

Lecture Notes on Compilation

15-417/817: HOT Compilation
Frank Pfenning

Lecture 4
January 23, 2025

1 Introduction

In the last lecture we introduced first-order functional programming (if such a thing exists) based on a computational interpretation of natural deduction for positive equirecursive types. For lack of a better name, we'll call this source language ND. In the lecture before we introduced an imperative intermediate language called SAX, based on a computational interpretation of the semi-axiomatic sequent calculus.

In this lecture we show how to compile the former to the latter. This echoes a proof-theoretic translation from natural deduction to the semi-axiomatic sequent calculus. We will also see our very first optimization!

We also address one of the most annoying parts of programming in ND, namely that, so far, we have to deconstruct data one constructor at a time. The solution is a form of matching with nested patterns.

2 Compiling from ND to SAX

Recall the two typing judgments. First, for ND:

$$\underbrace{x_1 : A_1, \dots, x_n : A_n}_{\text{value variables}} \vdash \underbrace{e : A}_{\text{computation}}$$

where the x_i stands for a large value of type A_i , and e computes a large value of type A . Next, for SAX:

$$\underbrace{x_1 : A_1, \dots, x_n : A_n}_{\text{source addresses}} \vdash P :: \underbrace{d : A}_{\text{destination address}}$$

where the x_i stand for addresses we may (and in the linear case, must) read and d is the destination command P must write a small value to.

It looks possible to keep variables the same, because the large value represented by x_i may be read off the heap at address x_i . On the other hand, the computation of e *returns* a value, while the command P *writes* a value to destination d . This suggests defining a translation taking an expression and a destination, returning a command

$$\llbracket e \rrbracket d = P$$

such that

$$\Gamma \vdash e : A \text{ implies } \Gamma \vdash \llbracket e \rrbracket d :: (d : A)$$

We now go through the few constructs we have one-by-one to see what the translation might be. We should keep in mind the typing requirement above, and the intuition:

If e evaluates to large value V then the heap at destination d after executing $\llbracket e \rrbracket d$ represents V .

Variables. This is easy:

$$\llbracket x \rrbracket d = \mathbf{id} \ d \ x$$

We can verify that this translation types correctly, and also follows our intuition about computation: if the heap at address x represents a value V , then the heap at d will do so afterwards.

Unit 1. The constructor is again straightforward.

$$\llbracket () \rrbracket d = \mathbf{write} \ d \ ()$$

The destructor for type **1** in ND is

$$\mathbf{match} \ e \ \mathbf{with} \ () \Rightarrow e'$$

We can translate this if we allocate a new destination x for the value of e , and then read from x matching against the unit value $()$.

$$\begin{aligned} \llbracket \mathbf{match} \ e \ () \Rightarrow e' \rrbracket d' &= \mathbf{cut} \ x : \mathbf{1} \\ &\quad \llbracket e \rrbracket x \\ &\quad \mathbf{read} \ x \ () \Rightarrow \llbracket e' \rrbracket d' \quad (x \text{ fresh}) \end{aligned}$$

Pairs $A \otimes B$. In this case, constructor expressions (e_1, e_2) require us to allocate two new destinations, one for the value of e_1 and one for the value of e_2 .

$$\begin{aligned} \llbracket (e_1, e_2) \rrbracket d &= \mathbf{cut} \ x_1 : A_1 \\ &\quad \llbracket e_1 \rrbracket x_1 \\ &\quad \mathbf{cut} \ x_2 : A_2 \\ &\quad \llbracket e_2 \rrbracket x_2 \\ &\quad \mathbf{write} \ d \ (x_1, x_2) \quad (x_1 \text{ and } x_2 \text{ fresh}) \end{aligned}$$

One problem we sweep under the rug here is how to obtain the types A_1 and A_2 for e_1 and e_2 . In order to understand this, let's verify that the translation is indeed type preserving as we claim. What we write below would be one case in proving the preservation of types under the translation, where $\llbracket \mathcal{D}_1 \rrbracket$ and $\llbracket \mathcal{D}_2 \rrbracket$ would be obtained by induction hypothesis.

