

MPRI, Typage

Didier Rémy
(With course material from François Pottier)

November 19, 2013



Plan of the course

Introduction

Simply-typed λ -calculus

Polymorphism and System F

Type reconstruction

Existential types

Overloading

Overloading

Contents

- Introduction
- Examples in Mini Haskell
- Mini Haskell
- Implicitly-typed terms
- Variations

What is overloading?

Overloading occurs when at some program point, several definitions for a same identifier are visible simultaneously.

An interpretation of the program (and a fortiori a run of the program) must choose the definition that applies at this point. This is called *overloading resolution*, which may use very different strategies and techniques.

All sorts of identifiers may be subject to overloading: variables, labels, constructors, types, etc.

Overloading must be distinguished from *shadowing* of identifiers by normal scoping rules, where in this case, a new definition may just shadow an older one and temporarily become the only one visible.

Why use overloading?

Naming convenience

It avoids name mangling, such as suffixing similar names by type information: printing functions, e.g. `print_int`, `print_string`, etc.; numerical operations, e.g. `+`, `+. .` etc.); or numerical constants e.g. `0`, `0. .` etc.

Modularity

To avoid clashing, the naming discipline (including name mangling conventions) must be known globally. Isolated identifiers with no particular naming convention may still interfere between different developments and cannot be used together unless fully qualified.

To think more abstractly

In terms of operations rather than of particular implementations. For instance, calling `to_string` conversion lets the system check whether one definition is available according to the type of the argument.

Why use overloading?

Type dependent functions

A function defined on $\tau[\alpha]$ for all α may have an implementation depending on the type of α . For instance, a marshaling function of type $\forall \alpha. \alpha \rightarrow \text{string}$ may execute different code for each base type α .

Ad hoc polymorphism

Overloading definitions may be ad hoc, *i.e.* completely unrelated for each type, or just share a same type schema.

For example 0 could mean either integer zero or the empty list. $\times$ could mean the either the integer product or string concatenation.

Why use overloading?

Type dependent functions

A function defined on $\tau[\alpha]$ for all α may have an implementation depending on the type of α . For instance, a marshaling function of type $\forall \alpha. \alpha \rightarrow \text{string}$ may execute different code for each base type α .

Polytypic polymorphism

Overloading definitions depend solely on the *type structure* (on whether it is a sum, a product, *etc.*) and can thus be derived mechanically for all types from their definitions on base types.

Typical examples of polytypic functions are marshaling functions or the generation of random values for arbitrary types, e.g. as used in [Quickcheck for Haskell](#).

Different forms of overloading

There are many variants of overloading, which can be classified by how overloading is *introduced* and *resolved*.

What are the restrictions on overloading definitions?

- None, *i.e.* arbitrary definitions can be overloaded!
- Can just functions or any definition be overloaded? *e.g.* can numerical values be overloaded?
- Are all overloaded definitions of the same name instances of a common type scheme? Are these type schemes arbitrary?
- Are overloaded definitions primitive (pre-existing), automatic (generated mechanically from other definitions), or user-defined?
- Can overloaded definitions overlap?
- Can overloaded definitions have a local scope?

How is overloading tamed?

How is overloading resolution defined?

- up to subtyping?
- *static* or *dynamic*?

Static resolution (rather simple)

- Overloaded symbols can/must be statically replaced by their implementations at the appropriate types.
- This does not increase expressiveness, but may still significantly reduce verbosity.

How is overloading tamed?

How is overloading resolution defined?

- up to subtyping?
- *static* or *dynamic*?

Dynamic resolution (more involved)

This is required when the choice of the implementation depends on the dynamic of the program execution. For example, the resolution at a program point in a polymorphic function may depend on the type of its argument so that different calls can make different choices.

The resolution is driven by information made available at runtime:

- it can be full or partial type information, or extra values (tags, dictionaries, *etc.*) correlated to types instead of types themselves.
- it can be attached to normal values or passed as extra arguments.

Static resolution

Examples

In SML

Overloaded definitions are primitive (for numerical operators), and automatic (for record accesses).

Typechecking fails if overloading cannot be resolved at outermost let-definitions. For example, `let twice x = x + x` is rejected in SML, at toplevel, as `+` could be the addition on either integers or floats.

Static resolution

Examples

In SML

Overloaded definitions are primitive (for numerical operators), and automatic (for record accesses).

Typechecking fails if overloading cannot be resolved at outermost let-definitions. For example, `let twice x = x + x` is rejected in SML, at toplevel, as `+` could be the addition on either integers or floats.

In Java?

Static resolution

Examples

In Java

Overloading is not primitive but automatically generated by subtyping. When a class extends another one and a method is redefined, the older definition is still visible, hence the method is overloaded.

Overloading is resolved at compile time by choosing the most specific definition. There is always a best choice—according to static knowledge.

An argument may have a runtime type that is a subtype of the best known compile-time type, and perhaps a more specific definition could have been used if overloading were resolved dynamically.

This is often a source of confusion for Java programmers.

Static resolution

Limits

It does not fit well with first-class functions and polymorphism:

For example, $\lambda x. x + x$ is rejected when $+$ is overloaded, as it cannot be statically resolved. The function must be specialized at some type at which $+$ is defined.

This argues in favor of some form of dynamic overloading: dynamic overloading allows to delay resolution of overloaded symbols until polymorphic functions have been sufficiently specialized.

How is dynamic resolution implemented?

Three main techniques for dynamic resolution

- Pass types at runtime and dispatch on the runtime type, using a general typecase construct.
- Tag values with their types—or, usually, an approximation of their types—and dispatch on these tags.
(This is one possible approach to object-orientation where objects may be tagged with the class they belong to.)
- Pass the appropriate implementations at runtime as extra arguments, usually grouped in *dictionaries* of implementations.

Dynamic resolution

Type passing semantics

Runtime type dispatch

- Use an explicitly-typed calculus (e.g. System F)
- Add a typecase function.
- The runtime cost of typecase may be high, unless type patterns are significantly restricted.
- By default, one pays even when overloading is not used.
- Monomorphization may be used to reduce type matching statically.
- Ensuring exhaustiveness of type matching is difficult.

ML& (Castagna)

- System F + intersection types + subtyping + type matching
- An expressive type system that keeps track of exhaustiveness; type matching functions are first-class and can be extended or overridden.
- Allows overlapping definitions with a best match resolution strategy.

Dynamic resolution

Type erasing semantics

Passing unresolved implementations as extra arguments

- Abstract over unresolved overloaded symbols and pass them around as extra arguments.

Hopefully, overloaded symbols can be resolved when their types are sufficiently specialized and before they are actually needed.

In short, *let* $f = \lambda x. x + x$ *in* a can be elaborated into

let $f = \lambda(+). \lambda x. x + x$ *in* a . Then, the application of f to to a float in a e.g. $f\ 1.0$ can be elaborated into $f\ (+.)\ 1.0$.

- This can be done based on the typing derivation.
- After elaboration, types are no longer needed and can be erased.
- Monomorphization or other simplifications may reduce the number of abstractions and applications introduced by overloading resolution.

Dynamic resolution

Type erasing semantics

This has been explored under different facets in the context of ML:

- Type classes, introduced in [1989] by Wadler and Blott are the most popular and widely explored framework of this kind.
- Other contemporary proposals were proposed by Rouaix [1990] and Kaes [1992].
- Tentative simplifications of type classes have been made [Odersky et al., 1995] but did not take over, because of their restrictions.
- Recent works on type classes is still going [Morris and Jones, 2010]

We present Mini-Haskell that contains the essence of Haskell.

Contents

- Introduction
- Examples in Mini Haskell
- Mini Haskell
- Implicitly-typed terms
- Variations

Mini-Haskell

Mini Haskell is a simplification of Haskell to avoid most of the difficulties of type classes while keeping their essence:

- single parameter type classes
- no overlapping instance definitions

It is close to [A second look at overloading](#) by Odersky et al. in terms of expressiveness and simplicity—but closer to Haskell in style: it can be easily generalized by lifting restrictions without changing the framework.

Our version of Mini-Haskell is explicitly typed. We present:

- Some examples in Mini-Haskell.
- Elaboration of Mini-Haskell into (the ML subset of) System F.
- An implicitly-typed version with type inference.

Mini-Haskell Example

Implicitly Typed

Mini-Haskell class declarations and instance definitions

```

class Eq X { equal : X → X → Bool }
inst Eq Int { equal = (==) }
inst Eq Char { equal = (==) }
inst      Eq X ⇒ Eq (List (X))
    { equal = λ(l1      ) λ(l2      ) match l1, l2 with
      | [],[] → true | [], - | [], - → false
      | h1::t1, h2::t2 → equal    h1 h2 && equal          t1 t2 }

```

This code:

- declares a class (dictionary) of type *Eq*(X) that contains definitions for *equal* : X → X → X,
- creates two concrete instances (dictionaries) of type *Eq* Int and *Eq* Char,
- creates a function that given a dictionary for *Eq* X builds a dictionary for *Eq* (List(X)).

