

SIROP: A Small IR for HLS with Parallel Patterns

Louis Hildebrand

McGill University

Montreal, Canada

louis.hildebrand@mail.mcgill.ca

Christophe Dubach

McGill University & Mila

Montreal, Canada

christophe.dubach@mcgill.ca

Abstract

Designers of custom streaming accelerators traditionally use HDLs (Hardware Description Languages), but this is time-consuming and requires advanced hardware expertise. C-based HLS (High-Level Synthesis) offers a higher level of abstraction and faster design time, but still requires some hardware expertise and performance is often left on the table. A promising direction is to use HLS with high-level functional parallel patterns such as map and reduce. Prior works have shown that high performance is achievable this way. However, designing such compiler systems is challenging because the optimizer must handle a large number of language primitives and interactions between them.

This paper introduces a minimal functional IR (Intermediate Representation) for hardware design, SIROP, which can express both pipelining and spatial parallelism with just five primitives. High-level operators from prior works are represented as syntax sugar and lowered to the core language. This simplifies hardware generation and optimization.

SIROP is compared with existing compilers on a set of image processing and linear algebra benchmarks. The SIROP designs use 61% fewer ALMs (Adaptive Logic Modules) than Aetherling, 68% fewer ALMs than SHIR, and 76% fewer ALMs than the Intel HLS compiler, all for the same throughput.

CCS Concepts: • **Hardware** → **High-level and register-transfer level synthesis**; • **Software and its engineering** → **Compilers**.

Keywords: intermediate representation, high-level synthesis, functional programming

ACM Reference Format:

Louis Hildebrand and Christophe Dubach. 2026. SIROP: A Small IR for HLS with Parallel Patterns. In *Proceedings of the 27th ACM SIGPLAN/SIGBED International Conference on Languages, Compilers, and Tools for Embedded Systems (LCTES '26)*, June 15–16, 2026, Boulder, CO, USA. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/3814943.3816175>



This work is licensed under a Creative Commons Attribution 4.0 International License.

LCTES '26, Boulder, CO, USA

© 2026 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-2721-4/2026/06

<https://doi.org/10.1145/3814943.3816175>

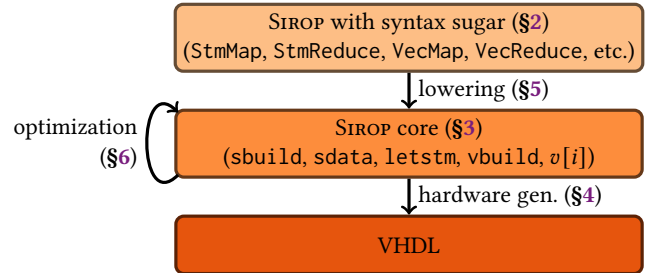


Figure 1. Steps in the SIROP compilation flow.

1 Introduction

Designing specialized accelerators is common in many domains from image processing to machine learning. Such accelerators offer higher performance and power efficiency than general-purpose processors. This is achieved by exploiting the fixed computational patterns found in the application and mapping these to spatial and pipeline parallelism.

Although specialized accelerators have significant benefits, they are difficult to design. Low-level HDLs such as VHDL and Verilog require deep hardware expertise. Even higher-level HDLs like Chisel [3] and Spatial [23], which simplify the process, demand hardware expertise.

C-based HLS lets designers define their hardware in C, with annotations manually added to guide the optimizer. This reduces design time but this still requires some expertise and the hardware may not reach peak performance [19, 29, 35].

A promising approach is to use HLS with parallel patterns from the functional programming paradigm, such as map and reduce [11, 19, 24, 25, 36]. Rewrite rules are used to lower operators step by step through different abstraction layers and apply optimizations such as operator fusion. This approach can produce high performance designs. However, implementing and testing a hardware template for each operator requires significant effort. At the time of writing, SHIR [36] includes over 40 stream and vector operators and Aetherling [11] has over 30. It is even harder to define rewrite rules to optimize all combinations of operators. Moreover, some rewrite rules cannot be defined because the result would not be expressible with the available language primitives.

