

RISE

A functional pattern-based language in MLIR

Martin Lücke | Michel Steuwer | Aaron Smith



THE UNIVERSITY
of EDINBURGH

supported by a Google Faculty Award sponsored by Jacques Pienaar and Albert Cohen

Why RISE?

Machine Learning Systems are Stuck in a Rut

[HotOS'19]

Paul Barham
Google Brain

Michael Isard
Google Brain

Abstract

In this paper we argue that systems for numerical computing are stuck in a local basin of performance and programmability. Systems researchers are doing an excellent job improving the performance of 5-year-old benchmarks, but gradually making it harder to explore innovative machine learning research ideas.

We explain how the evolution of hardware accelerators favors compiler back ends that hyper-optimize large monolithic kernels, show how this reliance on high-performance but inflexible kernels reinforces the dominant style of programming model, and argue these programming abstractions lack expressiveness, maintainability, and modularity; all of which hinders research progress.

We conclude by noting promising directions in the field, and advocate steps to advance progress towards high-performance general purpose numerical computing systems on modern accelerators.

ACM Reference Format:

Paul Barham and Michael Isard. 2019. Machine Learning Systems are Stuck in a Rut. In *Workshop on Hot Topics in Operating Systems (HotOS '19)*, May 13–15, 2019, Bertinoro, Italy. ACM, New York, NY, USA, 7 pages. <https://doi.org/10.1145/3317550.3321441>

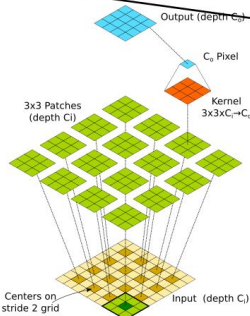


Figure 1. Conv2D operation with 3x3 kernel, stride=2

with 16 times fewer training parameters than the convolutional neural network (CNN) we were comparing it to, implementations in both TensorFlow[2] and PyTorch[3] were much slower and ran out of memory with much smaller models. We wanted to understand why.

1.1 New ideas often require new primitives

We won't discuss the full details of Capsule networks in this paper¹, but for our purposes it is sufficient to consider a simplified form of the inner loop, which is

Original authors
of TensorFlow



Why RISE?

Machine Learning Systems are Stuck in a Rut

Paul Barham
Google Brain

Michael Isard
Google Brain

Abstract

In this paper we argue that systems for numerical computing are stuck in a local basin of performance and programmability. Systems researchers are doing an excellent job improving the performance of 5-year-old benchmarks, but gradually making it harder to explore innovative machine learning research ideas.

We explain how the evolution of hardware accelerators favors compiler back ends that **hyper-optimize large monolithic kernels**, show how this reliance on high-performance but **inflexible** kernels reinforces the dominant style of programming model, and argue these programming abstractions lack expressiveness, maintainability, and **modularity**; all of which **hinders research progress**.

We conclude by noting promising directions in the field, and advocate steps to advance progress towards high-performance general purpose numerical computing systems on modern accelerators.

ACM Reference Format:

Paul Barham and Michael Isard. 2019. Machine Learning Systems are Stuck in a Rut. In *Workshop on Hot Topics in Operating Systems (HotOS '19)*, May 13–15, 2019, Bertinoro, Italy. ACM, New York, NY, USA, 7 pages. <https://doi.org/10.1145/3317550.3321441>

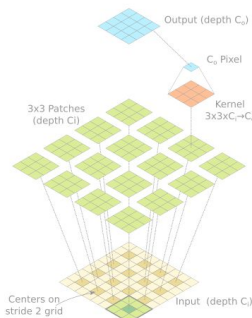


Figure 1. Conv2D operation with 3×3 kernel, stride=2

with 16 times fewer training parameters than the convolutional neural network (CNN) we were comparing it to, implementations in both TensorFlow[2] and PyTorch[3] were much slower and ran out of memory with much smaller models. We wanted to understand why.

1.1 New ideas often require new primitives

We won't discuss the full details of Capsule networks in this paper¹, but for our purposes it is sufficient to consider a simplified form of the inner loop, which is

Why RISE?

Machine Learning Systems are Stuck in a Rut

Paul Barham
Google Brain

Michael Isard
Google Brain

Abstract

In this paper we argue that systems for numerical computing are stuck in a local basin of performance and programmability. Systems researchers are doing an excellent job improving the performance of 5-year-old benchmarks, but gradually making it harder to explore innovative machine learning research ideas.

We explain how the evolution of hardware accelerators favors compiler back ends that **hyper-optimize large monolithic kernels**, show how this reliance on high-performance but **inflexible** kernels reinforces the dominant style of programming model, and argue these programming abstractions lack expressiveness, maintainability, and **modularity**; all of which **hinders research progress**.

We conclude by noting promising directions in the field and advocate steps to advance progress towards

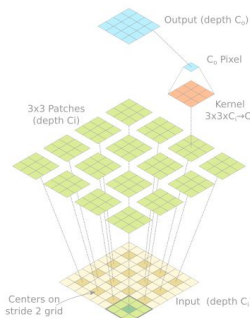


Figure 1. Conv2D operation with 3x3 kernel, stride=2 with 16 times fewer training parameters than the convo-

We should aim for more principled higher level intermediate representations

Paul Barham and Michael Isard. 2019. Machine Learning Systems are Stuck in a Rut. In *Workshop on Hot Topics in Operating Systems (HotOS '19)*, May 13–15, 2019, Bertinoro, Italy. ACM, New York, NY, USA, 7 pages. <https://doi.org/10.1145/3317550.3321441>

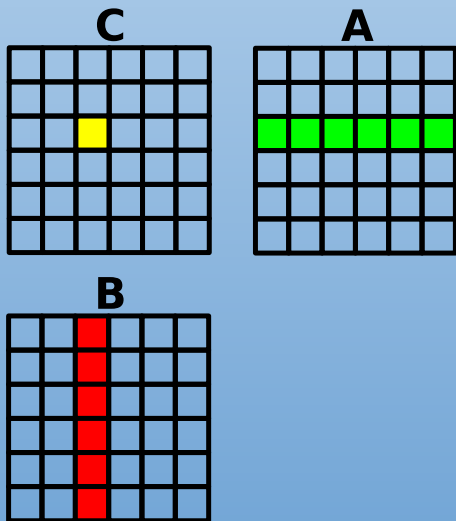
smaller models. we wanted to understand why.

1.1 New ideas often require new primitives

We won't discuss the full details of Capsule networks in this paper¹, but for our purposes it is sufficient to consider a simplified form of the inner loop, which is

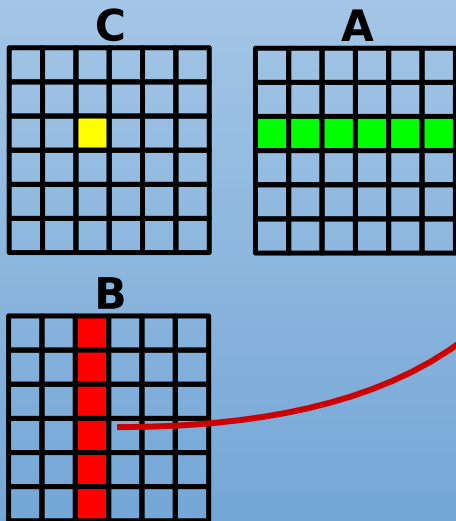
RISE by example: Matrix Multiplication

```
fun(A : N.K.float ⇒ fun(B : K.M.float ⇒  
  A ▷ map(fun(arow ⇒  
    B ▷ transpose ▷ map(fun(bcol ⇒  
      zip(arow, bcol) ▷ map(*) ▷ reduce(+, 0) )) )) ))
```



RISE by example: Matrix Multiplication

```
fun(A : N.K.float => fun(B : K.M.float =>
  A ▷ map(fun(arow =>
    B ▷ transpose ▷ map(fun(bcol =>
      zip(arow, bcol) ▷ map(*) ▷ reduce(+, 0) )) )) ))
```

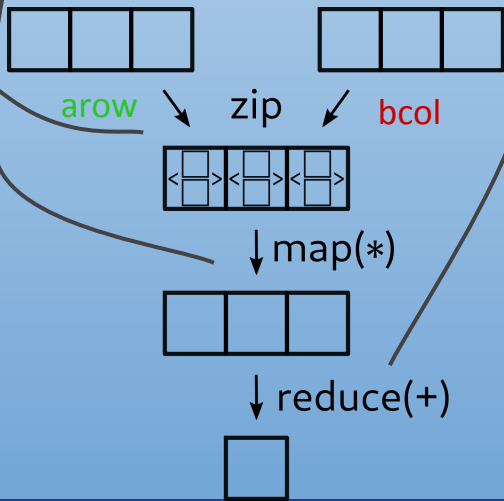
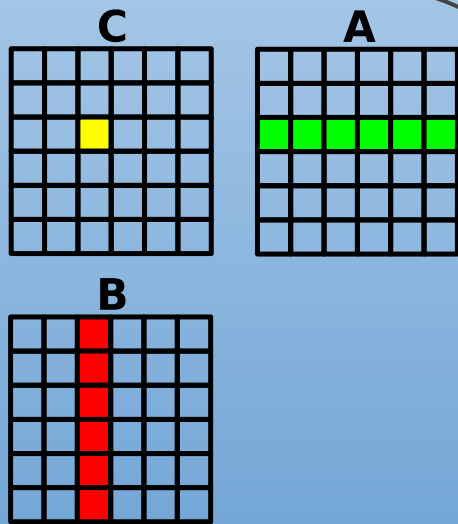


dot product computation:

$$\sum arow_i * bcol_i$$

RISE by example: Matrix Multiplication

```
fun(A : N.K.float ⇒ fun(B : K.M.float ⇒  
  A ▷ map(fun(arow ⇒  
    B ▷ transpose ▷ map(fun(bcol ⇒  
      zip(arow, bcol) ▷ map(*) ▷ reduce(+, 0) )) )) ))
```



