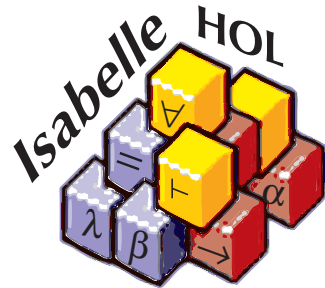


Tobias Nipkow

Programming and Proving in Isabelle/HOL



February 17, 2016

Contents

1	Introduction	1
2	Programming and Proving	3
2.1	Basics	3
2.2	Types <i>bool</i> , <i>nat</i> and <i>list</i>	5
2.3	Type and Function Definitions	13
2.4	Induction Heuristics	17
2.5	Simplification	19
3	Logic and Proof Beyond Equality	25
3.1	Formulas	25
3.2	Sets	26
3.3	Proof Automation	27
3.4	Single Step Proofs	30
3.5	Inductive Definitions	33
4	Isar: A Language for Structured Proofs	41
4.1	Isar by Example	42
4.2	Proof Patterns	44
4.3	Streamlining Proofs	46
4.4	Case Analysis and Induction	49
	References	59

Introduction

Isabelle is a generic system for implementing logical formalisms, and Isabelle/HOL is the specialization of Isabelle for HOL, which abbreviates Higher-Order Logic. We introduce HOL step by step following the equation

$$\text{HOL} = \text{Functional Programming} + \text{Logic}.$$

We assume that the reader is used to logical and set-theoretic notation and is familiar with the basic concepts of functional programming.

[Chapter 2](#) introduces HOL as a functional programming language and explains how to write simple inductive proofs of mostly equational properties of recursive functions. [Chapter 3](#) introduces the rest of HOL: the language of formulas beyond equality, automatic proof tools, single-step proofs, and inductive definitions, an essential specification construct. [Chapter 4](#) introduces Isar, Isabelle's language for writing structured proofs.

This introduction to the core of Isabelle is intentionally concrete and example-based: we concentrate on examples that illustrate the typical cases without explaining the general case if it can be inferred from the examples. We cover the essentials (from a functional programming point of view) as quickly and compactly as possible.

For a comprehensive treatment of all things Isabelle we recommend the *Isabelle/Isar Reference Manual* [6], which comes with the Isabelle distribution. The tutorial by Nipkow, Paulson and Wenzel [5] (in its updated version that comes with the Isabelle distribution) is still recommended for the wealth of examples and material, but its proof style is outdated. In particular it does not cover the structured proof language Isar.

If you want to apply what you have learned about Isabelle we recommend you download and read the book [Concrete Semantics](#) [4], a guided tour of the wonderful world of programming language semantics formalised in Isabelle. In fact, *Programming and Proving in Isabelle/HOL* constitutes part I of

[Concrete Semantics](#). The web pages for [Concrete Semantics](#) also provide a set of L^AT_EX-based slides and Isabelle demo files for teaching *Programming and Proving in Isabelle/HOL*.

Acknowledgements

I wish to thank the following people for their comments on this document: Florian Haftmann, Peter Johnson, René Thiemann, Sean Seefried, Christian Sternagel and Carl Witty.

Programming and Proving

This chapter introduces HOL as a functional programming language and shows how to prove properties of functional programs by induction.

2.1 Basics

2.1.1 Types, Terms and Formulas

HOL is a typed logic whose type system resembles that of functional programming languages. Thus there are

base types, in particular *bool*, the type of truth values, *nat*, the type of natural numbers ($\mathbb{N}$), and *int*, the type of mathematical integers ($\mathbb{Z}$).

type constructors, in particular *list*, the type of lists, and *set*, the type of sets. Type constructors are written postfix, i.e., after their arguments. For example, *nat list* is the type of lists whose elements are natural numbers.

function types, denoted by $\Rightarrow$.

type variables, denoted by *'a*, *'b*, etc., like in ML.

Note that *'a* $\Rightarrow$ *'b list* means *'a* $\Rightarrow$ (*'b list*), not (*'a* $\Rightarrow$ *'b*) *list*: postfix type constructors have precedence over $\Rightarrow$.

Terms are formed as in functional programming by applying functions to arguments. If *f* is a function of type $\tau_1 \Rightarrow \tau_2$ and *t* is a term of type τ_1 then *f t* is a term of type τ_2 . We write *t* :: τ to mean that term *t* has type τ .

 There are many predefined infix symbols like + and $\leq$. The name of the corresponding binary function is *op* +, not just +. That is, *x* + *y* is nice surface syntax (“syntactic sugar”) for *op* + *x y*.

HOL also supports some basic constructs from functional programming:

$(if\ b\ then\ t_1\ else\ t_2)$
 $(let\ x = t\ in\ u)$
 $(case\ t\ of\ pat_1 \Rightarrow t_1 \mid \dots \mid pat_n \Rightarrow t_n)$



The above three constructs must always be enclosed in parentheses if they occur inside other constructs.

Terms may also contain λ -abstractions. For example, $\lambda x. x$ is the identity function.

Formulas are terms of type *bool*. There are the basic constants *True* and *False* and the usual logical connectives (in decreasing order of precedence):

$\neg, \wedge, \vee, \longrightarrow$.

Equality is available in the form of the infix function $=$ of type $'a \Rightarrow 'a \Rightarrow bool$. It also works for formulas, where it means “if and only if”.

Quantifiers are written $\forall x. P$ and $\exists x. P$.

Isabelle automatically computes the type of each variable in a term. This is called **type inference**. Despite type inference, it is sometimes necessary to attach an explicit **type constraint** (or **type annotation**) to a variable or term. The syntax is $t :: \tau$ as in $m + (n::nat)$. Type constraints may be needed to disambiguate terms involving overloaded functions such as $+$.

Finally there are the universal quantifier $\bigwedge$ and the implication $\Longrightarrow$. They are part of the Isabelle framework, not the logic HOL. Logically, they agree with their HOL counterparts $\forall$ and $\longrightarrow$, but operationally they behave differently. This will become clearer as we go along.



Right-arrows of all kinds always associate to the right. In particular, the formula $A_1 \Longrightarrow A_2 \Longrightarrow A_3$ means $A_1 \Longrightarrow (A_2 \Longrightarrow A_3)$. The (Isabelle-specific¹) notation $\llbracket A_1; \dots; A_n \rrbracket \Longrightarrow A$ is short for the iterated implication $A_1 \Longrightarrow \dots \Longrightarrow A_n \Longrightarrow A$.

Sometimes we also employ inference rule notation:
$$\frac{A_1 \quad \dots \quad A_n}{A}$$

2.1.2 Theories

Roughly speaking, a **theory** is a named collection of types, functions, and theorems, much like a module in a programming language. All Isabelle text needs to go into a theory. The general format of a theory *T* is

```

theory T
imports T1 ... Tn
begin
  definitions, theorems and proofs
end

```

¹ To display implications in this style in Isabelle/jEdit you need to set Plugins > Plugin Options > Isabelle/General > Print Mode to “brackets” and restart.

where $T_1 \dots T_n$ are the names of existing theories that T is based on. The T_i are the direct **parent theories** of T . Everything defined in the parent theories (and their parents, recursively) is automatically visible. Each theory T must reside in a **theory file** named $T.thy$.

 HOL contains a theory *Main*, the union of all the basic predefined theories like arithmetic, lists, sets, etc. Unless you know what you are doing, always include *Main* as a direct or indirect parent of all your theories.

In addition to the theories that come with the Isabelle/HOL distribution (see <http://isabelle.in.tum.de/library/HOL/>) there is also the *Archive of Formal Proofs* at <http://afp.sourceforge.net>, a growing collection of Isabelle theories that everybody can contribute to.

2.1.3 Quotation Marks

The textual definition of a theory follows a fixed syntax with keywords like **begin** and **datatype**. Embedded in this syntax are the types and formulas of HOL. To distinguish the two levels, everything HOL-specific (terms and types) must be enclosed in quotation marks: "...". Quotation marks around a single identifier can be dropped. When Isabelle prints a syntax error message, it refers to the HOL syntax as the **inner syntax** and the enclosing theory language as the **outer syntax**.

2.1.4 Proof State

 By default Isabelle/jEdit does not show the proof state in the output window. You should enable this by ticking the "Proof state" box.

2.2 Types *bool*, *nat* and *list*

These are the most important predefined types. We go through them one by one. Based on examples we learn how to define (possibly recursive) functions and prove theorems about them by induction and simplification.

2.2.1 Type *bool*

The type of boolean values is a predefined datatype

datatype *bool* = *True* | *False*

with the two values *True* and *False* and with many predefined functions: $\neg$, $\wedge$, $\vee$, $\longrightarrow$, etc. Here is how conjunction could be defined by pattern matching:

```
fun conj :: "bool  $\Rightarrow$  bool  $\Rightarrow$  bool" where
  "conj True True = True" |
  "conj _ _ = False"
```

Both the datatype and function definitions roughly follow the syntax of functional programming languages.

2.2.2 Type *nat*

Natural numbers are another predefined datatype:

```
datatype nat = 0 | Suc nat
```

All values of type *nat* are generated by the constructors *0* and *Suc*. Thus the values of type *nat* are *0*, *Suc 0*, *Suc (Suc 0)*, etc. There are many predefined functions: $+$, $*$, $\leq$, etc. Here is how you could define your own addition:

```
fun add :: "nat  $\Rightarrow$  nat  $\Rightarrow$  nat" where
  "add 0 n = n" |
  "add (Suc m) n = Suc(add m n)"
```

And here is a proof of the fact that $\text{add } m \ 0 = m$:

```
lemma add_02: "add m 0 = m"
  apply(induction m)
  apply(auto)
  done
```

The **lemma** command starts the proof and gives the lemma a name, *add_02*. Properties of recursively defined functions need to be established by induction in most cases. Command **apply**(*induction m*) instructs Isabelle to start a proof by induction on *m*. In response, it will show the following proof state:

1. $\text{add } 0 \ 0 = 0$
2. $\wedge m. \text{add } m \ 0 = m \implies \text{add } (\text{Suc } m) \ 0 = \text{Suc } m$

The numbered lines are known as *subgoals*. The first subgoal is the base case, the second one the induction step. The prefix $\wedge m.$ is Isabelle's way of saying "for an arbitrary but fixed *m*". The $\implies$ separates assumptions from the conclusion. The command **apply**(*auto*) instructs Isabelle to try and prove all subgoals automatically, essentially by simplifying them. Because both subgoals are easy, Isabelle can do it. The base case $\text{add } 0 \ 0 = 0$ holds by definition of *add*, and the induction step is almost as simple: $\text{add } (\text{Suc } m) \ 0 = \text{Suc}(\text{add } m \ 0) = \text{Suc } m$ using first the definition of *add* and then the

induction hypothesis. In summary, both subproofs rely on simplification with function definitions and the induction hypothesis. As a result of that final **done**, Isabelle associates the lemma just proved with its name. You can now inspect the lemma with the command

```
thm add_02
```

which displays

```
add ?m 0 = ?m
```

The free variable m has been replaced by the **unknown** $?m$. There is no logical difference between the two but there is an operational one: unknowns can be instantiated, which is what you want after some lemma has been proved.

Note that there is also a proof method *induct*, which behaves almost like *induction*; the difference is explained in [Chapter 4](#).

! Terminology: We use **lemma**, **theorem** and **rule** interchangeably for propositions that have been proved.

! Numerals (0, 1, 2, ...) and most of the standard arithmetic operations (+, −, *, ≤, <, etc.) are overloaded: they are available not just for natural numbers but for other types as well. For example, given the goal $x + 0 = x$, there is nothing to indicate that you are talking about natural numbers. Hence Isabelle can only infer that x is of some arbitrary type where 0 and + exist. As a consequence, you will be unable to prove the goal. In this particular example, you need to include an explicit type constraint, for example $x+0 = (x::nat)$. If there is enough contextual information this may not be necessary: $Suc\ x = x$ automatically implies $x::nat$ because *Suc* is not overloaded.

An Informal Proof

Above we gave some terse informal explanation of the proof of $add\ m\ 0 = m$. A more detailed informal exposition of the lemma might look like this:

Lemma $add\ m\ 0 = m$

Proof by induction on m .

- Case 0 (the base case): $add\ 0\ 0 = 0$ holds by definition of *add*.
- Case $Suc\ m$ (the induction step): We assume $add\ m\ 0 = m$, the induction hypothesis (IH), and we need to show $add\ (Suc\ m)\ 0 = Suc\ m$. The proof is as follows:

$$\begin{aligned} add\ (Suc\ m)\ 0 &= Suc\ (add\ m\ 0) && \text{by definition of } add \\ &= Suc\ m && \text{by IH} \end{aligned}$$

Throughout this book, **IH** will stand for “induction hypothesis”.

We have now seen three proofs of $\text{add } m \ 0 = 0$: the Isabelle one, the terse four lines explaining the base case and the induction step, and just now a model of a traditional inductive proof. The three proofs differ in the level of detail given and the intended reader: the Isabelle proof is for the machine, the informal proofs are for humans. Although this book concentrates on Isabelle proofs, it is important to be able to rephrase those proofs as informal text comprehensible to a reader familiar with traditional mathematical proofs. Later on we will introduce an Isabelle proof language that is closer to traditional informal mathematical language and is often directly readable.