Mini-Haskell Example

Explicitly Typed

Mini-Haskell class declarations and instance definitions

```

class Eq X { equal : X → X → Bool }
inst Eq Int { equal = (==) }
inst Eq Char { equal = (==) }
inst  $\Lambda(X)$  Eq X  $\Rightarrow$  Eq (List (X))
  { equal =  $\lambda(l_1 : \text{List } X) \lambda(l_2 : \text{List } X)$  match  $l_1, l_2$  with
    | [], []  $\rightarrow$  true | [], - | -, -  $\rightarrow$  false
    |  $h_1::t_1, h_2::t_2 \rightarrow$  equal X  $h_1 h_2$  && equal (List X)  $t_1 t_2$  }

```

This code:

- declares a class (dictionary) of type $\text{Eq}(X)$ that contains definitions for $\text{equal} : X \rightarrow X \rightarrow \text{Bool}$,
- creates two concrete instances (dictionaries) of type Eq Int and Eq Char ,
- creates a function that given a dictionary for $\text{Eq } X$ builds a dictionary for Eq (List(X)) .

Example

Elaboration into explicit dictionaries

```

class Eq X { equal : X → X → Bool }
inst Eq Int { equal = (==) }
inst Eq Char { equal = (==) }
inst  $\Lambda(X)$  Eq X  $\Rightarrow$  Eq (List (X))
  { equal =  $\lambda(l_1 : \text{List } X) \lambda(l_2 : \text{List } X)$  match  $l_1, l_2$  with
    | [], []  $\rightarrow$  true | [], - | -, -  $\rightarrow$  false
    |  $h_1::t_1, h_2::t_2 \rightarrow$  equal  $X$   $h_1$   $h_2$  && equal (List  $X$ )  $t_1$   $t_2$  }

```

Becomes:

```

type Eq (X) = { equal : X → X → Bool }
let equal X (EqX : Eq X) : X → X → Bool = EqX.equal

let EqInt : Eq Int = { equal = ( (==) : int → int → bool ) }
let EqChar : Eq Char = { equal = printEqChar }
let EqList  $X$  (EqX : Eq X) : Eq (List X)
  { equal =  $\lambda(l_1 : \text{List } X) \lambda(l_2 : \text{List } X)$  match  $l_1, l_2$  with
    | [], []  $\rightarrow$  true | [], - | -, -  $\rightarrow$  false
    |  $h_1::t_1, h_2::t_2 \rightarrow$ 
      equal  $X$  EqX  $h_1$   $h_2$  && equal (List  $X$ ) (EqList  $X$  EqX)  $t_1$   $t_2$  }

```

Example

Class Inheritance

Classes may themselves depend on other classes (called superclasses):

```
class Eq X ⇒ Ord (X) { lt : X → X → Bool }
inst Ord Int { lt = (<) }
```

This declares a new class (dictionary) *Ord* X that depends on a dictionary *Eq* X and contains a method *lt* : X → X → Bool.

The instance definition builds a dictionary *Ord* Int from the existing dictionary *Eq* Int and the primitive (<) for *lt*.

The two declarations are elaborated into:

```
type Ord X = { Eq : Eq X; lt : X → X → Bool }
let EqOrd X (OrdX : Ord X) : Eq X = OrdX.eq
let lt X (OrdX : Ord X) : X → X → Bool = OrdX.lt

let OrdInt : Ord Int = { Eq = EqInt; lt = (<) }
```

Mini Haskell

Overloading

An overloaded function `search` is defined as follows:

```
let rec search :  $\forall (X) \text{ Ord } X \Rightarrow X \rightarrow \text{List } X \rightarrow \text{Bool} =$ 
   $\Lambda(X) \lambda(x : X) \lambda(l : \text{List } X)$ 
  match l with []  $\rightarrow$  false | h::t  $\rightarrow$  equal x h || search x t

let b = search Int 1 [1; 2; 3];;
```

This elaborates into:

```
let rec search X (OrdX : Ord X) (x : X) (l : List X) : Bool =
  match l with
  | []  $\rightarrow$  false
  | h::t  $\rightarrow$  equal X (EqOrd X OrdX) x h || search X OrdX x t

let b = search Int OrdInt 1 [1; 2; 3];;
```

That is, the code in **green** is inferred.

Contents

- Introduction
- Examples in Mini Haskell
- Mini Haskell
- Implicitly-typed terms
- Variations

Mini Haskell

We restrict to single parameter classes.

Class and instance declarations are restricted to the toplevel.

Their scope is the whole program.

Mini Haskell

In practice, a program is composed of *interleaved*

- class declarations,
- instance definitions,
- function definitions,

given *in any order* and

- ending with an expression.

Instance and function definitions are interpreted recursively.

Hence, their definition order does not matter.

For simplification, we assume that instance definitions do not depend on function definitions, which may then come last as part of the expression in a recursive let-binding.

Mini Haskell

In practice, a program is composed of *sequences of*

- class declarations,
- instance definitions,

given *in this order* and

- ending with an expression.

Instance definitions are interpreted recursively; their order does not matter.

We may assume, *w.l.o.g.*, that instance definitions come after all class declarations.

The order of class declaration matters, since they may only refer to other class constructors that have been previously defined.

Mini Haskell

Source programs p are of the form:

$$p ::= H_1 \dots H_p \ h_1 \dots h_q \ M$$

$$\begin{aligned} H &::= \text{class } \vec{P} \Rightarrow K \alpha \ \{\rho\} \\ \rho &::= u_1 : \tau_1, \dots, u_m : \tau_m \end{aligned}$$

$$\begin{aligned} h &::= \text{inst } \forall \vec{\beta}. \vec{P} \Rightarrow K (G \vec{\beta}) \ \{r\} \\ r &::= u_1 = M_1, \dots, u_k = M_k \end{aligned}$$

$$P ::= K \alpha \qquad Q ::= K \tau \qquad T ::= \tau \mid Q \qquad \sigma ::= \forall \vec{\alpha}. \vec{Q} \Rightarrow T$$

Letter u ranges over overloaded symbols.

Class constructors K may appear in Q but not in τ .

Only regular type constructors G may appear in τ .

We write $\forall \vec{\alpha}. Q_1 \Rightarrow \dots Q_m \Rightarrow T$ for $\forall \vec{\alpha}. Q_1, \dots, Q_m \Rightarrow T$ and see $\Rightarrow$ as an annotated version of $\rightarrow$.

Mini Haskell

Source programs p are of the form:

$$p ::= H_1 \dots H_p \ h_1 \dots h_q \ M$$

$$H ::= \text{class } \vec{P} \Rightarrow K \alpha \{ \rho \}$$

$$\rho ::= u_1 : \tau_1, \dots u_m : \tau_m$$

$$h ::= \text{inst } \forall \vec{\beta}. \vec{P} \Rightarrow K (G \vec{\beta}) \{ r \}$$

$$r ::= u_1 = M_1, \dots u_k = M_k$$

$$P ::= K \alpha$$

$$Q ::= K \tau$$

$$T ::= \tau \mid Q$$

$$\sigma ::= \forall \vec{\alpha}. \vec{Q} \Rightarrow T$$

The sequence $\vec{P}$ in class and instance definitions is a *typing context*. Each clause $\vec{P}$ is of the form $K' \alpha'$ and can be read as an assumption “*given a dictionary of type $\alpha \dots$* ”

The restriction to types of the form $K' \alpha'$ in typing contexts and class declarations, and to types of the form $K' (G' \vec{\alpha}')$ in instances are for simplicity. Generalization are discussed later.

Target language

System F, extended with record types, let-bindings, and let-rec.

Records are provided as data types. They are used to represent dictionaries. Record labels represented overloaded symbols u .

We may also use overloaded symbols u as variables.

This amounts to reserving a subset of variables x_u indexed by overloaded symbols, but just writing u as a shortcut for x_u .

We use letter N instead of M for elaborated terms, to distinguish them from source terms.

Class declarations

$$H \triangleq \text{class } K_1 \alpha_1, \dots K_p \alpha_p \Rightarrow K \alpha \{ \rho \}$$

A class declaration H defines a class constructor K .

Every class (constructor) K must be defined by one and only one class declaration. So we may say that H is the declaration of K .

Classes K_i 's are superclasses of K and we write $K_i < K$.

Class definitions must respect the order $<$ (*acyclic*)

The dictionary of K will contain a sub-dictionary for each K_i .

All K_i 's are independent in a *typing context*: there does not exist i and j such that $K_j < K_i$.

(If $K_j < K_i$, then K_i dictionary would contain a sub-dictionary for K_j , to which K has access via K_i so K does not itself need dictionary for K_j .)

Class declarations

$$H \triangleq \text{class } K_1 \alpha_1, \dots K_p \alpha_p \Rightarrow K \alpha \{\rho\}$$

The row type ρ is of the form

$$u_1 : \tau_1, \dots u_m : \tau_m$$

and declares *overloaded symbols* u_i (also called *methods*) of class K .

