

LINALG AND THE TRANSFORM DIALECT

Maximilian Bartel



EXERCISES

ALL EXERCISES AND SOLUTIONS ARE IN THE REPOSITORY!

Exercises `git checkout exercises/day-4-session-4`

Solutions `git checkout solutions/day-4-session-4`

AFTER THIS SESSION YOU CAN WRITE AND SCHEDULE LINEAR ALGEBRA ALGORITHMS USING LINALG AND TRANSFORM

- 1 Understand Linalg dialect and why this is a very helpful abstraction in a tensor compiler
- 2 Apply transformations to Linalg to make a tensor program run fast
- 3 Define your own transformations with the transform dialect
- 4 Get to know projects that use these concepts on real workloads in production

The running example for this lecture is going to be a matmul followed by an elementwise add

CLASSICAL KERNELS ARE WRITTEN TO GET THE BEST PERFORMANCE FOR A GIVEN TASK AND COMMIT TO A SPECIFIC SCHEDULE

```

1  #include <stddef.h>
2
3  void mm(int m, int n, int k, const float *restrict A,
4          const float *restrict B,
5          float *restrict C) {
6      for (int i = 0; i < m; i++)
7          for (int j = 0; j < n; j++) {
8              float s = 0;
9              for (int p = 0; p < k; p++)
10                 s += A[(size_t)i * k + p] * B[(size_t)p * n + j];
11             C[(size_t)i * n + j] = s;
12         }
13     }

```

Algorithm

$$C[i,j] += A[i,k] * B[k,j]$$

The mathematical update is independent of the loop nesting.

Schedule

loop order
work granularity
locality
parallelism

i -> j -> k
scalar
bad
single threaded

KERNELS ARE MADE OF THE ALGORITHM AND THE SCHEDULE OF THE ALGORITHM

Algorithm

$C[i,j] += A[i,k] * B[k,j]$

Schedule A

- plain loops
- row-major traversal
- scalar update

Schedule B

- tile M,N to 64
- parallelize outer tiles
- tile again for registers
- vectorize inner work

While the semantics of the computation stays the same, how the computation is executed to reach your optimization target vastly differs based on the hardware.

The idea is to split the what from the how

LINALG DESCRIBES LINEAR ALGEBRA DECLARATIVELY WITHOUT DECIDING HOW IT WILL BE EXECUTED

```
1 #map = affine_map<(d0, d1) -> (d0, d1)>
2 %0 = linalg.generic {
3     indexing_maps = [#map, #map, #map],
4     iterator_types = ["parallel", "parallel"]
5 }
6 ins(%arg0, %arg1 : tensor<?x?xf32>, tensor<?x?xf32>)
7 outs(%arg2 : tensor<?x?xf32>) {
8     ^bb0(%in: f32, %in_0: f32, %out: f32):
9         %1 = arith.addf %in, %in_0 : f32
10        linalg.yield %1 : f32
11    } -> tensor<?x?xf32>
12
```

`linalg.generic`

- The most important operation in the Linalg Dialect is `linalg.generic`.
- It describes linear algebra operations consisting of perfectly nested loops when executed naively.
- It is designed to make reasoning about the loop structure easy.

LINALG DESCRIBES LINEAR ALGEBRA DECLARATIVELY WITHOUT DECIDING HOW IT WILL BE EXECUTED

```
1  #map = affine_map<(d0, d1) -> (d0, d1)>
2  %0 = linalg.generic {
3      indexing_maps = [#map, #map, #map],
4      iterator_types = ["parallel", "parallel"]
5  }
6  ins(%arg0, %arg1 : tensor<?x?xf32>, tensor<?x?xf32>)
7  outs(%arg2 : tensor<?x?xf32>) {
8      ^bb0(%in: f32, %in_0: f32, %out: f32):
9          %1 = arith.addf %in, %in_0 : f32
10         linalg.yield %1 : f32
11     } -> tensor<?x?xf32>
12
```

Iteration space

- The iteration space for a loop is fully defined by the dimension sizes of inputs and outputs, even when they are dynamic.
- The indexing_maps then define which element is picked in which iteration of the innermost loop.

