

A Denotational Semantics of Inheritance and its Correctness

William Cook¹
Jens Palsberg²

Technical Report No. CS-90-30

September 1990

¹Computer Science Dept., Brown University, Providence, RI 02912

²Computer Science Dept, Aarhus University, Ny Munkegade, DK-8000, Aarhus C,
Denmark

A Denotational Semantics of Inheritance and its Correctness

William Cook*

Department of Computer Science
Box 1910 Brown University
Providence, RI 02912, USA
wrc@cs.brown.edu

Jens Palsberg

Computer Science Department
Aarhus University
Ny Munkegade, DK-8000
Aarhus C, Denmark
palsberg@daimi.dk

Abstract

This paper presents a denotational model of inheritance. The model is based on an intuitive motivation of the purpose of inheritance. The correctness of the model is demonstrated by proving it equivalent to an operational semantics of inheritance based upon the method-lookup algorithm of object-oriented languages. Although it was originally developed to explain inheritance in object-oriented languages, the model shows that inheritance is a general mechanism that may be applied to any form of recursive definition.

1 Introduction

Inheritance is one of the central concepts in object-oriented programming. Despite its importance, there seems to be a lack of consensus on the proper way to describe inheritance. This is evident from the following review of various formalizations of inheritance that have been proposed.

The concept of *prefixing* in Simula [5], which evolved into the modern concept of inheritance, was defined in terms of textual concatenation of program blocks. However, this definition was informal, and only partially accounted for more sophisticated aspects of prefixing like the pseudo-variable *this* and virtual operations.

The most precise and widely used definition of inheritance is given by the operational semantics of object-oriented languages. The canonical operational semantics is the “method lookup” algorithm of Smalltalk:

When a message is sent, the methods in the receiver’s class are searched for one with a matching selector. If none is found, the methods in that class’s superclass are searched next. The search continues up the superclass chain until a matching method is found. ...

When a method contains a message whose receiver is *self*, the search for the method for that message begins in the instance’s class, regardless of which class contains the method containing *self*. ...

When a message is sent to *super*, the search for a method ... begins in the superclass of the class containing the method. The use of *super* allows a method to access methods defined in a superclass even if the methods have been overridden in the subclasses. [6, pp. 61–64]

Unfortunately, such operational definitions do not necessarily foster intuitive understanding. As a result, insight into the proper use and purpose of inheritance is often gained only through an “Aha!” experience [1].

Cardelli [2] identifies inheritance with the subtype relation on record types: “a record type τ is a subtype (written $\leq$) of a record type τ' if τ has all the fields of τ' , and possibly more, and the common fields of τ and τ' are in the $\leq$ relation.” His work shows that a sound type-checking algorithm exists for strongly-typed, statically-scoped languages with inheritance, but it doesn’t give their dynamic semantics.

More recently, McAllister and Zabih [9] suggested a system of “boolean classes” similar to inheritance as used in knowledge representation. Stein [16] focused on shared attributes and methods. Minsky and Rosenzstein [10] characterized inheritance by “laws” regulating message sending. Although they express various aspects of inheritance, none of these presentations are convincing because they provide no verifiable evidence that the formal model corresponds to the form of inheritance

* *Current address:* Hewlett-Packard Laboratories, P.O. Box 10490
Palo Alto, CA 94303-0969, cook@hplabs.hp.com.

actually used in object-oriented programming.

This paper presents a denotational model of inheritance. The model is based upon an intuitive explanation of the proper use and purpose of inheritance. It is well-known that inheritance is a mechanism for “differential programming” by allowing a new class to be defined by incremental modification of an existing class. We show that self-reference complicates the mechanism of incremental programming. In order for derivation to have the same conceptual effect as direct modification, self-reference in the original definition must be changed to refer to the modified definition. This conceptual argument is useful for explaining the complex functionality of the pseudovariables *self* and *super* in Smalltalk.

Although the model was originally developed to describe inheritance in object-oriented languages, it shows that inheritance is a general mechanism that is applicable to any kind of recursive definition.

Essentially the same technical interpretation of inheritance was discovered independently by Reddy [12]. A closely related model was presented by Kamin [8]. However, Kamin describes inheritance as a global operation on programs, a formulation that blurs scope issues and inheritance.

These duplications, by themselves, are evidence for the validity of the model. This paper provides, in addition, a formal proof that the inheritance model is equivalent to the operational definition of inheritance quoted above.

In Section 2 we develop an intuitive motivation of inheritance. In Section 3 this intuition is formalized as a denotational model of inheritance. In Section 4 we demonstrate the correctness of the model by proving equivalence of two semantics of object-oriented systems, one based on the operational model and the other based upon the denotational model.

2 Motivating Inheritance

Inheritance is a mechanism for differential, or incremental, programming. Incremental programming is the construction of new program components by specifying how they differ from existing components. Incremental programming may be achieved by text editing, but this approach has a number of obvious disadvantages. A more disciplined approach to incremental programming is based upon using a form of “filter” to modify the external behavior of the original component. For example, to define a modified version of a function one simply defines a new function that performs some special computations and possibly calls the original function. This simple form of derivation is illustrated in Figure 1, where *P* is the original function, *M* is the modification, and the arrows represent invocation.

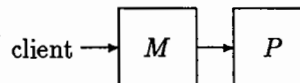


Figure 1: Derivation.

Incremental programming by explicit derivation is obviously more restrictive than text-editing; changes can be made either to the input passed to the original module or the output it returns, but the way in which the original works cannot be changed. Thus this form of derivation does not violate encapsulation [15]: the original structure can be replaced with an equivalent implementation and the derivation will have the same effect (text-editing is inherently unencapsulated).

