ implies $y \sim x$ for all $x, y \in A$
3. Transitive if $x \sim y$ and $y \sim z$, then $x \sim z$ for all $x, y, z \in A$

A relation is called an **equivalence relation** if it is reflexive, symmetric, and transitive. For example, let A be a set of people and $P(x, y)$ be the statement “ x has the same birthday (month and day) as y .” Then, we can define $\sim$ such that $a \sim b$ holds if and only if $P(x, y)$ is true. In this case, the set E is given by $E = \{(x, y) \in A \times A | P(x, y)\}$. One can verify that this is an equivalence relation by checking that it is reflexive, symmetric, and transitive.

One important characteristic of an equivalence relation is that it partitions the entire set A into disjoint **equivalence classes**. The equivalence class associated with $a \in A$ is given by $[a] = \{x \in A | x \sim a\}$. In the birthday example, there is a natural equivalence class associated with each day of the year. The set of all equivalence classes is called the **quotient set** and is denoted $A/\sim = \{[a] | a \in A\}$.

In fact, there is a natural equivalence relation defined by any disjoint partition of a set. For example, let $A_{i,j}$ be the set of people in A whose birthday was on the j -th day of the i -th month. It follows that $x \sim y$ if and only if there exists a unique pair i, j such that $x, y \in A_{i,j}$. In this case, the days of year are used as equivalence classes to define the equivalence relation.

Example 1.4.11. Consider the set $\mathbb{N} \times \mathbb{N} = \{(a, b) | a, b \in \mathbb{N}\}$ of ordered pairs of natural numbers. If one associates the element (a, b) with the fraction a/b , then the entire set is associated with the set of (possibly reducible) fractions. Now, consider the equivalence relation $(a, b) \sim (c, d)$ if $ad = bc$. In this case, two ordered pairs are equivalent if their associated fractions evaluate to the same real number. The quotient set $\mathbb{N}/\sim$ can therefore be associated with the set of reduced fractions.

Unfortunately, this section will not end on a happy note by saying that the ZFC axiomatic formulation of set theory is consistent. Instead, we observe that Kurt Gödel's Incompleteness Theorems imply that, if ZFC is consistent, then this cannot be proven using statements in ZFC and, moreover, it cannot be complete. On the other hand, if ZFC is inconsistent, then it contains a paradox and one can prove anything using statements in ZFC. Since ZFC manages to avoid all known paradoxes and no contradictions have been so far, it is still the most popular formal system in which to define mathematics.

1.5 Functions

In elementary mathematics, functions are typically described in terms of graphs and formulas. The drawback of this approach is that one tends to picture only “nice” functions. In fact, Cauchy himself published in 1821 an incorrect proof of the false assertion that “a sequence of continuous functions that converges everywhere has a continuous limit function.” Nowadays, every teacher warns their students that one must be careful because the world is filled with “not so nice” functions.

The modern approach to defining functions is based on set theory. A **function** $f: X \rightarrow Y$ is a rule that assigns a single value $f(x) \in Y$ to each element $x \in X$. The notation $f: X \rightarrow Y$ is used to emphasize the role of the **domain** X and the **codomain** Y . The **range** of f is the subset of Y which is actually achieved by f , $\{f(x) \in Y | x \in X\}$. Since the term codomain is somewhat uncommon, people

often use the term range instead of codomain either intentionally (for simplicity) or unintentionally (due to confusion).

Definition 1.5.1. Formally, a **function** $f: X \rightarrow Y$ from X to Y is defined by a subset $F \subset X \times Y$ such that $A_x = \{y \in Y \mid (x, y) \in F\}$ has exactly one element for each $x \in X$. The **value** of f at $x \in X$, denoted $f(x)$, is the unique element of Y contained in A_x .

Two functions are said to be equal if they have the same domain, codomain, and value for all elements of the domain. A function f is called:

1. **one-to-one** or **injective** if, for all $x, x' \in X$, if $f(x) = f(x')$ then $x = x'$;
2. **onto** or **surjective** if its range $\{f(x) \mid x \in X\}$ equals Y ;
3. a **one-to-one correspondence** or **bijective** if it is both one-to-one and onto.