This paper addresses these challenges by introducing a small IR, SIROP, with just five primitives for stream and vector operations. All the familiar parallel patterns can be expressed as syntax sugar and lowered to the core SIROP language. Optimization and hardware generation are then performed on the smaller core language, as shown in figure 1.

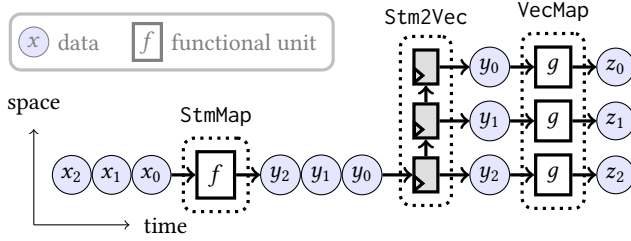


Figure 2. Stream and vector functional primitives.

The main contribution of this paper is the minimal IR for functional HLS (section 3). Its usefulness in practice is demonstrated by an optimizing compiler, which is described in sections 4 to 6. Specifically, the compiler demonstrates that the language has the following properties:

- It is easily translated to hardware (section 4).
- It can express familiar functional patterns for HLS, such as `StmMap` and `StmReduce` (section 5).
- It enables some more general optimizations, such as fusion of any two streaming operators (section 6).

Moreover, section 7 evaluates the compiler on a set of image processing and linear algebra kernels for an FPGA. The SIROP compiler generates more resource-efficient designs than prior works, including a commercial compiler (Intel HLS) and two state-of-the-art functional HLS compilers (SHIR and Aetherling). This is thanks to its ability to fuse any streaming operators and its finer control over each operator’s implementation during lowering and optimization.

2 Background and Overview

2.1 Functional Programming Primitives

Some prior works have used data-parallel patterns [15] from functional programming for hardware design. LIFTHLS [24], Aetherling [11], and SHIR [36] are recent examples. At the algorithmic level, operators are typically applied on sequences of data, which can be nested. Examples of such operators are shown below. Their semantics should be self-explanatory given their names and types.

```
count : (n: Int) → Seq[Int, n]
map : Seq[A, n] → (A → B) → Seq[B, n]
reduce : Seq[A, n] → ((A, A) → A) → Seq[A, 1]  (n ≥ 1)
concat : Seq[A, n] → Seq[A, m] → Seq[A, n + m]
zip : Seq[A, n] → Seq[B, n] → Seq[(A, B), n]
split : (m: Int) → Seq[A, nm] → Seq[Seq[A, m], n]
join : Seq[Seq[A, m], n] → Seq[A, nm]
```

Here, n, m are natural numbers; A, B are types; and $\text{Seq}[A, n]$ is the type of a sequence of n elements of type A .

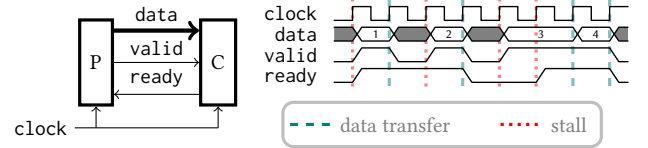


Figure 3. Streaming handshake protocol. The producer raises the valid signal when it has data to send and the consumer raises the ready signal when it can receive data.

```
1 (I0: Stm[u16, 4]) => (I1: Stm[u16, 4]) =>
2   StmZip(I0, I1)
3   .StmMap( @(x: u16, y: u16) => x * y )
4   .StmReduce( @(sum: u16, x: u16) => sum + x )
```

Figure 4. Functional-style dot product of two streams. In this paper, the notation $x.f(y, z, \dots)$ is equivalent to $f(x, y, z, \dots)$.

2.2 Stream and Vector Hardware Primitives

A key advantage of hardware acceleration is the ability to customize the level of spatial and pipeline parallelism. Therefore, functional HLS tools typically expose at least two sequence types: *streams* and *vectors*. This paper uses the “stream” and “vector” terminology from LIFTHLS [24] and SHIR [36]. Aetherling’s TSeq and SSeq (time and space sequences, respectively) are similar [11].