Lessons from Lift

- **RISE** (<https://rise-lang.org/>) is a spiritual successor to the Lift project
- **Lift has demonstrated:**
 - optimizing by rewriting
 - achieving high performance
 - performance portability

Lift: Performance Portability

```
fun(A : N.K.float ⇒ fun(B : K.M.float ⇒
  A ▷ map(fun(arow ⇒
    B ▷ transpose ▷ map(fun(bcol ⇒
      zip(arow, bcol) ▷ map(*) ▷ reduce(+, 0) )) )) ))
```

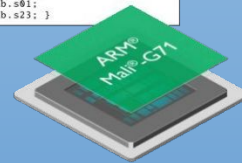
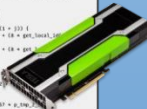


Lift compiler

```
kernel = void __global float * A, B, C;
const float L_row_stride = 1;
private float row_stride;
private float row_stride;
...
for (int i = 0; i < N; i++) {
  for (int j = 0; j < M; j++) {
    float sum = 0.0f;
    for (int k = 0; k < K; k++) {
      sum += A[i * L_row_stride + k] * B[k * M + j];
    }
    C[i * L_row_stride + j] = sum;
  }
}
```

```
kernel void mm(global float* const A,
               global float* const B,
               global float* const C,
               const float row_stride,
               const float col_stride) {
  int i = blockIdx.x * blockDim.x + threadIdx.x;
  int j = blockIdx.y * blockDim.y + threadIdx.y;
  int k = blockIdx.z * blockDim.z + threadIdx.z;
  if (i >= N || j >= M || k >= K) return;
  float sum = 0.0f;
  for (int l = 0; l < L; l++) {
    sum += A[i * L_row_stride + l] * B[l * M + j];
  }
  C[i * L_row_stride + j] = sum;
}
```

```
kernel void mm(global float4* const A,
               global float4* const B,
               global float4* const C,
               const float row_stride,
               const float col_stride) {
  int i = blockIdx.x * blockDim.x + threadIdx.x;
  int j = blockIdx.y * blockDim.y + threadIdx.y;
  int k = blockIdx.z * blockDim.z + threadIdx.z;
  if (i >= N || j >= M || k >= K) return;
  float4 sum = {0.0f, 0.0f, 0.0f, 0.0f};
  for (int l = 0; l < L; l++) {
    sum += A[i * L_row_stride + l] * B[l * M + j];
  }
  C[i * L_row_stride + j] = sum;
}
```



Lift: Optimization by Rewriting

```
fun(A : N.K.float ⇒ fun(B : K.M.float ⇒
  A ▷ map(fun(arow ⇒
    B ▷ transpose ▷ map(fun(bcol ⇒
      zip(arow, bcol) ▷ map(*) ▷ reduce(+, 0) )) )) ))
```

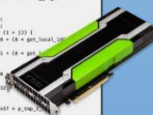


Lift compiler

```
kernel = void void global float * A, B, C;
local float (*row)(int i, int k);
private float row0000_0; ... row0000_7;
private float row0000_8; ... row0000_15;
...
for (int i = 0; i < N; i++) {
  for (int k = 0; k < M; k++) {
    row0000_0 = 0.0f;
    for (int j = 0; j < K; j++) {
      row0000_0 += A[i * K + j] * B[j * M + k];
    }
    C[i * M + k] = row0000_0;
  }
}
```



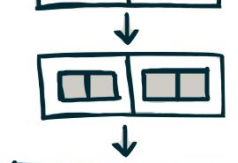
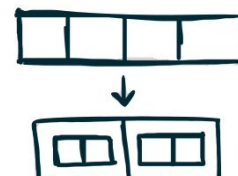
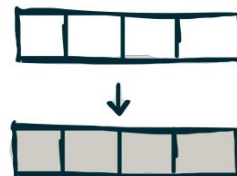
```
kernel void mm(global float* const A,
               global float* const B,
               global float* const C,
               int N, int M, int K) {
  local float (*row)(int i, int k);
  private float row0000_0; ... row0000_7;
  private float row0000_8; ... row0000_15;
  ...
  for (int i = 0; i < N; i++) {
    for (int k = 0; k < M; k++) {
      row0000_0 = 0.0f;
      for (int j = 0; j < K; j++) {
        row0000_0 += A[i * K + j] * B[j * M + k];
      }
      C[i * M + k] = row0000_0;
    }
  }
}
```



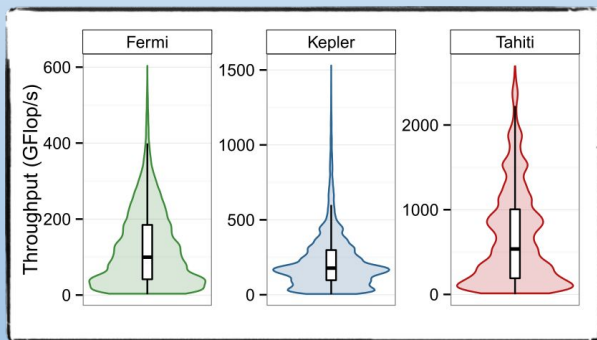
```
kernel void mm(global float* const A,
               global float* const B,
               global float* const C,
               int N, int M, int K) {
  local float (*row)(int i, int k);
  private float row0000_0; ... row0000_7;
  private float row0000_8; ... row0000_15;
  ...
  for (int i = 0; i < N; i++) {
    for (int k = 0; k < M; k++) {
      row0000_0 = 0.0f;
      for (int j = 0; j < K; j++) {
        row0000_0 += A[i * K + j] * B[j * M + k];
      }
      C[i * M + k] = row0000_0;
    }
  }
}
```

Rewrite Rules

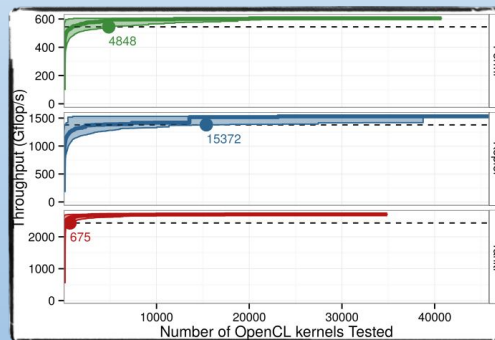
$\text{map}(f, A) \mapsto \text{join}(\text{map}(\text{map}(f), \text{split}(n, A)))$



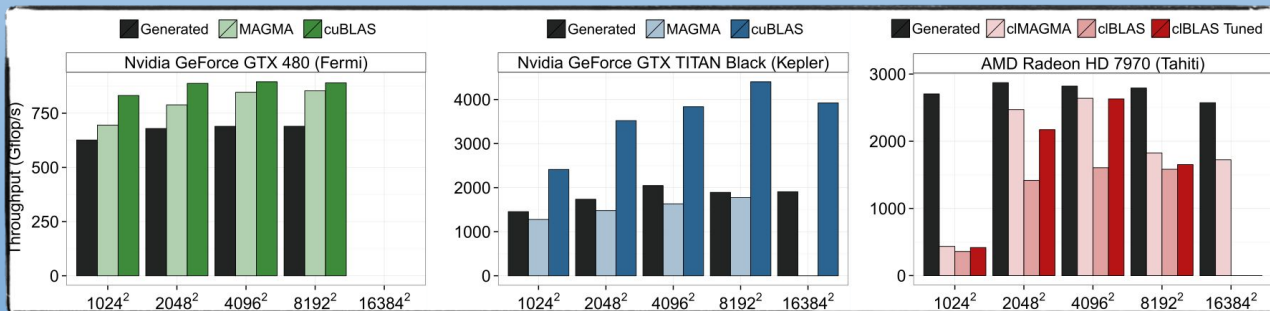
Lift High Performance Results for Matrix Multiplication



Only few generated code with very good performance



Still: One can expect to find a good performing kernel quickly!