A bijective function $f: X \rightarrow Y$ has a unique **inverse function** $f^{-1}: Y \rightarrow X$ such that $f^{-1}(f(x)) = x$ for all $x \in X$ and $f(f^{-1}(y)) = y$ for all $y \in Y$. In fact, any one-to-one function $f: X \rightarrow Y$ can be transformed into a bijective function $g: X \rightarrow R$ with $g(x) = f(x)$ by restricting its codomain Y to its range R .

Functions can also be applied to sets in a natural way. For a function $f: X \rightarrow Y$ and subset $A \subseteq X$, the **image** of A under f is

$$f(A) \triangleq \{y \in Y \mid \exists x \in A \text{ s.t. } f(x) = y\} = \{f(x) \mid x \in A\}.$$

Using this definition, we see that the range of f is simply $f(X)$. One benefit of allowing functions to have set-valued images is that a set-valued inverse function always exists. The **inverse image** or **preimage** of a subset $B \subseteq Y$ is

$$f^{-1}(B) \triangleq \{x \in X \mid f(x) \in B\}.$$

For a one-to-one function f , the inverse image of any singleton set $\{f(x)\}$ is the singleton set $\{x\}$. It is worth noting that the notation $f^{-1}(B)$ for the preimage of B can be somewhat misleading because, in some cases, $f^{-1}(f(A)) \neq A$. In general, a function gives rise to the following property, $f(f^{-1}(B)) \subseteq B$ and $f^{-1}(f(A)) \supseteq A$.

Example 1.5.2. Let the function $f: \mathbb{R} \rightarrow \mathbb{R}$ be defined by $f(x) = x^2$. Let $A = [1, 2]$ and notice that $B = f(A) = [1, 4]$. Then,

$$f^{-1}(B) = f^{-1}([1, 4]) = [-2, -1] \cup [1, 2] \supseteq A.$$

Example 1.5.3. Let the function $f: \mathbb{R} \rightarrow \mathbb{R}$ be defined by $f(x) = x^2 + 1$. Let $B = [0, 2]$ and notice that $A = f^{-1}(B) = [-1, 1]$. Then,

$$f(A) = f([-1, 1]) = [1, 2] \subseteq B.$$

Problem 1.5.4. For all $f: X \rightarrow Y$, $A \subseteq X$, and $B \subseteq Y$, we have the rules:

$$\begin{aligned} (a) \quad x \in A &\Rightarrow f(x) \in f(A) & (b) \quad y \in f(A) &\Rightarrow \exists x \in A \text{ s.t. } f(x) = y \\ (c) \quad x \in f^{-1}(B) &\Rightarrow f(x) \in B & (d) \quad f(x) \in B &\Rightarrow x \in f^{-1}(B). \end{aligned}$$

Use these rules to show that $f^{-1}(f(A)) \supseteq A$ and $f(f^{-1}(B)) \subseteq B$.

Solution 1.5.4. The first result follows from

$$x \in A \xRightarrow{(a)} f(x) \in f(A) \xRightarrow{(d)} x \in f^{-1}(f(A)),$$

and the definition of subset. The second result follows from

$$y \in f(f^{-1}(B)) \xRightarrow{(b)} \exists x \in f^{-1}(B) \text{ s.t. } f(x) = y \xRightarrow{(c)} y \in B,$$

and the definition of subset.

Problem 1.5.5. Let $f: X \rightarrow Y$, $A_i \subseteq X$ for all $i \in I$, and $B_i \subseteq Y$ for all $i \in I$. Show that the following expressions hold:

$$\begin{aligned} (1) \quad f\left(\bigcup_{i \in I} A_i\right) &= \bigcup_{i \in I} f(A_i) & (2) \quad f\left(\bigcap_{i \in I} A_i\right) &\subseteq \bigcap_{i \in I} f(A_i) \\ (3) \quad f^{-1}\left(\bigcup_{i \in I} B_i\right) &= \bigcup_{i \in I} f^{-1}(B_i) & (4) \quad f^{-1}\left(\bigcap_{i \in I} B_i\right) &= \bigcap_{i \in I} f^{-1}(B_i). \end{aligned}$$