2.2.3 Type *list*

Although lists are already predefined, we define our own copy for demonstration purposes:

```
datatype 'a list = Nil | Cons 'a "'a list"
```

- Type *'a list* is the type of lists over elements of type *'a*. Because *'a* is a type variable, lists are in fact **polymorphic**: the elements of a list can be of arbitrary type (but must all be of the same type).
- Lists have two constructors: *Nil*, the empty list, and *Cons*, which puts an element (of type *'a*) in front of a list (of type *'a list*). Hence all lists are of the form *Nil*, or *Cons x Nil*, or *Cons x (Cons y Nil)*, etc.
- **datatype** requires no quotation marks on the left-hand side, but on the right-hand side each of the argument types of a constructor needs to be enclosed in quotation marks, unless it is just an identifier (e.g., *nat* or *'a*).

We also define two standard functions, *append* and *reverse*:

```
fun app :: "'a list  $\Rightarrow$  'a list  $\Rightarrow$  'a list" where
  "app Nil ys = ys" |
  "app (Cons x xs) ys = Cons x (app xs ys)"
```

```
fun rev :: "'a list  $\Rightarrow$  'a list" where
  "rev Nil = Nil" |
  "rev (Cons x xs) = app (rev xs) (Cons x Nil)"
```

By default, variables *xs*, *ys* and *zs* are of *list* type.

Command **value** evaluates a term. For example,

```
value "rev(Cons True (Cons False Nil))"
```

yields the result *Cons False (Cons True Nil)*. This works symbolically, too:

```
value "rev(Cons a (Cons b Nil))"
```

yields *Cons b (Cons a Nil)*.

Figure 2.1 shows the theory created so far. Because *list*, *Nil*, *Cons*, etc. are already predefined, Isabelle prints qualified (long) names when executing this theory, for example, *MyList.Nil* instead of *Nil*. To suppress the qualified names you can insert the command `declare [[names_short]]`. This is not recommended in general but is convenient for this unusual example.

```
theory MyList
imports Main
begin

datatype 'a list = Nil | Cons 'a "'a list"

fun app :: "'a list => 'a list => 'a list" where
"app Nil ys = ys" |
"app (Cons x xs) ys = Cons x (app xs ys)"

fun rev :: "'a list => 'a list" where
"rev Nil = Nil" |
"rev (Cons x xs) = app (rev xs) (Cons x Nil)"

value "rev(Cons True (Cons False Nil))"

(* a comment *)

end
```

Fig. 2.1. A theory of lists

Structural Induction for Lists

Just as for natural numbers, there is a proof principle of induction for lists. Induction over a list is essentially induction over the length of the list, although the length remains implicit. To prove that some property *P* holds for all lists *xs*, i.e., *P xs*, you need to prove

1. the base case *P Nil* and
2. the inductive case *P (Cons x xs)* under the assumption *P xs*, for some arbitrary but fixed *x* and *xs*.

This is often called **structural induction** for lists.

2.2.4 The Proof Process

We will now demonstrate the typical proof process, which involves the formulation and proof of auxiliary lemmas. Our goal is to show that reversing a list twice produces the original list.

theorem *rev_rev* [*simp*]: " $rev(rev\ xs) = xs$ "

Commands **theorem** and **lemma** are interchangeable and merely indicate the importance we attach to a proposition. Via the bracketed attribute *simp* we also tell Isabelle to make the eventual theorem a **simplification rule**: future proofs involving simplification will replace occurrences of $rev\ (rev\ xs)$ by xs . The proof is by induction:

apply(*induction xs*)

As explained above, we obtain two subgoals, namely the base case (*Nil*) and the induction step (*Cons*):

1. $rev\ (rev\ Nil) = Nil$
2. $\bigwedge x\ xs.$
 $rev\ (rev\ xs) = xs \implies rev\ (rev\ (Cons\ x\ xs)) = Cons\ x\ xs$

Let us try to solve both goals automatically:

apply(*auto*)

Subgoal 1 is proved, and disappears; the simplified version of subgoal 2 becomes the new subgoal 1:

1. $\bigwedge x\ xs.$
 $rev\ (rev\ xs) = xs \implies$
 $rev\ (app\ (rev\ xs)\ (Cons\ x\ Nil)) = Cons\ x\ xs$

In order to simplify this subgoal further, a lemma suggests itself.

A First Lemma

We insert the following lemma in front of the main theorem:

lemma *rev_app* [*simp*]: " $rev(app\ xs\ ys) = app\ (rev\ ys)\ (rev\ xs)$ "

There are two variables that we could induct on: *xs* and *ys*. Because *app* is defined by recursion on the first argument, *xs* is the correct one:

apply(*induction xs*)

This time not even the base case is solved automatically:

apply(*auto*)

1. $rev\ ys = app\ (rev\ ys)\ Nil$
A total of 2 subgoals...

Again, we need to abandon this proof attempt and prove another simple lemma first.

A Second Lemma

We again try the canonical proof procedure:

```
lemma app_Nil2 [simp]: "app xs Nil = xs"
apply(induction xs)
apply(auto)
done
```

Thankfully, this worked. Now we can continue with our stuck proof attempt of the first lemma:

```
lemma rev_app [simp]: "rev(app xs ys) = app (rev ys) (rev xs)"
apply(induction xs)
apply(auto)
```

We find that this time *auto* solves the base case, but the induction step merely simplifies to

1. $\bigwedge x\ l\ xs.$
 $rev\ (app\ xs\ ys) = app\ (rev\ ys)\ (rev\ xs) \implies$
 $app\ (app\ (rev\ ys)\ (rev\ xs))\ (Cons\ x\ l\ Nil) =$
 $app\ (rev\ ys)\ (app\ (rev\ xs)\ (Cons\ x\ l\ Nil))$

The missing lemma is associativity of *app*, which we insert in front of the failed lemma *rev_app*.

Associativity of *app*

The canonical proof procedure succeeds without further ado:

```
lemma app_assoc [simp]: "app (app xs ys) zs = app xs (app ys zs)"
apply(induction xs)
apply(auto)
done
```

Finally the proofs of *rev_app* and *rev_rev* succeed, too.

Another Informal Proof

Here is the informal proof of associativity of *app* corresponding to the Isabelle proof above.

Lemma $app\ (app\ xs\ ys)\ zs = app\ xs\ (app\ ys\ zs)$

Proof by induction on *xs*.

- Case *Nil*: $app\ (app\ Nil\ ys)\ zs = app\ ys\ zs = app\ Nil\ (app\ ys\ zs)$ holds by definition of *app*.
- Case *Cons x xs*: We assume

$$app\ (app\ xs\ ys)\ zs = app\ xs\ (app\ ys\ zs) \quad (IH)$$

and we need to show

$$app\ (app\ (Cons\ x\ xs)\ ys)\ zs = app\ (Cons\ x\ xs)\ (app\ ys\ zs).$$

The proof is as follows:

$$\begin{aligned} & app\ (app\ (Cons\ x\ xs)\ ys)\ zs \\ &= app\ (Cons\ x\ (app\ xs\ ys))\ zs \quad \text{by definition of } app \\ &= Cons\ x\ (app\ (app\ xs\ ys)\ zs) \quad \text{by definition of } app \\ &= Cons\ x\ (app\ xs\ (app\ ys\ zs)) \quad \text{by IH} \\ &= app\ (Cons\ x\ xs)\ (app\ ys\ zs) \quad \text{by definition of } app \end{aligned}$$

Didn't we say earlier that all proofs are by simplification? But in both cases, going from left to right, the last equality step is not a simplification at all! In the base case it is $app\ ys\ zs = app\ Nil\ (app\ ys\ zs)$. It appears almost mysterious because we suddenly complicate the term by appending *Nil* on the left. What is really going on is this: when proving some equality $s = t$, both s and t are simplified until they "meet in the middle". This heuristic for equality proofs works well for a functional programming context like ours. In the base case both $app\ (app\ Nil\ ys)\ zs$ and $app\ Nil\ (app\ ys\ zs)$ are simplified to $app\ ys\ zs$, the term in the middle.

2.2.5 Predefined Lists

Isabelle's predefined lists are the same as the ones above, but with more syntactic sugar:

- $[]$ is *Nil*,
- $x \# xs$ is *Cons x xs*,
- $[x_1, \dots, x_n]$ is $x_1 \# \dots \# x_n \# []$, and
- $xs @ ys$ is $app\ xs\ ys$.

There is also a large library of predefined functions. The most important ones are the length function $length :: 'a\ list \Rightarrow nat$ (with the obvious definition), and the *map* function that applies a function to all elements of a list:

```
fun map :: "('a  $\Rightarrow$  'b)  $\Rightarrow$  'a list  $\Rightarrow$  'b list" where
  "map f Nil = Nil" |
  "map f (Cons x xs) = Cons (f x) (map f xs)"
```

From now on lists are always the predefined lists.

Exercises

Exercise 2.1. Use the **value** command to evaluate the following expressions: $"1 + (2::nat)"$, $"1 + (2::int)"$, $"1 - (2::nat)"$ and $"1 - (2::int)"$.

Exercise 2.2. Start from the definition of *add* given above. Prove that *add* is associative and commutative. Define a recursive function $double :: nat \Rightarrow nat$ and prove $double\ m = add\ m\ m$.

Exercise 2.3. Define a function $count :: 'a \Rightarrow 'a\ list \Rightarrow nat$ that counts the number of occurrences of an element in a list. Prove $count\ x\ xs \leq length\ xs$.

Exercise 2.4. Define a recursive function $snoc :: 'a\ list \Rightarrow 'a \Rightarrow 'a\ list$ that appends an element to the end of a list. With the help of *snoc* define a recursive function $reverse :: 'a\ list \Rightarrow 'a\ list$ that reverses a list. Prove $reverse\ (reverse\ xs) = xs$.

Exercise 2.5. Define a recursive function $sum :: nat \Rightarrow nat$ such that $sum\ n = 0 + \dots + n$ and prove $sum\ n = n * (n + 1) \div 2$.

2.3 Type and Function Definitions

Type synonyms are abbreviations for existing types, for example

```
type_synonym string = "char list"
```

Type synonyms are expanded after parsing and are not present in internal representation and output. They are mere conveniences for the reader.

2.3.1 Datatypes

The general form of a datatype definition looks like this:

```
datatype ('a1, ..., 'an) t = C1 "τ1,1" ... "τ1,n1"
                        | ...
                        | Ck "τk,1" ... "τk,nk"
```

It introduces the constructors $C_i :: \tau_{i,1} \Rightarrow \dots \Rightarrow \tau_{i,n_i} \Rightarrow ('a_1, \dots, 'a_n)t$ and expresses that any value of this type is built from these constructors in a unique manner. Uniqueness is implied by the following properties of the constructors:

- *Distinctness*: $C_i \dots \neq C_j \dots$ if $i \neq j$
- *Injectivity*: $(C_i x_1 \dots x_{n_i} = C_i y_1 \dots y_{n_i}) = (x_1 = y_1 \wedge \dots \wedge x_{n_i} = y_{n_i})$

The fact that any value of the datatype is built from the constructors implies the **structural induction** rule: to show $P\ x$ for all x of type $('a_1, \dots, 'a_n)t$, one needs to show $P(C_i x_1 \dots x_{n_i})$ (for each i) assuming $P(x_j)$ for all j where $\tau_{i,j} = ('a_1, \dots, 'a_n)t$. Distinctness and injectivity are applied automatically by *auto* and other proof methods. Induction must be applied explicitly.

Like in functional programming languages, datatype values can be taken apart with case expressions, for example

$(\text{case } xs \text{ of } [] \Rightarrow 0 \mid x \# _ \Rightarrow \text{Suc } x)$

Case expressions must be enclosed in parentheses.

As an example of a datatype beyond *nat* and *list*, consider binary trees:

datatype *'a tree* = *Tip* | *Node* "'a tree" 'a "'a tree"

with a mirror function:

```
fun mirror :: "'a tree  $\Rightarrow$  'a tree" where
  "mirror Tip = Tip" |
  "mirror (Node l a r) = Node (mirror r) a (mirror l)"
```

The following lemma illustrates induction:

```
lemma "mirror(mirror t) = t"
apply(induction t)
```

yields

1. *mirror* (*mirror Tip*) = *Tip*
2. $\bigwedge t1\ x2\ t2.$

$$\llbracket \text{mirror} (\text{mirror } t1) = t1; \text{mirror} (\text{mirror } t2) = t2 \rrbracket$$

$$\implies \text{mirror} (\text{mirror} (\text{Node } t1\ x2\ t2)) = \text{Node } t1\ x2\ t2$$

The induction step contains two induction hypotheses, one for each subtree. An application of *auto* finishes the proof.