An overloaded symbol cannot be declared twice in the same class and must be declared only in one class.

Types τ_i 's must be closed with respect to α .

Each class dictionary will contain a definition for each method of the class.

Class declarations

Elaboration

$$H \triangleq \text{class } K_1 \alpha_1, \dots K_p \alpha_p \Rightarrow K \alpha \{ \rho \}$$

Its elaboration consists in a record type declaration to represent the dictionary and the definition of accessors for each field of the record.

The row ρ only lists methods $u_1 : T_1, \dots u_m : T_m$. We extend it with sub-dictionary fields and define ρ^K to be $\rho, u_{K_1}^K : K_1 \alpha, \dots u_{K_p}^K : K_p \alpha$, which expanding ρ is of the form $u_1 : T_1, \dots u_n : T_n$. We introduce:

- a record type definition $K \alpha \approx \{u_1 : T_1, \dots u_n : T_n\}$,
- for each i in $1..n$ we define the accessor to field u_i :
 - let N_i be $\Lambda \alpha. \lambda z : K \alpha. (z.u_i)$.
 - let σ_i be $\forall \alpha. K \alpha \Rightarrow T_i$, i.e. the type of N_i
 - let $\mathcal{R}_i$ be the program context *let* $u_i : \sigma_i = N_i$ *in* $[]$.

Then, $\llbracket H \rrbracket$ is $\mathcal{R}_1 \circ \dots \mathcal{R}_n$ and we write Γ_H for the typing environment $u_1 : \sigma_1 \dots u_p : \sigma_p$ in the hole of $\llbracket H \rrbracket$.

Class declarations

Elaboration

The elaboration $\llbracket \vec{H} \rrbracket$ of the sequence of class definitions $\vec{H}$ is the composition of the elaboration of each.

$$\llbracket H_1 \dots H_p \rrbracket \stackrel{\Delta}{=} \llbracket H_1 \rrbracket \circ \dots \circ \llbracket H_p \rrbracket$$

Record type definitions are collected in the program prelude.

We write $\Gamma_{H_1 \dots H_p}$ for $\Gamma_{H_1}, \dots, \Gamma_{H_p}$.

Instance definitions

$$h \triangleq \text{inst } \forall \vec{\beta}. \vec{P} \Rightarrow K (G \vec{\beta}) \{r\}$$

It defines an instance of a class K .

The *typing context* $\vec{P}$ describes the dictionaries that must be available on type parameters $\vec{\beta}$ so that the dictionary $K (G \vec{\beta})$ can be built.

This is not related to the superclasses of the class K !

For example, class K' may be a superclass of K , so the creation of the dictionary $K \alpha$ requires a dictionary $K' \alpha$ while an instance declaration $K G$ (where G is nullary) need not request a dictionary of type $K' G$.

Indeed, either such a dictionary can already be built, hence the instance does not require it, or it will never be possible to build one (remember that instance definitions are recursively defined so all of them are already visible) and the program must be rejected.

Instance declarations

$$h \triangleq \text{inst } \forall \vec{\beta}. K_1 \beta_1, \dots K_k \beta_k \Rightarrow K(G \vec{\beta}) \{r\}$$

In summary, the typing context describes dictionaries that cannot yet be built because they depend on some unknown type β in $\vec{\beta}$.

We assume that the typing context is such that:

- each β_i is in $\vec{\beta}$
- β_i and β_j may be equal, except if $K_i < K_j$ or $K_j < K_i$ or $K_i = K_j$.

The reason for the latter condition is, as for class declarations, that it would be useless to require both dictionaries $K_i \beta$ and $K_j \beta$ when there are equal or one is contained in the other.

Such typing contexts are said to be *canonical*.

Instance declarations

Elaboration

$$h \triangleq \text{inst } \forall \vec{\beta}. K'_1 \beta_1, \dots K'_k \beta_k \Rightarrow K (G \vec{\beta}) \{r\}$$

This instance definition h is elaborated into a triple (z_h, N^h, σ_h) where z_h is an identifier to refer to the elaborated body N^h of type σ_h .

The type σ_h is $\forall \vec{\beta}. K'_1 \beta_1 \Rightarrow \dots K'_k \beta_k \Rightarrow K (G \vec{\beta})$

The expression N^h builds a dictionary of type $K (G \vec{\beta})$, given $k \geq 0$ dictionaries of respective types $K'_1 \beta_1, \dots K'_k \beta_k$:

$$\Lambda \vec{\beta}. \lambda(z_1 : K'_1 \beta_1). \dots \lambda(z_k : K'_k \beta_k). \\ \{u_1 = N_1^h, \dots u_m = N_m^h, u_{K'_1}^K = q_1, \dots u_{K'_p}^K = q_p\}$$

The types of fields are as prescribed by the class definition K :

- N_i^h is the elaboration of M_i where r is $u_1 = M_1, \dots u_m = M_m$.
- q_i is a dictionary of type $K_i (G \vec{\beta})$ (the i 'th subdictionary of K)

(We write z for a variable x that binds a dictionary.)

Elaboration of whole programs

The elaboration of all class instances $\llbracket \vec{h} \rrbracket$ is the program context

$$\text{let rec } (\vec{z}_h : \vec{\sigma}_h) = \vec{N}^h \text{ in } []$$

The elaboration of the whole program $\vec{H} \vec{h} M$ is

$$\llbracket \vec{H} \vec{h} M \rrbracket \triangleq \text{let } \vec{u} : \vec{\sigma}_u = \vec{N}_u \text{ in } \text{let rec } (\vec{z}_h : \vec{\sigma}_h) = \vec{N}^h \text{ in } N$$

Hence, the expression N and all expressions N^h are typed (and elaborated) in the environment Γ_0 equal to $\Gamma_{\vec{H}}, \Gamma_{\vec{h}}$ where

- $\Gamma_{\vec{H}}$ declares functions to access components of dictionaries (both sub-dictionaries and definitions of overloaded symbols).
- $\Gamma_{\vec{h}}$ equal to $(\vec{z}_h : \vec{\sigma}_h)$ declares functions to build dictionaries (i.e. all class instances).

Elaboration of expressions

The elaboration of expressions is defined by a judgment

$$\Gamma \vdash M \rightsquigarrow N : \sigma$$

where Γ is a System-F typing context, M is the source expression, N is the elaborated expression and σ its type in Γ .

In particular, $\Gamma \vdash M \rightsquigarrow N : \sigma$ implies $\Gamma \vdash N : \sigma$ in F .

We write q for dictionary terms, *i.e.* the following subset of F terms:

$$q ::= u \mid z \mid q \tau \mid q q$$

(u and z are just particular cases of x)

The elaboration of dictionaries is the judgment $\Gamma \vdash q : \sigma$ which is just typability in System F—but restricted to dictionary expressions.

Elaboration of expressions

$$\begin{array}{c}
 \text{VAR} \\
 \frac{x : \sigma \in \Gamma}{\Gamma \vdash x \rightsquigarrow x : \sigma}
 \end{array}
 \qquad
 \begin{array}{c}
 \text{INST} \\
 \frac{\Gamma \vdash M \rightsquigarrow N : \forall \alpha. \sigma}{\Gamma \vdash M \tau \rightsquigarrow N \tau : [\alpha \mapsto \tau] \sigma}
 \end{array}
 \qquad
 \begin{array}{c}
 \text{GEN} \\
 \frac{\Gamma, \alpha \vdash M \rightsquigarrow N : \sigma}{\Gamma \vdash \Lambda \alpha. M \rightsquigarrow \Lambda \alpha. N : \forall \alpha. \sigma}
 \end{array}$$

$$\begin{array}{c}
 \text{LET} \\
 \frac{\Gamma \vdash M_1 \rightsquigarrow N_1 : \sigma \quad \Gamma, x : \sigma \vdash M_2 \rightsquigarrow N_2 : \tau}{\Gamma \vdash \text{let } x : \sigma = M_1 \text{ in } M_2 \rightsquigarrow \text{let } x : \sigma = N_1 \text{ in } N_2 : \tau}
 \end{array}$$

$$\begin{array}{c}
 \text{APP} \\
 \frac{\Gamma \vdash M_1 \rightsquigarrow N_1 : \tau_2 \rightarrow \tau_1 \quad \Gamma \vdash M_2 \rightsquigarrow N_2 : \tau_2}{\Gamma \vdash M_1 M_2 \rightsquigarrow N_1 N_2 : \tau_1}
 \end{array}
 \qquad
 \begin{array}{c}
 \text{ABS} \\
 \frac{\Gamma, x : \tau' \vdash M \rightsquigarrow N : \tau}{\Gamma \vdash \lambda x : \tau'. M \rightsquigarrow \lambda x : \tau'. N : \tau' \rightarrow \tau}
 \end{array}$$

In rule [LET](#), we require σ to be canonical, i.e. of the form $\forall \vec{\alpha}. \vec{P} \Rightarrow T$ where $\vec{P}$ is itself empty or canonical. See also [this restriction](#).

Rules [APP](#) and [ABS](#) do not apply to overloaded expressions of type σ .