$$\begin{aligned} &\frac{\frac{\mathcal{D}_1}{\Gamma_1 \vdash e_1 \Leftarrow A_1} \quad \frac{\mathcal{D}_2}{\Gamma_2 \vdash e_2 \Leftarrow A_2}}{\Gamma_1, \Gamma_2 \vdash (e_1, e_2) \Leftarrow A_1 \otimes A_2} \otimes I \\ &\quad \rightsquigarrow \\ &\frac{\frac{\llbracket \mathcal{D}_1 \rrbracket}{\Gamma_1 \vdash \llbracket e_1 \rrbracket x_1 :: (x_1 : A_1)} \quad \frac{\frac{\llbracket \mathcal{D}_2 \rrbracket}{\Gamma_2 \vdash \llbracket e_2 \rrbracket x_2 :: (x_2 : A_2)} \quad \frac{}{x_1 : A_1, x_2 : A_2 \vdash \mathbf{write} \ d \ (x_1, x_2) :: (d : A_1 \otimes A_2)} \otimes X}{\frac{}{x_1, \Gamma_2 : A \vdash _ :: (d : A_1 \otimes A_2)} \text{ cut}} \text{ cut} \\ &\quad \Gamma_1, \Gamma_2 \vdash _ :: (d : A_1 \otimes A_2) \end{aligned}$$

Here we have elided some of the commands resulting from the translation for the sake of readability. We see that the translation may really be considered a translation of typing derivations instead of terms. If we follow the structure of bidirectional derivations, the types we need for the cuts (namely A_1 and A_2) will be available to us.

The translation of the destructor follows the previous idea: allocate a new destination x , evaluate the match subject to write a pair to x , and then read the result from x .

$$\begin{aligned} \llbracket \mathbf{match} \ e \ (x_1, x_2) \Rightarrow e' \rrbracket d' &= \mathbf{cut} \ x : A_1 \otimes A_2 \\ &\quad \llbracket e \rrbracket x \\ &\quad \mathbf{read} \ x \ (x_1, x_2) \Rightarrow \llbracket e' \rrbracket d' \quad (x \text{ fresh}) \end{aligned}$$

Again, we can map this onto a translation between typing derivations.

$$\begin{aligned} &\frac{\mathcal{D} \quad \mathcal{E}}{\Gamma \vdash e \Rightarrow A_1 \otimes A_2 \quad \Delta, x_1 : A_1, x_2 : A_2 \vdash e' \Leftarrow A'} \otimes E \\ &\quad \Gamma, \Delta \vdash \mathbf{match} \ e \ \mathbf{with} \ (x_1, x_2) \Rightarrow e' \Leftarrow A' \\ &\quad \rightsquigarrow \\ &\frac{\llbracket \mathcal{D} \rrbracket \quad \frac{\llbracket \mathcal{E} \rrbracket \quad \Delta, x_1 : A_1, x_2 : A_2 \vdash \llbracket e' \rrbracket d' :: (d' : A')}{\Delta, x : A_1 \otimes A_2 \vdash \mathbf{read} \ x \ (x_1, x_2) \Rightarrow \llbracket e' \rrbracket d' :: (d' : A')} \otimes L}{\Gamma \vdash \llbracket e \rrbracket x :: (x : A_1 \otimes A_2) \quad \Delta, x : A_1 \otimes A_2 \vdash \mathbf{read} \ x \ (x_1, x_2) \Rightarrow \llbracket e' \rrbracket d' :: (d' : A')} \text{cut} \\ &\quad \Gamma, \Delta \vdash _ :: (d' : A') \end{aligned}$$

Sums $\oplus\{\ell : A_\ell\}$. There isn't really anything new here, except that each branch of a match expression must be translated. We don't show how to compute the types, but it follows the reasoning for pairs.

$$\begin{aligned} \llbracket k(e) \rrbracket d &= \mathbf{cut} \ x : A_k \\ &\quad \llbracket e \rrbracket x \\ &\quad \mathbf{write} \ d \ k(x) \quad (x \text{ fresh}) \\ \llbracket \mathbf{match} \ e \ \mathbf{with} \ \{\ell(x_\ell) \Rightarrow e_\ell\}_{\ell \in L} \rrbracket d' &= \mathbf{cut} \ x : \oplus\{\ell : A_\ell\}_{\ell \in L} \\ &\quad \llbracket e \rrbracket x \\ &\quad \mathbf{read} \ x \ \{\ell(x_\ell) \Rightarrow \llbracket e_\ell \rrbracket d'\}_{\ell \in L} \quad (x \text{ fresh}) \end{aligned}$$

This already concludes our compiler. Even if we don't give a correctness proof, it doesn't seem like it should be difficult to construct.