However, there is one way in which this naive interpretation of derivation is radically different from text-editing: in the treatment of self-reference or recursion in the original structure. Figure 2 illustrates a naive derivation from a self-referential component.

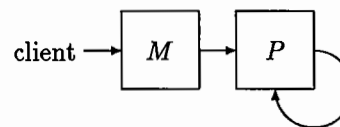


Figure 2: Naive derivation from a recursive structure.

Notice that the modification only affects external clients of the function — it does not modify the function’s recursive calls. Thus naive derivation does not represent a true modification of the original component. To achieve the effect of a true modification of the original component, self-reference in the original function must be changed to refer to the modification, as illustrated in Figure 3.

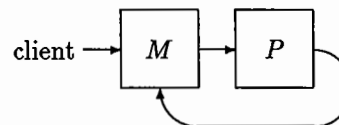


Figure 3: Inheritance.

This construction represents the essence of inheritance: it is a mechanism for deriving modified versions of recursive structures.

3 A Model of Inheritance

This section develops the informal account of inheritance into a formal model of inheritance in object-

oriented languages. Since manipulation of self-reference was identified as a central feature of inheritance, inheritance is explained within a traditional semantics of self-reference or recursion.

3.1 Fixed-Point Semantics

The fixed-point semantics of recursive programs, developed by Scott [14], provides the mathematical setting for the inheritance model. The technique of fixed point semantics is merely illustrated here; thorough introductory explanations are given by Stoy [17], Gordon [7], Scott [14], and Schmidt [13].

To illustrate the use of fixed points for the analysis of recursive programs, consider the following definition of the factorial function:

$fact = \lambda n. \text{if } n = 1 \text{ then } 1 \text{ else } n \times fact(n - 1)$

From a mathematical point of view, this definition is not very satisfactory because *fact* appears on both sides of the equation: the definition is merely an equation that *fact* must satisfy. There is no guarantee that any function satisfying this equation exists, and even less that there is a unique one.

This problem is solved by fixed-point analysis, which indicates how to construct the denotation of such self-referential definitions by “solving” the equation given above. To use fixed-point techniques, the recursive definition is transformed into a non-recursive form. First, the body of the function is converted into an explicit abstraction, or function, in which the parameter *f* is substituted for *fact*:

$FACT = \lambda f. \lambda n. \text{if } n = 1 \text{ then } 1 \text{ else } n \times f(n - 1)$

Functions derived in this way will be called *generators*. *FACT* is a *functional*, or mapping from functions to functions. Its definition is not recursive, because ‘*FACT*’ does not appear in its body. The formal parameter *f* represents the function to call in order to compute the factorial of numbers less than *n*, if needed. The original definition of *fact* is then given in terms of *FACT*:

$fact = FACT(fact)$

But now *fact* is defined as a value that is unchanged when *FACT* is applied. Such a value is called a *fixed point* of *FACT*. Under certain conditions, it is possible to define a particular *fixed point*, the *least fixed point*, of any function by using the *fixed-point function*, *fix*. The fixed-point function has the property that if $f = \text{fix}(F)$, then $F(f) = f$. Another way to express the least fixed point, which is utilized later in the paper, is as the limit of the series $\perp, F(\perp), F(F(\perp)), \dots$. For *fix* to work the function *F* must be *continuous* (our domains are cpos), a condition that is satisfied by our definitions in the rest of the paper.

3.2 Self-referential Objects

The following example illustrates how fixed-point semantics is used to describe the behavior of objects with mutually recursive methods. This is essentially the standard interpretation of mutually recursive procedures [18, 19]. The example involves a simple class of ‘points’ given in Figure 4. Points have *x* and *y* components to specify their location. The *distFromOrig* method computes their distance from the origin. *closerToOrg* is a method that takes another point object and returns “true” if the point is closer to the origin than the other point, and “false” otherwise.

```
class Point(a, b)
  method x = a
  method y = b
  method distFromOrig =
    sqrt(self.x2 + self.y2)
  method closerToOrg(p) =
    (self.distFromOrig < p.distFromOrig)
```

Figure 4: The class Point.

Objects are modeled as record values whose fields represent methods [12, 3]. The notation $\{l_1 \mapsto v_1, \dots, l_n \mapsto v_n\}$ represents a record associating the value v_i with label l_i . Records may in turn be viewed as finite functions from a domain of labels to a heterogeneous domain of values. Selection of the field *l* from a record *m* is achieved by applying the record to the label: $m.l$ or $m(l)$.

By using this model, the class Point is represented as a function that creates objects. References to *self* in Point are explained using fixed points. A function *MakeGenPoint* is defined to construct points and supply them with appropriate methods. Since points are self-referential, they are explained as the fixed point of the “generator” of the methods. *MakeGenPoint*, defined in Figure 5, is a function that takes the coordinates of the new point and returns a generator, whose fixed point is a point.

A point (3,4) is created as shown in Figure 6. The *closerToOrg* function takes a single argument which is assumed to be a point. Actually, all that is required is that it be a record with a *distFromOrig* component, whose value is a number.

3.3 Class Inheritance

Inheritance allows a new class to be defined by adding or replacing methods in an existing class. In the following example, the Point class is inherited to define a class of

```

MakeGenPoint(a, b) = λ self.
{ x ↦ a,
  y ↦ b,
  distFromOrig ↦
    sqrt(self.x2 + self.y2),
  closerToOrig ↦
    λ p. (self.distFromOrig < p.distFromOrig) }

```

Figure 5: The generator associated with Point.

```

p = fix(MakeGenPoint(3, 4))
= { x ↦ 3,
    y ↦ 4,
    distFromOrig ↦ 5,
    closerToOrig ↦
      λ p. (5 < p.distFromOrig) }

```

Figure 6: A point at location (3,4).

circles. Circles have a radius and thus a different notion of distance from the origin. The definition in Figure 7 gives only the differences between circles and points.