LINALG DESCRIBES LINEAR ALGEBRA DECLARATIVELY WITHOUT DECIDING HOW IT WILL BE EXECUTED

```
1  #map = affine_map<(d0, d1) -> (d0, d1)>
2  %0 = linalg.generic {
3      indexing_maps = [#map, #map, #map],
4      iterator_types = ["parallel", "parallel"]
5  }
6  ins(%arg0, %arg1 : tensor<?x?xf32>, tensor<?x?xf32>)
7  outs(%arg2 : tensor<?x?xf32>) {
8      ^bb0(%in: f32, %in_0: f32, %out: f32):
9          %1 = arith.addf %in, %in_0 : f32
10         linalg.yield %1 : f32
11     } -> tensor<?x?xf32>
12
```

Iterator_types

- The `iterator_types` tell you easily how many loops a naïve lowering would produce and what properties the loops has.
- Upstream MLIR implements parallel and reduction loops.

LINALG DESCRIBES LINEAR ALGEBRA DECLARATIVELY WITHOUT DECIDING HOW IT WILL BE EXECUTED

```
1  #map = affine_map<(d0, d1) -> (d0, d1)>
2  %0 = linalg.generic {
3      indexing_maps = [#map, #map, #map],
4      iterator_types = ["parallel", "parallel"]
5  }
6  ins(%arg0, %arg1 : tensor<?x?xf32>, tensor<?x?xf32>)
7  outs(%arg2 : tensor<?x?xf32>) {
8      ^bb0(%in: f32, %in_0: f32, %out: f32):
9          %1 = arith.addf %in, %in_0 : f32
10         linalg.yield %1 : f32
11     } -> tensor<?x?xf32>
12
```

Linalg operation body

- The body of a linalg operation then operates on scalar values and tell you what the operation does semantically.

LINALG DESCRIBES LINEAR ALGEBRA DECLARATIVELY WITHOUT DECIDING HOW IT WILL BE EXECUTED

```
1  #map = affine_map<(d0, d1) -> (d0, d1)>
2  %0 = linalg.generic {
3      indexing_maps = [#map, #map, #map],
4      iterator_types = ["parallel", "parallel"]
5  }
6  ins(%arg0, %arg1 : tensor<?x?xf32>, tensor<?x?xf32>)
7  outs(%arg2 : tensor<?x?xf32>) {
8      ^bb0(%in: f32, %in_0: f32, %out: f32):
9          %1 = arith.addf %in, %in_0 : f32
10         linalg.yield %1 : f32
11     } -> tensor<?x?xf32>
12
```

Destination passing style

- The `linalg.generic` operation is what we call a “destination passing style” operation.
- This is because we are passing in a tensor as an argument where the result should be written to.
- This breaks the tensor abstraction but makes bufferization later a lot easier.

EXERCISE 1

MAPPING LINEAR ALGEBRA TO A LINALG GENERIC NEEDS A BIT OF PRACTICE

```
1 %0 = linalg.generic {
2     indexing_maps = [???],
3     iterator_types = [???]
4 }
5 ins(???)
6 outs(???) {
7     ^bb0(???):
8         ???
9 } -> ???
10
```

- 1 $C = A * B$ Matmul
- 2 $D = A * B + C$ Matmul and elementwise addition
- 3 NHWC convolution
- 4 Bonus: NCHW convolution

This exercise should hopefully give you a bit of intuition on how to map real workloads onto linalg.generic operations.

EXERCISE 2

UNDERSTANDING THE ALGORITHM FROM THE LINALG OPERATION ALSO TAKES PRACTICE

```
1 %result = linalg.generic {
2     indexing_maps = [
3         affine_map<(b, m, n, k) -> (b, m, k)>,
4         affine_map<(b, m, n, k) -> (n, k)>,
5         affine_map<(b, m, n, k) -> (b, m, n)>
6     ],
7     iterator_types = [
8         "parallel", "parallel", "parallel", "reduction"
9     ]
10 }
11 ins(%A, %B : tensor<?x?x?xf32>, tensor<?x?xf32>)
12 outs(%C : tensor<?x?x?xf32>) {
13     ^bb0(%a: f32, %b: f32, %acc: f32):
14         %product = arith.mulf %a, %b : f32
15         %sum = arith.addf %acc, %product : f32
16         linalg.yield %sum : f32
17 } -> tensor<?x?x?xf32>
```