A *stream* can only be read once, one element at a time, in order. Operations on streams are processed in time, enabling pipeline parallelism. By contrast, *vectors* support random access, even access to multiple elements at once. Operations on vectors are processed in space, enabling spatial parallelism.

Most functional primitives have two versions: one for streams and one for vectors. There is `StmMap` and `VecMap`, `StmReduce` and `VecReduce`, and so on. There are also primitives `Stm2Vec` and `Vec2Stm` for converting between streams and vectors. These can be implemented in hardware as serial-in parallel-out and parallel-in serial-out shift registers. Figure 2 shows examples of stream and vector primitives.

Streams may be statically scheduled [11] or dynamically scheduled [20, 24, 36]. This paper uses dynamic scheduling with the handshake protocol illustrated in figure 3.

2.3 Overview

Figure 1 summarizes the SIROP compiler. The input is a program using stream and vector functional operators, like the dot product program in figure 4. Prior works [11, 36] have shown ways to lower operators like `map` to their stream and vector variants; this paper focuses on the IR instead.

The compiler lowers the program to the core language, which has only five stream and vector primitives: a vector constructor (`vbuid`) and extractor (`v[i]`), a stream constructor (`sbuid`) and extractor (`sdata`), and a `letstm` construct to share streams. These will be defined in the next section.

Type	T	$::= \text{bool} \mid i_{n^+} \mid u_n \mid \text{fix}_{n,k} \mid T_1 \rightarrow T_2$ $\mid (T_0, \dots, T_{n-1}) \mid \text{Vec}[T, n] \mid \text{Stm}[T, n]$	Booleans, numbers, functions Collections
Data type	D	$::= \text{bool} \mid i_{n^+} \mid u_n \mid \text{fix}_{n,k} \mid (D_0, \dots, D_{n-1}) \mid \text{Vec}[D, n]$	
Expression	e	$::= x \mid (x : T) \mapsto e \mid e_1(e_2)$ $\mid e_1 + e_2 \mid e_1 * e_2 \mid \text{if } e_1 \text{ then } e_2 \text{ else } e_3 \mid \dots$ $\mid (e_0, \dots, e_{n-1}) \mid e.n$ $\mid \text{vbuild}(n, (x : T) \mapsto e) \mid e_1[e_2]$ $\mid \text{sbuild}(n)(e_1, e_2) \{ \vec{a} \} \{ \vec{p} \} \mid \text{sdata}(e)$ $\mid \text{letstm}[n] \ x = e_1 \text{ in } e_2$	Variables, functions, function applications Standard logical, arithmetic operations Tuple constructor and extractor Vector constructor and extractor Stream constructor and extractor Stream sharing
Accumulator	a	$::= (x : D) = \{ \text{init} : e_1, \text{next} : e_2 \}$	
Producer	p	$::= (x : \text{Stm}[D, n]) = \{ \text{stm} : e_1, \text{ready} : e_2 \}$	

Figure 5. Abstract syntax of SIROP. $n, k \in \mathbb{N}$ and $n^+ \in \mathbb{N}^+$ are natural numbers. i_{n^+} and u_n are the types of n -bit signed and unsigned integers, respectively. $\text{fix}_{n,k}$ is for fixed-point numbers with n total bits and k fractional bits. $\vec{e}$ is an arbitrary number of expressions $e_1, e_2, \dots, e_n$. Most logical and arithmetic operations are omitted for brevity; they are standard.

The eager reader can find the lowered code for the dot product example in [figure 7](#), but it will be explained later.

Finally, after running several optimization passes, the compiler generates a VHDL description of a dynamically-scheduled streaming accelerator. [Section 7](#) will show that SIROP generates more resource-efficient hardware than the Intel HLS compiler, Aetherling [11], and SHIR [36].

3 The Core Language

The core language is a typed lambda calculus with expressions for booleans, integers, fixed-point numbers, tuples, vectors, and streams, as shown in [figure 5](#).

Vectors. Vectors have only one constructor:

$$\text{vbuild}(n, f) = [f(0), f(1), \dots, f(n-1)]$$

Vectors also have an extractor, $v[i]$, which returns the element at index i in vector v . The vbuild primitive was introduced by prior work on compiling functional array programs to imperative code [37].