3 Some Simple Examples

Say, we have in ND:

```
type nat =  $\oplus\{\underline{\mathbf{zero}} : 1, \underline{\mathbf{succ}} : \mathbf{nat}\}$ 
defn one : nat =  $\underline{\mathbf{succ}}(\underline{\mathbf{zero}}())$ 
```

The type will be the same in SAX, and the definition becomes

```
proc one (d : nat) =  $\llbracket \underline{\mathbf{succ}}(\underline{\mathbf{zero}}()) \rrbracket d$ 
```

It's easy to work out the definition, but we recommend you go through it.

```

[[succ(zero())]] d = cut x1 : nat
                     [[zero()]] x1
                     write d succ(x1)

= cut x1 : nat
  cut x2 : nat
  [()] x2
  write x1 zero(x2)
  write d succ(x1)

= cut x1 : nat
  cut x2 : nat
  write x2 ()
  write x1 zero(x2)
  write d succ(x1)

```

This looks like a quite plausible implementation: we allocate cells on the heap and fill them with the correct small values.

Let's consider the predecessor as something with an elimination rule.

```

defn pred (x : nat) : nat = match x with
| zero(u) => zero(u)
| succ(y) => y

```

This becomes

```

proc pred (d : nat) (x : nat) = [..] d

```

where ... is the match expression above. Following the translation we obtain

```

cut z : nat
  id z x
read z {
| zero(u) => cut w : 1
            id w u
            write d zero(w)
| succ(y) => id d y
}

```

Perhaps the only thing noteworthy here is that two of the three uses of the identity seems unnecessary. These are cuts followed by identities.

In proof theory, it has been observed that cut and identities are opposites and cancel each other out. A cut allows us to use a succedent as an antecedent and the identity allows us to use an antecedent as a succedent. These cancellation laws are

$$\frac{\frac{}{A \vdash A} \text{id} \quad \frac{\mathcal{E}}{\Delta, A \vdash C} \text{cut}}{\Delta, A \vdash C} \rightsquigarrow \frac{\mathcal{E}}{\Delta, A \vdash C}$$

$$\frac{\frac{\mathcal{D}}{\Gamma \vdash A} \quad \frac{}{A \vdash A} \text{id}}{\Gamma \vdash A} \text{cut} \rightsquigarrow \frac{\mathcal{D}}{\Gamma \vdash A}$$

Writing out the SAX commands that are assigned to these derivations, we see hidden renamings.

$$\text{cut } (x : A) (\text{id } x \ y) \ Q(y) \rightsquigarrow Q(x)$$

$$\text{cut } (x : A) \ P(x) (\text{id } y \ x) \rightsquigarrow P(y)$$

It is easy to verify that these optimization not only preserve the types, but also the computational behavior.

One can build these optimizations directly into the translation, or one can design a separate optimization pass that eliminates cut/identity pairs of these two forms. From general principles, we recommend the latter. Keep the translations as simple as possible (and therefore most likely to be correct) and introduce optimizations separately. An advantage of this approach is that it is also easier to measure the impact of an optimization. Our experience is that in the cut/identity optimization are significant. In Lab 2, we may be able to back this up with some numbers.

4 Weak Inversion

From proof search in logical systems we are familiar with the concept of *inversion*. Essentially, a rule is invertible if the premises of a rule are valid whenever the conclusion is. In goal-directed (that is, bottom-up) proof search we can always apply the an invertible rule, reducing the goal of proving a sequent to proving the premises of the rule *without having to consider any alternatives*.

Inversion is mostly formulated in the sequent calculus where all rules are read from the conclusion to the premises. There is also a version for natural deduction which turns out to be an excellent basis for deriving pattern matching that is robust in ways we explain later.

For the positive fragment (which is what we have), we can break down an assumption whenever it is made. Since assumptions are made in each elimination rule we get the following reconstruction of inversion.

$$\frac{\Gamma \vdash A \quad \Delta \vdash A \triangleright C}{\Gamma, \Delta \vdash C} \text{ match}$$

We do not add A to Δ directly, but put it on a “stoop” to be broken down before we return to the usual hypothetical judgment. Already the first rule requires us to generalize our viewpoint:

$$\frac{\Delta \vdash A \cdot B \triangleright C}{\Delta \vdash A \otimes B \triangleright C} \otimes L?$$

We should be able to continue to break down both A and B separately, so instead of a single proposition (which for us will be a type), we have a whole sequence of them. We’ll call these Ω . Then we get

$$\frac{\Delta \vdash A \cdot B \cdot \Omega \triangleright C}{\Delta \vdash A \otimes B \cdot \Omega \triangleright C} \otimes L$$

The unit just disappears.

$$\frac{\Delta \vdash \Omega \triangleright C}{\Delta \vdash \mathbf{1} \cdot \Omega \triangleright C} \mathbf{1}L$$