Elaboration of overloaded expressions

The interesting rules are the elaboration of missing abstractions and applications of dictionaries.

OABS

$$\frac{\Gamma, x : Q \vdash M \rightsquigarrow N : \sigma \quad x \# M}{\Gamma \vdash M \rightsquigarrow \lambda x : Q. N : Q \Rightarrow \sigma}$$

OAPP

$$\frac{\Gamma \vdash M \rightsquigarrow N : Q \Rightarrow \sigma \quad \Gamma \vdash q : Q}{\Gamma \vdash M \rightsquigarrow N \ q : \sigma}$$

Rule **OABS** pushes dictionary abstractions as prescribed by the expected type in the context Γ . These may then be used (in addition to dictionary accessors and instance definitions already in Γ) to elaborate dictionaries as described by the premise $\Gamma \vdash q : Q$ of rule **OAPP**.

Elaboration of overloaded expressions

The interesting rules are the elaboration of missing abstractions and applications of dictionaries.

OABS

$$\frac{\Gamma, x : Q \vdash M \rightsquigarrow N : \sigma \quad x \# M}{\Gamma \vdash M \rightsquigarrow \lambda x : Q. N : Q \Rightarrow \sigma}$$

OAPP

$$\frac{\Gamma \vdash M \rightsquigarrow N : Q \Rightarrow \sigma \quad \Gamma \vdash q : Q}{\Gamma \vdash M \rightsquigarrow N \ q : \sigma}$$

Judgment $\Gamma \vdash q : Q$ is just well-typedness in System F, but restricted to dictionary expressions. There is an algorithmic reading of the rule, described [further](#), where Γ and Q are given and q is inferred.

Elaboration of overloaded expressions

The interesting rules are the elaboration of missing abstractions and applications of dictionaries.

OABS

$$\frac{\Gamma, x : Q \vdash M \rightsquigarrow N : \sigma \quad x \# M}{\Gamma \vdash M \rightsquigarrow \lambda x : Q. N : Q \Rightarrow \sigma}$$

OAPP

$$\frac{\Gamma \vdash M \rightsquigarrow N : Q \Rightarrow \sigma \quad \Gamma \vdash q : Q}{\Gamma \vdash M \rightsquigarrow N \ q : \sigma}$$

By construction, elaboration produces well-typed expressions: that is $\Gamma_0 \vdash M \rightsquigarrow N : \tau$ implies that is $\Gamma_0 \vdash N : \tau$.

Resuming the elaboration

An instance declaration h of the form:

$$\text{inst } \forall \vec{\beta}. K'_1 \beta_1, \dots K'_k \beta_k \Rightarrow K (G \vec{\tau}) \{u_1 = M_1, \dots u_m = M_m\}$$

is translated into

$$\Lambda \vec{\beta}. \lambda(z_1 : K'_1 \beta_1). \dots \lambda(z_k : K'_k \beta_k). \\ \{u_1 = N_1^h, \dots u_m = N_m^h, u_{K_1}^K = q_1, \dots u_{K_p}^K = q_p\}$$

where:

- $u_{K_i}^K : Q_i$ are the superclasses fields, $u_i : \tau_i$ are the method fields
- Γ_h is $\vec{\beta}, K'_1 \beta_1, \dots K'_k \beta_k$
- $\Gamma_0, \Gamma_h \vdash q_i : Q_i$
- $\Gamma_0, \Gamma_h \vdash M_i \rightsquigarrow N_i : \tau_i$

Finally, given the program p equal to $\vec{H} \vec{h} M$, we elaborate M as N such that $\Gamma_0 \vdash M \rightsquigarrow N : \forall \bar{\alpha}. \tau$.

Notice that $\forall \bar{\alpha}. \tau$ is an unconstrained type scheme. Why?

Resuming the elaboration

Otherwise, N could elaborate into an abstraction over dictionaries, *i.e.* it would be a value and never applied!

Where else should we be careful that the *intended* semantics is preserved?

?

Let-monomorphization

Otherwise, N could elaborate into an abstraction over dictionaries, *i.e.* it would be a value and never applied!

Where else should we be careful that the *intended* semantics is preserved?

In a call-by-value setting, we must not elaborate applications into abstractions, since it would delay and perhaps duplicate the order of evaluations.

For that purpose, we must restrict rule **LET** so that either σ is of the form $\forall \bar{\alpha}. \tau$ or M_1 is a value or a variable.

What about call-by-name? and Haskell?

Let-monomorphization

In call-by-name, an application is not evaluated until it is needed. Hence, adding an abstraction in front of an application should not change the evaluation order $M_1 M_2$.

We must in fact compare:

$$\text{let } x_1 = \text{let } x_2 = \lambda y. V_1 V_2 \text{ in } [x_2 \mapsto x_2 q] M_2 \text{ in } M_1 \quad (1)$$

$$\text{let } x_1 = \lambda y. \text{let } x_2 = V_1 V_2 \text{ in } M_2 \text{ in } [x_1 \mapsto x_1 q] M_1 \quad (2)$$

The order of evaluation of $V_1 V_2$ is preserved.

However, Haskell is call-by-need, and not call-by-name!

Hence, applications are delayed as in call-by-name but their evaluation is shared and only reduced once.

The application $V_1 V_2$ will be reduced once in (2), but as many times as there are occurrences of x_2 in M_2 in (1).

Let-monomorphization

The final result will still be the same in both cases because Haskell is pure, but the intended semantics is changed regarding the efficiency.

Hence, Haskell may also use monomorphization in this case. This is a delicate design choice

(Of course, monomorphization reduces polymorphism, hence the set of typable programs.)

Resuming the elaboration

Sources of failures

The elaboration may fail for several reasons:

- The input expression does not obey one of the restrictions we have requested.
- A typing error may occur during elaboration of an expression.
- Some required dictionary cannot be built.

If elaboration fails, the program p is rejected, indeed.

When elaboration succeeds

When the elaboration of p succeeds it returns $\llbracket p \rrbracket$, well-typed in F .

Then, the semantics of p is given by that of $\llbracket p \rrbracket$.

When elaboration succeeds

When the elaboration of p succeeds it returns $\llbracket p \rrbracket$, well-typed in F .

Then, the semantics of p is given by that of $\llbracket p \rrbracket$.

Hum...

When elaboration succeeds

When the elaboration of p succeeds it returns $\llbracket p \rrbracket$, well-typed in F .

Then, the semantics of p is given by that of $\llbracket p \rrbracket$.

Hum... Although terms are explicitly-typed, their elaboration may not be unique! Indeed, there might be several ways to build dictionaries of some given type (see [below](#) for details).

In the worst case, a source program may elaborate to completely unrelated programs. In the best case, all possible elaborations are *equivalent* programs and we say that the elaboration is *coherent*: the programs has a deterministic semantics given by elaboration.

But what does it mean for programs be equivalent?

On program equivalence

There are several notions of program equivalence:

- If programs have a denotational semantics, the equivalence of programs should be the equality of their denotations.
- As a subcase, two programs having a common reduct should definitely be equivalent. However, this will in general not be complete: values may contain functions that are not identical, but perhaps would reduce to the same value whenever applied to the same arguments.
- This leads to the notion of *observational equivalence*. Two expressions are observationally equivalent (at some observable type, such as integers) if their are indistinguishable whenever they are put in arbitrary (well-typed) contexts of the observable type.

On program equivalence

For instance, two different elaborations that would just consistently change the representation of dictionaries (e.g. by ordering records in reverse order), would be equivalent if we cannot observe the representation of dictionaries.

Sufficient conditions for coherence

Since terms are explicitly typed, the only source of non-determinism is the elaboration of dictionaries.

One way to ensure coherence is that two dictionary *values* of the same type are always equal. This does not mean that there is a unique way of building dictionaries, but that all ways are equivalent as they eventually return the same dictionary.

Elaboration of dictionaries

Elaboration of dictionaries is just typing in System F.

More precisely, it infers a dictionary q given Γ and Q so that $\Gamma \vdash q : Q$.

The relevant subset of rules for dictionary expressions are:

$$\begin{array}{c}
 \text{D-OVAR} \\
 \hline
 x : \sigma \in \Gamma \\
 \hline
 \Gamma \vdash x : \sigma
 \end{array}
 \qquad
 \begin{array}{c}
 \text{D-INST} \\
 \hline
 \Gamma \vdash q : \forall \alpha. \sigma \\
 \hline
 \Gamma \vdash q \tau : [\alpha \mapsto \tau] \sigma
 \end{array}
 \qquad
 \begin{array}{c}
 \text{D-APP} \\
 \hline
 \Gamma \vdash q_1 : Q_1 \Rightarrow Q_2 \quad \Gamma \vdash q_2 : Q_1 \\
 \hline
 \Gamma \vdash q_1 q_2 : Q_2
 \end{array}$$

Can we give a type-directed presentation?