Stream Constructor. The idea of having just one constructor and one extractor per collection type can be extended to streams. The expression

$$\text{sbuild}(n)(\text{data}, \text{valid}) \{ \vec{a} \} \{ \vec{p} \}$$

defines a stream of length n whose output is given by data and valid . The stream has a set of accumulators $\vec{a}$ modeled as recurrence equations. The stream also has a set of input producing streams $\vec{p}$, each with a corresponding ready expression. The next, ready, data, and valid expressions can depend on any accumulators and producers.

sbuild can implement any of the higher-level streaming operators (e.g., StmMap , StmReduce). This will be discussed further in [section 5](#). Intuitively, it has the required properties:

- It can choose when its output is valid. For example, StmReduce only has valid output after reading its entire input.

- It has some internal state, like the counter in StmCount .
- It has zero or more input streams and, at any given time, it can choose whether to read a given input. For example, StmConcat only reads its second input after exhausting its first input.

Stream Extractor. To read the current data from a producer, the consumer must use the sdata primitive. sdata can only be used within a stream and its value is undefined if the corresponding ready expression evaluates to false. All producers for which the ready expression evaluates to true must have valid data before the consumer will proceed.

Stream Sharing. A stream e_1 can be shared by many consumers using the construct $\text{letstm}[n] \ x = e_1 \text{ in } e_2$. x can appear multiple times in e_2 . letstm has a FIFO (First-In First-Out) queue whose capacity is at least n elements, which allows consumers to read the shared stream at different time steps. However, no consumer can read element $i + n$ before all consumers have read element i . For example, suppose s is a stream of length $n \geq 2$. Then $\text{letstm}[1] \ x = s \text{ in } \text{StmConcat}(x, x)$ will get stuck, but $\text{letstm}[n] \ x = s \text{ in } \text{StmConcat}(x, x)$ is valid.

Example. [Figure 7](#) shows how the dot product program from [figure 4](#) is expressed in the core language. The three distinct high-level operators StmReduce , StmMap and StmZip are all represented by sbuild and sdata . [Figure 7](#) also shows its behaviour over time.

Type Checking. Type checking is straightforward; most rules are standard. [Figure 6](#) shows typing rules for vbuild , sbuild , and letstm .

There are restrictions on the contents of collections in the core language. Collections can only contain “data types” D as defined in [figure 5](#). This means streams cannot be nested, which greatly simplifies hardware generation. [Section 5](#) will show that expressions with nested streams can be lowered to expressions in the core language that satisfy this constraint.

$$\begin{array}{c}
\boxed{\Gamma \vdash e : T} \text{ in context } \Gamma, \quad \boxed{\Gamma|\Gamma' \vdash a : D} \text{ in external context } \Gamma \text{ and internal context } \Gamma', a \text{ has type } D, \quad \boxed{\Gamma|\Gamma' \vdash p : T} \text{ in external context } \Gamma \text{ and internal context } \Gamma', p \text{ has type } T \\
\\
\Gamma' = \Gamma, a_i.x : D_i, p_i.x : T_i \quad \forall i. \Gamma|\Gamma' \vdash a_i : D_i \quad \forall i. \Gamma|\Gamma' \vdash p_i : T_i \\
\Gamma' \vdash d : D \quad \Gamma' \vdash v : \text{bool} \\
\hline
\Gamma \vdash \text{sbuild}(n)(d, v) \{ \vec{a} \} \{ \vec{p} \} : \text{Stm}[D, n] \\
a.T = D \quad \Gamma \vdash a.\text{init} : D \quad \Gamma' \vdash a.\text{next} : D \\
\hline
\Gamma|\Gamma' \vdash a : D \\
\\
p.T = T \quad \Gamma \vdash p.\text{stm} : T \quad \Gamma' \vdash p.\text{ready} : \text{bool} \\
\hline
\Gamma|\Gamma' \vdash p : T
\end{array}$$

$$\frac{\Gamma \vdash f : u_k \rightarrow D \quad 2^k - 1 \geq n - 1}{\Gamma \vdash \text{vbuild}(n, f) : \text{Vec}[D, n]}$$

$$\frac{\Gamma \vdash e_1 : \text{Stm}[T_2, m] \quad \Gamma, x : \text{Stm}[T_2, m] \vdash e_2 : \text{Stm}[T_1, n]}{\Gamma \vdash \text{letstm}[k] x = e_1 \text{ in } e_2 : \text{Stm}[T_1, n]}$$

Figure 6. Typing judgments and rules for sbuild, vbuild, and letstm. Γ maps variables to types. For sbuild, there is an “internal” context Γ' which includes the variables defined in the sbuild and an “external” context Γ which does not.