A very simple but also very useful datatype is the predefined

datatype *'a option* = *None* | *Some* 'a

Its sole purpose is to add a new element *None* to an existing type *'a*. To make sure that *None* is distinct from all the elements of *'a*, you wrap them

up in *Some* and call the new type *'a option*. A typical application is a lookup function on a list of key-value pairs, often called an association list:

```
fun lookup :: "('a * 'b) list ⇒ 'a ⇒ 'b option" where
  "lookup [] x = None" |
  "lookup ((a,b) # ps) x = (if a = x then Some b else lookup ps x)"
```

Note that $\tau_1 * \tau_2$ is the type of pairs, also written $\tau_1 \times \tau_2$. Pairs can be taken apart either by pattern matching (as above) or with the projection functions *fst* and *snd*: *fst* $(x, y) = x$ and *snd* $(x, y) = y$. Tuples are simulated by pairs nested to the right: (a, b, c) is short for $(a, (b, c))$ and $\tau_1 \times \tau_2 \times \tau_3$ is short for $\tau_1 \times (\tau_2 \times \tau_3)$.

2.3.2 Definitions

Non-recursive functions can be defined as in the following example:

```
definition sq :: "nat ⇒ nat" where
  "sq n = n * n"
```

Such definitions do not allow pattern matching but only $f x_1 \dots x_n = t$, where f does not occur in t .

2.3.3 Abbreviations

Abbreviations are similar to definitions:

```
abbreviation sq' :: "nat ⇒ nat" where
  "sq' n ≡ n * n"
```

The key difference is that *sq'* is only syntactic sugar: after parsing, *sq' t* is replaced by $t * t$; before printing, every occurrence of $u * u$ is replaced by *sq' u*. Internally, *sq'* does not exist. This is the advantage of abbreviations over definitions: definitions need to be expanded explicitly (Section 2.5.5) whereas abbreviations are already expanded upon parsing. However, abbreviations should be introduced sparingly: if abused, they can lead to a confusing discrepancy between the internal and external view of a term.

The ASCII representation of $\equiv$ is `==` or `\<equiv>`.

2.3.4 Recursive Functions

Recursive functions are defined with **fun** by pattern matching over datatype constructors. The order of equations matters, as in functional programming languages. However, all HOL functions must be total. This simplifies the logic

— terms are always defined — but means that recursive functions must terminate. Otherwise one could define a function $f\ n = f\ n + 1$ and conclude $0 = 1$ by subtracting $f\ n$ on both sides.

Isabelle's automatic termination checker requires that the arguments of recursive calls on the right-hand side must be strictly smaller than the arguments on the left-hand side. In the simplest case, this means that one fixed argument position decreases in size with each recursive call. The size is measured as the number of constructors (excluding 0-ary ones, e.g., *Nil*). Lexicographic combinations are also recognized. In more complicated situations, the user may have to prove termination by hand. For details see [2].

Functions defined with **fun** come with their own induction schema that mirrors the recursion schema and is derived from the termination order. For example,

```
fun div2 :: "nat ⇒ nat" where
  "div2 0 = 0" |
  "div2 (Suc 0) = 0" |
  "div2 (Suc (Suc n)) = Suc (div2 n)"
```

does not just define *div2* but also proves a customized induction rule:

$$\frac{P\ 0 \quad P\ (Suc\ 0) \quad \bigwedge n. P\ n \implies P\ (Suc\ (Suc\ n))}{P\ m}$$

This customized induction rule can simplify inductive proofs. For example,

```
lemma "div2(n) = n div 2"
apply(induction n rule: div2.induct)
```

(where the infix *div* is the predefined division operation) yields the subgoals

1. *div2* 0 = 0 *div* 2
2. *div2* (Suc 0) = Suc 0 *div* 2
3. $\bigwedge n. \text{div2 } n = n \text{ div } 2 \implies$
 $\text{div2 } (Suc\ (Suc\ n)) = Suc\ (Suc\ n) \text{ div } 2$

An application of *auto* finishes the proof. Had we used ordinary structural induction on *n*, the proof would have needed an additional case analysis in the induction step.

This example leads to the following induction heuristic:

Let f be a recursive function. If the definition of f is more complicated than having one equation for each constructor of some datatype, then properties of f are best proved via $f.induct$.

The general case is often called **computation induction**, because the induction follows the (terminating!) computation. For every defining equation

$$f(e) = \dots f(r_1) \dots f(r_k) \dots$$

where $f(r_i)$, $i=1\dots k$, are all the recursive calls, the induction rule $f.induct$ contains one premise of the form

$$P(r_1) \implies \dots \implies P(r_k) \implies P(e)$$

If $f :: \tau_1 \Rightarrow \dots \Rightarrow \tau_n \Rightarrow \tau$ then $f.induct$ is applied like this:

`apply(induction  $x_1 \dots x_n$  rule:  $f.induct$ )`

where typically there is a call $f\ x_1 \dots x_n$ in the goal. But note that the induction rule does not mention f at all, except in its name, and is applicable independently of f .

Exercises

Exercise 2.6. Starting from the type `'a tree` defined in the text, define a function `contents :: 'a tree  $\Rightarrow$  'a list` that collects all values in a tree in a list, in any order, without removing duplicates. Then define a function `treesum :: nat tree  $\Rightarrow$  nat` that sums up all values in a tree of natural numbers and prove `treesum t = listsum (contents t)`.

Exercise 2.7. Define a new type `'a tree2` of binary trees where values are also stored in the leaves of the tree. Also reformulate the `mirror` function accordingly. Define two functions `pre_order` and `post_order` of type `'a tree2  $\Rightarrow$  'a list` that traverse a tree and collect all stored values in the respective order in a list. Prove `pre_order (mirror t) = rev (post_order t)`.

Exercise 2.8. Define a function `intersperse :: 'a  $\Rightarrow$  'a list  $\Rightarrow$  'a list` such that `intersperse a [x1, ..., xn] = [x1, a, x2, a, ..., a, xn]`. Now prove that `map f (intersperse a xs) = intersperse (f a) (map f xs)`.

2.4 Induction Heuristics

We have already noted that theorems about recursive functions are proved by induction. In case the function has more than one argument, we have followed the following heuristic in the proofs about the `append` function:

*Perform induction on argument number i
if the function is defined by recursion on argument number i.*

The key heuristic, and the main point of this section, is to *generalize the goal before induction*. The reason is simple: if the goal is too specific, the induction hypothesis is too weak to allow the induction step to go through. Let us illustrate the idea with an example.

Function *rev* has quadratic worst-case running time because it calls *append* for each element of the list and *append* is linear in its first argument. A linear time version of *rev* requires an extra argument where the result is accumulated gradually, using only *#*:

```
fun itrev :: "'a list ⇒ 'a list ⇒ 'a list" where
  "itrev [] ys = ys" |
  "itrev (x#xs) ys = itrev xs (x#ys)"
```

The behaviour of *itrev* is simple: it reverses its first argument by stacking its elements onto the second argument, and it returns that second argument when the first one becomes empty. Note that *itrev* is tail-recursive: it can be compiled into a loop; no stack is necessary for executing it.

Naturally, we would like to show that *itrev* does indeed reverse its first argument provided the second one is empty:

```
lemma "itrev xs [] = rev xs"
```

There is no choice as to the induction variable:

```
apply(induction xs)
apply(auto)
```

Unfortunately, this attempt does not prove the induction step:

1. $\bigwedge a \text{ xs. } \text{itrev xs []} = \text{rev xs} \implies \text{itrev xs [a]} = \text{rev xs} @ [a]$

The induction hypothesis is too weak. The fixed argument, *[]*, prevents it from rewriting the conclusion. This example suggests a heuristic:

Generalize goals for induction by replacing constants by variables.

Of course one cannot do this naively: *itrev xs ys = rev xs* is just not true. The correct generalization is

```
lemma "itrev xs ys = rev xs @ ys"
```

If *ys* is replaced by *[]*, the right-hand side simplifies to *rev xs*, as required. In this instance it was easy to guess the right generalization. Other situations can require a good deal of creativity.

Although we now have two variables, only *xs* is suitable for induction, and we repeat our proof attempt. Unfortunately, we are still not there:

1. $\bigwedge a \text{ xs. } \text{itrev xs ys} = \text{rev xs} @ \text{ys} \implies$
 $\text{itrev xs (a \# ys)} = \text{rev xs} @ \text{a \# ys}$

The induction hypothesis is still too weak, but this time it takes no intuition to generalize: the problem is that the *ys* in the induction hypothesis is fixed,

but the induction hypothesis needs to be applied with $a \# ys$ instead of ys . Hence we prove the theorem for all ys instead of a fixed one. We can instruct induction to perform this generalization for us by adding *arbitrary*: ys .

```
apply(induction xs arbitrary: ys)
```

The induction hypothesis in the induction step is now universally quantified over ys :

1. $\bigwedge ys. \text{itrev } [] \text{ } ys = \text{rev } [] \text{ } @ \text{ } ys$
2. $\bigwedge a \text{ } xs \text{ } ys.$
 $(\bigwedge ys. \text{itrev } xs \text{ } ys = \text{rev } xs \text{ } @ \text{ } ys) \implies$
 $\text{itrev } (a \# xs) \text{ } ys = \text{rev } (a \# xs) \text{ } @ \text{ } ys$

Thus the proof succeeds:

```
apply auto
done
```

This leads to another heuristic for generalization:

Generalize induction by generalizing all free variables
(except the induction variable itself).

Generalization is best performed with *arbitrary*: $y_1 \dots y_k$. This heuristic prevents trivial failures like the one above. However, it should not be applied blindly. It is not always required, and the additional quantifiers can complicate matters in some cases. The variables that need to be quantified are typically those that change in recursive calls.

Exercises

Exercise 2.9. Write a tail-recursive variant of the *add* function on *nat*: *itadd*. Tail-recursive means that in the recursive case, *itadd* needs to call itself directly: $\text{itadd } (\text{Suc } m) \text{ } n = \text{itadd } \dots$. Prove $\text{itadd } m \text{ } n = \text{add } m \text{ } n$.

2.5 Simplification

So far we have talked a lot about simplifying terms without explaining the concept. **Simplification** means

- using equations $l = r$ from left to right (only),
- as long as possible.

To emphasize the directionality, equations that have been given the *simp* attribute are called **simplification rules**. Logically, they are still symmetric,

but proofs by simplification use them only in the left-to-right direction. The proof tool that performs simplifications is called the **simplifier**. It is the basis of *auto* and other related proof methods.

The idea of simplification is best explained by an example. Given the simplification rules

$$0 + n = n \quad (1)$$

$$\text{Suc } m + n = \text{Suc } (m + n) \quad (2)$$

$$(\text{Suc } m \leq \text{Suc } n) = (m \leq n) \quad (3)$$

$$(0 \leq m) = \text{True} \quad (4)$$

the formula $0 + \text{Suc } 0 \leq \text{Suc } 0 + x$ is simplified to *True* as follows:

$$(0 + \text{Suc } 0 \leq \text{Suc } 0 + x) \quad \stackrel{(1)}{=}$$

$$(\text{Suc } 0 \leq \text{Suc } 0 + x) \quad \stackrel{(2)}{=}$$

$$(\text{Suc } 0 \leq \text{Suc } (0 + x)) \quad \stackrel{(3)}{=}$$

$$(0 \leq 0 + x) \quad \stackrel{(4)}{=}$$

True

Simplification is often also called **rewriting** and simplification rules **rewrite rules**.

2.5.1 Simplification Rules

The attribute *simp* declares theorems to be simplification rules, which the simplifier will use automatically. In addition, **datatype** and **fun** commands implicitly declare some simplification rules: **datatype** the distinctness and injectivity rules, **fun** the defining equations. Definitions are not declared as simplification rules automatically! Nearly any theorem can become a simplification rule. The simplifier will try to transform it into an equation. For example, the theorem $\neg P$ is turned into $P = \text{False}$.

Only equations that really simplify, like $\text{rev } (\text{rev } xs) = xs$ and $xs @ [] = xs$, should be declared as simplification rules. Equations that may be counterproductive as simplification rules should only be used in specific proof steps (see [Section 2.5.4](#) below). Distributivity laws, for example, alter the structure of terms and can produce an exponential blow-up.

2.5.2 Conditional Simplification Rules

Simplification rules can be conditional. Before applying such a rule, the simplifier will first try to prove the preconditions, again by simplification. For example, given the simplification rules

$$p\ 0 = \text{True}$$

$$p\ x \implies f\ x = g\ x,$$

the term $f\ 0$ simplifies to $g\ 0$ but $f\ 1$ does not simplify because $p\ 1$ is not provable.

2.5.3 Termination

Simplification can run forever, for example if both $f\ x = g\ x$ and $g\ x = f\ x$ are simplification rules. It is the user's responsibility not to include simplification rules that can lead to nontermination, either on their own or in combination with other simplification rules. The right-hand side of a simplification rule should always be "simpler" than the left-hand side — in some sense. But since termination is undecidable, such a check cannot be automated completely and Isabelle makes little attempt to detect nontermination.