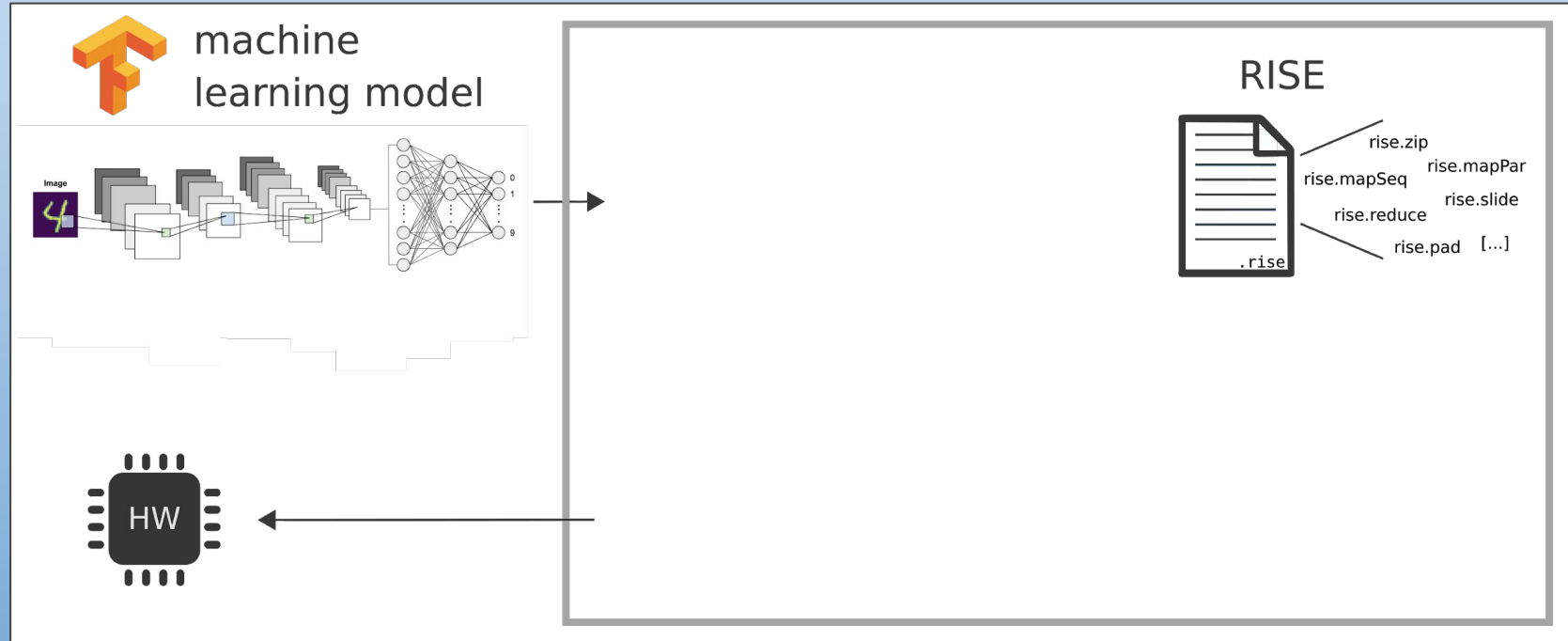
Performance close or better than hand-tuned library code

[GPGPU'16]

Lift - The Caveats

- Does everyone have to write functional programs now?
- Academic work written in Scala
- Does not integrate well with existing compiler infrastructures

How can we achieve end-to-end integration with existing solutions?





MLIR - Multi-Level Intermediate Representation

- Extensible infrastructure to define compiler intermediate representations (dialects)
- Dialects can capture different levels of abstraction:
 - High-level domain specific ----- Hardware specific backend
- Existing dialects available for:
 - TensorFlow / TensorFlow Lite
 - Targeting GPUs
 - ...
 - Performing polyhedral optimizations
 - LLVM IR



MLIR - Martin Lücke Intermediate Representation

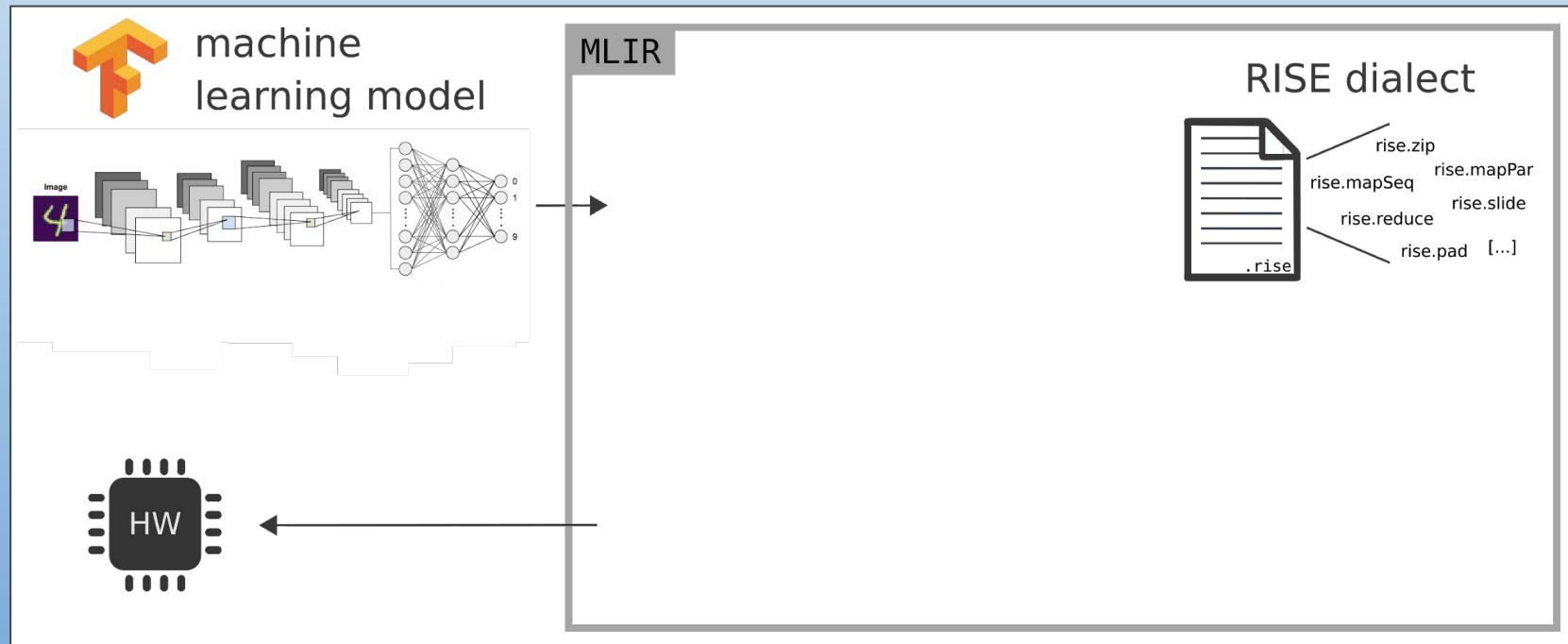
- Extensible infrastructure to define compiler intermediate representations (dialects)
- Dialects can capture different levels of abstraction:
 - High-level domain specific ----- Hardware specific backend
- Existing dialects available for:
 - TensorFlow / TensorFlow Lite
 - Targeting GPUs
 - ...
 - Performing polyhedral optimizations
 - LLVM IR



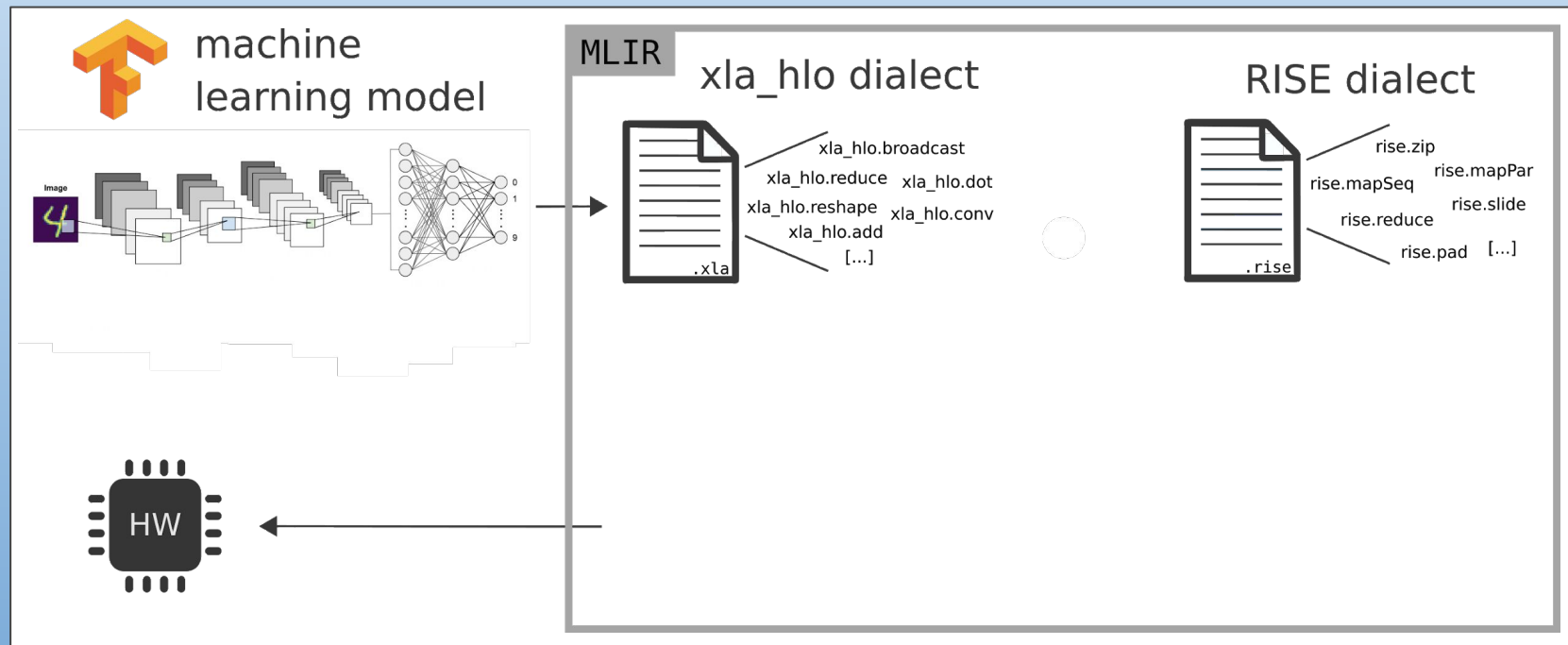
MLIR - Multi-Level Intermediate Representation