This form of inheritance is modeled as an operation on generators. There are three aspects to this process: (1) the addition or replacement of methods, (2) the redirection of self-reference in the original generator to refer to the modified methods, and (3) the binding of *super* in the modification to refer to the original methods.

The modifications effected during class inheritance are naturally expressed as a record of methods to be combined with the inherited methods. The new methods M and the original methods O are combined into a new record $M \oplus O$ such that any method defined in M replaces the corresponding method in O .

The modifications, however, are also defined in terms of the original methods (via *super*). In addition, the modifications refer to the resulting structure (via *self*). Thus a modification is naturally expressed as a function of two arguments, one representing *self* and the other representing *super*, that returns a record of locally defined methods. Such functions will be called *wrappers*. A wrapper contains just the information in the subclass definition. The wrapper for the subclass Circle is given in Figure 8.

The other aspect of inheritance that must be formalized is the change of self-reference in inherited methods. The methods to be inherited are contained in a generator, a function whose argument is used for self-reference in the methods. The result of inheritance should be a new class definition, that is, a new generator.

This mechanism is provided by *wrapper application*.

```

class Circle(a, b, r) inherit Point(a, b)
  method radius = r
  method distFromOrig =
    max(super.distFromOrig - self.radius, 0)

```

Figure 7: The class Circle.

```

CircleWrapper = λ a, b, r. λ self. λ super.
{ radius ↦ r,
  distFromOrig ↦
    max(super.distFromOrig - self.radius, 0) }

```

Figure 8: The wrapper associated with Circle.

A wrapper is applied to a generator to produce a new generator by first distributing *self* to both the wrapper and the original generator. Then the modifications defined by the wrapper are applied to the original record definition to produce a modification record. This is then combined with the original record using $\oplus$. The mechanism of wrapper application is defined by the infix inheritance operator $\triangleright$ as follows:

$$W \triangleright P = \lambda \text{self}. (W(\text{self})(P(\text{self}))) \oplus P(\text{self})$$

Note that two occurrences of P refer to the same behavioral denotation, and do not indicate that the parent is instantiated twice. The mechanism of wrapper application is illustrated in Figure 9.

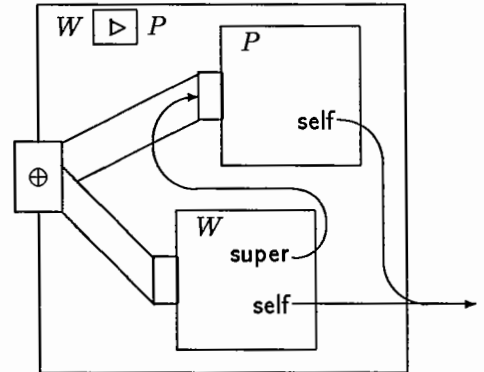


Figure 9: Wrapper application.

The generator which represents the class Circle can now be defined by wrapper application of CircleWrapper to MakeGenPoint, as shown in Figure 10. The figure also shows the partial expansion of the expression into a form that represents what one might write if circles

```

MakeGenCircle =  $\lambda a, b, r.$ 
  CircleWrapper( $a, b, r$ )  $\triangleright$  MakeGenPoint( $a, b$ )
=  $\lambda a, b, r. \lambda \text{self}.$ 
  {
     $x \mapsto a,$ 
     $y \mapsto b,$ 
     $\text{radius} \mapsto r,$ 
     $\text{distFromOrig} \mapsto$ 
       $\max(0, \text{sqrt}(\text{self}.x^2 + \text{self}.y^2) - \text{self}.radius),$ 
     $\text{closerToOrig} \mapsto$ 
       $\lambda p. (\text{self}.distFromOrig < p.distFromOrig) \}$ 

```

Figure 10: The generator associated with Circle.

had been defined without using inheritance. Note that `distFromOrig` has changed in such a way that `closerToOrig` uses the notion of distance for circles, instead of the original one for points. Thus inheritance has achieved a consistent modification of the point class.

4 Correctness of the Model

To show the correctness of the inheritance model, we prove that it is equivalent to the definition of inheritance provided by the operational semantics of an object-oriented language. We introduce method systems as a useful framework in which to prove correctness. Two different semantics for method systems are then defined, based on the operational and denotational definitions of inheritance. Finally we prove the equivalence of the two semantics.

4.1 Method Systems

Method systems are a simple formalization of object-oriented programming that support semantics based upon both the operational and the denotational models of inheritance. Method systems encompass only those aspects of object-oriented programming that are directly related to inheritance or method determination. As such, many important aspects are omitted, including instance variables, assignment, and object creation.

A method system may be understood as part of a snapshot of an object-oriented system. It consists of all the objects and relationships that exist at a given point during execution of an object-oriented program. The basic ontology for method systems includes instances, classes, and method descriptions, which are mappings from message keys to method expressions. Each object is an instance of a class. Classes have an associated method description and may inherit methods from other classes. These (flat) domains and their interconnections

Method System Domains			
Instances	ρ	$\in$	Instance
Classes	κ	$\in$	Class
Messages	m	$\in$	Key
Primitives	f	$\in$	Primitive
Methods	e	$\in$	Exp := self super arg $e_1 \ m \ e_2$ $f(e_1, \dots, e_q)$

Method System Operations	
<i>class</i>	: Instance $\rightarrow$ Class
<i>parent</i>	: Class $\rightarrow$ (Class + ?)
<i>methods</i>	: Class $\rightarrow$ Key $\rightarrow$ (Exp + ?)