When conditional simplification rules are applied, their preconditions are proved first. Hence all preconditions need to be simpler than the left-hand side of the conclusion. For example

$$n < m \implies (n < \text{Suc}\ m) = \text{True}$$

is suitable as a simplification rule: both $n < m$ and True are simpler than $n < \text{Suc}\ m$. But

$$\text{Suc}\ n < m \implies (n < m) = \text{True}$$

leads to nontermination: when trying to rewrite $n < m$ to True one first has to prove $\text{Suc}\ n < m$, which can be rewritten to True provided $\text{Suc}\ (\text{Suc}\ n) < m$, *ad infinitum*.

2.5.4 The *simp* Proof Method

So far we have only used the proof method *auto*. Method *simp* is the key component of *auto*, but *auto* can do much more. In some cases, *auto* is overeager and modifies the proof state too much. In such cases the more predictable *simp* method should be used. Given a goal

$$1. \llbracket P_1; \dots; P_m \rrbracket \implies C$$

the command

$$\text{apply}(\text{simp add: } th_1 \dots th_n)$$

simplifies the assumptions P_i and the conclusion C using

- all simplification rules, including the ones coming from **datatype** and **fun**,
- the additional lemmas $th_1 \dots th_n$, and

- the assumptions.

In addition to or instead of *add* there is also *del* for removing simplification rules temporarily. Both are optional. Method *auto* can be modified similarly:

`apply(auto simp add: ... simp del: ...)`

Here the modifiers are *simp add* and *simp del* instead of just *add* and *del* because *auto* does not just perform simplification.

Note that *simp* acts only on subgoal 1, *auto* acts on all subgoals. There is also *simp_all*, which applies *simp* to all subgoals.

2.5.5 Rewriting with Definitions

Definitions introduced by the command **definition** can also be used as simplification rules, but by default they are not: the simplifier does not expand them automatically. Definitions are intended for introducing abstract concepts and not merely as abbreviations. Of course, we need to expand the definition initially, but once we have proved enough abstract properties of the new constant, we can forget its original definition. This style makes proofs more robust: if the definition has to be changed, only the proofs of the abstract properties will be affected.

The definition of a function *f* is a theorem named *f_def* and can be added to a call of *simp* like any other theorem:

`apply(simp add: f_def)`

In particular, let-expressions can be unfolded by making *Let_def* a simplification rule.

2.5.6 Case Splitting With *simp*

Goals containing if-expressions are automatically split into two cases by *simp* using the rule

$$P \text{ (if } A \text{ then } s \text{ else } t) = ((A \longrightarrow P \ s) \wedge (\neg A \longrightarrow P \ t))$$

For example, *simp* can prove

$$(A \wedge B) = (\text{if } A \text{ then } B \text{ else } \text{False})$$

because both $A \longrightarrow (A \wedge B) = B$ and $\neg A \longrightarrow (A \wedge B) = \text{False}$ simplify to *True*.

We can split case-expressions similarly. For *nat* the rule looks like this:

$$P \text{ (case } e \text{ of } 0 \Rightarrow a \mid \text{Suc } n \Rightarrow b \ n) = \\ ((e = 0 \longrightarrow P \ a) \wedge (\forall n. e = \text{Suc } n \longrightarrow P \ (b \ n)))$$

Case expressions are not split automatically by *simp*, but *simp* can be instructed to do so:

```
apply(simp split: nat.split)
```

splits all case-expressions over natural numbers. For an arbitrary datatype *t* it is *t.split* instead of *nat.split*. Method *auto* can be modified in exactly the same way. The modifier *split:* can be followed by multiple names. Splitting if or case-expressions in the assumptions requires *split: if_splits* or *split: t.splits*.

Exercises

Exercise 2.10. Define a datatype *tree0* of binary tree skeletons which do not store any information, neither in the inner nodes nor in the leaves. Define a function *nodes* :: *tree0* $\Rightarrow$ *nat* that counts the number of all nodes (inner nodes and leaves) in such a tree. Consider the following recursive function:

```
fun explode :: "nat  $\Rightarrow$  tree0  $\Rightarrow$  tree0" where
  "explode 0 t = t" |
  "explode (Suc n) t = explode n (Node t t)"
```

Find an equation expressing the size of a tree after exploding it (*nodes* (*explode* *n* *t*)) as a function of *nodes* *t* and *n*. Prove your equation. You may use the usual arithmetic operators, including the exponentiation operator “ $\wedge$ ”. For example, $2 \wedge 2 = 4$.

Hint: simplifying with the list of theorems *algebra_simps* takes care of common algebraic properties of the arithmetic operators.

Exercise 2.11. Define arithmetic expressions in one variable over integers (type *int*) as a data type:

```
datatype exp = Var | Const int | Add exp exp | Mult exp exp
```

Define a function *eval* :: *exp* $\Rightarrow$ *int* $\Rightarrow$ *int* such that *eval* *e* *x* evaluates *e* at the value *x*.

A polynomial can be represented as a list of coefficients, starting with the constant. For example, $[4, 2, -1, 3]$ represents the polynomial $4 + 2x - x^2 + 3x^3$. Define a function *evalp* :: *int* *list* $\Rightarrow$ *int* $\Rightarrow$ *int* that evaluates a polynomial at the given value. Define a function *coeffs* :: *exp* $\Rightarrow$ *int* *list* that transforms an expression into a polynomial. This may require auxiliary functions. Prove that *coeffs* preserves the value of the expression: *evalp* (*coeffs* *e*) *x* = *eval* *e* *x*. Hint: consider the hint in Exercise 2.10.

Logic and Proof Beyond Equality

3.1 Formulas

The core syntax of formulas (*form* below) provides the standard logical constructs, in decreasing order of precedence:

$$\begin{aligned} \text{form} ::= & (\text{form}) \mid \text{True} \mid \text{False} \mid \text{term} = \text{term} \\ & \mid \neg \text{form} \mid \text{form} \wedge \text{form} \mid \text{form} \vee \text{form} \mid \text{form} \longrightarrow \text{form} \\ & \mid \forall x. \text{form} \mid \exists x. \text{form} \end{aligned}$$

Terms are the ones we have seen all along, built from constants, variables, function application and λ -abstraction, including all the syntactic sugar like infix symbols, *if*, *case*, etc.

! Remember that formulas are simply terms of type *bool*. Hence $=$ also works for formulas. Beware that $=$ has a higher precedence than the other logical operators. Hence $s = t \wedge A$ means $(s = t) \wedge A$, and $A \wedge B = B \wedge A$ means $A \wedge (B = B) \wedge A$. Logical equivalence can also be written with $\longleftrightarrow$ instead of $=$, where $\longleftrightarrow$ has the same low precedence as $\longrightarrow$. Hence $A \wedge B \longleftrightarrow B \wedge A$ really means $(A \wedge B) \longleftrightarrow (B \wedge A)$.

! Quantifiers need to be enclosed in parentheses if they are nested within other constructs (just like *if*, *case* and *let*).

The most frequent logical symbols and their ASCII representations are shown in Fig. 3.1. The first column shows the symbols, the other columns ASCII representations. The $\langle \dots \rangle$ form is always converted into the symbolic form by the Isabelle interfaces, the treatment of the other ASCII forms depends on the interface. The ASCII forms $/\wedge$ and $\backslash/$ are special in that they are merely keyboard shortcuts for the interface and not logical symbols by themselves.

$\forall$	<code>\<forall></code>	ALL
$\exists$	<code>\<exists></code>	EX
λ	<code>\<lambda></code>	%
$\longrightarrow$	<code>--></code>	
$\longleftrightarrow$	<code><-></code>	
$\wedge$	<code>\&</code>	&
$\vee$	<code>\ </code>	
$\neg$	<code>\<not></code>	~
$\neq$	<code>\<noteq></code>	~=

Fig. 3.1. Logical symbols and their ASCII forms

! The implication $\implies$ is part of the Isabelle framework. It structures theorems and proof states, separating assumptions from conclusions. The implication $\longrightarrow$ is part of the logic HOL and can occur inside the formulas that make up the assumptions and conclusion. Theorems should be of the form $\llbracket A_1; \dots; A_n \rrbracket \implies A$, not $A_1 \wedge \dots \wedge A_n \longrightarrow A$. Both are logically equivalent but the first one works better when using the theorem in further proofs.

3.2 Sets

Sets of elements of type *'a* have type *'a set*. They can be finite or infinite. Sets come with the usual notation:

- $\{\}, \{e_1, \dots, e_n\}$
- $e \in A, A \subseteq B$
- $A \cup B, A \cap B, A - B, -A$

(where $A - B$ and $-A$ are set difference and complement) and much more. *UNIV* is the set of all elements of some type. Set comprehension is written $\{x. P\}$ rather than $\{x \mid P\}$.

! In $\{x. P\}$ the x must be a variable. Set comprehension involving a proper term t must be written $\{t \mid x y. P\}$, where $x y$ are those free variables in t that occur in P . This is just a shorthand for $\{v. \exists x y. v = t \wedge P\}$, where v is a new variable. For example, $\{x + y \mid x. x \in A\}$ is short for $\{v. \exists x. v = x + y \wedge x \in A\}$.

Here are the ASCII representations of the mathematical symbols:

$\in$	<code>\<in></code>	:
$\subseteq$	<code>\<subsetq></code>	<code><=</code>
$\cup$	<code>\<union></code>	Un
$\cap$	<code>\<inter></code>	Int

Sets also allow bounded quantifications $\forall x \in A. P$ and $\exists x \in A. P$.

For the more ambitious, there are also $\bigcup$ and $\bigcap$:

$$\bigcup A = \{x. \exists B \in A. x \in B\} \quad \bigcap A = \{x. \forall B \in A. x \in B\}$$

The ASCII forms of $\bigcup$ are `\<Union>` and `Union`, those of $\bigcap$ are `\<Inter>` and `Inter`. There are also indexed unions and intersections:

$$\begin{aligned} (\bigcup_{x \in A} B \ x) &= \{y. \exists x \in A. y \in B \ x\} \\ (\bigcap_{x \in A} B \ x) &= \{y. \forall x \in A. y \in B \ x\} \end{aligned}$$

The ASCII forms are `UN x:A. B` and `INT x:A. B` where x may occur in B . If A is `UNIV` you can write `UN x. B` and `INT x. B`.

Some other frequently useful functions on sets are the following:

<code>set :: 'a list $\Rightarrow$ 'a set</code>	converts a list to the set of its elements
<code>finite :: 'a set $\Rightarrow$ bool</code>	is true iff its argument is finite
<code>card :: 'b set $\Rightarrow$ nat</code>	is the cardinality of a finite set and is 0 for all infinite sets
<code>f ' A = {y. $\exists x \in A. y = f \ x$}</code>	is the image of a function over a set

See [3] for the wealth of further predefined functions in theory *Main*.

Exercises

Exercise 3.1. Start from the data type of binary trees defined earlier:

```
datatype 'a tree = Tip | Node "'a tree" 'a "'a tree"
```

Define a function `set :: 'a tree  $\Rightarrow$  'a set` that returns the elements in a tree and a function `ord :: int tree  $\Rightarrow$  bool` that tests if an *int tree* is ordered.

Define a function `ins` that inserts an element into an ordered *int tree* while maintaining the order of the tree. If the element is already in the tree, the same tree should be returned. Prove correctness of `ins`: `set (ins x t) = {x}  $\cup$  set t` and `ord t  $\implies$  ord (ins i t)`.

3.3 Proof Automation

So far we have only seen *simp* and *auto*: Both perform rewriting, both can also prove linear arithmetic facts (no multiplication), and *auto* is also able to prove simple logical or set-theoretic goals:

```
lemma " $\forall x. \exists y. x = y$ "
by auto
```

```
lemma " $A \subseteq B \cap C \implies A \subseteq B \cup C$ "
by auto
```

where

by *proof-method*

is short for

apply *proof-method*
done

The key characteristics of both *simp* and *auto* are

- They show you where they got stuck, giving you an idea how to continue.
- They perform the obvious steps but are highly incomplete.

A proof method is **complete** if it can prove all true formulas. There is no complete proof method for HOL, not even in theory. Hence all our proof methods only differ in how incomplete they are.

A proof method that is still incomplete but tries harder than *auto* is *fastforce*. It either succeeds or fails, it acts on the first subgoal only, and it can be modified like *auto*, e.g., with *simp add*. Here is a typical example of what *fastforce* can do:

lemma "[$\forall x s \in A. \exists y s. xs = ys @ ys; us \in A$]
 $\implies \exists n. \text{length } us = n + n$ "
 by *fastforce*

This lemma is out of reach for *auto* because of the quantifiers. Even *fastforce* fails when the quantifier structure becomes more complicated. In a few cases, its slow version *force* succeeds where *fastforce* fails.

The method of choice for complex logical goals is *blast*. In the following example, *T* and *A* are two binary predicates. It is shown that if *T* is total, *A* is antisymmetric and *T* is a subset of *A*, then *A* is a subset of *T*:

lemma
 "[$\forall x y. T x y \vee T y x;$
 $\forall x y. A x y \wedge A y x \longrightarrow x = y;$
 $\forall x y. T x y \longrightarrow A x y$]
 $\implies \forall x y. A x y \longrightarrow T x y$ "
 by *blast*

We leave it to the reader to figure out why this lemma is true. Method *blast*

- is (in principle) a complete proof procedure for first-order formulas, a fragment of HOL. In practice there is a search bound.
- does no rewriting and knows very little about equality.
- covers logic, sets and relations.
- either succeeds or fails.