- Extensible infrastructure to define compiler intermediate representations (dialects)
- Dialects can capture different levels of abstraction:
 - High-level domain specific ----- Hardware specific backend
- Existing dialects available for:
 - TensorFlow / TensorFlow Lite
 - Targeting GPUs
 - ...
 - Performing polyhedral optimizations
 - LLVM IR

How can we achieve end-to-end integration with existing solutions?



How can we achieve end-to-end integration with existing solutions?



Lowering TensorFlow models to RISE

1. Lower TensorFlow model to XLA_HLO dialect
2. Match for supported operations
3. Replace operations with corresponding RISE operations

```
func @mnist_predict(%input: tensor<1x28x28x1xf32>) → tensor<1x10xf32> {  
  %1 = hlo.reshape(%input) : (tensor<1x28x28x1xf32>) → tensor<1x784xf32>  
  %2 = hlo.dot(%1, %kernel) : (tensor<1x784xf32>, tensor<784x128xf32>) → tensor<1x128xf32>  
  %3 = hlo.add(%2, %bias) : (tensor<1x128xf32>, tensor<128xf32>) → tensor<1x128xf32>  
  [...]   
  %4 = hlo.dot(%3, %kernel_2) : (tensor<1x128xf32>, tensor<128x10xf32>) → tensor<1x10xf32>  
  %5 = hlo.add(%4, %bias_2) : (tensor<1x10xf32>, tensor<10xf32>) → tensor<1x10xf32>  
  [...]   
  return %5 : tensor<1x10xf32>  
}
```

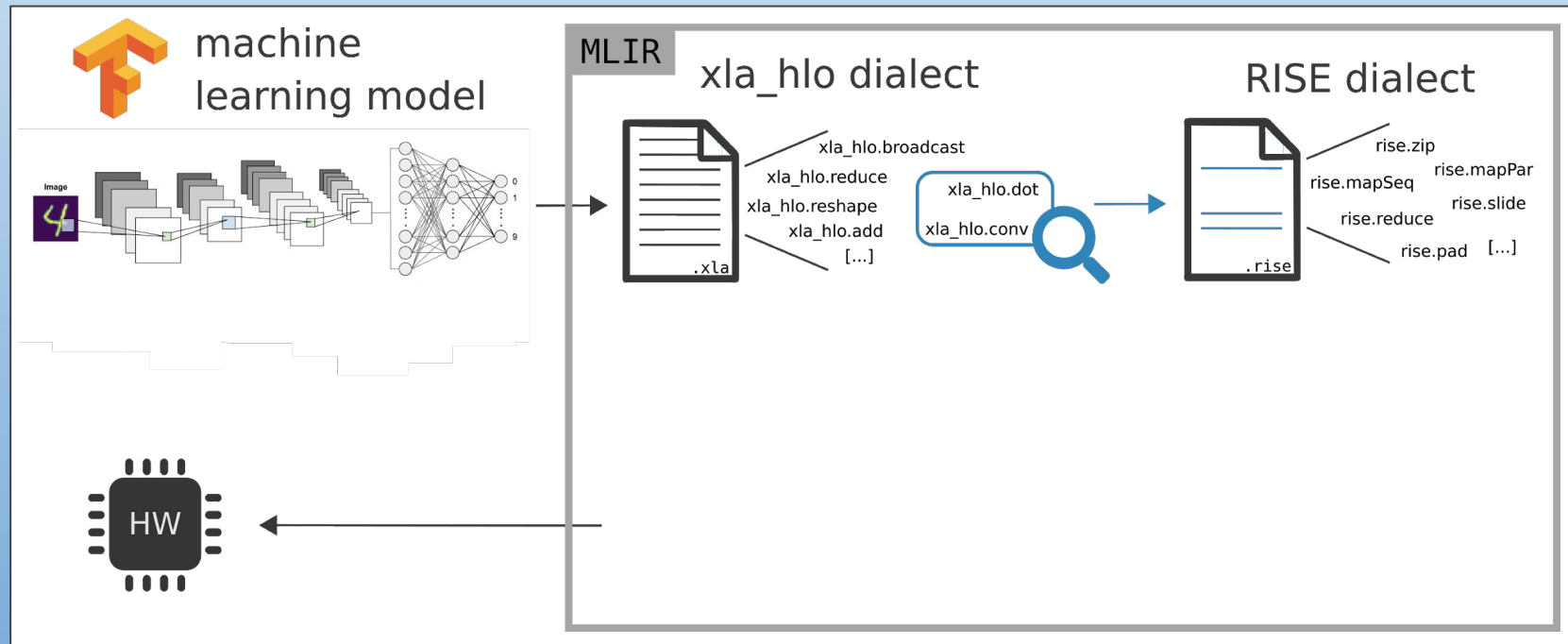
xla_hlo.dot

xla_hlo.conv

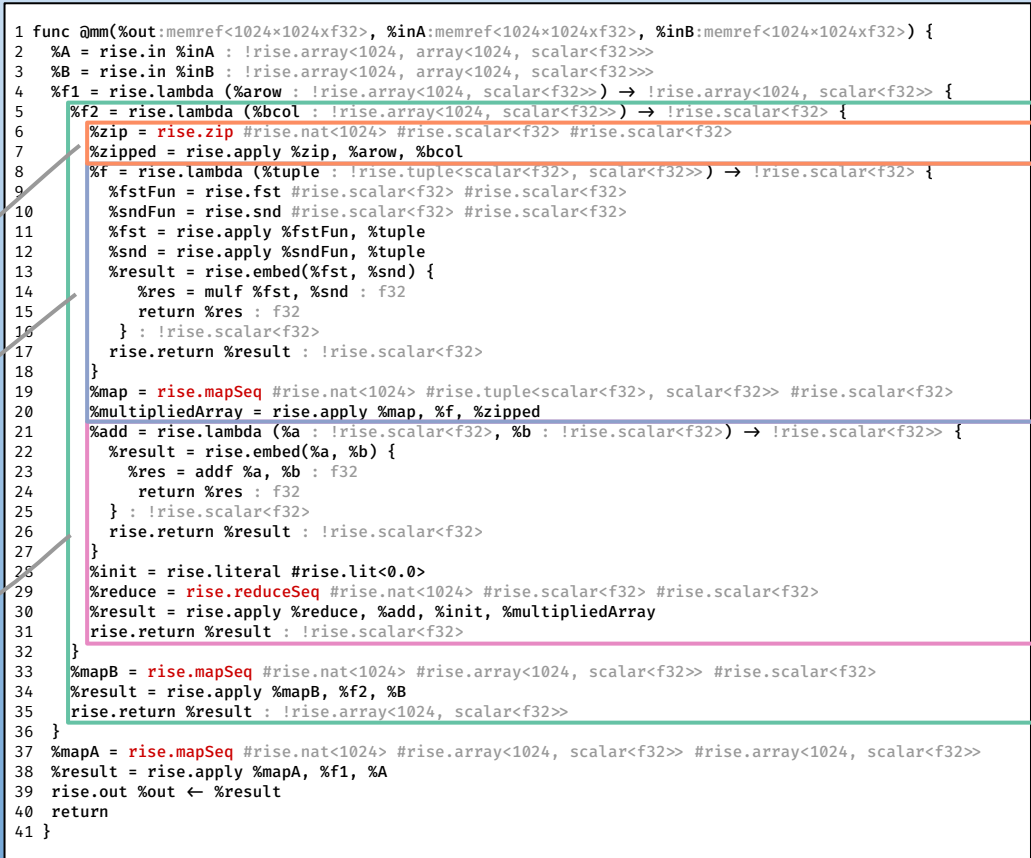
%1 ▷ map(**fun**(arow ⇒
 %kernel ▷ transpose ▷ map(**fun**(bcol ⇒
 zip(arow, bcol) ▷ map(*) ▷ reduce(+, 0)))))