Figure 11: Syntactic domains and interconnections.

are defined in Figure 11. Figure 12 illustrates a method system.

The syntax of method expressions is defined by the **Exp** domain which defines a restricted language used to implement the behavior of objects. For simplicity, methods all have exactly one argument, referenced by the symbol **arg** within the body of the method. Self-reference is denoted by the symbol **self**, which may be returned as the value of a method, passed as an actual argument, or sent additional messages. A subclass method may invoke the previous definition of a redefined method with the expression **super**. Message-passing is represented by the expression $e_1 \ m \ e_2$, in which the message consisting of the key m and the argument e_2 is sent to the object e_1 . Finally, primitive values and computations are represented by the expression $f(e_1, \dots, e_q)$. If $q = 0$, then the primitive represents a constant.

class gives the class of an instance. Every instance has exactly one class, although a class may have many instances.

parent defines the inheritance hierarchy, which is required to be a tree. For any class κ , the value of $\text{parent}(\kappa)$ is the parent class of κ , or else $\perp?$ if κ is the root. $?$ is a one-point domain consisting of only $\perp?$. The use of (**Class** + ?) allows us to test monotonically whether a class is the root. Note that + denotes “separated” sum, so that the elements of (**Class** + ?) are (distinguished copies of) the elements of **Class**, the element $\perp?$, and a new bottom element. We omit the injections into sum domains; the meaning of expressions, in particular $\perp?$, is always unambiguously implied by the context.

methods specifies the local method expressions defined by a class. For any class κ and any message key m , the value of $\text{methods}(\kappa)m$ is either an expression or $\perp?$ if κ doesn’t define an expression for m . Let us assume that the root of the inheritance hierarchy doesn’t

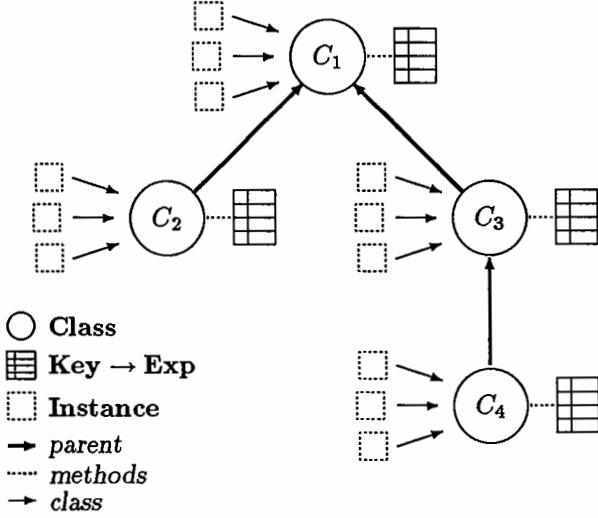


Figure 12: A method system.

define any methods. Note that inheritance allows instances of a class to respond to more than the locally defined methods.

In the following two sections we give the method system both a conventional method lookup semantics and a denotational semantics. Both define the result of sending a message to an instance.

4.2 Method Lookup Semantics

The method lookup semantics given in Figure 14 closely resembles the implementation of method lookup in object-oriented languages like Smalltalk [6]. It is given in a denotational style due to the abstract nature of method systems. A more traditional operational semantics is not needed because of the absence of updatable storage.

The domains used to represent the behavior of an instance are defined in Figure 13. A behavior is a mapping from message keys to *functions* or $\perp$. This is clearly contrasted with the methods of a class, which are given by a mapping from message keys to *expressions* or $\perp$. Thus a behavior is a semantic entity, while methods are syntactic. Another difference between the behavior of an instance and its class's methods is that the behavior contains a function for every message the class handles, while methods associate an expression only with messages that are different from the class's parent. In the rest of this paper, $\perp$ (without subscript) denotes the bottom element of **Behavior**.

The semantics also uses an auxiliary function *root*, defined in Figure 13, that determines whether a class is the root of the inheritance hierarchy. **Boolean** is the

Semantic Domains			
Number			
α	$\in$	Value	$= \text{Behavior} + \text{Number}$
σ, π	$\in$	Behavior	$= \text{Key} \rightarrow (\text{Fun} + ?)$
ϕ	$\in$	Fun	$= \text{Value} \rightarrow \text{Value}$

root : **Class** $\rightarrow$ **Boolean**
root(κ) = $[\lambda \kappa' \in \text{Class} . \text{false},$
 $\lambda v \in ? . \text{true}$
 $](\text{parent}(\kappa))$

Figure 13: Semantic domains and *root*.

send : **Instance** $\rightarrow$ **Behavior**
send(ρ) = *lookup*(*class*(ρ)) ρ

lookup : **Class** $\rightarrow$ **Instance** $\rightarrow$ **Behavior**
lookup(κ) ρ = $\lambda m \in \text{Key} .$
 $[\lambda e \in \text{Exp} . \text{do}[e]\rho\kappa,$
 $\lambda v \in ? . \text{if } \text{root}(\kappa)$
 $\text{then } \perp?$
 $\text{else } \text{lookup}(\text{parent}(\kappa))\rho m$
 $](\text{methods}(\kappa)m)$

do : **Exp** $\rightarrow$ **Instance** $\rightarrow$ **Class** $\rightarrow$ **Fun**
do[**self**] $\rho\kappa$ = $\lambda \alpha . \text{send}(\rho)$
do[**super**] $\rho\kappa$ = $\lambda \alpha . \text{lookup}(\text{parent}(\kappa))\rho$
do[**arg**] $\rho\kappa$ = $\lambda \alpha . \alpha$
do[$e_1 \ m \ e_2$] $\rho\kappa$ = $\lambda \alpha . (\text{do}[e_1]\rho\kappa\alpha)m(\text{do}[e_2]\rho\kappa\alpha)$
do[$f(e_1, \dots, e_q)$] $\rho\kappa$ =
 $\lambda \alpha . f(\text{do}[e_1]\rho\kappa\alpha, \dots, \text{do}[e_q]\rho\kappa\alpha)$