Interpreting linalg.generic

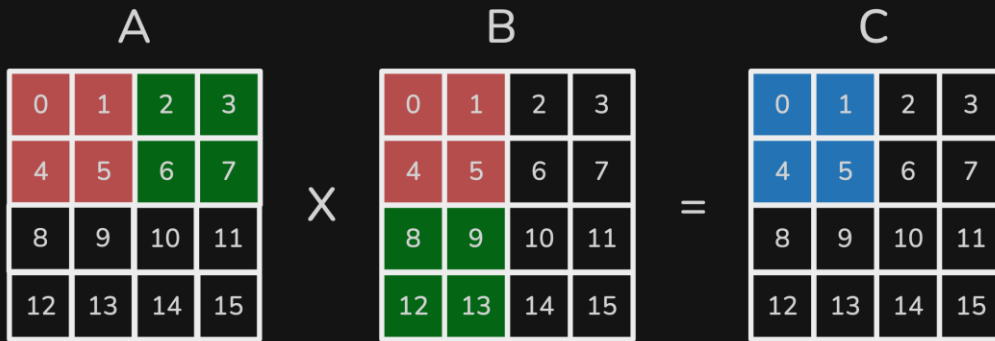
- Interpreting linalg.generic operations is also not always trivial.
- Let's practice this a few times! The output should be a pseudo loop nest like

```
for b in 0..B:           // parallel
  for i in 0..M:         // parallel
    for j in 0..N:       // parallel
      for k in 0..K:     // reduction
        C[b,i,j] += A[b,i,k] * B[b,k,j]
```

TILING

TILING IS A LOOP TRANSFORMATION THAT SPLITS LOOPS TO GET BETTER MEMORY ACCESS PATTERNS

Logical layout



Fast algorithm, but it generates cache misses:

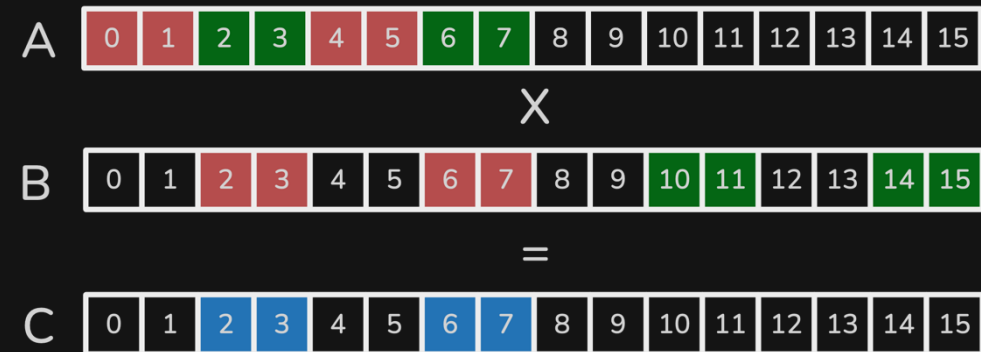
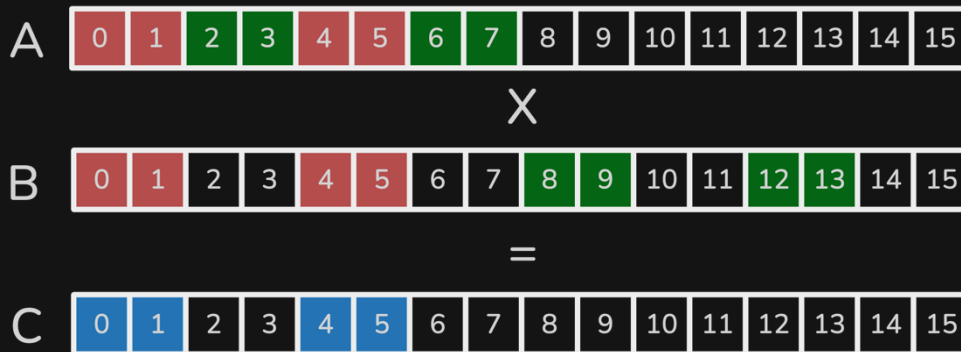