%3 ▷ map(**fun**(arow ⇒
 %kernel_2 ▷ transpose ▷ map(**fun**(bcol ⇒
 zip(arow, bcol) ▷ map(*) ▷ reduce(+, 0)))))

How can we achieve end-to-end integration with existing solutions?



```
%A > map(fun(arow =>
  %B > transpose > map(fun(bcol =>
    zip(arow, bcol) > map(*) > reduce(+, 0) )) ))
```



RISE dialect by example: Matrix Multiplication

```
1 func @mm(%out:memref<1024x1024xf32>, %inA:memref<1024x1024xf32>, %inB:memref<1024x1024xf32>) {  
2   %A = rise.in %inA : !rise.array<1024, array<1024, scalar<f32>>>  
3   %B = rise.in %inB : !rise.array<1024, array<1024, scalar<f32>>>  
4   %f1 = rise.lambda (%arow : !rise.array<1024, scalar<f32>>) → !rise.array<1024, scalar<f32>> {  
5     %f2 = rise.lambda (%bcol : !rise.array<1024, scalar<f32>>) → !rise.scalar<f32> {  
6       %zip = rise.zip #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>  
7       %zipped = rise.apply %zip, %arow, %bcol  
8       %f = rise.lambda (%tuple : !rise.tuple<scalar<f32>, scalar<f32>>) → !rise.scalar<f32> {  
9         %fstFun = rise.fst #rise.scalar<f32> #rise.scalar<f32>  
10        %sndFun = rise.snd #rise.scalar<f32> #rise.scalar<f32>  
11        %fst = rise.apply %fstFun, %tuple  
12        %snd = rise.apply %sndFun, %tuple  
13        %result = rise.embed(%fst, %snd) {  
14          %res = mulf %fst, %snd : f32  
15          return %res : f32  
16        } : !rise.scalar<f32>  
17        rise.return %result : !rise.scalar<f32>  
18      }  
19      %map = rise.mapSeq #rise.nat<1024> #rise.tuple<scalar<f32>, scalar<f32>> #rise.scalar<f32>  
20      %multipliedArray = rise.apply %map, %f, %zipped  
21      %add = rise.lambda (%a : !rise.scalar<f32>, %b : !rise.scalar<f32>) → !rise.scalar<f32>> {  
22        %result = rise.embed(%a, %b) {  
23          %res = addf %a, %b : f32  
24          return %res : f32  
25        } : !rise.scalar<f32>  
26        rise.return %result : !rise.scalar<f32>  
27      }  
28      %init = rise.literal #rise.lit<0.0>  
29      %reduce = rise.reduceSeq #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>  
30      %result = rise.apply %reduce, %add, %init, %multipliedArray  
31      rise.return %result : !rise.scalar<f32>  
32    }  
33    %mapB = rise.mapSeq #rise.nat<1024> #rise.array<1024, scalar<f32>> #rise.scalar<f32>  
34    %result = rise.apply %mapB, %f2, %B  
35    rise.return %result : !rise.array<1024, scalar<f32>>  
36  }  
37  %mapA = rise.mapSeq #rise.nat<1024> #rise.array<1024, scalar<f32>> #rise.array<1024, scalar<f32>>  
38  %result = rise.apply %mapA, %f1, %A  
39  rise.out %out ← %result  
40  return  
41 }
```

Data Types

Function Types

Types

Data Types

```
2  %A = rise.in %inA : !rise.array<1024, array<1024, scalar<f32>>>
```

- Rise data types: array types, tuple types, scalar types
- Nested array types represent higher dimensional data
- Scalar wraps an arbitrary scalar MLIR type

Function Types

```
8  %f = rise.lambda (%t : !rise.tuple<scalar<f32>, scalar<f32>>) → scalar<f32>  
   // %f : !rise.fun<tuple<scalar<f32>, scalar<f32>> → scalar<f32>>
```

- Rise function types: types of lambda expressions and functional patterns
- Type system prevents mixing of function and data types

Patterns

```
1 func @mm(%out:memref<1024×1024xf32>, %inA:memref<1024×1024xf32>, %inB:memref<1024×1024xf32>) {
2   %A = rise.in %inA : !rise.array<1024, array<1024, scalar<f32>>>
3   %B = rise.in %inB : !rise.array<1024, array<1024, scalar<f32>>>
4   %f1 = rise.lambda (%arow : !rise.array<1024, scalar<f32>>) → !rise.array<1024, scalar<f32>> {
5     %f2 = rise.lambda (%bcol : !rise.array<1024, scalar<f32>>) → !rise.scalar<f32> {
6       %zip = rise.zip #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>
7       %zipped = rise.apply %zip, %arow, %bcol
8       %f = rise.lambda (%tuple : !rise.tuple<scalar<f32>, scalar<f32>>) → !rise.scalar<f32> {
9         %fstFun = rise.fst #rise.scalar<f32> #rise.scalar<f32>
10        %sndFun = rise.snd #rise.scalar<f32> #rise.scalar<f32>
11        %fst = rise.apply %fstFun, %tuple
12        %snd = rise.apply %sndFun, %tuple
13        %result = rise.embed(%fst, %snd) {
14          %res = mulf %fst, %snd : f32
15          return %res : f32
16        } : !rise.scalar<f32>
17        rise.return %result : !rise.scalar<f32>
18      }
19      %map = rise.mapSeq #rise.nat<1024> #rise.tuple<scalar<f32>, scalar<f32>> #rise.scalar<f32>
20      %multipliedArray = rise.apply %map, %f, %zipped
21      %add = rise.lambda (%a : !rise.scalar<f32>, %b : !rise.scalar<f32>) → !rise.scalar<f32> {
22        %result = rise.embed(%a, %b) {
23          %res = addf %a, %b : f32
24          return %res : f32
25        } : !rise.scalar<f32>
26        rise.return %result : !rise.scalar<f32>
27      }
28      %init = rise.literal #rise.lit<0.0>
29      %reduce = rise.reduceSeq #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>
30      %result = rise.apply %reduce, %add, %init, %multipliedArray
31      rise.return %result : !rise.scalar<f32>
32    }
33    %mapB = rise.mapSeq #rise.nat<1024> #rise.array<1024, scalar<f32>> #rise.scalar<f32>
34    %result = rise.apply %mapB, %f2, %B
35    rise.return %result : !rise.array<1024, scalar<f32>>
36  }
37  %mapA = rise.mapSeq #rise.nat<1024> #rise.array<1024, scalar<f32>> #rise.array<1024, scalar<f32>>
38  %result = rise.apply %mapA, %f1, %A
39  rise.out %out ← %result
40  return
41 }
```

Patterns: zip

```
6 %zip = rise.zip #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>
```

Patterns: zip

```
6 %zip = rise.zip #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>

// type: () → !rise.fun<array<1024, scalar<f32>> →
                                fun<array<1024, scalar<f32>> →
                                    array<1024, tuple<scalar<f32>, scalar<f32>>>>>>

// type: () → 
```

- Each **RISE** pattern is implemented as an MLIR operation
- Operations customized with attributes specifying information for the type
- Patterns encoded as operations have a **RISE** function type

Function application

4

5

6

7

```
%zip = rise.zip #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>
```

Function application

```
4 %f1 = rise.lambda (%arow : !rise.array<1024, scalar<f32>>) →  
                                !rise.array<1024, scalar<f32>> {  
5   %f2 = rise.lambda (%bcol : !rise.array<1024, scalar<f32>>) →  
                                !rise.scalar<f32> {  
6     %zip = rise.zip #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>  
7     %zipped = rise.apply %zip, %arow, %bcol
```