Elaboration of dictionaries

Elaboration is driven by the type of the expected dictionary and the bindings available in the typing environment, which may be:

- a dictionary constructor z_h given by an instance definition h ;
- a dictionary accessor $u_K^{K'}$ given by a class declaration K' ;
- a dictionary argument z , given by the local typing context.

Hence, the typing rules may be reorganized as follows:

D-OVAR-INST

$$\frac{z_h : \forall \vec{\beta}. P_1 \Rightarrow \dots P_n \Rightarrow K (G \vec{\beta}) \in \Gamma \quad \Gamma \vdash q_i : [\vec{\beta} \mapsto \vec{\tau}] P_i}{\Gamma \vdash z_h \vec{\tau} \vec{q} : K (G \vec{\tau})}$$

D-PROJ

$$\frac{u_K^{K'} : \forall \alpha. K' \alpha \Rightarrow K \alpha \in \Gamma \quad \Gamma \vdash q : K' \tau}{\Gamma \vdash u_K^{K'} \tau q : K \tau}$$

D-VAR

$$\frac{z : K \alpha \in \Gamma}{\Gamma \vdash z : K \alpha}$$

Elaboration of dictionary values

Dictionary **values** are typed in Γ_0 , which does not contain free type variables, hence, only the first rule applies:

D-OVAR-INST

$$\frac{z_h : \forall \vec{\beta}. P_1 \Rightarrow \dots P_n \Rightarrow K(G \vec{\beta}) \in \Gamma \quad \Gamma \vdash q_i : [\vec{\beta} \mapsto \vec{\tau}] P_i}{\Gamma \vdash z_h \vec{\tau} \vec{q} : K(G \vec{\tau})}$$

This rule for the judgment $\Gamma \vdash q : \tau$ can be read as an algorithm where Γ and Q are inputs (and Γ is constant) and q is an output.

Provided there is no choice in finding $z_h : \forall \vec{\beta}. P_1 \Rightarrow \dots P_n \Rightarrow K(G \vec{\beta}) \in \Gamma$

Since each such clause is coming from an instance definition h , there is no choice in the application of this rule if *instance definitions never overlap*.

This assumption ensures coherence.

Overlapping instances

Two instances $inst \forall \vec{\beta}_i. \vec{P} \Rightarrow K (G_i \vec{\beta}_i) \{r_i\}$ for i in $\{1, 2\}$ of a class K **overlap** if the type schemes $\forall \vec{\beta}_i. K (G_i \vec{\tau}_i)$ have a common instance, *i.e.* in the present setting, if G_1 and G_2 are equal.

Overlapping instances are an inherent source of incoherence: it means that for some type Q (in the common instance), a dictionary of type Q may (possibly) be built using two different implementations.

Elaboration of dictionary arguments

Dictionary expressions, as opposed to dictionary values, will also be built by extracting dictionaries from other dictionaries.

Why?

Elaboration of dictionary arguments

Dictionary expressions, as opposed to dictionary values, will also be built by extracting dictionaries from other dictionaries.

Indeed, in overloaded code, the exact type is not fully known at compile time, hence dictionaries must be passed as arguments, from which superclass dictionaries may (and must, as we forbid to pass both a class and one of its super class dictionary simultaneously) be extracted.

Technically, they are typed in an extension of the typing context Γ_0 which may contain typing assumptions $z : K' \beta$ about dictionaries received as arguments. Hence rules [D-PROJ](#) and [D-VAR](#) may also apply.

Elaboration of dictionary arguments

The elaboration of dictionaries uses the three rules (reminder):

D-OVAR-INST

$$\frac{z : \forall \vec{\beta}. P_1 \Rightarrow \dots P_n \Rightarrow K(G \vec{\beta}) \in \Gamma \quad \Gamma \vdash q_i : [\vec{\beta} \mapsto \vec{\tau}] P_i}{\Gamma \vdash z \vec{\tau} \vec{q} : K(G \vec{\tau})}$$

D-PROJ

$$\frac{u : \forall \alpha. K' \alpha \Rightarrow K \alpha \in \Gamma \quad \Gamma \vdash q : K' \tau}{\Gamma \vdash z \tau q : K \tau}$$

D-VAR

$$\frac{z : K \alpha \in \Gamma}{\Gamma \vdash z : K \alpha}$$

They can be read as a backtracking algorithm.

Elaboration of dictionary arguments

Termination

The proof search always terminates, since premises have smaller Q than the conclusion when using the lexicographic order of first the height of τ , then the reverse order of class inheritance.

If no rule applies, we fail. If rule **D-VAR** applies, the derivation ends with success.

If rule **D-PROJ** applies, the premise is called with a smaller problem since the height is unchanged and $K' \vec{\tau}$ with $K' < K$.

If **D-OVAR-INST** applies, the premises are called at type $K_i \tau_j$ where τ_j is subtype of $\vec{\tau}$, hence of a strictly smaller height.

Elaboration of dictionary arguments

Non determinism

For instance, in the introduction, we defined two instances `eqInt` and `ordInt`, while the later contains an instance of the former.

Hence, a dictionary of type `eqInt` may be obtained:

- directly as `EqInt`, or
- indirectly as `OrdInt.Eq`, by projecting the `Eq` sub-dictionary of class `Ord Int`

In fact, the latter choice could then be reduced at compile time and be equivalent to the first one.

One may enforce determinism by fixing a simple and sensible strategy for elaboration. Restrict the use of rule `D-PROJ` to cases where Q is P —when `D-OVAR-INST` does not apply. However, the extra flexibility is harmless and perhaps useful freedom for the compiler.

Typing dictionaries

Example

In the introductory example Γ_0 is:

$$\begin{array}{lll}
 \text{equal} & \triangleq & u_{\text{equal}} : \forall \alpha. \text{Eq } \alpha \Rightarrow \alpha \rightarrow \alpha \rightarrow \text{bool}, \\
 \text{EqInt} & \triangleq & z_{\text{Eq}}^{\text{Int}} : \text{Eq } \text{int} \\
 \text{EqList} & \triangleq & z_{\text{Eq}}^{\text{List}} : \forall \alpha. \text{Eq } \alpha \Rightarrow \text{Eq } (\text{list } \alpha) \\
 \\
 \text{EqOrd} & \triangleq & u_{\text{Eq}}^{\text{Ord}} : \forall \alpha. \text{Ord } \alpha \Rightarrow \text{Eq } \alpha \\
 \text{lt} & \triangleq & u_{\text{lt}} : \forall \alpha. \text{Ord } \alpha \Rightarrow \alpha \rightarrow \alpha \rightarrow \text{bool}
 \end{array}$$

When elaborating the body of `search`, we have to infer a dictionary for `EqOrd X OrdX` in the local context $X, \text{Ord}X : \text{Ord } X$. Thus, Γ is $\Gamma_0, \alpha, z : \text{Ord } \alpha$ and `EqOrd` is $u_{\text{Eq}}^{\text{Ord}}$. We have:

$$\text{D-PROJ} \frac{\begin{array}{c} \text{D-OVAR-INST} \\ \Gamma \vdash u_{\text{Eq}}^{\text{Ord}} \alpha : \text{Ord } \alpha \rightarrow \text{Eq } \alpha \end{array} \quad \begin{array}{c} \text{D-VAR} \\ \Gamma \vdash z : \text{Ord } \alpha \end{array}}{\Gamma \vdash u_{\text{Eq}}^{\text{Ord}} \alpha z : \text{Eq } \alpha}$$

Contents

- Introduction
- Examples in Mini Haskell
- Mini Haskell
- Implicitly-typed terms
- Variations

What can be left implicit?

Class declarations?

What can be left implicit?

Class declarations must remain explicit:

- They define the structure of dictionaries: a record type definition and its accessors.
- They define the type scheme of overloaded symbols and the class they belong to.

The type of instance declarations?

What can be left implicit?

Class declarations must remain explicit:

- They define the structure of dictionaries: a record type definition and its accessors.
- They define the type scheme of overloaded symbols and the class they belong to.

The type of instance declarations must also remain explicit:

- These are polymorphic recursive definitions, hence their types are mandatory.

What can be left implicit?

Class declarations must remain explicit:

- They define the structure of dictionaries: a record type definition and its accessors.
- They define the type scheme of overloaded symbols and the class they belong to.

The type of instance declarations must also remain explicit:

- These are polymorphic recursive definitions, hence their types are mandatory.

However, all core language expressions (in instance declarations and the final one) can be left implicit.