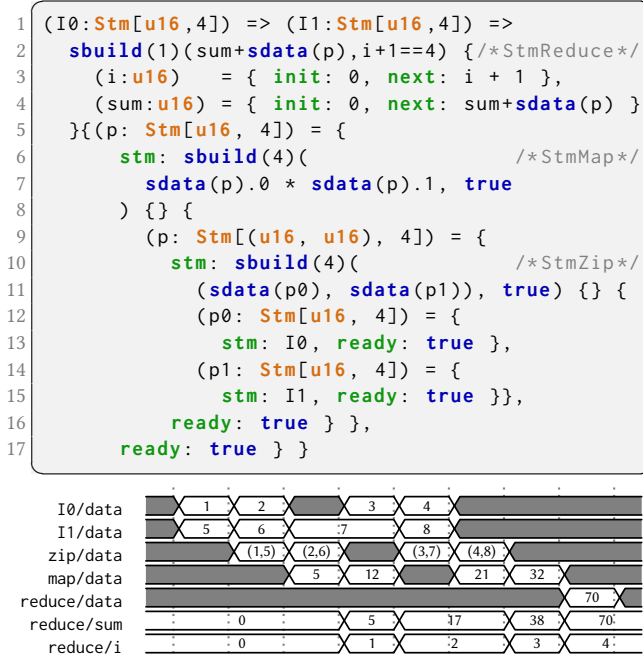


Figure 7. The dot product of two streams, expressed in the core language. The outermost sbuild (StmReduce) adds up the products using accumulator sum. It uses counter i to know when to stop. In the timing diagram, grey sections show where valid is false. reduce/data is the final dot product. reduce/sum and reduce/ i are the accumulators in StmReduce. When the StmMap does not have valid data (because I0 stalled earlier), the accumulators are not updated.

Summary. The core language of SIROP includes only five main primitives for vectors and streams. There is a constructor and extractor for vectors, vbuild and $v[i]$, borrowed from prior work. There is also a constructor and extractor for streams, sbuild and sdata. Finally, letstm allows a producer to have multiple consumers. The three stream-related primitives are one of the main contributions of this work.

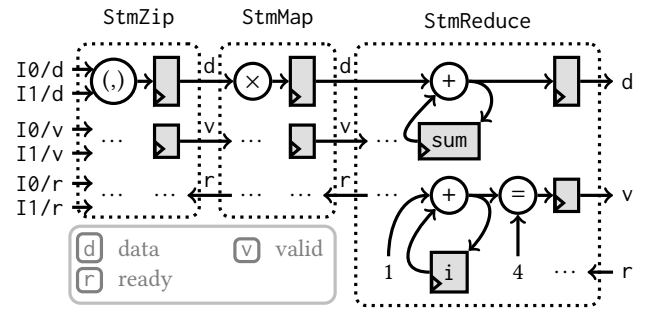


Figure 8. Hardware for the dot product example from figure 7. Handshake logic and register initial values are omitted.

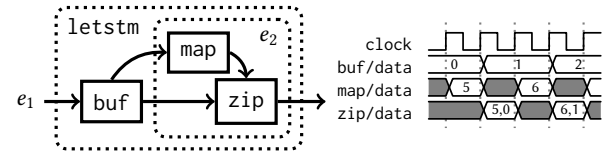


Figure 9. Block and timing diagrams for letstm[1] $x = e_1$ in StmZip(StmMap(x , +5), x), where e_1 is the stream 0, 1, 2, ... Each element of e_1 will be read at cycle t by map and then at cycle $t + 1$ by zip.