Reads from A are good

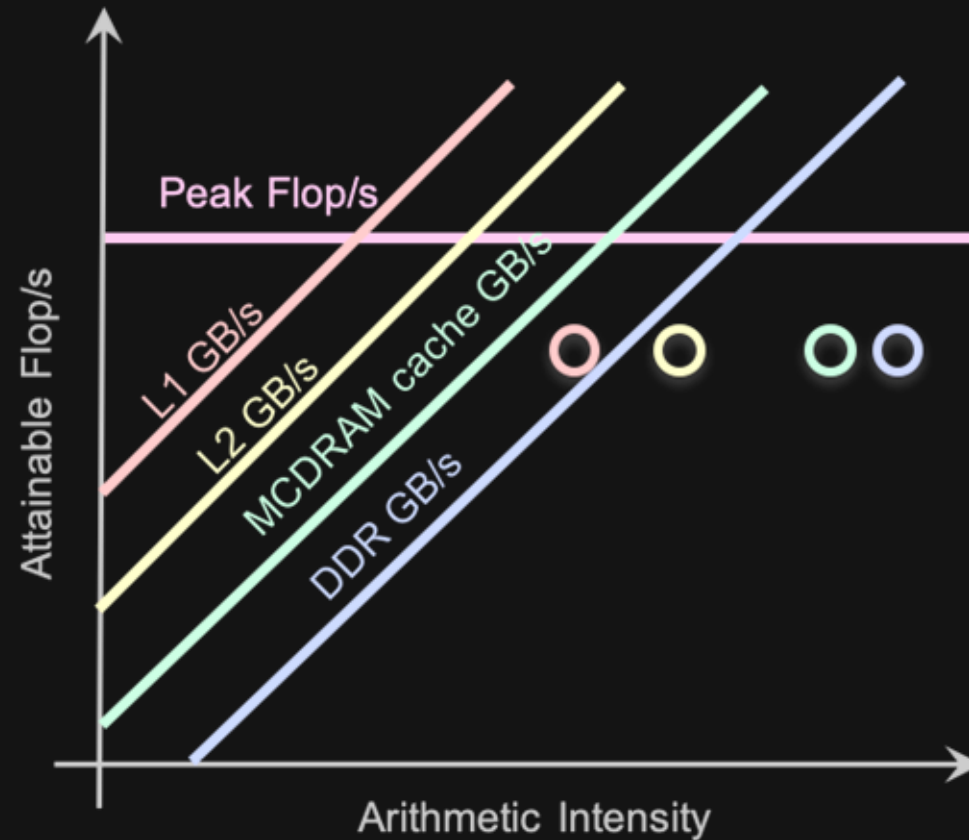


Reads from B and C are not

In-memory layout



TILING INCREASES THE PERFORMANCE BECAUSE WE ARE REUSING DATA THAT IS AVAILABLE IN FASTER MEMORIES



TILING A LINALG OPERATION GIVES YOU A NEW LINALG OPERATION TO CREATE YOUR SCHEDULE STEP BY STEP

```
1 %0 = linalg.matmul
2   ins(%arg0, %arg1 : tensor<?x?xf32>, tensor<?x?xf32>)
3   outs(%arg2 : tensor<?x?xf32>) -> tensor<?x?xf32>
```

- Tiling a linalg operation will create loops and a new linalg operation.
- You can tile in multiple steps or stages and create your schedule with that.
- While this transformation is very mechanical, composing them in a meaningful way is very complex.



Tiling

```
1 %0 = scf.for %arg6 = %c0 to %dim step %c8
2   iter_args(%arg7 = %arg2) -> (tensor<?x?xf32>) {
3     %1 = scf.for %arg8 = %c0 to %dim_0 step %c16
4       iter_args(%arg9 = %arg7) -> (tensor<?x?xf32>) {
5         %2 = affine.min #map(%arg6)[%dim]
6         %3 = affine.min #map1(%arg8)[%dim_0]
7         %extracted_slice = tensor.extract_slice ...
8         %extracted_slice_2 = tensor.extract_slice ...
9         %extracted_slice_3 = tensor.extract_slice ...
10        %4 = linalg.matmul
11          ins(%extracted_slice, %extracted_slice_2 :
12            tensor<?x?xf32>, tensor<?x?xf32>)
13          outs(%extracted_slice_3 : tensor<?x?xf32>)
14          -> tensor<?x?xf32>
15        %inserted_slice = tensor.insert_slice %4 into ...
16        scf.yield %inserted_slice : tensor<?x?xf32>
17      }
18    scf.yield %1 : tensor<?x?xf32>
19  }
```