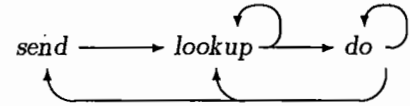


Figure 14: The method lookup semantics.

flat three-point domain of truth values. $[f, g]$ denotes the case analysis of two functions $f \in D_f \rightarrow t$ and $g \in D_g \rightarrow t$ with result in type t , mapping $x \in D_f + D_g$ to $f(x)$ if $x \in D_f$ or to $g(x)$ if $x \in D_g$.

Sending a message m to an instance ρ is performed by looking up the message in the instance's class. The lookup process yields a function that takes a message key and an actual argument and computes the value of the message send.

Performing message m in a class κ on behalf of an instance ρ involves searching the sequence of class par-

ents until a method is found to handle the message. This method is then evaluated. In *lookup*, the instance and message remain constant, while the class argument is recursively bound to each of the parents in sequence. At each stage there are two possibilities: (1) the message key has an associated method expression in class κ , in which case it is evaluated, and (2) the method is not defined, in which case a recursive call is made to *lookup* after computing the parent of the class. The tail-recursion in *lookup* would be replaced by iteration in a real interpreter.

Evaluation of methods is complicated by the need to interpret occurrences of **self** and **super**. The *do* function has three extra arguments, besides the expression being evaluated: the original instance ρ that received the message whose method is being evaluated, the class κ in which the method was found, and an actual argument α . The expression **self** evaluates to the behavior of the original instance. The expression **super** requires a continuation of the method search starting from the superclass of the class in which the method occurs. The expression **arg** evaluates to α . The expression $e_1 \ m \ e_2$ evaluates to the result of applying the behavior of the object denoted by e_1 to m and the meaning of the argument e_2 .

One important aspect of the method-lookup semantics is that the functions are mutually recursive, because *do* contains calls to *send* and *lookup*.

4.3 Denotational Semantics

The denotational semantics based on generator modification given in Figure 16 uses two additional domains representing behavior generators and wrappers, defined in Figure 15. A formal definition of $\oplus$ is also given in Figure 15.

The behavior of an instance is defined as the fixed point of the generator associated with its class. The generator specifies a self-referential behavior, and its fixed point is that behavior. The generator of the root class produces a behavior in which all messages are undefined.

The generator of a class that isn't the root is created by modifying the generator of the class's parent. The modifications to be made are found in the wrapper of the class, which is a semantic entity derived from the block of syntactic method expressions defined by the class. These modifications are effected by the inheritance operator $\triangleright$.

The function *wrap* computes the wrapper of a class as a mapping from messages to the evaluation of the corresponding method, or to $\perp$. A wrapper has two behavioral arguments, one used for self-reference, and the other for reference to the parent behavior (i.e. the behavior being 'wrapped'). These arguments may be understood as representing the behavior of **self** and the

Generator Semantics Domains		
Generator	=	Behavior $\rightarrow$ Behavior
Wrapper	=	Behavior $\rightarrow$ Generator

$$\begin{aligned} \oplus : (\mathbf{Behavior} \times \mathbf{Behavior}) &\rightarrow \mathbf{Behavior} \\ r_1 \oplus r_2 &= \lambda m \in \mathbf{Key} . [\lambda \phi \in \mathbf{Fun} . \phi, \\ &\quad \lambda v \in ? . r_2(m) \\ &\quad]r_1(m) \end{aligned}$$

Figure 15: Semantic domains and $\oplus$.

```

behave : Instance  $\rightarrow$  Behavior
behave( $\rho$ ) = fix(gen(class( $\rho$ )))

gen : Class  $\rightarrow$  Generator
gen( $\kappa$ ) = if root( $\kappa$ )
      then  $\lambda \sigma \in \mathbf{Behavior} . \lambda m \in \mathbf{Key} . \perp?$ 
      else wrap( $\kappa$ )  $\triangleright$  gen(parent( $\kappa$ ))

wrap : Class  $\rightarrow$  Wrapper
wrap( $\kappa$ ) =  $\lambda \sigma . \lambda \pi . \lambda m \in \mathbf{Key} .$ 
      [ $\lambda e \in \mathbf{Exp} . \text{eval}[e] \sigma \pi$ 
        $\lambda v \in ? . \perp?$ 
       ]methods( $\kappa$ )m

eval : Exp  $\rightarrow$  Behavior  $\rightarrow$  Behavior  $\rightarrow$  Fun
eval[self] $\sigma \pi = \lambda \alpha . \sigma$ 
eval[super] $\sigma \pi = \lambda \alpha . \pi$ 
eval[arg] $\sigma \pi = \lambda \alpha . \alpha$ 
eval[ $e_1 \ m \ e_2$ ] $\sigma \pi =$ 
       $\lambda \alpha . (\text{eval}[e_1] \sigma \pi \alpha) m (\text{eval}[e_2] \sigma \pi \alpha)$ 
eval[ $f(e_1, \dots, e_q)$ ] $\sigma \pi =$ 
       $\lambda \alpha . f(\text{eval}[e_1] \sigma \pi \alpha, \dots, \text{eval}[e_q] \sigma \pi \alpha)$ 

```

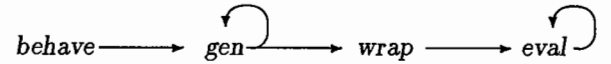


Figure 16: The denotational semantics.

behavior of **super**. In the definitions, the behavior for **self** is named σ and the one for **super** is named π .