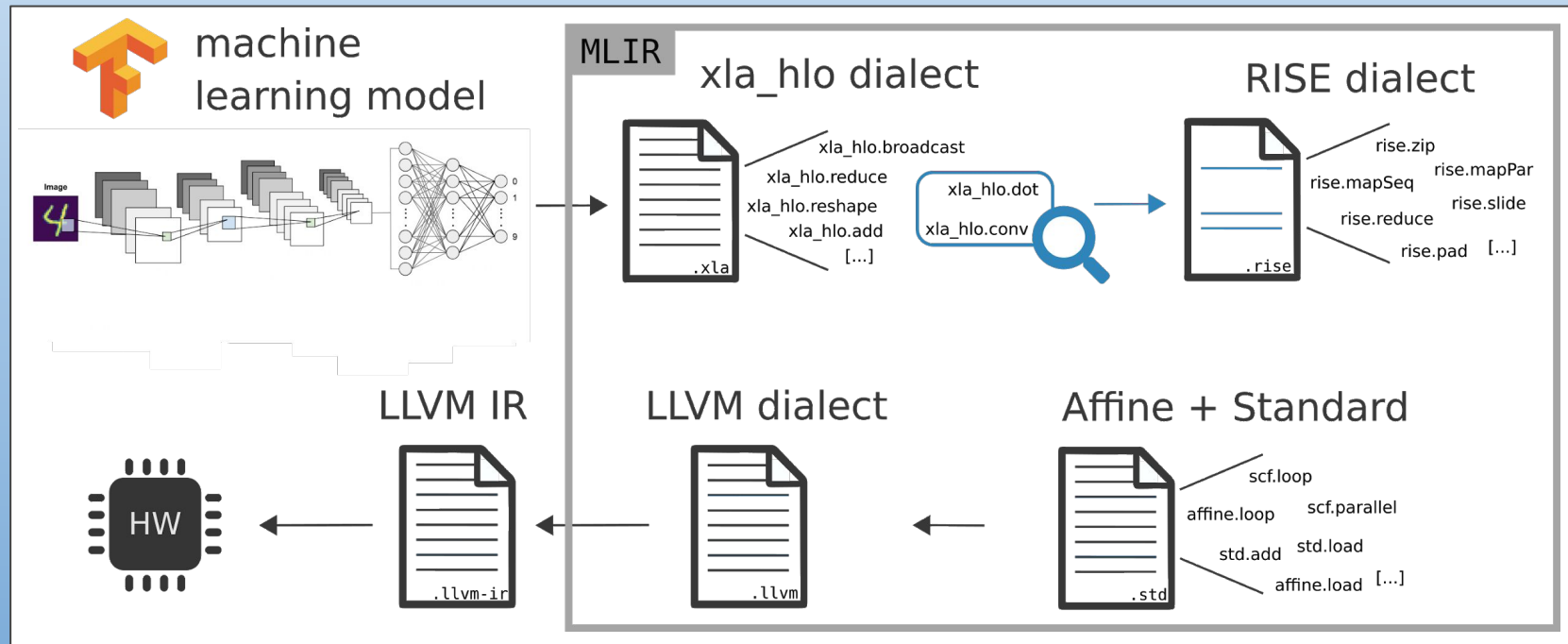
- **rise.apply** models function application:
expects an SSA value with a **RISE function type** (%zip)
and arguments to the function (%arow, %bcol)

Function application

```
4 %f1 = rise.lambda (%arow : !rise.array<1024, scalar<f32>>) →  
                                     !rise.array<1024, scalar<f32>> {  
5   %f2 = rise.lambda (%bcol : !rise.array<1024, scalar<f32>>) →  
                                     !rise.scalar<f32> {  
6     %zip = rise.zip #rise.nat<1024> #rise.scalar<f32> #rise.scalar<f32>  
7     %zipped = rise.apply %zip, %arow, %bcol  
  
//type: (%zip.type, array<1024, scalar<f32>>, array<1024, scalar<f32>>) →  
                                     array<1024, tuple<scalar<f32>, scalar<f32>>>
```

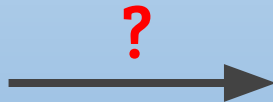
- **rise.apply** models function application:
expects an SSA value with a **Rise function type** (%zip)
and arguments to the function (%arow, %bcol)

How can we achieve end-to-end integration with existing solutions?



Lowering of the RISE dialect

```
1 func @mm(%outArg, %inA, %inB) {
2   %A = in %inA
3   %B = in %inB
4   %m1fun = lambda(%arow) -> array<2048, scalar<f32>> {
5     %m2fun = lambda(%bcol) -> scalar<f32> {
6       %zipFun = zip #nat<2048> #scalar<f32> #scalar<f32>
7       %zippedArrays = apply %zipFun, %arow, %bcol
8       %reductionLambda = lambda(%tuple, %acc) -> scalar<f32>{
9         %fstFun = fst #scalar<f32> #scalar<f32>
10        %sndFun = snd #scalar<f32> #scalar<f32>
11        %first = apply %fstFun, %tuple
12        %second = apply %sndFun, %tuple
13        %result = embed(%first, %second, %acc) {
14          %product = mulf %first, %second : f32
15          %result = addf %product, %acc : f32
16          return %result : f32
17        }
18        return %result : scalar<f32>
19      }
20      %init = literal #lit<0.0>
21      %reduceFun = reduceSeq #nat<2048>
22        #tuple<scalar<f32>, scalar<f32>> #scalar<f32>
23      %result = apply %reduceFun, %reductionLambda,
24        %init, %zippedArrays
25      return %result : scalar<f32>
26    }
27    %m2 = mapSeq #nat<2048> #array<2048, scalar<f32>>
28      #scalar<f32>
29    %result = apply %m2, %m2fun, %B
30    return %result : array<2048, scalar<f32>>
31  }
32  %m1 = mapSeq #nat<2048> #array<2048, scalar<f32>>
33    #array<2048, scalar<f32>>
34  %result = apply %m1, %m1fun, %A
35  out %outArg <- %result
36  return
37 }
```



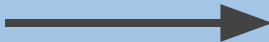
```
1 func @mm(%outArg, %inA, %inB) {
2   %init = constant 0.000000e+00 : f32
3   affine.for %i = 0 to 2048 {
4     affine.for %j = 0 to 2048 {
5       affine.store %init, %outArg[%i, %j]
6       affine.for %k = 0 to 2048 {
7         %a = affine.load %inA[%i, %k]
8         %b = affine.load %inB[%k, %j]
9         %c = affine.load %outArg[%i, %j]
10        %1 = mulf %a, %b : f32
11        %2 = addf %1, %c : f32
12        affine.store %2, %outArg[%i, %j]
13      }
14    }
15  }
16  return
17 }
```

Matrix Multiplication in Affine + Standard

Lowering of the RISE dialect

```
1 func @mm(%outArg, %inA, %inB) {
2   %A = in %inA
3   %B = in %inB
4   %m1fun = lambda(%arow) -> array<2048, scalar<f32>> {
5     %m2fun = lambda(%bcol) -> scalar<f32> {
6       %zipFun = zip #nat<2048> #scalar<f32> #scalar<f32>
7       %zippedArrays = apply %zipFun, %arow, %bcol
8       %reductionLambda = lambda(%tuple, %acc) -> scalar<f32>{
9         %fstFun = fst #scalar<f32> #scalar<f32>
10        %sndFun = snd #scalar<f32> #scalar<f32>
11        %first = apply %fstFun, %tuple
12        %second = apply %sndFun, %tuple
13        %result = embed(%first, %second, %acc) {
14          %product = mulf %first, %second :f32
15          %result = addf %product, %acc : f32
16          return %result : f32
17        }
18        return %result : scalar<f32>
19      }
20      %init = literal #lit<0.0>
21      %reduceFun = reduceSeq #nat<2048>
22                  #tuple<scalar<f32>, scalar<f32>> #scalar<f32>
23      %result = apply %reduceFun, %reductionLambda,
24                    %init, %zippedArrays
25      return %result : scalar<f32>
26    }
27    %m2 = mapSeq #nat<2048> #array<2048, scalar<f32>>
28          #scalar<f32>
29    %result = apply %m2, %m2fun, %B
30    return %result : array<2048, scalar<f32>>
31  }
32  %m1 = mapSeq #nat<2048> #array<2048, scalar<f32>>
33        #array<2048, scalar<f32>>
34  %result = apply %m1, %m1fun, %A
35  out %outArg <- %result
36  return
37 }
```

1. Lowering functional
to imperative
representation



```
1 func @mm_codegen(%outArg, %inA, %inB){
2   %A = codegen.cast(%inA)
3   %B = codegen.cast(%inB)
4   %out = codegen.cast(%outArg)
5   affine.for %i = 0 to 2048 {
6     %3 = codegen.idx(%A, %i)
7     %4 = codegen.idx(%out, %i)
8     affine.for %j = 0 to 2048 {
9       %5 = codegen.idx(%B, %j)
10      %6 = codegen.idx(%4, %j)
11      %7 = codegen.zip(%3, %5)
12      %init = embed() {
13        %cst = constant 0.0 : f32
14        return(%cst) : (f32) -> ()
15      }
16      codegen.assign(%init, %6)
17      affine.for %kk = 0 to 2048 {
18        %9 = codegen.idx(%7, %kk)
19        %10 = codegen.fst(%9)
20        %11 = codegen.snd(%9)
21        %12 = embed(%10, %11, %6) {
22          %13 = mulf %arg6, %arg7 : f32
23          %14 = addf %13, %arg8 : f32
24          return(%14) : (f32) -> ()
25        }
26        codegen.assign(%12, %6)
27      }
28    }
29    return
30  }
```

Matrix Multiplication in
imperative RISE



```
1 func @mm(%outA
2   %init = const
3   affine.for %i
4     affine.for %
5       affine.st
6       affine.fo
7       %a = af
8       %b = af
9       %c = af
10      %1 = mu
11      %2 = ad
12      affine.
13    }
14  }
15  }
16  return
17 }
```