TRANSFORM DIALECT CAN REPRESENT TRANSFORMATIONS AS IR WITHOUT THE NEED TO CHANGE THE COMPILER

```

1 func.func @foo(
2   %arg0: tensor<?x?xf32>, %arg1: tensor<?x?xf32>, %arg2: tensor<?x?xf32>)
3   -> tensor<?x?xf32> {
4   %0 = linalg.matmul {__producer__}
5     ins(%arg0, %arg1 : tensor<?x?xf32>, tensor<?x?xf32>)
6     outs(%arg2 : tensor<?x?xf32>) -> tensor<?x?xf32>
7   return %0 : tensor<?x?xf32>
8 }
9 module attributes {transform.with_named_sequence} {
10   transform.named_sequence @__transform_main(
11     %arg0: !transform.any_op {transform.readonly}) {
12     %0 = transform.structured.match
13       attributes {__producer__} in %arg0 : (!transform.any_op)
14     -> !transform.any_op
15     %tiled_linalg_op, %loops:2 = transform.structured.tile_using_for %0
16     tile_sizes [8, 16] : (!transform.any_op)
17     -> (!transform.any_op, !transform.any_op, !transform.any_op)
18     transform.yield
19   }
20 }

```

Transform dialect

- The transform dialect lets you describe transformations of IR as IR.
- The IR that you are going to transform is called payload IR.
- The transformation instructions transform IR.

THE TRANSFORMS ARE DRIVEN BY AN INTERPRETER ON OPERATIONS THAT IMPLEMENT SPECIFIC INTERFACES

Writing transform ops

- To write your own transform operation you need to define the operation and then write implementations for a few interfaces.
- Calling `--transform-interpreter` runs the code then.

```

1 // Define the new operation. By convention, prefix its name with the name of the dialect
2 // extension, "my.". The full operation name will be further prefixed with "transform.".
3 def ChangeCallTargetOp : Op<Transform_Dialect, "my.change_call_target",
4   // Indicate that the operation implements the required TransformOpInterface and
5   // MemoryEffectsOpInterface.
6   [DeclareOpInterfaceMethods<TransformOpInterface>,
7     DeclareOpInterfaceMethods<MemoryEffectsOpInterface>]> {
8   // Provide a brief and a full description. It is recommended that the latter describes
9   // the effects on the operands and how the operation processes various failure modes.
10  let summary = "Changes the callee of a call operation to the specified one";
11  let description = [{ ...
12  }];

```

```

13
14 // The arguments include the handle to the payload operations and the attribute that
15 // specifies the new callee. The handle must implement TransformHandleTypeInterface.
16 // We use a string attribute as the symbol may not exist in the transform IR so the
17 // verification may fail.
18 let arguments = (ins
19   TransformHandleTypeInterface:$call,
20   StrAttr:$new_target);
21
22 // The results are empty as the transformation does not produce any new payload.
23 let results = (outs);
24
25 // Provide nice syntax.
26 let assemblyFormat = "$call `, ` $new_target attr-dict `:` type($call)";
27 }

```

EXERCISE 3

WHEN WRITING COMPLEX TRANSFORMATION CHAINS, YOU WANT TO INSPECT THE CODE SOMETIMES



Everyone makes mistakes. Let's implement a transform operation that lets us `printf` debug the IR during the execution of a transform script!

TRANSFORM DIALECT CAN BE USED VERY SIMILAR TO HALIDE TO SCHEDULE LINALG OPERATIONS

```

1 Func blur_3x3(Func input) {
2   Func blur_x, blur_y;
3   Var x, y, xi, yi;
4
5   // The algorithm - no storage or order
6   blur_x(x, y) = (input(x-1, y) + input(x, y) + input(x+1, y))/3;
7   blur_y(x, y) = (blur_x(x, y-1) + blur_x(x, y) + blur_x(x, y+1))/3;
8
9   // The schedule - defines order, locality; implies storage
10  blur_y.tile(x, y, xi, yi, 256, 32)
11    .vectorize(xi, 8).parallel(y);
12  blur_x.compute_at(blur_y, x).vectorize(x, 8);
13
14  return blur_y;
15 }

```

Halide, transform & schedules

- Halide and transform dialect follow the same mental model of separation of concerns.
- Schedules are explicit, reviewable, testable, and can be change independently of the payload.
- Transform IR goes even a step beyond that because to change the schedule there is no need to recompile the compiler itself. It is an input that can be written in its own tool.