Because of its strength in logic and sets and its weakness in equality reasoning, it complements the earlier proof methods.

3.3.1 Sledgehammer

Command `sledgehammer` calls a number of external automatic theorem provers (ATPs) that run for up to 30 seconds searching for a proof. Some of these ATPs are part of the Isabelle installation, others are queried over the internet. If successful, a proof command is generated and can be inserted into your proof. The biggest win of `sledgehammer` is that it will take into account the whole lemma library and you do not need to feed in any lemma explicitly. For example,

```
lemma "[ xs @ ys = ys @ xs; length xs = length ys ] ==> xs = ys"
```

cannot be solved by any of the standard proof methods, but `sledgehammer` finds the following proof:

```
by (metis append_eq_conv_conj)
```

We do not explain how the proof was found but what this command means. For a start, Isabelle does not trust external tools (and in particular not the translations from Isabelle's logic to those tools!) and insists on a proof that it can check. This is what *metis* does. It is given a list of lemmas and tries to find a proof using just those lemmas (and pure logic). In contrast to using *simp* and friends who know a lot of lemmas already, using *metis* manually is tedious because one has to find all the relevant lemmas first. But that is precisely what `sledgehammer` does for us. In this case lemma *append_eq_conv_conj* alone suffices:

$$(xs @ ys = zs) = (xs = take (length xs) zs \wedge ys = drop (length xs) zs)$$

We leave it to the reader to figure out why this lemma suffices to prove the above lemma, even without any knowledge of what the functions *take* and *drop* do. Keep in mind that the variables in the two lemmas are independent of each other, despite the same names, and that you can substitute arbitrary values for the free variables in a lemma.

Just as for the other proof methods we have seen, there is no guarantee that `sledgehammer` will find a proof if it exists. Nor is `sledgehammer` superior to the other proof methods. They are incomparable. Therefore it is recommended to apply *simp* or *auto* before invoking `sledgehammer` on what is left.

3.3.2 Arithmetic

By arithmetic formulas we mean formulas involving variables, numbers, $+$, $-$, $=$, $<$, $\leq$ and the usual logical connectives $\neg$, $\wedge$, $\vee$, $\longrightarrow$, $\longleftrightarrow$. Strictly speaking, this is known as **linear arithmetic** because it does not involve multiplication, although multiplication with numbers, e.g., $2*n$, is allowed. Such formulas can be proved by *arith*:

lemma "[(a::nat) ≤ x + b; 2*x < c] ⇒ 2*a + 1 ≤ 2*b + c"
by *arith*

In fact, *auto* and *simp* can prove many linear arithmetic formulas already, like the one above, by calling a weak but fast version of *arith*. Hence it is usually not necessary to invoke *arith* explicitly.

The above example involves natural numbers, but integers (type *int*) and real numbers (type *real*) are supported as well. As are a number of further operators like *min* and *max*. On *nat* and *int*, *arith* can even prove theorems with quantifiers in them, but we will not enlarge on that here.

3.3.3 Trying Them All

If you want to try all of the above automatic proof methods you simply type
try

You can also add specific simplification and introduction rules:

try simp: ... **intro:** ...

There is also a lightweight variant **try0** that does not call sledgehammer.

3.4 Single Step Proofs

Although automation is nice, it often fails, at least initially, and you need to find out why. When *fastforce* or *blast* simply fail, you have no clue why. At this point, the stepwise application of proof rules may be necessary. For example, if *blast* fails on $A \wedge B$, you want to attack the two conjuncts A and B separately. This can be achieved by applying *conjunction introduction*

$$\frac{?P \quad ?Q}{?P \wedge ?Q} \text{ conjI}$$

to the proof state. We will now examine the details of this process.

3.4.1 Instantiating Unknowns

We had briefly mentioned earlier that after proving some theorem, Isabelle replaces all free variables x by so called **unknowns** $?x$. We can see this clearly in rule *conjI*. These unknowns can later be instantiated explicitly or implicitly:

- By hand, using *of*. The expression *conjI*[*of* " $a=b$ " "*False*"] instantiates the unknowns in *conjI* from left to right with the two formulas $a=b$ and *False*, yielding the rule

$$\frac{a = b \quad False}{a = b \wedge False}$$

In general, $th[of\ string_1 \dots string_n]$ instantiates the unknowns in the theorem th from left to right with the terms $string_1$ to $string_n$.

- By unification. **Unification** is the process of making two terms syntactically equal by suitable instantiations of unknowns. For example, unifying $?P \wedge ?Q$ with $a = b \wedge False$ instantiates $?P$ with $a = b$ and $?Q$ with $False$.

We need not instantiate all unknowns. If we want to skip a particular one we can write $_$ instead, for example $conjI[of\ _ "False"]$. Unknowns can also be instantiated by name using *where*, for example $conjI[where\ ?P = "a=b"$ and $?Q = "False"]$.

3.4.2 Rule Application

Rule application means applying a rule backwards to a proof state. For example, applying rule *conjI* to a proof state

$$1. \dots \Longrightarrow A \wedge B$$

results in two subgoals, one for each premise of *conjI*:

$$\begin{aligned} 1. \dots &\Longrightarrow A \\ 2. \dots &\Longrightarrow B \end{aligned}$$

In general, the application of a rule $\llbracket A_1; \dots; A_n \rrbracket \Longrightarrow A$ to a subgoal $\dots \Longrightarrow C$ proceeds in two steps:

1. Unify A and C , thus instantiating the unknowns in the rule.
2. Replace the subgoal C with n new subgoals A_1 to A_n .

This is the command to apply rule *xyz*:

apply(rule *xyz*)

This is also called **backchaining** with rule *xyz*.

3.4.3 Introduction Rules

Conjunction introduction (*conjI*) is one example of a whole class of rules known as **introduction rules**. They explain under which premises some logical construct can be introduced. Here are some further useful introduction rules:

$$\frac{?P \Longrightarrow ?Q}{?P \longrightarrow ?Q} \text{ impI} \qquad \frac{\bigwedge x. ?P\ x}{\forall x. ?P\ x} \text{ allI}$$

$$\frac{?P \Rightarrow ?Q \quad ?Q \Rightarrow ?P}{?P = ?Q} \text{ iffI}$$

These rules are part of the logical system of **natural deduction** (e.g., [1]). Although we intentionally de-emphasize the basic rules of logic in favour of automatic proof methods that allow you to take bigger steps, these rules are helpful in locating where and why automation fails. When applied backwards, these rules decompose the goal:

- *conjI* and *iffI* split the goal into two subgoals,
- *impI* moves the left-hand side of a HOL implication into the list of assumptions,
- and *allI* removes a $\forall$ by turning the quantified variable into a fixed local variable of the subgoal.

Isabelle knows about these and a number of other introduction rules. The command

apply rule

automatically selects the appropriate rule for the current subgoal.

You can also turn your own theorems into introduction rules by giving them the *intro* attribute, analogous to the *simp* attribute. In that case *blast*, *fastforce* and (to a limited extent) *auto* will automatically backchain with those theorems. The *intro* attribute should be used with care because it increases the search space and can lead to nontermination. Sometimes it is better to use it only in specific calls of *blast* and friends. For example, *le_trans*, transitivity of $\leq$ on type *nat*, is not an introduction rule by default because of the disastrous effect on the search space, but can be useful in specific situations:

```
lemma "[ (a::nat) ≤ b; b ≤ c; c ≤ d; d ≤ e ] ⇒ a ≤ e"
by(blast intro: le_trans)
```

Of course this is just an example and could be proved by *arith*, too.

3.4.4 Forward Proof

Forward proof means deriving new theorems from old theorems. We have already seen a very simple form of forward proof: the *of* operator for instantiating unknowns in a theorem. The big brother of *of* is *OF* for applying one theorem to others. Given a theorem $A \Rightarrow B$ called *r* and a theorem A' called *r'*, the theorem $r[OF\ r']$ is the result of applying *r* to *r'*, where *r* should be viewed as a function taking a theorem *A* and returning *B*. More precisely, *A* and *A'* are unified, thus instantiating the unknowns in *B*, and the result is the instantiated *B*. Of course, unification may also fail.

! Application of rules to other rules operates in the forward direction: from the premises to the conclusion of the rule; application of rules to proof states operates in the backward direction, from the conclusion to the premises.

In general r can be of the form $\llbracket A_1; \dots; A_n \rrbracket \Longrightarrow A$ and there can be multiple argument theorems r_1 to r_m (with $m \leq n$), in which case $r[OF\ r_1 \dots r_m]$ is obtained by unifying and thus proving A_i with r_i , $i = 1 \dots m$. Here is an example, where *refl* is the theorem $?t = ?t$:

```
thm conjI[OF refl[of "a"] refl[of "b"]]
```

yields the theorem $a = a \wedge b = b$. The command **thm** merely displays the result.

Forward reasoning also makes sense in connection with proof states. Therefore *blast*, *fastforce* and *auto* support a modifier *dest* which instructs the proof method to use certain rules in a forward fashion. If r is of the form $A \Longrightarrow B$, the modifier *dest: r* allows proof search to reason forward with r , i.e., to replace an assumption A' , where A' unifies with A , with the correspondingly instantiated B . For example, *Suc_leD* is the theorem $Suc\ m \leq n \Longrightarrow m \leq n$, which works well for forward reasoning:

```
lemma "Suc(Suc(Suc a)) ≤ b ⟹ a ≤ b"
by(blast dest: Suc_leD)
```

In this particular example we could have backchained with *Suc_leD*, too, but because the premise is more complicated than the conclusion this can easily lead to nontermination.

! To ease readability we will drop the question marks in front of unknowns from now on.

3.5 Inductive Definitions

Inductive definitions are the third important definition facility, after datatypes and recursive function.

3.5.1 An Example: Even Numbers

Here is a simple example of an inductively defined predicate:

- 0 is even
- If n is even, so is $n + 2$.

The operative word “inductive” means that these are the only even numbers. In Isabelle we give the two rules the names *ev0* and *evSS* and write

inductive $ev :: "nat \Rightarrow bool"$ **where**

$ev0:$ $"ev\ 0"$ |

$evSS:$ $"ev\ n \Longrightarrow ev\ (n + 2)"$

To get used to inductive definitions, we will first prove a few properties of ev informally before we descend to the Isabelle level.

How do we prove that some number is even, e.g., $ev\ 4$? Simply by combining the defining rules for ev :

$$ev\ 0 \Longrightarrow ev\ (0 + 2) \Longrightarrow ev\ ((0 + 2) + 2) = ev\ 4$$

Rule Induction

Showing that all even numbers have some property is more complicated. For example, let us prove that the inductive definition of even numbers agrees with the following recursive one:

fun $evn :: "nat \Rightarrow bool"$ **where**

$"evn\ 0 = True"$ |

$"evn\ (Suc\ 0) = False"$ |

$"evn\ (Suc\ (Suc\ n)) = evn\ n"$

We prove $ev\ m \Longrightarrow evn\ m$. That is, we assume $ev\ m$ and by induction on the form of its derivation prove $evn\ m$. There are two cases corresponding to the two rules for ev :

Case $ev0$: $ev\ m$ was derived by rule $ev\ 0$:

$$\Longrightarrow m = 0 \Longrightarrow evn\ m = evn\ 0 = True$$

Case $evSS$: $ev\ m$ was derived by rule $ev\ n \Longrightarrow ev\ (n + 2)$:

$$\Longrightarrow m = n + 2 \text{ and by induction hypothesis } evn\ n$$

$$\Longrightarrow evn\ m = evn\ (n + 2) = evn\ n = True$$

What we have just seen is a special case of **rule induction**. Rule induction applies to propositions of this form

$$ev\ n \Longrightarrow P\ n$$

That is, we want to prove a property $P\ n$ for all even n . But if we assume $ev\ n$, then there must be some derivation of this assumption using the two defining rules for ev . That is, we must prove

Case $ev0$: $P\ 0$

Case $evSS$: $\llbracket ev\ n; P\ n \rrbracket \Longrightarrow P\ (n + 2)$

The corresponding rule is called $ev.induct$ and looks like this:

$$\frac{ev\ n \quad P\ 0 \quad \bigwedge n. \llbracket ev\ n; P\ n \rrbracket \Longrightarrow P\ (n + 2)}{P\ n}$$

The first premise $ev\ n$ enforces that this rule can only be applied in situations where we know that n is even.

Note that in the induction step we may not just assume $P\ n$ but also $ev\ n$, which is simply the premise of rule $evSS$. Here is an example where the local assumption $ev\ n$ comes in handy: we prove $ev\ m \implies ev\ (m - 2)$ by induction on $ev\ m$. Case $ev0$ requires us to prove $ev\ (0 - 2)$, which follows from $ev\ 0$ because $0 - 2 = 0$ on type nat . In case $evSS$ we have $m = n + 2$ and may assume $ev\ n$, which implies $ev\ (m - 2)$ because $m - 2 = (n + 2) - 2 = n$. We did not need the induction hypothesis at all for this proof (it is just a case analysis of which rule was used) but having $ev\ n$ at our disposal in case $evSS$ was essential. This case analysis of rules is also called “rule inversion” and is discussed in more detail in [Chapter 4](#).