For sums, we get as many premises as there are summands.

$$\frac{(\Delta \vdash A_\ell \triangleright C) \quad (\forall \ell \in L)}{\Delta \vdash \oplus\{\ell : A_\ell\}_{\ell \in L} \triangleright C} \oplus L$$

When the list of possibly invertible propositions is empty, we revert back to the usual judgment.

$$\frac{\Delta \vdash C}{\Delta \vdash (\cdot) \triangleright C} \text{ empty}$$

As given, the system does not work for recursive proposition (think recursive type like `nat`) because we would have to continue ad infinitum. So we can cut off the inversion phase at any point at our discretion, moving the leftmost proposition into the collection of hypotheses.

$$\frac{\Delta, A \vdash \Omega \triangleright C}{\Delta \vdash A \cdot \Omega \triangleright C} \text{ stop}$$

So, inversion is allowed, but not forced, which is why we call this *weak inversion*.

5 Weak Inversion and Pattern Matching

It turns out that allowing (weak) inversion corresponds to allowing nested patterns in match expressions. This has been made explicit and investigated in depth by [Zeilberger \[2009\]](#). We use a somewhat different notation, more suited to our programming language.

The pattern themselves are symmetric to large values, but can stop at variables.

$$\begin{array}{ll} \text{Patterns} & p, q ::= (p_1, p_2) \mid () \mid k(p) \mid x \\ \text{Pattern Sequences} & \bar{p} ::= p \cdot \bar{p} \mid () \end{array}$$

We need pattern sequences to match the sequences of types Ω (formerly interpreted as propositions). A single branch in a pattern matching expression then has the form

$$p_1 \cdots p_n \Rightarrow e$$

We check a sequence of such branches against a sequence of types, in the judgment

$$\Delta \vdash A_1 \cdots A_n \triangleright (p_1 \cdot p_n \Rightarrow e)^* \Leftarrow C$$

The $()^*$ here indicates that we have a finite sequence of such branches, where each may have different pattern sequences and expressions, but each pattern sequence must have the same length.

Just as the logical inversion judgment distinguishes cases on A_1 , we will do the same, projection out the relevant patterns. We write K for a branch, and K^* for a sequence of branches. filter^* applies the operation to each branch and collects the results.

$$\frac{\Delta \vdash A \cdot B \cdot \Omega \triangleright (\text{filter}^* (_, _) K^*) \Leftarrow C}{\Delta \vdash (A \otimes B) \cdot \Omega \triangleright K^* \Leftarrow C} \otimes L$$

where

$$\text{filter} (_, _) ((x, y) \cdot \bar{q} \Rightarrow e) = x \cdot y \cdot \bar{q} \Rightarrow e$$

The filter operation for pairs results in an error in all other cases.

$$\frac{\Delta \vdash \Omega \triangleright (\text{filter}^* () K^*) \Leftarrow C}{\Delta \vdash \mathbf{1} \cdot \Omega \triangleright K^* \Leftarrow C} \mathbf{1}L$$

where

$$\text{filter } () ((\cdot) \cdot \bar{q} \Rightarrow e) = \bar{q} \Rightarrow e$$

The filter operation for unit results in an error in all other cases.

Finally, the filter for sums.

$$\frac{(\Delta \vdash A_\ell \cdot \Omega \triangleright (\text{filter}^* (\ell(_)) K^*) \Leftarrow C) \quad (\forall \ell \in L)}{\Delta \vdash \oplus \{\ell : A_\ell\}_{\ell \in L} \cdot \Omega \triangleright K^* \Leftarrow C} \oplus L$$

where each application of $\text{filter}^* \ell(_)$ returns either a single branch or none. These are then concatenated.

$$\begin{aligned} \text{filter } \ell(_) (k(p) \cdot \bar{q} \Rightarrow e) &= p \cdot \bar{q} \Rightarrow e \quad \text{for } \ell = k \\ \text{filter } \ell(_) (k(p) \cdot \bar{q} \Rightarrow e) &= (\cdot) \quad \text{for } \ell \neq k \end{aligned}$$

The filter operation for labeled patterns results in an error in all other cases.