A method is always evaluated in the context of a behavior for **self** (represented by σ) and **super** (represented by π). The evaluation of the corresponding expressions, **self** and **super**, is therefore simple. The evaluation of the other expressions is essentially the same as in the method lookup semantics.

Note that each of the functions in the denotational semantics is recursive only within itself: there is no mu-

tual recursion among the functions, except that which is achieved by the explicit fixed point.

4.4 Equivalence

The method lookup semantics and the denotational semantics are equivalent because they assign the same behavior to an instance. This proposition is captured by theorem 1.

Theorem 1 $send = behave$

In the proof of the theorem we use an “intermediate semantics” defined in Figure 17 and inspired by the one used by Mosses and Plotkin [11] in their proof of limiting completeness. The semantics uses $n \in \text{Nat}$, the flat domain of natural numbers.

The intermediate semantics resembles the method-lookup semantics but differs in that each of the syntactic domains of instances, classes, and expressions has a whole family of semantic equations, indexed by natural numbers. The intuition behind the definition is that $send'_n i$ allows $(n-1)$ evaluations of **self** before it stops and gives $\perp$. $send'_n i$ is defined in terms of $send'_{n-1} i$ via $lookup'_n$ and do'_n because the **self** expression evaluates to the result of $send'_{n-1} i$, which allows one less evaluation of **self**. (The values of $(lookup'_n c \ i)$ and $(do'_n e \ i \ c \ a)$ are irrelevant; let them be $\perp$.)

The following four lemmas state useful properties of the intermediate semantics. Here we only outline their proofs, leaving the full proofs to the appendix.

Lemma 1 *If $n > 0$ then*

$$do'_n \llbracket e \rrbracket \rho \kappa = eval \llbracket e \rrbracket (send'_{n-1} \rho) (lookup'_n (parent(\kappa)) \rho)$$

PROOF: By induction on the structure of e , using the definitions of do' and $eval$.

Lemma 2 *If $n > 0$ then*

$$lookup'_n \kappa \rho = gen(\kappa)(send'_{n-1} \rho)$$

PROOF: By induction on the number of ancestors of κ , using the definitions of gen , $\boxed{\triangleright}$, $\oplus$, and $wrap$, Lemma 1, and the definition of $lookup'$.

Lemma 3 $send'_n \rho = (gen(class(\rho)))^n(\perp)$

PROOF: By induction on n , using Lemma 2 and the definition of $send'$.

Lemma 4 $send'$, $lookup'$, and do' are monotone functions of the natural numbers with the usual ordering.

PROOF: Immediate from Lemma 1–3.

Lemma 4 expresses that the family of $send'_n$'s is an increasing sequence of functions.

$send' : \text{Nat} \rightarrow \text{Instance} \rightarrow \text{Behavior}$
 $send'_0(\rho) = \perp$
 $send'_n(\rho) = lookup'_n(class(\rho))\rho \quad n > 0$

$lookup' : \text{Nat} \rightarrow \text{Class} \rightarrow \text{Instance} \rightarrow \text{Behavior}$
 $lookup'_0 \kappa \rho = \perp$
if $n > 0$ then
 $lookup'_n \kappa \rho = \lambda m \in \text{Key} .$
 $\quad [\lambda e \in \text{Exp} . do'_n \llbracket e \rrbracket \rho \kappa,$
 $\quad \lambda v \in ? . \text{if } root(\kappa)$
 $\quad \quad \text{then } \perp?$
 $\quad \quad \text{else } lookup'_n(parent(\kappa))\rho m$
 $\quad](methods(\kappa)m)$

$do' : \text{Nat} \rightarrow \text{Exp} \rightarrow \text{Instance} \rightarrow \text{Class} \rightarrow \text{Fun}$
 $do'_0 \llbracket e \rrbracket \rho \kappa = \lambda \alpha . \perp$
if $n > 0$ then
 $do'_n \llbracket \text{self} \rrbracket \rho \kappa = \lambda \alpha . send'_{n-1} \rho$
 $do'_n \llbracket \text{super} \rrbracket \rho \kappa = \lambda \alpha . lookup'_n(parent(\kappa))\rho$
 $do'_n \llbracket \text{arg} \rrbracket \rho \kappa = \lambda \alpha . \alpha$
 $do'_n \llbracket e_1 \ m \ e_2 \rrbracket \rho \kappa =$
 $\quad \lambda \alpha . (do'_n \llbracket e_1 \rrbracket \rho \kappa \alpha) m (do'_n \llbracket e_2 \rrbracket \rho \kappa \alpha)$
 $do'_n \llbracket f(e_1, \dots, e_q) \rrbracket \rho \kappa =$
 $\quad \lambda \alpha . f(do'_n \llbracket e_1 \rrbracket \rho \kappa \alpha, \dots, do'_n \llbracket e_q \rrbracket \rho \kappa \alpha)$

Figure 17: The intermediate semantics.

Definition 1

$interpret : \text{Instance} \rightarrow \text{Behavior}$
 $interpret = \bigsqcup_n (send'_n)$

The following three propositions express the relations among the method lookup semantics, the intermediate semantics, and the denotational semantics.