In Isabelle

Let us now recast the above informal proofs in Isabelle. For a start, we use Suc terms instead of numerals in rule $evSS$:

$$ev\ n \implies ev\ (Suc\ (Suc\ n))$$

This avoids the difficulty of unifying $n+2$ with some numeral, which is not automatic.

The simplest way to prove $ev\ (Suc\ (Suc\ (Suc\ (Suc\ 0))))$ is in a forward direction: $evSS[OF\ evSS[OF\ ev0]]$ yields the theorem $ev\ (Suc\ (Suc\ (Suc\ (Suc\ 0))))$. Alternatively, you can also prove it as a lemma in backwards fashion. Although this is more verbose, it allows us to demonstrate how each rule application changes the proof state:

lemma "ev(Suc(Suc(Suc(Suc 0)))"

1. ev (Suc (Suc (Suc (Suc 0))))

apply(rule evSS)

1. ev (Suc (Suc 0))

apply(rule evSS)

1. ev 0

apply(rule ev0)

done

Rule induction is applied by giving the induction rule explicitly via the *rule:* modifier:

```

lemma "ev m  $\implies$  evn m"
apply(induction rule: ev.induct)
by(simp_all)

```

Both cases are automatic. Note that if there are multiple assumptions of the form $ev\ t$, method *induction* will induct on the leftmost one.

As a bonus, we also prove the remaining direction of the equivalence of ev and evn :

```

lemma "evn n  $\implies$  ev n"
apply(induction n rule: evn.induct)

```

This is a proof by computation induction on n (see [Section 2.3.4](#)) that sets up three subgoals corresponding to the three equations for evn :

1. $evn\ 0 \implies ev\ 0$
2. $evn\ (Suc\ 0) \implies ev\ (Suc\ 0)$
3. $\bigwedge n. \llbracket evn\ n \implies ev\ n; evn\ (Suc\ (Suc\ n)) \rrbracket \implies ev\ (Suc\ (Suc\ n))$

The first and third subgoals follow with $ev0$ and $evSS$, and the second subgoal is trivially true because $evn\ (Suc\ 0)$ is *False*:

```

by (simp_all add: ev0 evSS)

```

The rules for ev make perfect simplification and introduction rules because their premises are always smaller than the conclusion. It makes sense to turn them into simplification and introduction rules permanently, to enhance proof automation. They are named *ev.intros* by Isabelle:

```

declare ev.intros[simp,intro]

```

The rules of an inductive definition are not simplification rules by default because, in contrast to recursive functions, there is no termination requirement for inductive definitions.

Inductive Versus Recursive

We have seen two definitions of the notion of evenness, an inductive and a recursive one. Which one is better? Much of the time, the recursive one is more convenient: it allows us to do rewriting in the middle of terms, and it expresses both the positive information (which numbers are even) and the negative information (which numbers are not even) directly. An inductive definition only expresses the positive information directly. The negative information, for example, that 1 is not even, has to be proved from it (by induction or rule inversion). On the other hand, rule induction is tailor-made for proving $ev\ n \implies P\ n$ because it only asks you to prove the positive cases. In the proof of $evn\ n \implies P\ n$ by computation induction via *evn.induct*, we are also

presented with the trivial negative cases. If you want the convenience of both rewriting and rule induction, you can make two definitions and show their equivalence (as above) or make one definition and prove additional properties from it, for example rule induction from computation induction.

But many concepts do not admit a recursive definition at all because there is no datatype for the recursion (for example, the transitive closure of a relation), or the recursion would not terminate (for example, an interpreter for a programming language). Even if there is a recursive definition, if we are only interested in the positive information, the inductive definition may be much simpler.

3.5.2 The Reflexive Transitive Closure

Evenness is really more conveniently expressed recursively than inductively. As a second and very typical example of an inductive definition we define the reflexive transitive closure.

The reflexive transitive closure, called *star* below, is a function that maps a binary predicate to another binary predicate: if r is of type $\tau \Rightarrow \tau \Rightarrow \text{bool}$ then $\text{star } r$ is again of type $\tau \Rightarrow \tau \Rightarrow \text{bool}$, and $\text{star } r \ x \ y$ means that x and y are in the relation $\text{star } r$. Think r^* when you see $\text{star } r$, because $\text{star } r$ is meant to be the reflexive transitive closure. That is, $\text{star } r \ x \ y$ is meant to be true if from x we can reach y in finitely many r steps. This concept is naturally defined inductively:

```
inductive star :: "('a  $\Rightarrow$  'a  $\Rightarrow$  bool)  $\Rightarrow$  'a  $\Rightarrow$  'a  $\Rightarrow$  bool" for r where
  refl: "star r x x" |
  step: "r x y  $\Longrightarrow$  star r y z  $\Longrightarrow$  star r x z"
```

The base case *refl* is reflexivity: $x = y$. The step case *step* combines an r step (from x to y) and a $\text{star } r$ step (from y to z) into a $\text{star } r$ step (from x to z). The “for r ” in the header is merely a hint to Isabelle that r is a fixed parameter of *star*, in contrast to the further parameters of *star*, which change. As a result, Isabelle generates a simpler induction rule.

By definition $\text{star } r$ is reflexive. It is also transitive, but we need rule induction to prove that:

```
lemma star_trans: "star r x y  $\Longrightarrow$  star r y z  $\Longrightarrow$  star r x z"
  apply(induction rule: star.induct)
```

The induction is over $\text{star } r \ x \ y$ (the first matching assumption) and we try to prove $\text{star } r \ y \ z \Longrightarrow \text{star } r \ x \ z$, which we abbreviate by $P \ x \ y$. These are our two subgoals:

1. $\bigwedge x. \text{star } r \ x \ z \Longrightarrow \text{star } r \ x \ z$

2. $\bigwedge u x y.$

$$\llbracket r u x; \text{star } r x y; \text{star } r y z \implies \text{star } r x z; \text{star } r y z \rrbracket$$

$$\implies \text{star } r u z$$

The first one is $P x x$, the result of case *refl*, and it is trivial:

`apply(assumption)`

Let us examine subgoal 2, case *step*. Assumptions $r u x$ and $\text{star } r x y$ are the premises of rule *step*. Assumption $\text{star } r y z \implies \text{star } r x z$ is $P x y$, the IH coming from $\text{star } r x y$. We have to prove $P u y$, which we do by assuming $\text{star } r y z$ and proving $\text{star } r u z$. The proof itself is straightforward: from $\text{star } r y z$ the IH leads to $\text{star } r x z$ which, together with $r u x$, leads to $\text{star } r u z$ via rule *step*:

`apply(metis step)`
`done`

3.5.3 The General Case

Inductive definitions have approximately the following general form:

inductive $I :: "\tau \Rightarrow \text{bool}"$ **where**

followed by a sequence of (possibly named) rules of the form

$$\llbracket I a_1; \dots; I a_n \rrbracket \implies I a$$

separated by $|$. As usual, n can be 0. The corresponding rule induction principle $I.\text{induct}$ applies to propositions of the form

$$I x \implies P x$$

where P may itself be a chain of implications.



Rule induction is always on the leftmost premise of the goal. Hence $I x$ must be the first premise.

Proving $I x \implies P x$ by rule induction means proving for every rule of I that P is invariant:

$$\llbracket I a_1; P a_1; \dots; I a_n; P a_n \rrbracket \implies P a$$

The above format for inductive definitions is simplified in a number of respects. I can have any number of arguments and each rule can have additional premises not involving I , so-called **side conditions**. In rule inductions, these side conditions appear as additional assumptions. The **for** clause seen in the definition of the reflexive transitive closure simplifies the induction rule.

Exercises

Exercise 3.2. Formalize the following definition of palindromes

- The empty list and a singleton list are palindromes.
- If xs is a palindrome, so is $a \# xs @ [a]$.

as an inductive predicate $palindrome :: 'a list \Rightarrow bool$ and prove that $rev\ xs = xs$ if xs is a palindrome.

Exercise 3.3. We could also have defined $star$ as follows:

```
inductive star' :: "('a \Rightarrow 'a \Rightarrow bool) \Rightarrow 'a \Rightarrow 'a \Rightarrow bool" for r where
  refl': "star' r x x" |
  step': "star' r x y \Longrightarrow r y z \Longrightarrow star' r x z"
```

The single r step is performed after rather than before the $star'$ steps. Prove $star' r x y \Longrightarrow star\ r\ x\ y$ and $star\ r\ x\ y \Longrightarrow star' r x y$. You may need lemmas. Note that rule induction fails if the assumption about the inductive predicate is not the first assumption.

Exercise 3.4. Analogous to $star$, give an inductive definition of the n -fold iteration of a relation r : $iter\ r\ n\ x\ y$ should hold if there are $x_0, \dots, x_n$ such that $x = x_0$, $x_n = y$ and $r\ x_i\ x_{i+1}$ for all $i < n$. Correct and prove the following claim: $star\ r\ x\ y \Longrightarrow iter\ r\ n\ x\ y$.

Exercise 3.5. A context-free grammar can be seen as an inductive definition where each nonterminal A is an inductively defined predicate on lists of terminal symbols: $A(w)$ means that w is in the language generated by A . For example, the production $S \rightarrow aSb$ can be viewed as the implication $S\ w \Longrightarrow S\ (a \# w @ [b])$ where a and b are terminal symbols, i.e., elements of some alphabet. The alphabet can be defined like this: **datatype** $alpha = a \mid b \mid \dots$

Define the two grammars (where ε is the empty word)

$$\begin{array}{lcl} S & \rightarrow & \varepsilon \mid aSb \mid SS \\ T & \rightarrow & \varepsilon \mid TaTb \end{array}$$

as two inductive predicates. If you think of a and b as “(” and “)”, the grammar defines strings of balanced parentheses. Prove $T\ w \Longrightarrow S\ w$ and $S\ w \Longrightarrow T\ w$ separately and conclude $S\ w = T\ w$.

Isar: A Language for Structured Proofs

Apply-scripts are unreadable and hard to maintain. The language of choice for larger proofs is **Isar**. The two key features of Isar are:

- It is structured, not linear.
- It is readable without its being run because you need to state what you are proving at any given point.

Whereas apply-scripts are like assembly language programs, Isar proofs are like structured programs with comments. A typical Isar proof looks like this:

```
proof
  assume "formula0"
  have "formula1" by simp
  ⋮
  have "formulan" by blast
  show "formulan+1" by ...
qed
```

It proves $formula_0 \implies formula_{n+1}$ (provided each proof step succeeds). The intermediate **have** statements are merely stepping stones on the way towards the **show** statement that proves the actual goal. In more detail, this is the Isar core syntax:

```
proof = by method
      | proof [method] step* qed

step = fix variables
      | assume proposition
      | [from fact+] (have | show) proposition proof

proposition = [name:] "formula"

fact = name | ...
```

A proof can either be an atomic **by** with a single proof method which must finish off the statement being proved, for example *auto*, or it can be a **proof-qed** block of multiple steps. Such a block can optionally begin with a proof method that indicates how to start off the proof, e.g., (*induction xs*).

A step either assumes a proposition or states a proposition together with its proof. The optional **from** clause indicates which facts are to be used in the proof. Intermediate propositions are stated with **have**, the overall goal is stated with **show**. A step can also introduce new local variables with **fix**. Logically, **fix** introduces $\bigwedge$ -quantified variables, **assume** introduces the assumption of an implication ($\implies$) and **have/show** introduce the conclusion.

Propositions are optionally named formulas. These names can be referred to in later **from** clauses. In the simplest case, a fact is such a name. But facts can also be composed with *OF* and *of* as shown in Section 3.4.4 — hence the ... in the above grammar. Note that assumptions, intermediate **have** statements and global lemmas all have the same status and are thus collectively referred to as **facts**.

Fact names can stand for whole lists of facts. For example, if *f* is defined by command **fun**, *f.simps* refers to the whole list of recursion equations defining *f*. Individual facts can be selected by writing *f.simps*(2), whole sublists by writing *f.simps*(2–4).

4.1 Isar by Example

We show a number of proofs of Cantor's theorem that a function from a set to its powerset cannot be surjective, illustrating various features of Isar. The constant *surj* is predefined.

```
lemma "¬ surj (f :: 'a ⇒ 'a set)"
proof
  assume 0: "surj f"
  from 0 have 1: "∀ A. ∃ a. A = f a" by (simp add: surj_def)
  from 1 have 2: "∃ a. {x. x ∉ f x} = f a" by blast
  from 2 show "False" by blast
qed
```

The **proof** command lacks an explicit method by which to perform the proof. In such cases Isabelle tries to use some standard introduction rule, in the above case for $\neg$:

$$\frac{P \implies False}{\neg P}$$

In order to prove $\neg P$, assume *P* and show *False*. Thus we may assume *surj f*. The proof shows that names of propositions may be (single!) digits —

meaningful names are hard to invent and are often not necessary. Both **have** steps are obvious. The second one introduces the diagonal set $\{x. x \notin f x\}$, the key idea in the proof. If you wonder why 2 directly implies *False*: from 2 it follows that $(a \notin f a) = (a \in f a)$.