Example

```

class Eq X { equal : X → X → Bool }
inst Eq Int { equal = (==) }
inst Eq Char { equal = (==) }
inst  $\Lambda(X)$  Eq X  $\Rightarrow$  Eq (List (X))
  { eq =  $\lambda(l_1 : \text{List } X) \lambda(l_2 : \text{List } X)$  match  $l_1, l_2$  with
    | [], []  $\rightarrow$  true | [], - | -, []  $\rightarrow$  false
    |  $h_1::t_1, h_2::t_2 \rightarrow$  eq  $X$   $h_1$   $h_2$  && eq ( $\text{List } X$ )  $t_1$   $t_2$  }

```

```

class Eq (X)  $\Rightarrow$  Ord (X) { lt : X → X → Bool }
inst Ord (Int) { lt = (<) }

```

```

let rec search :  $\forall(X)$  Ord X  $\Rightarrow$  X  $\rightarrow$  List X  $\rightarrow$  Bool =
   $\Lambda(X)$   $\lambda(x : X) \lambda(l : \text{List } X)$ 
    match l with []  $\rightarrow$  false |  $h::t \rightarrow$  eq  $X$  x h || search X x t

```

```

let b = search Int 1 [1; 2; 3];;

```

Example

```

class Eq X { equal : X → X → Bool }
inst Eq Int { equal = (==) }
inst Eq Char { equal = (==) }
inst      Eq X ⇒ Eq (List (X))
  { eq = λ(l1      ) λ(l2      ) match l1, l2 with
    | [],[] → true | [], - | [], - → false
    | h1::t1, h2::t2 → eq    h1 h2 && eq      t1 t2 }

```

```

class Eq (X) ⇒ Ord (X) { lt : X → X → Bool }
inst Ord (Int) { lt = (<) }

```

```

let rec search      =
  λ(x      ) λ(l      )
  match l with [] → false | h::t → equal    x h || search    x t

```

```

let b = search Int 1 [1; 2; 3];;

```

Type inference

The idea is to see dictionary types $K\tau$, which can only appear in type schemes and not in types, as a type constraint to mean “*there exists a dictionary of type $K\alpha$* ”.

Just read $\forall \vec{\alpha}. \vec{P} \Rightarrow \tau$ as the constraint type scheme $\forall \vec{\alpha} [\vec{P}]. \tau$.

We extend constraints with dictionary predicates:

$$C ::= \dots \mid K\tau$$

On ground types a constraint Kt is satisfied if one can build a dictionary of type Kt in the initial environment Γ_0 (that contains all class and instance declarations), *i.e.* formally, if there exists a dictionary expression q such that $\Gamma_0 \vdash q : Kt$.

The satisfiability of class-membership constraints is thus:

$$\frac{K\phi\tau}{\phi \vdash K\tau}$$

Reasoning with class-membership constraints

For every class declaration $class\ K_1\ \alpha_1, \dots K_n\ \alpha_n \Rightarrow K\ \alpha\ \{\rho\}$,

$$K\ \alpha \Vdash K_1\ \alpha_1 \wedge \dots K_n\ \alpha_n \quad (1)$$

This rule allows to decompose any set of simple constraints into a canonical one.

Proof of (1).

Reasoning with class-membership constraints

$$\begin{array}{l} \text{For every class declaration } \textit{class } K_1 \alpha_1, \dots K_n \alpha_n \Rightarrow K \alpha \{ \rho \}, \\ K \alpha \Vdash K_1 \alpha_1 \wedge \dots K_n \alpha_n \end{array} \quad (1)$$

This rule allows to decompose any set of simple constraints into a canonical one.

Proof of (1). Assume $\phi \vdash K \alpha$, i.e. $\Gamma_0 \vdash q : K (\phi \alpha)$ for some q .

From the class declaration, we know that $K \alpha$ is a record type definition that contains fields $u_{K_i}^K$ of type $K_i \alpha_i$. Hence, the dictionary value q contains field values of types $K_i (\phi \alpha)$. Therefore, we have $\phi \vdash K_i \alpha$ for all i in $1..n$, which implies $\phi \vdash K_1 \alpha \wedge \dots K_n \alpha$. □

Reasoning with class-membership constraints

$$\begin{aligned} \text{For every instance definition } inst \forall \vec{\beta}. K_1 \beta_1, \dots K_p \beta_p \Rightarrow K(G \beta) \{r\} \\ K(G \vec{\beta}) \equiv K_1 \beta_1 \wedge \dots K_p \beta_p \end{aligned} \quad (2)$$

This rule allows to decompose all class constraints into simple constraints of the form $K \alpha$.

Proof of (2) ($\dashv$ direction). Assume $\phi \vdash K_i \beta_i$. There exists dictionaries q_i such that $\Gamma_0 \vdash q_i : K_i (\phi \beta_i)$. Hence, $\Gamma_0 \vdash x_h \vec{\beta} q_1 \dots q_p : K(G(\phi \vec{\beta}))$, i.e. $\phi \vdash K(G(\phi \vec{\beta}))$.

($\Vdash$ direction). Assume, $\phi \vdash K(G(\phi \vec{\beta}))$. i.e. there exists a dictionary q such that $\Gamma_0 \vdash q : K(G \phi \vec{\beta})$. By *non-overlapping of instance declarations*, the only way to build such a dictionary is by an application of x_h . Hence, q must be of the form $x_h \vec{\beta} q_1 \dots q_p$ with $\Gamma_0 \vdash q_i : K_i (\phi \beta_i)$, that is, $\phi \vdash K_i \beta_i$ for every i , which implies $\phi \vdash K_1 \beta_1 \wedge \dots K_p \beta_p$.

Reasoning with class-membership constraints

For every instance definition $inst \forall \vec{\beta}. K_1 \beta_1, \dots K_p \beta_p \Rightarrow K(G \beta) \{r\}$

$$K(G \vec{\beta}) \equiv K_1 \beta_1 \wedge \dots K_p \beta_p \quad (2)$$

This rule allows to decompose all class constraints into simple constraints of the form $K \alpha$.

Notice that the equivalence still holds in an open-world assumption where new instance clauses may be added later, because another future instance definition cannot overlap with existing ones.

If overlapping of instances were allowed, the $\Vdash$ direction would not hold. Then, the rewriting rule:

$$K(G \vec{\beta}) \longrightarrow K_1 \beta_1 \wedge \dots K_p \beta_p$$

would still be sound (the right-hand side entails the left-hand side, and thus type inference would infer sound typings), i.e. but not complete (type inference could miss some typings).

Reasoning with class-membership constraints

For every class K and type constructor G for which there is no instance of K ,

$$K (G \vec{\beta}) \equiv \text{false} \quad (3)$$

This rule allows failure to be reported as soon as constraints of the form $K (G \vec{\tau})$ appear and there is not instance of K for G .

Proof of (3). The $\dashv\!\!\vdash$ direction is a tautology, so it suffices to prove the $\Vdash$ direction. By contradiction. Assume $\phi \vdash K (G \vec{\beta})$. This implies the existence of a dictionary q such that $\Gamma_0 \vdash q : K (G (\phi \vec{\beta}))$. Then, there must be some x_h in Γ whose type scheme is of the form $\forall \vec{\beta}. \vec{P} \Rightarrow K (G \vec{\beta})$, i.e. there must be an instance of class K for G .

Notice that this rule does not work in an open world assumption. The rewriting rule

$$K (G \vec{\beta}) \longrightarrow \text{false}$$

would still remain sound but incomplete.

Typing constraints

Constraint generation is as in ML.

A constraint type scheme can always be decomposed into one of the form $\forall \bar{\alpha}[P_1 \wedge P_2].\tau$ where $\text{ftv}(P_1) \in \bar{\alpha}$ and $\text{ftv}(P_2) \# \bar{\alpha}$.

The constraints P_2 can then be extruded in the enclosing context if any, so we are in general left with just P_1 .

Checking well-typedness

To check well-typedness of the program p equal to $\vec{H} \vec{h} a$, we must check that: each expression a_i^h and the expression a are well-typed, in the environment used to elaborate them:

This amounts to checking:

- $\Gamma_0, \Gamma_h \vdash a_i^h : \tau_i^h$ where τ_i^h is given.
Thus, we check that the constraints *def* Γ_0, Γ_h *in* $\langle a_i^h \rangle \leq \tau_i^h \equiv \text{true}$.
- $\Gamma_0 \vdash a : \tau$ for some τ .
Thus, we check that *def* Γ_0 *in* $\exists \alpha. \langle a \rangle \leq \alpha \equiv \text{true}$.

However, ...

Checking well-typedness

To check well-typedness of the program p equal to $\vec{H} \vec{h} a$, we must check that: each expression a_i^h and the expression a are well-typed, in the environment used to elaborate them:

This amounts to checking:

- $\Gamma_0, \Gamma_h \vdash a_i^h : \tau_i^h$ where τ_i^h is given.

Thus, we check that the constraints *def* Γ_0, Γ_h *in* $\langle a_i^h \rangle \leq \tau_i^h \equiv \text{true}$.

- $\Gamma_0 \vdash a : \tau$ for some τ .

Thus, we check that *def* Γ_0 *in* $\exists \alpha. \langle a \rangle \leq \alpha \equiv \text{true}$.

However, ... Typechecking is not sufficient!

Type reconstruction should also return an explicitly-typed term M that can then be elaborated into N .

Type reconstruction

As for ML the resolution strategy for constraints may be tuned to keep persistent constraints from which an explicitly typed term M can be read back.

Back to coherence

When the source language is implicitly-typed, the elaboration from the source language into System F code is the composition of type reconstruction with elaboration of explicitly-typed terms.