Filtering based on sums is quite lenient in the sense that extraneous branches in the pattern match are ignored. Based on the types of the subject of the match, these are unreachable and therefore ignoring them will not lead to a runtime error. On the other hand, it might seem to be a likely error if the programmer has written extraneous branches. A brief further remark on that in the final section of this lecture's notes. When we reach the empty sequence of types, we also must reach the empty sequence of patterns, and revert back to our usual typing rules. Considering we are in the bidirectional system, this is a checking judgment.

$$\frac{\Delta \vdash e \Leftarrow C}{\Delta \vdash (\cdot) \triangleright (\cdot) \Rightarrow e \Leftarrow C} \text{ empty}$$

It is an error if the empty sequence of types is not matched by an empty sequence of patterns. Also note that there can only be a single branch in order to avoid nondeterminism (or redundant patterns, however, one wants to look at it). So we are less lenient here as in the case of extraneous patterns.

Finally, we come to variable patterns which represent a kind of special case because they are not driven by the structure of the type.

$$\frac{\Delta, x : A \vdash \Omega \triangleright \text{filter}^* x K^* \Leftarrow C \quad (x \text{ fresh})}{\Delta \vdash A \cdot \Omega \triangleright K^* \Leftarrow C} \text{ stop}$$

where

$$\text{filter } x (y \cdot \bar{q} \Rightarrow e(y)) = \bar{q} \Rightarrow e(x)$$

Note that the bound variable name y could be different from x , but that we substitute the fresh name x for y . There may be multiple such branches (which are disambiguated later based on Ω as it matches $\bar{q}$), but the pattern in each of them must be a variable.

Overall, this form of pattern matching has a distinct left-to-right flavor in the way it analyzes the match. Moreover, it does not allow some catch-all patterns that can often be useful. The justification for this is three-fold. First, this form of nested checking of patterns can be compiled without contortions to the one-level-at-a-time matching provided by read commands in SAX. Second, when patterns are matched sequentially (including catch-alls for patterns as yet unmatched), assessing linearity of the resulting code is less clear. Finally, under the parallel semantics permitted by SAX, it matters whether cells are read. For example, the two simple matches

```
match a with zero(u) => zero(u)
match a with zero() => zero()
```

have different behaviors! The first can proceed as soon as a small value `zero(b)` is written into the cell with address `a`, while the second must also wait for the unit value to be written into `b`.

At some later point in the course we might consider more general forms of pattern matching. Some investigations in this direction have been made by [Stock \[2020\]](#). For Lab 2, when you implement nested pattern matching, we are lenient as explained above. We do not write out how the definition of projection carries over to compilation. This could be done either as a translation from ND to ND, eliminating nested pattern matches, followed by the translation we gave, or it could be built into a generalized translation from ND to SAX.

6 Subtyping and Intersection Types

We did not cover this in lecture, but one reason one might want to consider allowing extraneous branches is to have nice properties for subtyping. In generally, we would want the following metatheorems to be true:

1. If $\Gamma \vdash e : A$ and $A \leq B$ then $\Gamma \vdash e : B$
2. If $A \leq B$ and $\Gamma, x : B \vdash e : C$ then $\Gamma, x : A \vdash e : C$

The first just transports to expressions the idea that any large value of type A should also have to B if $A \leq B$. The second exploits the same definition, but in reverse because it is used on variables that stand for values. However, the second part will be false if we do not allow extraneous branches. For example, with

```
type nat = +{'zero : 1, 'succ : nat}
type pos = +{'succ : nat}
```

we have $\text{pos} \leq \text{nat}$. The expression

```
match x with (
| 'zero() => 'zero()
| 'succ(y) => y
)
```

checks with $x : \text{nat}$. But unless we allow extraneous branches, it would not check at $x : \text{pos}$, even though $\text{pos} \leq \text{nat}$.

More generally, we might want to assign more than one possible type to a given definition. For example,

```
type nat = +{'zero : 1, 'succ : nat}
type pos = +{'succ : nat}
type zero = +{'zero : 1}
```

```
decl pred (x : nat) : nat
decl pred (x : pos) : nat
decl pred (x : zero) : zero
```

```
defn pred x = match x with (  
  | 'zero() => 'zero()  
  | 'succ(y) => y  
)
```

When checking the second type, we do not consider the zero branch, when checking the third type, we do not consider the succ branch. And yet, it is perfectly valid to assign all three types to the given definition of `pred`.

Assigning multiple types to the same expression is the domain of *intersection types*. If they are considered only at the top level as here, we do not form any explicit intersections, but some of the issues and solutions remain the same.

References

- Benedikt Stock. General pattern matching for session-typed concurrent programs. Bachelor Thesis, Jacobs University, Bremen, Germany, May 2020.
- Noam Zeilberger. *The Logical Basis of Evaluation Order and Pattern-Matching*. PhD thesis, Department of Computer Science, Carnegie Mellon University, May 2009. Available as Technical Report CMU-CS-09-122.