4.1.1 *this*, then, hence and thus

Labels should be avoided. They interrupt the flow of the reader who has to scan the context for the point where the label was introduced. Ideally, the proof is a linear flow, where the output of one step becomes the input of the next step, piping the previously proved fact into the next proof, like in a UNIX pipe. In such cases the predefined name *this* can be used to refer to the proposition proved in the previous step. This allows us to eliminate all labels from our proof (we suppress the **lemma** statement):

```
proof
  assume "surj f"
  from this have " $\exists a. \{x. x \notin f x\} = f a$ " by(auto simp: surj_def)
  from this show "False" by blast
qed
```

We have also taken the opportunity to compress the two **have** steps into one.

To compact the text further, Isar has a few convenient abbreviations:

```
then = from this
thus = then show
hence = then have
```

With the help of these abbreviations the proof becomes

```
proof
  assume "surj f"
  hence " $\exists a. \{x. x \notin f x\} = f a$ " by(auto simp: surj_def)
  thus "False" by blast
qed
```

There are two further linguistic variations:

```
(have|show) prop using facts = from facts (have|show) prop
with facts = from facts this
```

The **using** idiom de-emphasizes the used facts by moving them behind the proposition.

4.1.2 Structured Lemma Statements: fixes, assumes, shows

Lemmas can also be stated in a more structured fashion. To demonstrate this feature with Cantor's theorem, we rephrase $\neg \text{surj } f$ a little:

lemma

```
fixes f :: "'a ⇒ 'a set"
assumes s: "surj f"
shows "False"
```

The optional **fixes** part allows you to state the types of variables up front rather than by decorating one of their occurrences in the formula with a type constraint. The key advantage of the structured format is the **assumes** part that allows you to name each assumption; multiple assumptions can be separated by **and**. The **shows** part gives the goal. The actual theorem that will come out of the proof is $\text{surj } f \implies \text{False}$, but during the proof the assumption $\text{surj } f$ is available under the name s like any other fact.

proof —

```
have "∃ a. {x. x ∉ f x} = f a" using s
  by(auto simp: surj_def)
thus "False" by blast
```

qed

! Note the hyphen after the **proof** command. It is the null method that does nothing to the goal. Leaving it out would be asking Isabelle to try some suitable introduction rule on the goal *False* — but there is no such rule and **proof** would fail.

In the **have** step the assumption $\text{surj } f$ is now referenced by its name s . The duplication of $\text{surj } f$ in the above proofs (once in the statement of the lemma, once in its proof) has been eliminated.

Stating a lemma with **assumes-shows** implicitly introduces the name *assms* that stands for the list of all assumptions. You can refer to individual assumptions by *assms*(1), *assms*(2), etc., thus obviating the need to name them individually.

4.2 Proof Patterns

We show a number of important basic proof patterns. Many of them arise from the rules of natural deduction that are applied by **proof** by default. The patterns are phrased in terms of **show** but work for **have** and **lemma**, too.

We start with two forms of **case analysis**: starting from a formula P we have the two cases P and $\neg P$, and starting from a fact $P \vee Q$ we have the two cases P and Q :

<pre> show "R" proof cases assume "P" ⋮ show "R" ... next assume "¬ P" ⋮ show "R" ... qed </pre>	<pre> have "P ∨ Q" ... then show "R" proof assume "P" ⋮ show "R" ... next assume "Q" ⋮ show "R" ... qed </pre>
--	--

How to prove a logical equivalence:

```

show "P ⟷ Q"
proof
  assume "P"
  ⋮
  show "Q" ...
next
  assume "Q"
  ⋮
  show "P" ...
qed

```

Proofs by contradiction (*ccontr* stands for “classical contradiction”):

<pre> show "¬ P" proof assume "P" ⋮ show "False" ... qed </pre>	<pre> show "P" proof (rule ccontr) assume "¬ P" ⋮ show "False" ... qed </pre>
---	---

How to prove quantified formulas:

<pre> show "∀ x. P(x)" proof fix x ⋮ show "P(x)" ... qed </pre>	<pre> show "∃ x. P(x)" proof ⋮ show "P(witness)" ... qed </pre>
---	---

In the proof of $\forall x. P(x)$, the step `fix x` introduces a locally fixed variable x into the subproof, the proverbial “arbitrary but fixed value”. Instead of x

we could have chosen any name in the subproof. In the proof of $\exists x. P(x)$, *witness* is some arbitrary term for which we can prove that it satisfies P .

How to reason forward from $\exists x. P(x)$:

```
have "∃ x. P(x)" ...
then obtain x where p: "P(x)" by blast
```

After the **obtain** step, x (we could have chosen any name) is a fixed local variable, and p is the name of the fact $P(x)$. This pattern works for one or more x . As an example of the **obtain** command, here is the proof of Cantor's theorem in more detail:

```
lemma "¬ surj(f :: 'a ⇒ 'a set)"
proof
  assume "surj f"
  hence "∃ a. {x. x ∉ f x} = f a" by(auto simp: surj_def)
  then obtain a where "{x. x ∉ f x} = f a" by blast
  hence "a ∉ f a ⟷ a ∈ f a" by blast
  thus "False" by blast
qed
```

Finally, how to prove set equality and subset relationship:

```
show "A = B"          show "A ⊆ B"
proof                proof
  show "A ⊆ B" ...      fix x
next                  assume "x ∈ A"
  show "B ⊆ A" ...      ⋮
qed                    show "x ∈ B" ...
                        qed
```

4.3 Streamlining Proofs

4.3.1 Pattern Matching and Quotations

In the proof patterns shown above, formulas are often duplicated. This can make the text harder to read, write and maintain. Pattern matching is an abbreviation mechanism to avoid such duplication. Writing

```
show formula (is pattern)
```

matches the pattern against the formula, thus instantiating the unknowns in the pattern for later use. As an example, consider the proof pattern for $\longleftrightarrow$:

```
show "formula1 ⟷ formula2" (is "?L ⟷ ?R")
```

```

proof
  assume "?L"
  ⋮
  show "?R" ...
next
  assume "?R"
  ⋮
  show "?L" ...
qed

```

Instead of duplicating $formula_i$ in the text, we introduce the two abbreviations $?L$ and $?R$ by pattern matching. Pattern matching works wherever a formula is stated, in particular with **have** and **lemma**.

The unknown $?thesis$ is implicitly matched against any goal stated by **lemma** or **show**. Here is a typical example:

```

lemma "formula"
proof —
  ⋮
  show ?thesis ...
qed

```

Unknowns can also be instantiated with **let** commands

```
let ?t = "some-big-term"
```

Later proof steps can refer to $?t$:

```
have "... ?t ..."
```



Names of facts are introduced with *name*: and refer to proved theorems. Unknowns $?X$ refer to terms or formulas.

Although abbreviations shorten the text, the reader needs to remember what they stand for. Similarly for names of facts. Names like 1, 2 and 3 are not helpful and should only be used in short proofs. For longer proofs, descriptive names are better. But look at this example:

```

have x_gr_0: "x > 0"
⋮
from x_gr_0 ...

```

The name is longer than the fact it stands for! Short facts do not need names; one can refer to them easily by quoting them:

```

have "x > 0"
:
from 'x>0' ...

```

Note that the quotes around $x>0$ are **back quotes**. They refer to the fact not by name but by value.

4.3.2 moreover

Sometimes one needs a number of facts to enable some deduction. Of course one can name these facts individually, as shown on the right, but one can also combine them with **moreover**, as shown on the left:

have "P ₁ " ...	have lab ₁ : "P ₁ " ...
moreover have "P ₂ " ...	have lab ₂ : "P ₂ " ...
moreover	:
:	have lab _n : "P _n " ...
moreover have "P _n " ...	from lab ₁ lab ₂ ...
ultimately have "P" ...	have "P" ...

The **moreover** version is no shorter but expresses the structure more clearly and avoids new names.

4.3.3 Raw Proof Blocks

Sometimes one would like to prove some lemma locally within a proof, a lemma that shares the current context of assumptions but that has its own assumptions and is generalized over its locally fixed variables at the end. This is what a **raw proof block** does:

```

{ fix x1 ... xn
  assume A1 ... Am
  :
  have B
}

```

proves $\llbracket A_1; \dots; A_m \rrbracket \implies B$ where all x_i have been replaced by unknowns $?x_i$.



The conclusion of a raw proof block is *not* indicated by **show** but is simply the final **have**.

As an example we prove a simple fact about divisibility on integers. The definition of *dvd* is $(b \text{ dvd } a) = (\exists k. a = b * k)$.

```

lemma fixes a b :: int assumes "b dvd (a+b)" shows "b dvd a"
proof -
  { fix k assume k: "a+b = b*k"
    have "∃ k'. a = b*k'"
    proof
      show "a = b*(k - 1)" using k by(simp add: algebra_simps)
    qed }
  then show ?thesis using assms by(auto simp add: dvd_def)
qed

```

Note that the result of a raw proof block has no name. In this example it was directly piped (via **then**) into the final proof, but it can also be named for later reference: you simply follow the block directly by a **note** command:

```
note name = this
```

This introduces a new name *name* that refers to *this*, the fact just proved, in this case the preceding block. In general, **note** introduces a new name for one or more facts.

Exercises

Exercise 4.1. Give a readable, structured proof of the following lemma:

```

lemma assumes T: "∀ x y. T x y ∨ T y x"
  and A: "∀ x y. A x y ∧ A y x ⟶ x = y"
  and TA: "∀ x y. T x y ⟶ A x y" and "A x y"
  shows "T x y"

```

Exercise 4.2. Give a readable, structured proof of the following lemma:

```

lemma "(∃ ys zs. xs = ys @ zs ∧ length ys = length zs)
  ∨ (∃ ys zs. xs = ys @ zs ∧ length ys = length zs + 1)"

```

Hint: There are predefined functions *take* :: *nat* ⇒ '*a list* ⇒ '*a list* and *drop* :: *nat* ⇒ '*a list* ⇒ '*a list* such that *take* *k* [*x*₁,...] = [*x*₁,...,*x*_{*k*}] and *drop* *k* [*x*₁,...] = [*x*_{*k*+1},...]. Let sledgehammer find and apply the relevant *take* and *drop* lemmas for you.

4.4 Case Analysis and Induction

4.4.1 Datatype Case Analysis

We have seen case analysis on formulas. Now we want to distinguish which form some term takes: is it 0 or of the form *Suc n*, is it [] or of the form *x # xs*, etc. Here is a typical example proof by case analysis on the form of *xs*:

```

lemma "length (tl xs) = length xs - 1"
proof (cases xs)
  assume "xs = []"
  thus ?thesis by simp
next
  fix y ys assume "xs = y # ys"
  thus ?thesis by simp
qed

```

Function *tl* ("tail") is defined by $tl [] = []$ and $tl (x1.0 \# x22.0) = x22.0$. Note that the result type of *length* is *nat* and $0 - 1 = 0$.

This proof pattern works for any term *t* whose type is a datatype. The goal has to be proved for each constructor *C*:

```
fix x1 ... xn assume "t = C x1 ... xn"
```

Each case can be written in a more compact form by means of the **case** command:

```
case (C x1 ... xn)
```

This is equivalent to the explicit **fix-assume** line but also gives the assumption $"t = C x_1 \dots x_n"$ a name: *C*, like the constructor. Here is the **case** version of the proof above:

```

proof (cases xs)
  case Nil
  thus ?thesis by simp
next
  case (Cons y ys)
  thus ?thesis by simp
qed

```

Remember that *Nil* and *Cons* are the alphanumeric names for $[]$ and $\#$. The names of the assumptions are not used because they are directly piped (via **thus**) into the proof of the claim.

4.4.2 Structural Induction

We illustrate structural induction with an example based on natural numbers: the sum ($\sum$) of the first *n* natural numbers ($\{0..n::nat\}$) is equal to $n * (n + 1) \text{ div } 2$. Never mind the details, just focus on the pattern:

```

lemma "∑ {0..n::nat} = n*(n+1) div 2"
proof (induction n)
  show "∑ {0..0::nat} = 0*(0+1) div 2" by simp

```

```

next
  fix n assume "∑ {0..n::nat} = n*(n+1) div 2"
  thus "∑ {0..Suc n} = Suc n*(Suc n+1) div 2" by simp
qed

```

Except for the rewrite steps, everything is explicitly given. This makes the proof easily readable, but the duplication means it is tedious to write and maintain. Here is how pattern matching can completely avoid any duplication:

```

lemma "∑ {0..n::nat} = n*(n+1) div 2" (is "?P n")
proof (induction n)
  show "?P 0" by simp
next
  fix n assume "?P n"
  thus "?P (Suc n)" by simp
qed

```

The first line introduces an abbreviation $?P\ n$ for the goal. Pattern matching $?P\ n$ with the goal instantiates $?P$ to the function $\lambda n. \sum \{0..n\} = n * (n + 1) \text{ div } 2$. Now the proposition to be proved in the base case can be written as $?P\ 0$, the induction hypothesis as $?P\ n$, and the conclusion of the induction step as $?P\ (Suc\ n)$.