Proposition 1 $interpret = behave$

PROOF: The following proof uses the definition of $interpret$, Lemma 3, the fixed-point theorem, and the definition of $behave$.

$$\begin{aligned} interpret(\rho) &= \bigsqcup_n (send'_n(\rho)) \\ &= \bigsqcup_n (gen(class(\rho)))^n(\perp) \\ &= \text{fix}(gen(class(\rho))) \\ &= behave(\rho) \end{aligned}$$

QED

Proposition 2 $send \sqsubseteq behave$

PROOF: The following facts have proofs analogous to those of Lemma 1–2 (we omit the proofs).

1. $do \llbracket e \rrbracket \rho \kappa = eval \llbracket e \rrbracket (send(\rho)) (lookup(parent(\kappa))\rho)$

$$2. \text{lookup}(\kappa)\rho = \text{gen}(\kappa)(\text{send}(\rho))$$

From the definition of send and the second fact we get $\text{send}(\rho) = \text{lookup}(\text{class}(\rho))\rho = \text{gen}(\text{class}(\rho))(\text{send}(\rho))$. Hence $\text{send}(\rho)$ is a fixed point of $\text{gen}(\text{class}(\rho))$. The definition of behave expresses that $\text{behave}(\rho)$ is the least fixed point of $\text{gen}(\text{class}(\rho))$; thus $\text{send}(\rho) \sqsubseteq \text{behave}(\rho)$. QED

Proposition 3 $\text{send} \sqsubseteq \text{interpret}$

PROOF: The functions defined in the method-lookup semantics are mutually recursive. Their meaning is the least fixed point of the generator g defined in the obvious way, as outlined below.

$$\begin{aligned} D = & (\text{Instance} \rightarrow \text{Behavior}) \\ & \times (\text{Class} \rightarrow \text{Instance} \rightarrow \text{Behavior}) \\ & \times (\text{Exp} \rightarrow \text{Instance} \rightarrow \text{Class} \rightarrow \text{Fun}) \end{aligned}$$

Let $g : D \rightarrow D$ be defined by

$$g(s, l, d) = (\lambda i \in \text{Instance}. l(\text{class}(\rho))\rho, \dots, \dots)$$

Now we can prove by induction on n that

$$g^n(\perp_D) \sqsubseteq (\text{send}'_n, \text{lookup}'_n, \text{do}'_n)$$

In the base case, where $n = 0$, the inequality holds trivially. Then assume that the inequality holds for $(n-1)$, where $n > 0$. The following proof of the induction step uses the associativity of function composition, the induction hypothesis, and Lemma 4.

$$\begin{aligned} g^n(\perp_D) &= g(g^{n-1}(\perp_D)) \\ &\sqsubseteq g(\text{send}'_{n-1}, \text{lookup}'_{n-1}, \text{do}'_{n-1}) \\ &\sqsubseteq (\text{send}'_n, \text{lookup}'_n, \text{do}'_n) \end{aligned}$$

Now

$$\begin{aligned} (\text{send}, \text{lookup}, \text{do}) &= \text{fix}(g) \\ &= \bigsqcup_n g^n(\perp_D) \\ &\sqsubseteq \bigsqcup_n (\text{send}'_n, \text{lookup}'_n, \text{do}'_n) \end{aligned}$$

and in particular $\text{send} \sqsubseteq \bigsqcup_n (\text{send}'_n) = \text{interpret}$
QED

PROOF of Theorem 1: Combine Propositions 1–3. QED

5 Conclusion

A denotational semantics of inheritance was presented, using a general notation that is applicable to the analysis of different object-oriented languages [4]. The semantics was supported by an intuitive explanation of

inheritance as a mechanism for incremental programming that simulate destructive modification. An explanation of the binding of self- and super-reference was given at this conceptual level. To provide evidence for the correctness of the model, it was proven equivalent to the most widely accepted definition of inheritance, the operational method-lookup semantics used in object-oriented languages.

In comparing the denotational semantics with the operational semantics, the denotational one does not seem to be much simpler. It may even be argued that it is a great deal more complex, because it requires an understanding of fixed points. The primary advantage of the denotational semantics is the intuitive explanation it provides. It suggests that inheritance may be useful for other kinds of recursive structures, like types and functions, in addition to classes. Another important advantage of the denotational semantics is a demonstration that inheritance, while a natural extension of existing mechanisms, does provide expressive power not found in conventional languages by allowing more flexible use of the fixed-point function.

Acknowledgement. The authors would like to thank Peter Wegner for the original motivation for an equivalence proof, Peter Mosses for helpful comments on the second part of the paper, and John Mitchell for comments on an earlier draft.

References

- [1] Alan H. Borning and Tim O'Shea. Deltatalk: An empirically and aesthetically motivated simplification of the Smalltalk-80 language. In *European Conf. on Object-Oriented Programming*, pages 1–10, 1987.
- [2] Luca Cardelli. A semantics of multiple inheritance. In *Semantics of Data Types, LNCS 173*, pages 51–68. Springer-Verlag, 1984.
- [3] Luca Cardelli and Peter Wegner. On understanding types, data abstraction, and polymorphism. *Computing Surveys*, 17(4):471–522, 1985.
- [4] William Cook. *A Denotational Semantics of Inheritance*. PhD thesis, Brown University, 1989.
- [5] Ole-Johan Dahl and Kristen Nygaard. *The SIMULA 67 Common Base Language*. 1970.
- [6] Adele Goldberg and Dave Robson. *Smalltalk-80: the Language and Its Implementation*. Addison-Wesley, 1983.
- [7] Michael J. C. Gordon. *The Denotational Description of Programming Languages*. Springer-Verlag, 1979.