Lowering of the RISE dialect

```
<f32>> {  
  scalar<f32>  
  ol  
  > scalar<f32>{
```

1. Lowering functional to imperative representation



```
#scalar<f32>  
oda,
```

```
<f32>>
```

```
<f32>>  
scalar<f32>>
```

```
1 func @mm_codegen(%outArg, %inA, %inB){  
2   %A = codegen.cast(%inA)  
3   %B = codegen.cast(%inB)  
4   %out = codegen.cast(%outArg)  
5   affine.for %i = 0 to 2048 {  
6     %3 = codegen.idx(%A, %i)  
7     %4 = codegen.idx(%out, %i)  
8     affine.for %j = 0 to 2048 {  
9       %5 = codegen.idx(%B, %j)  
10      %6 = codegen.idx(%4, %j)  
11      %7 = codegen.zip(%3, %5)  
12      %init = embed() {  
13        %cst = constant 0.0 : f32  
14        return(%cst) : (f32) -> ()  
15      }  
16      codegen.assign(%init, %6)  
17      affine.for %k = 0 to 2048 {  
18        %9 = codegen.idx(%7, %k)  
19        %10 = codegen.fst(%9)  
20        %11 = codegen.snd(%9)  
21        %12 = embed(%10, %11, %6) {  
22          %13 = mulf %arg6, %arg7 : f32  
23          %14 = addf %13, %arg8 : f32  
24          return(%14) : (f32) -> ()  
25        }  
26        codegen.assign(%12, %6)  
27      }  
28    }  
29  }  
30  return  
31 }
```

Matrix Multiplication in
imperative RISE

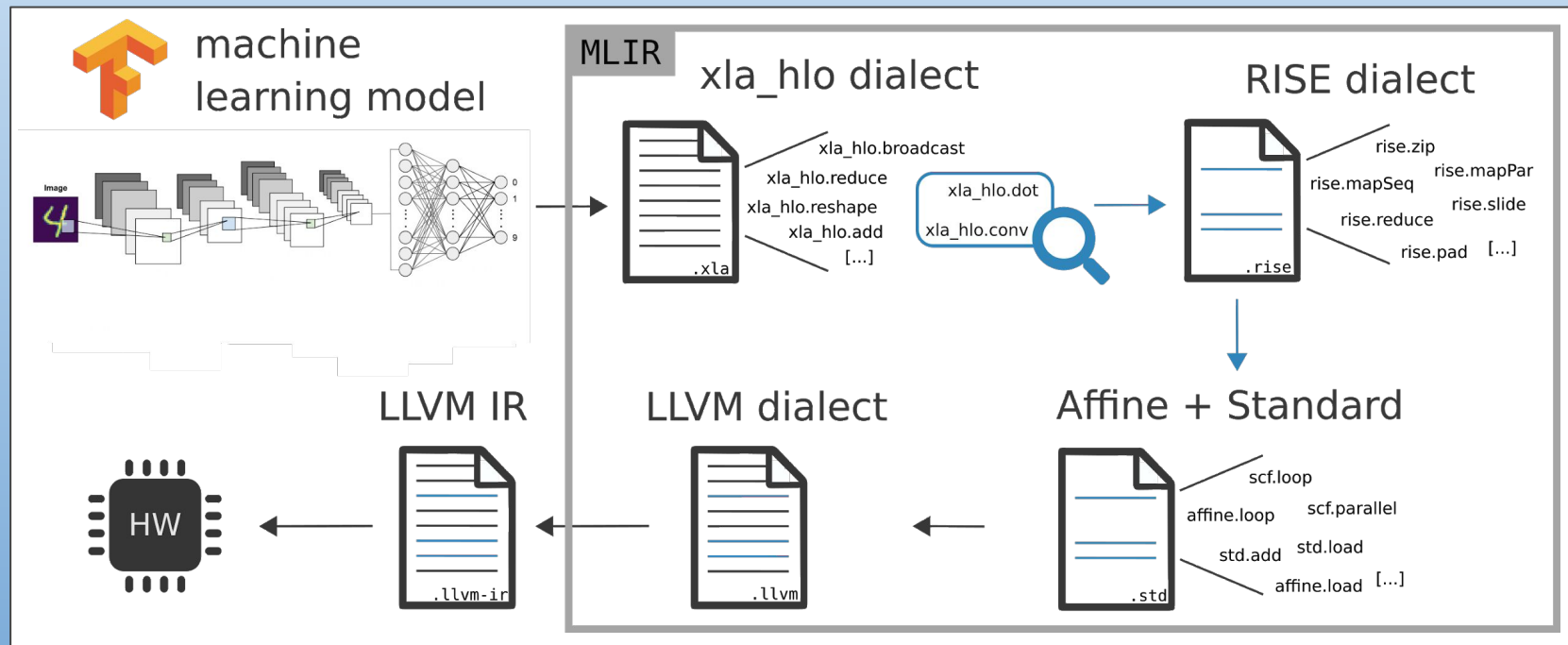
2. Lowering imperative to target representation



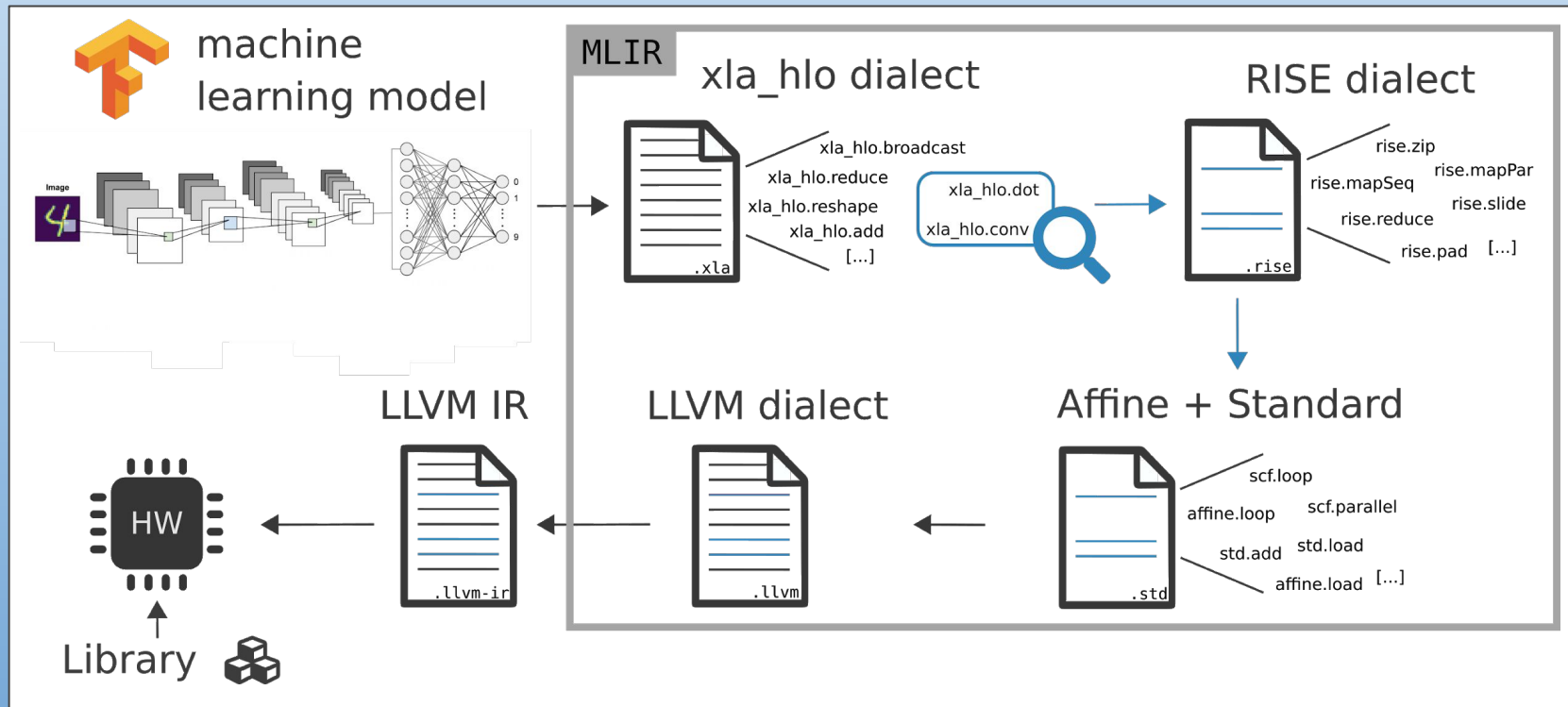
```
1 func @mm(%outArg, %inA, %inB) {  
2   %init = constant 0.000000e+00 : f32  
3   affine.for %i = 0 to 2048 {  
4     affine.for %j = 0 to 2048 {  
5       affine.store %init, %outArg[%i, %j]  
6       affine.for %k = 0 to 2048 {  
7         %a = affine.load %inA[%i, %k]  
8         %b = affine.load %inB[%k, %j]  
9         %c = affine.load %outArg[%i, %j]  
10        %1 = mulf %a, %b : f32  
11        %2 = addf %1, %c : f32  
12        affine.store %2, %outArg[%i, %j]  
13      }  
14    }  
15  }  
16  return  
17 }
```