That is, a elaborates to N if $\Gamma \vdash a \rightsquigarrow M : \tau$ and $\Gamma \vdash M \rightsquigarrow N : \tau$.

Hence, even if the elaboration is coherent for explicitly-typed terms, this may not be true for implicitly-typed terms.

There are two potential problems:

- The language has principal constrained type schemes, but the elaboration requests unconstrained type schemes.
- Ambiguities could be hidden (and missed) by non principal type reconstructions.

Coherence

Toplevel unresolved constraints

Thanks to the several restrictions on class declarations and instance definitions, the type system has principal constrained schemes (and principal typing reconstructions). However, this does not imply that there are principal *unconstrained* type schemes.

Indeed, assume that the principal constrained type scheme is $\forall \alpha [K \alpha]. \alpha \rightarrow \alpha$ and the typing environment contains two instances of $K \ G_1$ and $K \ G_2$ of class K . Constraint-free instances of this type scheme are $G_1 \rightarrow G_1$ and $G_2 \rightarrow G_2$ but $\forall \alpha. \alpha \rightarrow \alpha$ is certainly not one.

Not only neither choice is principal, but the two choices would elaborate into expressions with different (non-equivalent) semantics.

We must fail in such cases.

Coherence

Toplevel unresolved constraints

This problem may appear while typechecking the final expression a in Γ_0 that request an unconstrained type scheme $\forall \alpha. \tau$

It may also occur when typechecking the body of an instance definition, which requests an explicit type scheme $\forall \vec{\alpha}[\vec{Q}]. \tau$ in Γ_0 or equivalently that request a type τ in $\Gamma_0, \vec{\alpha}, \vec{Q}$.

Coherence

Example of unresolved constraints

```
class Num (X) { 0 : X, (+) : X → X → X }  
inst Num Int { 0 = Int.(0), (+) = Int.(+) }  
inst Num Float { 0 = Float.(0), (+) = Float.(+) }  
let zero = 0 + 0;
```

The type of *zero* or *zero + zero* is $\forall \alpha [Num\ \alpha]. \alpha$ —and several classes are possible for *Num X*. The semantics of the program is undetermined.

```
class Readable (X) { read : descr → X }  
inst Readable (Int) { read = read_int }  
inst Readable (Char) { read = read_char }  
let x = read (open_in())
```

The type of *x* is $\forall \alpha [Readable\ \alpha]. unit \rightarrow \alpha$ —and several classes are possible for *Readable α*. The program is rejected.

Coherence

Inaccessible constraint variables

In the previous examples, the incoherence comes from the obligation to infer type schemes without constraints. A similar problem may occur with isolated constraints in a type scheme.

Assume, for instance, that the elaboration of *let* $x = a_1$ *in* a_2 is *let* $x : \forall \alpha [K \alpha]. \text{int} \rightarrow \text{int} = N_1$ *in* N_2 .

All applications of x in N_2 will lead to an unresolved constraint $K \alpha$ since neither the argument nor the context of this application can determine the value of the type parameter α . Still, a dictionary of type $K \tau$ must be given before N_1 can be executed.

Although x may not be used in N_2 , in which case, all elaborations of the expression may be coherent, we may still raise an error, since an unusable local definition is certainly useless, hence probably a programmer's mistake. The error may then be raised immediately, at the definition site, instead of at every use of x .

Coherence

The open-world view

When there is a single instance $K\ G$ for a class K that appears in an unresolved or isolated constraint $K\ \alpha$, the problem formally disappears, as all possible type reconstructions are coherent.

However, we may still not accept this situation, for modularity reasons, as an extension of the program with another non-overlapping *correct* instance declaration would make the program become ambiguous.

Formally, this amounts to saying that the program must be coherent in its current form, but also in all possible extensions with well-typed class definitions. This is taking an *open-world* view.

On the importance of principal type reconstruction

In the source of incoherence we have seen, some class constraint remained undetermined.

As noticed earlier some (usually arbitrary) less general elaboration would cover the problem—but the source program would remain incoherent.

Hence, in order to detect incoherent (i.e. ambiguous) programs it is essential that type reconstruction is principal.

Once a program has been checked coherent, *i.e.* with no undetermined constraint, based on a principal type reconstruction, can we still use another non principal type reconstruction for its elaboration?

On the importance of principal type reconstruction

In the source of incoherence we have seen, some class constraint remained undetermined.

As noticed earlier some (usually arbitrary) less general elaboration would cover the problem—but the source program would remain incoherent.

Hence, in order to detect incoherent (i.e. ambiguous) programs it is essential that type reconstruction is principal.

Once a program has been checked coherent, *i.e.* with no undetermined constraint, based on a principal type reconstruction, can we still use another non principal type reconstruction for its elaboration?

Yes, indeed, this will preserve the semantics.

This freedom may actually be very useful for optimizations.

On the importance of principal type reconstruction

Consider the program

```
let twice =  $\lambda(x)$  x + x in twice (twice 1)
```

Its principal type reconstruction is:

```
let twice :  $\forall(X)$  [ Num X ] X  $\rightarrow$  X =  $\Lambda(X)$  [Num X]  $\lambda(x)$  x + x in  
twice Int (twice Int) 1
```

which elaborates into

On the importance of principal type reconstruction

Consider the program

```
let twice =  $\lambda(x)$  x + x in twice (twice 1)
```

Its principal type reconstruction is:

```
let twice :  $\forall(X)$  [ Num X ] X  $\rightarrow$  X =  $\Lambda(X)$  [Num X]  $\lambda(x)$  x + x in  
twice Int (twice Int) 1
```

which elaborates into

```
let twice X numX =  $\lambda(x : X)$  x (plus numX) x in  
twice Int NumInt (twice Int NumInt 1)
```

while, avoiding the generalization of twice, it would elaborate into:

```
let twice =  $\lambda(x : \text{Int})$  x (plus NumInt) x in twice (twice 1)
```

where moreover, the *plus NumInt* can be statically reduced.

Overloading by return types

All previous ambiguous examples are overloaded by return types:

- $0 : X$.

The value 0 has an overloaded type that is not constraint by the argument.

- $\text{read} : \text{desc} \rightarrow X$.

The function read applied to some ground type argument will be under specified.

Odersky et al. [1995] suggested to prevent overloading by return types by requesting that overloaded symbols of a class $K\alpha$ have types of the form $\alpha \rightarrow \tau$.

The above examples are indeed rejected by this definition.

Overloading by return types

In fact, disallowing overloading by return types suffices to ensure that all well-typed programs are coherent.

Moreover, untyped programs can then be given a semantics directly (which of course coincides with the semantics obtained by elaboration).

Many interesting examples of overloading fits in this schema.

However, overloading by returns types is also found useful in several cases and Haskell allows it, using default rules to resolve ambiguities.

This is still an arguable design choice in the Haskell community.

Contents

- Introduction
- Examples in Mini Haskell
- Mini Haskell
- Implicitly-typed terms
- Variations

Changing the representation of dictionaries

An overloaded method call u of a class K is elaborated into an application $u\ q$ of u to a dictionary expression q of class K . The function u and the representation of the dictionary are both defined in the elaboration of the class K and need not be known at the call site.

This leaves quite a lot of flexibility in the representation of dictionaries.

For example, we used record data-type definitions to represent dictionaries, but tuples would have been sufficient.

An alternative compilation of type classes

The dictionary passing semantics is quite intuitive and very easy to type in the target language.

However, dictionaries may be replaced by a derivation tree that proves the existence of the dictionary. This derivation tree can be passed around instead of the dictionary and at the call site be used to dispatch to the appropriate implementation of the method.

This has been studied in [Furuse, 2003b].

This can also elegantly be explained as a type preserving compilation of dictionaries called concretization and described in [Pottier and Gauthier, 2006]. It is somehow similar to defunctionalization and also requires that the target language is equipped with GADT (Guarded Abstract Data Types). See the following course.

Multi-parameter type classes

Multi-parameter type classes are of the form

$$\text{class } \vec{P} \Rightarrow K \vec{\alpha} \{ \rho \}$$

where free variables of $\vec{P}$ are in $\vec{\alpha}$.

The current framework can easily be extended to handle multi-parameter type classes.

Example: Collections represented by type C whose elements are of type E can be defined as follows:

```
class Collection C E { find : C → E → Option(E), add : C → E → C }  
inst Collection (List X) X { find = List.find, add = λ(c)λ(e) e::c }  
inst Collection (Set X) X { ... }
```

Type dependencies

However, the class `Collection` does not provide the intended intuition that collections be homogeneous:

```
let add2 c x y = add (add c x) y
add2 :  $\forall (C, E, E')$  Collection C E, Collection C E'  $\Rightarrow C \rightarrow E \rightarrow E' \rightarrow C$ 
```

This definition assumes that collections may be heterogeneous. This may not be intended, and perhaps no instance of heterogeneous collections will ever be provided.