Induction also provides the **case** idiom that abbreviates the **fix-assume** step. The above proof becomes

```

proof (induction n)
  case 0
  show ?case by simp
next
  case (Suc n)
  thus ?case by simp
qed

```

The unknown $?case$ is set in each case to the required claim, i.e., $?P\ 0$ and $?P\ (Suc\ n)$ in the above proof, without requiring the user to define a $?P$. The general pattern for induction over nat is shown on the left-hand side:

```

show "P(n) "
proof (induction n)
  case 0                let ?case = "P(0) "
  :
  show ?case ...
next
  case (Suc n)          fix n assume Suc: "P(n) "
  :                      let ?case = "P(Suc n) "
  show ?case ...
qed

```

On the right side you can see what the `case` command on the left stands for.

In case the goal is an implication, induction does one more thing: the proposition to be proved in each case is not the whole implication but only its conclusion; the premises of the implication are immediately made assumptions of that case. That is, if in the above proof we replace `show "P(n) "` by `show "A(n)  $\implies$  P(n) "` then `case 0` stands for

```

assume 0: "A(0) "
let ?case = "P(0) "

```

and `case (Suc n)` stands for

```

fix n
assume Suc: "A(n)  $\implies$  P(n) "
           "A(Suc n) "
let ?case = "P(Suc n) "

```

The list of assumptions `Suc` is actually subdivided into `Suc.IH`, the induction hypotheses (here $A(n) \implies P(n)$), and `Suc.premis`, the premises of the goal being proved (here $A(Suc\ n)$).

Induction works for any datatype. Proving a goal $\llbracket A_1(x); \dots; A_k(x) \rrbracket \implies P(x)$ by induction on x generates a proof obligation for each constructor C of the datatype. The command `case (C x1 ... xn)` performs the following steps:

1. `fix x1 ... xn`
2. `assume` the induction hypotheses (calling them `C.IH`) and the premises $A_i(C\ x_1 \dots x_n)$ (calling them `C.premis`) and calling the whole list `C`
3. `let ?case = "P(C x1 ... xn) "`

4.4.3 Rule Induction

Recall the inductive and recursive definitions of even numbers in [Section 3.5](#):

```

inductive ev :: "nat  $\Rightarrow$  bool" where
  ev0: "ev 0" |
  evSS: "ev n  $\implies$  ev (Suc (Suc n))"

```

```

fun evn :: "nat  $\Rightarrow$  bool" where
  "evn 0 = True" |
  "evn (Suc 0) = False" |
  "evn (Suc (Suc n)) = evn n"

```

We recast the proof of $ev\ n \implies evn\ n$ in Isar. The left column shows the actual proof text, the right column shows the implicit effect of the two `case` commands:

<pre> lemma "ev n $\implies$ evn n" proof(induction rule: ev.induct) case ev0 show ?case by simp next case evSS </pre>	<pre> let ?case = "evn 0" fix n assume evSS: "ev n" "evn n" let ?case = "evn (Suc (Suc n))" </pre>
<pre> thus ?case by simp qed </pre>	

The proof resembles structural induction, but the induction rule is given explicitly and the names of the cases are the names of the rules in the inductive definition. Let us examine the two assumptions named *evSS*: *ev n* is the premise of rule *evSS*, which we may assume because we are in the case where that rule was used; *evn n* is the induction hypothesis.

 Because each case command introduces a list of assumptions named like the case name, which is the name of a rule of the inductive definition, those rules now need to be accessed with a qualified name, here *ev.ev0* and *ev.evSS*.

In the case *evSS* of the proof above we have pretended that the system fixes a variable *n*. But unless the user provides the name *n*, the system will just invent its own name that cannot be referred to. In the above proof, we do not need to refer to it, hence we do not give it a specific name. In case one needs to refer to it one writes

```
case (evSS m)
```

like `case (Suc n)` in earlier structural inductions. The name *m* is an arbitrary choice. As a result, case *evSS* is derived from a renamed version of rule

evSS: $ev\ m \implies ev(Suc(Suc\ m))$. Here is an example with a (contrived) intermediate step that refers to *m*:

```
lemma "ev n  $\implies$  evn n"
proof(induction rule: ev.induct)
  case ev0 show ?case by simp
next
  case (evSS m)
  have "evn(Suc(Suc m)) = evn m" by simp
  thus ?case using 'evn m' by blast
qed
```

In general, let *I* be a (for simplicity unary) inductively defined predicate and let the rules in the definition of *I* be called *rule*₁, ..., *rule*_{*n*}. A proof by rule induction follows this pattern:

```
show "I x  $\implies$  P x"
proof(induction rule: I.induct)
  case rule1
  :
  show ?case ...
next
  :
next
  case rulen
  :
  show ?case ...
qed
```

One can provide explicit variable names by writing *case* (*rule*_{*i*} *x*₁ ... *x*_{*k*}), thus renaming the first *k* free variables in rule *i* to *x*₁ ... *x*_{*k*}, going through rule *i* from left to right.

4.4.4 Assumption Naming

In any induction, *case name* sets up a list of assumptions also called *name*, which is subdivided into three parts:

name.IH contains the induction hypotheses.

name.hyps contains all the other hypotheses of this case in the induction rule. For rule inductions these are the hypotheses of rule *name*, for structural inductions these are empty.

name.premis contains the (suitably instantiated) premises of the statement being proved, i.e., the *A*_{*i*} when proving $\llbracket A_1; \dots; A_n \rrbracket \implies A$.

! Proof method *induct* differs from *induction* only in this naming policy: *induct* does not distinguish *IH* from *hyps* but subsumes *IH* under *hyps*.

More complicated inductive proofs than the ones we have seen so far often need to refer to specific assumptions — just *name* or even *name.prem*s and *name.IH* can be too unspecific. This is where the indexing of fact lists comes in handy, e.g., *name.IH*(2) or *name.prem*s(1–2).

4.4.5 Rule Inversion

Rule inversion is case analysis of which rule could have been used to derive some fact. The name **rule inversion** emphasizes that we are reasoning backwards: by which rules could some given fact have been proved? For the inductive definition of *ev*, rule inversion can be summarized like this:

$$ev\ n \implies n = 0 \vee (\exists k. n = Suc\ (Suc\ k) \wedge ev\ k)$$

The realisation in Isabelle is a case analysis. A simple example is the proof that $ev\ n \implies ev\ (n - 2)$. We already went through the details informally in [Section 3.5.1](#). This is the Isar proof:

```

assume "ev n"
from this have "ev(n - 2)"
proof cases
  case ev0 thus "ev(n - 2)" by (simp add: ev.ev0)
next
  case (evSS k) thus "ev(n - 2)" by (simp add: ev.evSS)
qed

```

The key point here is that a case analysis over some inductively defined predicate is triggered by piping the given fact (here: *from this*) into a proof by *cases*. Let us examine the assumptions available in each case. In case *ev0* we have $n = 0$ and in case *evSS* we have $n = Suc\ (Suc\ k)$ and *ev k*. In each case the assumptions are available under the name of the case; there is no fine-grained naming schema like there is for induction.

Sometimes some rules could not have been used to derive the given fact because constructors clash. As an extreme example consider rule inversion applied to *ev (Suc 0)*: neither rule *ev0* nor rule *evSS* can yield *ev (Suc 0)* because *Suc 0* unifies neither with 0 nor with *Suc (Suc n)*. Impossible cases do not have to be proved. Hence we can prove anything from *ev (Suc 0)*:

```

assume "ev(Suc 0)" then have P by cases

```

That is, *ev (Suc 0)* is simply not provable:

```

lemma "¬ ev(Suc 0)"

```

```

proof
  assume "ev (Suc 0) " then show False by cases
qed

```

Normally not all cases will be impossible. As a simple exercise, prove that $\neg \text{ev } (\text{Suc } (\text{Suc } (\text{Suc } 0)))$.

4.4.6 Advanced Rule Induction

So far, rule induction was always applied to goals of the form $I\ x\ y\ z \implies \dots$ where I is some inductively defined predicate and x, y, z are variables. In some rare situations one needs to deal with an assumption where not all arguments r, s, t are variables:

```

lemma "I r s t  $\implies \dots$ "

```

Applying the standard form of rule induction in such a situation will lead to strange and typically unprovable goals. We can easily reduce this situation to the standard one by introducing new variables x, y, z and reformulating the goal like this:

```

lemma "I x y z  $\implies x = r \implies y = s \implies z = t \implies \dots$ "

```

Standard rule induction will work fine now, provided the free variables in r, s, t are generalized via *arbitrary*.

However, induction can do the above transformation for us, behind the curtains, so we never need to see the expanded version of the lemma. This is what we need to write:

```

lemma "I r s t  $\implies \dots$ "
proof(induction "r" "s" "t" arbitrary: ... rule: I.induct)

```

Like for rule inversion, cases that are impossible because of constructor clashes will not show up at all. Here is a concrete example:

```

lemma "ev (Suc m)  $\implies \neg \text{ev } m$ "
proof(induction "Suc m" arbitrary: m rule: ev.induct)
  fix n assume IH: " $\bigwedge m. n = \text{Suc } m \implies \neg \text{ev } m$ "
  show " $\neg \text{ev } (\text{Suc } n)$ "
  proof — contradiction
    assume "ev (Suc n)"
    thus False
  proof cases — rule inversion
    fix k assume "n = Suc k" "ev k"
    thus False using IH by auto
qed

```

qed
qed

Remarks:

- Instead of the `case` and `?case` magic we have spelled all formulas out. This is merely for greater clarity.
- We only need to deal with one case because the `ev0` case is impossible.
- The form of the *IH* shows us that internally the lemma was expanded as explained above: $ev\ x \implies x = Suc\ m \implies \neg\ ev\ m$.
- The goal $\neg\ ev\ (Suc\ n)$ may surprise. The expanded version of the lemma would suggest that we have a `fix m assume Suc (Suc n) = Suc m` and need to show $\neg\ ev\ m$. What happened is that Isabelle immediately simplified $Suc\ (Suc\ n) = Suc\ m$ to $Suc\ n = m$ and could then eliminate m . Beware of such nice surprises with this advanced form of induction.

! This advanced form of induction does not support the *IH* naming schema explained in [Section 4.4.4](#): the induction hypotheses are instead found under the name *hyps*, as they are for the simpler *induct* method.

Exercises

Exercise 4.3. Give a structured proof by rule inversion:

lemma assumes $a: "ev(Suc(Suc\ n))"$ shows $"ev\ n"$

Exercise 4.4. Give a structured proof of $\neg\ ev\ (Suc\ (Suc\ (Suc\ 0)))$ by rule inversions. If there are no cases to be proved you can close a proof immediately with `qed`.

Exercise 4.5. Recall predicate *star* from [Section 3.5.2](#) and *iter* from [Exercise 3.4](#). Prove $iter\ r\ n\ x\ y \implies star\ r\ x\ y$ in a structured style; do not just sledgehammer each case of the required induction.

Exercise 4.6. Define a recursive function $elems :: 'a\ list \Rightarrow 'a\ set$ and prove $x \in elems\ xs \implies \exists\ ys\ zs. xs = ys @ x \# zs \wedge x \notin elems\ ys$.

Exercise 4.7. Extend [Exercise 3.5](#) with a function that checks if some *alpha list* is a balanced string of parentheses. More precisely, define a recursive function $balanced :: nat \Rightarrow alpha\ list \Rightarrow bool$ such that $balanced\ n\ w$ is true iff (informally) $S\ (a^n @ w)$. Formally, prove that $balanced\ n\ w = S\ (replicate\ n\ a @ w)$ where $replicate :: nat \Rightarrow 'a \Rightarrow 'a\ list$ is predefined and $replicate\ n\ x$ yields the list $[x, \dots, x]$ of length n .

References

1. Michael Huth and Mark Ryan. *Logic in Computer Science*. Cambridge University Press, 2004.
2. Alexander Krauss. *Defining Recursive Functions in Isabelle/HOL*. <http://isabelle.in.tum.de/doc/functions.pdf>.
3. Tobias Nipkow. *What's in Main*. <http://isabelle.in.tum.de/doc/main.pdf>.
4. Tobias Nipkow and Gerwin Klein. *Concrete Semantics with Isabelle/HOL*. Springer, 2014. 298 pp. <http://concrete-semantics.org>.
5. Tobias Nipkow, Lawrence Paulson, and Markus Wenzel. *Isabelle/HOL — A Proof Assistant for Higher-Order Logic*, volume 2283 of *Lect. Notes in Comp. Sci.* Springer-Verlag, 2002.
6. Makarius Wenzel. *The Isabelle/Isar Reference Manual*. <http://isabelle.in.tum.de/doc/isar-ref.pdf>.