WHILE EXECUTING A TRANSFORM SCRIPT YOU ARE CHANGING THE IR SIMULTANEOUSLY

```
1 %mm = transform.structured.match ...
2
3 %tiled, %loops:2 =
4     transform.structured.tile_using_for %mm
5     tile_sizes [8, 16, 0]
6
7 // Error: %mm was consumed.
8 transform.debug.emit_remark_at %mm, "old matmul"
9
10 // Correct: %tiled is the new tiled matmul operation.
11 transform.debug.emit_remark_at %tiled, "new matmul"
```

Handle invalidation

- The transform interpreter changes the payload IR on the fly.
- This means that if a transformation is destructive, e.g., erasing an operation, then handle is also invalid.
- There is `--transform-dialect-check-uses` to make sure the script is valid. With dynamic behavior this can get quite complicated.

Rule of thumb: After a consuming transform, use the new handles it returns.

YOU CAN USE MULTIPLE TRANSFORM OPERATIONS AFTER ONE ANOTHER TO EXECUTE MULTIPLE TRANSFORMATIONS

```
1 transform.named_sequence @__transform_main(%arg0: !transform.any_op {transform.readonly}) {
2     %0 = transform.structured.match attributes {__root__} in %arg0
3     : (!transform.any_op) -> !transform.any_op
4     %1 = transform.structured.match attributes {__producer__} in %arg0
5     : (!transform.any_op) -> !transform.any_op
6     %tiled_linalg_op, %loops:2 = transform.structured.tile_using_for %0 tile_sizes [8, 16]
7     : (!transform.any_op) -> (!transform.any_op, !transform.any_op, !transform.any_op)
8     %fused_op, %new_containing_op = transform.structured.fuse_into_containing_op %1 into %loops#0
9     : (!transform.any_op, !transform.any_op) -> (!transform.any_op, !transform.any_op)
10    transform.yield
11 }
```

EXERCISE 4

WRITING SCHEDULES FOR LINALG WITH THE TRANSFORM DIALECT IS NOT HARD

1. Part: Let's get a bit of practice writing longer transform scripts! On the right you see the transformations we want to apply.

2. Part: After you finished the script, try to lower the IR to a lower-level abstraction to see if your mental model of the computations are still correct.

```
mlir-opt exercise.mlir \  
  --transform-interpreter \  
  --linalg-generalize-named-ops \  
  --one-shot-bufferize="bufferize-function-boundaries" \  
  --convert-linalg-to-loops
```

Reference for a 2048×2048 float32 Matmul on Apple M3 MAX:

- Naïve loops: 8.082 s
- Accelerate Blas: 10.747 ms

Start from the matmul + add payload. Complete the transform sequence.

- 1 Match add with `__root__`
- 2 Match fill + matmul with `__producer__`
- 3 Tile the root with [8, 16]
- 4 Fuse producers into one generated loop
- 5 Tile the root again with [4, 4]
- 6 print or diff the resulting payload

UPSTREAM MLIR USES TRANSFORM AND LINALG FOR END-TO-END TESTING OF SOME HW FEATURES

Upstream MLIR SME tests

<https://github.com/llvm/llvm-project/blob/main/mlir/test/Integration/Dialect/Linalg/CPU/ArmSME/multi-tile-matmul.mlir>

PROJECT LIGHTHOUSE AND IREE ARE PROMINENT USERS OF THE TRANSFORM DIALECT

MLIR Lighthouse Project

https://github.com/llvm/lighthouse/blob/5d6431fe66f80469589b0f987727d4d80baf1f9e/examples/llama/test_llama3.py#L61

IREE

https://github.com/iree-org/iree/tree/13d9ad9501d878ff61ca1a0f753419f6d78ce85c/samples/transform_dialect

ABOUT US

WE BRING THE EXPERTISE TO ENABLE YOUR AI DEPLOYMENT

Founded by AI compiler and AI system engineers from RWTH Aachen University.

Built around a core team of ex-AMD engineers, researchers and strategists.

Multi-million funding secured from public institutions and private investors.



CEO
Dr.-Ing. Moritz Joseph

joseph@roofline.ai



www.roofline.ai



CTO
Max Bartel

bartel@roofline.ai



CFO & COO
Dr. Thomas Zimmermann

zimmermann@roofline.ai

Trusted and supported by:



Co-funded by
the European Union



Federal Ministry
for Economic Affairs
and Climate Action