Matrix Multiplication in
Affine+Standard

RISE end-to-end integration



RISE end-to-end integration

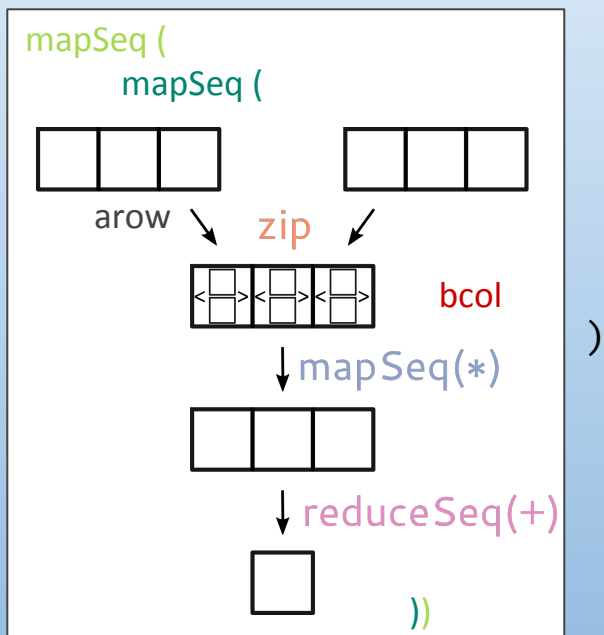


Lowering RISE to library calls

[illegible]

```
match_for(
```

matchSuccess

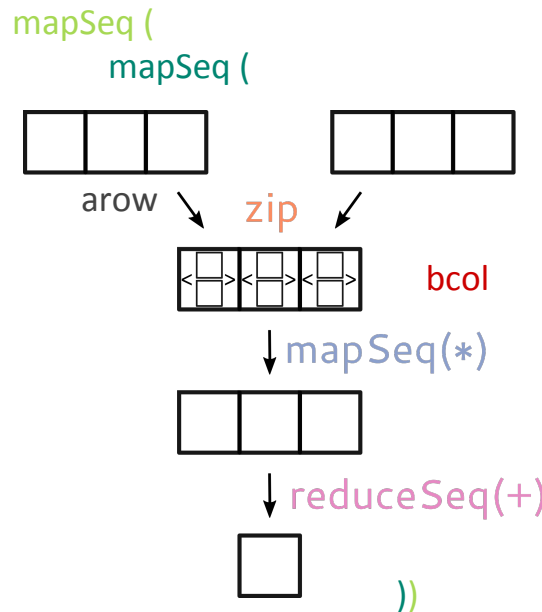


Lowering RISE to library calls



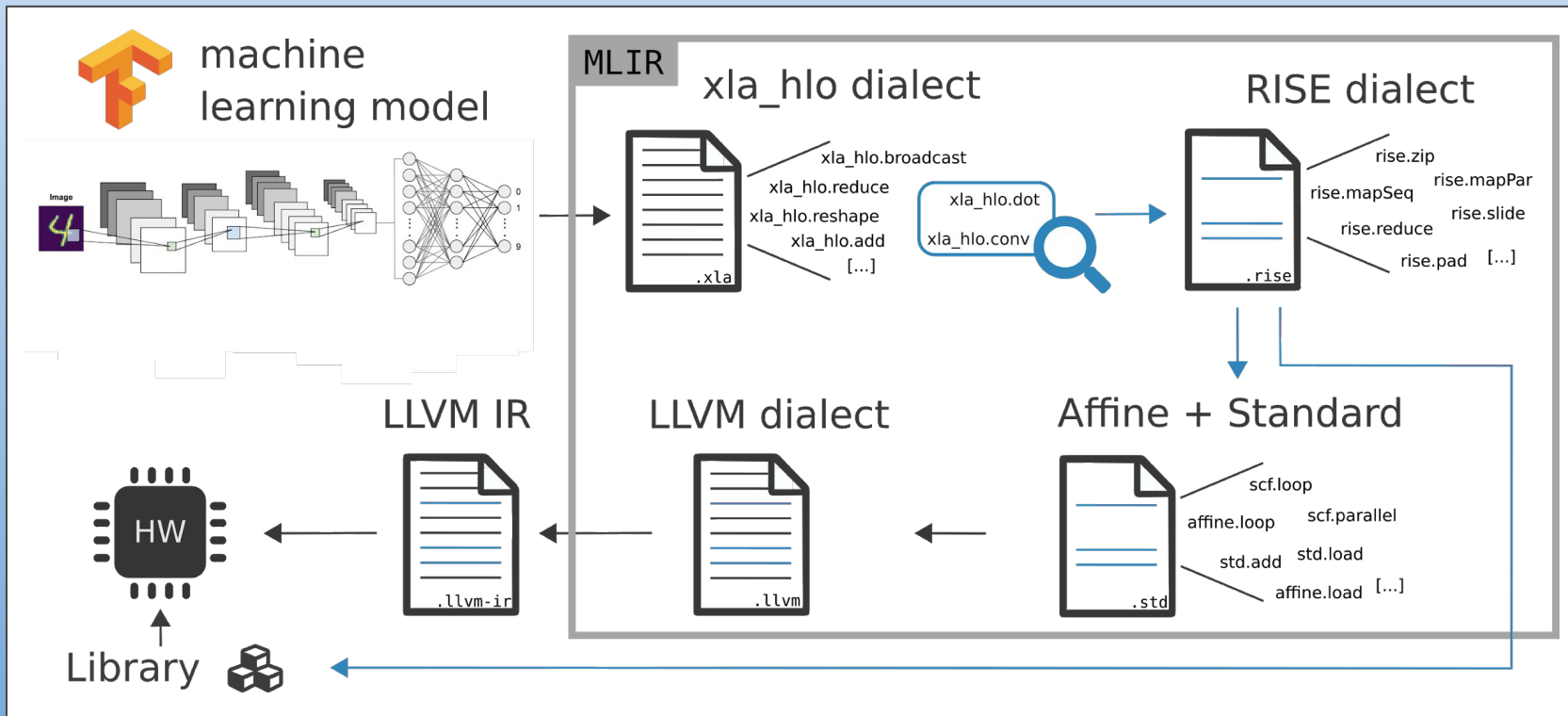
match_for(

matchSuccess



```
1 func @mm(%out:memref<1024x1024xf32>,%A:memref<1024x1024xf32>,%B:memref<1024x1024xf32>) {
2   %CblasRowMajor = constant 101 : i32
3   %CblasNoTrans = constant 111 : i32
4   %M = constant 1024 : i32
5   %N = constant 1024 : i32
6   %K = constant 1024 : i32
7   %LDA = constant 1024 : i32
8   %LDB = constant 1024 : i32
9   %LDC = constant 1024 : i32
10  %alpha = constant 1.0 : f32
11  %beta = constant 1.0 : f32
12  call @cblas_sgemm_wrapper(%CblasRowMajor, %CblasNoTrans, %CblasNoTrans, %M, %N, %K, %alpha, %A, %LDA, %B, %LDB, %beta, %out, %LDC) :
13    (i32, i32, i32, i32, i32, i32, f32, memref<1024x1024xf32>, i32, memref<1024x1024xf32>, i32, f32, memref<1024x1024xf32>, i32) → ()
14  return
15 }
```

RISE end-to-end integration



RISE end-to-end demonstration

```
tools git:(master) head mnist.mlir -n 4
func @mnist_predict(%input: tensor<1x28x28x1xf32>) -> tensor<1x10xf32> {
    %1 = hlo.reshape (%input) : (tensor<1x28x28x1xf32>) -> tensor<1x784xf32>
    %2 = hlo.dot (%1, %kernel) : (tensor<1x784xf32>, tensor<784x128xf32>) -> tensor<1x128xf32>
    %3 = hlo.add (%2, %bias) : (tensor<1x128xf32>, tensor<128xf32>) -> tensor<1x128xf32>
tools git:(master) run-mlir.sh -target-backends=llvm-ir mnist.mlir -input-value="1x28x28x1xf32"
Lowering XLA HLO -> RISE
Lowering RISE -> Affine
Lowering Affine -> LLVM-IR
Compiling for target backend 'llvm-ir*'
Evaluating all functions in module for driver 'llvm'
Creating driver and device for 'llvm'
EXECUTE @mnist_predict

result[0]: Buffer<float32[1x10]> 1x10xf32=[0.0 0.0 0.9 0.0 0.0 0.0 0.0 0.1 0.0 0.0]
```

RISE

A functional Pattern-based language in MLIR

We are Open Source!

<https://rise-lang.org/mlir>

<https://github.com/rise-lang/mlir>

Martin Lücke | Michel Steuwer | Aaron Smith



THE UNIVERSITY
of EDINBURGH