- [8] Samuel Kamin. Inheritance in Smalltalk-80: A denotational definition. In *Proc. of Conf. on Principles of Programming Languages*, pages 80–87, 1988.
- [9] D. McAllester and R. Zabih. Boolean classes. In *Proc. ACM Conf. on Object-Oriented Programming: Systems, Languages and Applications*, pages 417–423, 1987.
- [10] N. Minsky and D. Rozenshtein. A law-based approach to object-oriented programming. In *Proc. ACM Conf. on Object-Oriented Programming: Systems, Languages and Applications*, pages 482–493, 1987.
- [11] Peter Mosses and Gordon Plotkin. On proving limiting completeness. *SIAM Journal of Computing*, 16:179–194, 1987.
- [12] Uday S. Reddy. Objects as closures: Abstract semantics of object-oriented languages. In *Proc. ACM conf. on Lisp and Functional Programming*, pages 289–297, 1988.
- [13] David A. Schmidt. *Denotational Semantics: A Methodology for Language Development*. Allyn & Bacon, 1986.
- [14] Dana S. Scott. Data types as lattices. *SIAM Journal*, 5(3):522–586, 1976.
- [15] Alan Snyder. Encapsulation and inheritance in object-oriented programming languages. In *Proc. ACM Conf. on Object-Oriented Programming: Systems, Languages and Applications*, pages 38–45, 1986.
- [16] Lynn Andrea Stein. Delegation is inheritance. In *Proc. ACM Conf. on Object-Oriented Programming: Systems, Languages and Applications*, pages 138–146, 1987.
- [17] Joseph Stoy. *Denotational Semantics: The Scott-Strachey Approach to Programming Language Semantics*. MIT Press, 1977.
- [18] Robert D. Tennent. On a new approach to representation-independent data classes. *Acta Informatica*, 8:315–324, 1977.
- [19] Robert D. Tennent. Denotational semantics of Hoare classes. Unpublished manuscript, 1982.

Appendix: Proofs of Lemmas

In this appendix we give the full proofs of Lemma 1–4.

PROOF of Lemma 1: Recall that we want to prove, for $n > 0$, that

$do'_n[e]\rho\kappa = eval[e](send'_{n-1}\rho)(lookup'_n(parent(\kappa))\rho)$
by induction on the structure of e . The base case is proved as follows.

$$\begin{aligned}
do'_n[\mathbf{self}]\rho\kappa\alpha &= send'_{n-1}\rho \\
&= eval[\mathbf{self}](send'_{n-1}\rho)(lookup'_n(parent(\kappa))\rho)\alpha \\
do'_n[\mathbf{super}]\rho\kappa\alpha &= lookup'_n(parent(\kappa))\rho \\
&= eval[\mathbf{super}](send'_{n-1}\rho)(lookup'_n(parent(\kappa))\rho)\alpha \\
do'_n[\mathbf{arg}]\rho\kappa\alpha &= \alpha \\
&= eval[\mathbf{arg}](send'_{n-1}\rho)(lookup'_n(parent(\kappa))\rho)\alpha
\end{aligned}$$

The induction step is proven below using the abbreviation $\pi = lookup'_n(parent(\kappa))\rho$.

$$\begin{aligned}
do'_n[e_1 \ m \ e_2]\rho\kappa\alpha &= do'_n[e_1]\rho\kappa\alpha m (do'_n[e_2]\rho\kappa\alpha) \\
&= eval[e_1](send'_{n-1}\rho)\pi\alpha m (eval[e_2](send'_{n-1}\rho)\pi\alpha) \\
&= eval[e_1 \ m \ e_2](send'_{n-1}\rho)\pi\alpha \\
do'_n[f(e_1, \dots, e_q)]\rho\kappa\alpha &= f(do'_n[e_1]\rho\kappa\alpha, \dots, do'_n[e_q]\rho\kappa\alpha) \\
&= f(eval[e_1](send'_{n-1}\rho)\pi\alpha, \dots, \\
&\quad eval[e_q](send'_{n-1}\rho)\pi\alpha) \\
&= eval[f(e_1, \dots, e_q)](send'_{n-1}\rho)\pi\alpha
\end{aligned}$$

QED

PROOF of Lemma 2: Recall that we want to prove $lookup'_n\kappa\rho = gen(\kappa)(send'_{n-1}\rho)$, where $n > 0$, by induction on the number of ancestors of κ . In the base case, where κ is the root, both sides evaluate to $(\lambda m \in \mathbf{Key}. \perp?)$ because κ doesn't define any methods. Then assume that the lemma holds for $parent(\kappa)$. The proof of the induction step given below uses the definition of gen (κ isn't the root), the definition of $\boxed{\triangleright}$, the induction hypothesis, the definitions of $\oplus$ and $wrap$, the properties of case analysis, Lemma 1, and the definition of $lookup'$ (κ isn't the root). Note that we use the abbreviation $\pi = lookup'_n(parent(\kappa))\rho$.

$$\begin{aligned}
gen(\kappa)(send'_{n-1}\rho) &= (wrap(\kappa) \boxed{\triangleright} gen(parent(\kappa)))(send'_{n-1}\rho) \\
&= (wrap(\kappa)(send'_{n-1}\rho)(gen(parent(\kappa))(send'_{n-1}\rho))) \\
&\quad \oplus (gen(parent(\kappa))(send'_{n-1}\rho)) \\
&= (wrap(\kappa)(send'_{n-1}\rho)\pi) \oplus \pi
\end{aligned}$$