To statically enforce collections to be homogeneous in types, the definition can add a clause to say that the parameter `C` determines the parameter `E`:

```
class Collection C E |  $C \rightarrow E$  { ... }
```

Then, `add2` would enforce E and E' to be equal.

Type dependencies

Type dependencies also reduce overlapping between class declarations.

Hence they allow examples that would have to be rejected if type dependencies could not be expressed.

Associated types

Functional dependencies are being replaced by the notion of associated types.

Associated types allow a class to declare its own type function.

Correspondingly, instance definitions must provide a definition for associated type (in addition to values for overloaded symbols as before).

For example, the `Collection` class becomes a single parameter class with an associated type definition:

```
class Collection E {  
  type C : * → *  
  find : C → E → Option E  
  add : C → E → C  
}  
inst Collection Eq X ⇒ Collection X {type C = List E, ... }  
inst Collection Eq X ⇒ Collection X {type C = Set E, ... }
```

Associated types

Associated types increase the expressiveness of type classes.

Overlapping instances

Example

In practice, overlapping instances may be desired! For example, one could provide a generic implementation of sets provided an ordering relation on elements, but also provide a more efficient version for bit sets.

If overlapping instances are allowed, further rules are needed to disambiguate the resolution of overloading, such as giving priority to rules, or using the most specific match.

However, the semantics depend on some particular resolution strategy and becomes more fragile. See [Jones et al., 1997] for a discussion.

See also [Morris and Jones, 2010] for a recent new proposal.

Overlapping instances

Example

```
inst Eq(X) { equal = (=) }  
inst Eq(Int) { equal = (==) }
```

This elaborates into the creation of a generic dictionary

```
let Eq X : Eq X = { equal = (=) }  
let EqInt : Eq Int = { equal = (==) }
```

Then, *EqInt* or *Eq Int* are two dictionaries of type *Eq Int* but with different implementations.

Restriction that are harder to lift

One may consider removing other restrictions on the class declarations or instance definitions. While some of these generalizations would make sense in theory, they may raise serious difficulties in practice.

For example:

- If constrained type schemes are of the form $K\tau$ instead of $K\alpha$? (which affects all aspects of the language), then it becomes difficult to control the termination of constrained resolution and of the elaboration of dictionaries.
- If a class instance $inst\ \forall\vec{\beta}. \vec{P} \Rightarrow K\tau\ \{\rho\}$ could be such that τ is $G\vec{\tau}$ and not $G\vec{\beta}$, then it would be more difficult to check non-overlapping of class instances.

Methods as overloading functions

One approach to object-orientation is to see methods as over as overloaded functions.

Object then carry class tags that can be used at runtime to find the best matching definition.

This approach has been studied in detail by [Millstein and Chambers, 1999]. See also [Bonniot, 2002, 2005].

Summary

Static overloading is not a solution for polymorphic languages.
Dynamics overloading must be used instead.

Dynamics overloading is a powerful mechanism.

Haskell type classes are a practical, general and powerful solution to dynamic overloading,

Dynamic overloading works relatively well in combination with ML-like type inference.

However, besides the simplest case where every one agrees, useful extensions often come with some drawbacks, and there is not yet an agreement on the best design choices.

The design decisions are often in favor of expressiveness, but losing some of the properties and the canonicity of the minimal design.

Summary

Dynamic overloading is a typical and very elegant use of elaboration.

The programmer could in principle write the elaborated program, build and pass dictionaries explicitly, but this would be cumbersome, tricky, error prone, and it would obfuscate the code.

The elaboration does this automatically, without arbitrary choices (in the minimal design) and with only local transformations that preserve the structure of the source.

Further Reading

For an all-in-one explanation of Haskell-like overloading, see *The essence of Haskell* by Odersky et al.

See also the Jones's monograph *Qualified types: theory and practice*.

For a calculus of overloading see ML& [Castagna, 1997]

Type classes have also been added to Coq [Sozeau and Oury, 2008].
Interestingly, the elaboration of proof terms need not be coherent which makes it a simpler situation for overloading.

Bibliography I

(Most titles have a clickable mark “▷” that links to online versions.)

- ▷ Lennart Augustsson. [Implementing Haskell overloading](#). In *FPCA '93: Proceedings of the conference on Functional programming languages and computer architecture*, pages 65–73, New York, NY, USA, 1993. ACM. ISBN 0-89791-595-X.
- Daniel Bonniot. [Typage modulaire des multi-méthodes](#). PhD thesis, École des Mines de Paris, November 2005.
- ▷ Daniel Bonniot. [Type-checking multi-methods in ML \(a modular approach\)](#). In *Workshop on Foundations of Object-Oriented Languages (FOOL)*, January 2002.
- Giuseppe Castagna. [Object-Oriented Programming: A Unified Foundation](#). Progress in Theoretical Computer Science Series. Birkäuser, Boston, 1997.

Bibliography II

Jun Furuse. [Extensional polymorphism by flow graph dispatching](#). In Ohori [2003], pages 376–393. ISBN 3-540-20536-5.

▷ Jun Furuse. [Extensional polymorphism by flow graph dispatching](#). In *Asian Symposium on Programming Languages and Systems (APLAS)*, volume 2895 of *Lecture Notes in Computer Science*. Springer, November 2003b.

▷ Mark P. Jones. [Simplifying and improving qualified types](#). In *FPCA '95: Proceedings of the seventh international conference on Functional programming languages and computer architecture*, pages 160–169, New York, NY, USA, 1995a. ACM. ISBN 0-89791-719-7.

Mark P. Jones. [Typing Haskell in Haskell](#). In *In Haskell Workshop*, 1999.

Mark P. Jones. [Qualified types: theory and practice](#). Cambridge University Press, New York, NY, USA, 1995b. ISBN 0-521-47253-9.

Bibliography III

- ▷ Simon Peyton Jones, Mark Jones, and Erik Meijer. [Type classes: an exploration of the design space](#). In *Haskell workshop*, 1997.
- Stefan Kaes. [Type inference in the presence of overloading, subtyping and recursive types](#). In *LFP '92: Proceedings of the 1992 ACM conference on LISP and functional programming*, pages 193–204, New York, NY, USA, 1992. ACM. ISBN 0-89791-481-3. doi: <http://doi.acm.org/10.1145/141471.141540>.
- Todd D. Millstein and Craig Chambers. [Modular statically typed multimethods](#). In *ECOOP '99: Proceedings of the 13th European Conference on Object-Oriented Programming*, pages 279–303, London, UK, 1999. Springer-Verlag. ISBN 3-540-66156-5.

Bibliography IV

- J. Garrett Morris and Mark P. Jones. [Instance chains: type class programming without overlapping instances](#). In *ICFP '10: Proceedings of the 15th ACM SIGPLAN international conference on Functional programming*, pages 375–386, New York, NY, USA, 2010. ACM. ISBN 978-1-60558-794-3. doi: <http://doi.acm.org/10.1145/1863543.1863596>.
- ▷ Matthias Neubauer, Peter Thiemann, Martin Gasbichler, and Michael Sperber. [Functional logic overloading](#). pages 233–244, 2002. doi: <http://doi.acm.org/10.1145/565816.503294>.
- ▷ Martin Odersky, Philip Wadler, and Martin Wehr. [A second look at overloading](#). In *FPCA '95: Proceedings of the seventh international conference on Functional programming languages and computer architecture*, pages 135–146, New York, NY, USA, 1995. ACM. ISBN 0-89791-719-7.

Bibliography V

Atsushi Ohori, editor. *Programming Languages and Systems, First Asian Symposium, APLAS 2003, Beijing, China, November 27-29, 2003, Proceedings*, volume 2895 of *Lecture Notes in Computer Science*, 2003. Springer. ISBN 3-540-20536-5.

▷ François Pottier and Nadji Gauthier. *Polymorphic typed defunctionalization and concretization*. *Higher-Order and Symbolic Computation*, 19:125–162, March 2006.

François Rouaix. *Safe run-time overloading*. In *Proceedings of the 17th ACM Conference on Principles of Programming Languages*, pages 355–366, 1990. doi: <http://doi.acm.org/10.1145/96709.96746>.

Geoffrey S. Smith. *Principal type schemes for functional programs with overloading and subtyping*. In *Science of Computer Programming*, 1994.

Bibliography VI

- ▷ Matthieu Sozeau and Nicolas Oury. [First-Class Type Classes](#). In Sofiène Tahar, Otmane Ait-Mohamed, and César Muñoz, editors, *TPHOLs 2008: Theorem Proving in Higher Order Logics, 21th International Conference*, Lecture Notes in Computer Science. Springer, August 2008.
- ▷ Peter J. Stuckey and Martin Sulzmann. [A theory of overloading](#). In *ICFP '02: Proceedings of the seventh ACM SIGPLAN international conference on Functional programming*, pages 167–178, New York, NY, USA, 2002. ACM. ISBN 1-58113-487-8.
- ▷ Philip Wadler and Stephen Blott. [How to make ad-hoc polymorphism less ad-hoc](#). In *ACM Symposium on Principles of Programming Languages (POPL)*, pages 60–76, January 1989.