4 Hardware Generation

An expression in the core IR can be straightforwardly translated to a streaming accelerator expressed in an HDL. The SIROP VHDL generator makes no optimization decisions; instead, all optimization decisions are made on the core IR. This will be discussed in section 6.

Combinational Logic. Expressions with no streams correspond to combinational logic. Tuples and vectors are represented by VHDL records and arrays, respectively.

Streams. Streams are translated to sequential logic. Each sbuild is mapped to one VHDL component. For example, the dot product program from figure 7 would be translated to a three-stage pipeline as shown in figure 8. Each accumulator

$$\begin{aligned}
\mathcal{L}[\text{StmCount}(n : T)] &= \text{sbuild}(n)(i, \text{true}) \{ (i : T) = \{ \text{init} : 0, \text{next} : 1 + i \} \} \{ \} \\
\mathcal{L}[\text{StmReduce}(s : \text{Stm}[D, n], f)] &= \text{sbuild}(1)(\text{if } (\text{first}) \text{ then } \text{sdata}(p) \text{ else } f(\text{acc}, \text{sdata}(p)), 1 + i = n) \\
&\quad \{ (\text{acc} : D) = \{ \text{init} : _, \text{next} : \text{if } (\text{first}) \text{ then } \text{sdata}(p) \text{ else } f(\text{acc}, \text{sdata}(p)) \} \}, \\
&\quad \{ (\text{first} : \text{bool}) = \{ \text{init} : \text{true}, \text{next} : \text{false} \}, (i : u_{\lceil \log_2(n+1) \rceil}) = \{ \text{init} : 0, \text{next} : 1 + i \} \} \\
&\quad \{ (p : \text{Stm}[D, n]) = \{ \text{stm} : s, \text{ready} : \text{true} \} \} \quad \text{where } n \geq 1 \\
\mathcal{L}[\text{StmConcat}(s_0 : \text{Stm}[D, n], s_1 : \text{Stm}[D, m])] &= \text{sbuild}(n + m)(\text{if } i = n \text{ then } \text{sdata}(p_1) \text{ else } \text{sdata}(p_0), \text{true}) \\
&\quad \{ (i : u_{\lceil \log_2(n+1) \rceil}) = \{ \text{init} : 0, \text{next} : \text{if } i = n \text{ then } i \text{ else } 1 + i \} \} \\
&\quad \{ (p_0 : \text{Stm}[D, n]) = \{ \text{stm} : s_0, \text{ready} : i \neq n \}, (p_1 : \text{Stm}[D, m]) = \{ \text{stm} : s_1, \text{ready} : i = n \} \} \\
\mathcal{L}[\text{StmZip}(s_0 : \text{Stm}[D_0, n], s_1 : \text{Stm}[D_1, n])] &= \text{sbuild}(n)(\text{sdata}(p_0), \text{sdata}(p_1), \text{true}) \{ \} \{ \\
&\quad (p_0 : \text{Stm}[D_0, n]) = \{ \text{stm} : s_0, \text{ready} : \text{true} \}, (p_1 : \text{Stm}[D_1, n]) = \{ \text{stm} : s_1, \text{ready} : \text{true} \} \} \\
\mathcal{L}[\text{StmMap}(s, (x : T) \mapsto e)] &= [s/x] \mathcal{R}[\mathcal{L}[e]] \quad \mathcal{L}[\text{StmSplit}(s, m)] = s \quad \mathcal{L}[\text{StmJoin}(s)] = s
\end{aligned}$$

Figure 10. Rules for lowering high-level stream operations into the core SIROP language. $\mathcal{L}[\cdot]$ is the lowering transformation. $\mathcal{S}[\cdot]$ and $\mathcal{R}[\cdot]$ are the streamification and resetting transformations, respectively. $[e_1/x]e_2$ denotes substitution.

is represented by a register with the given initial value. The output (data and valid) is also stored in a register to avoid long combinational paths in the pipeline. Finally, for each input producer p , the ready signal is set using $p.\text{ready}$.

Algorithm 1 shows the pseudocode for converting an sbuild to VHDL. Notably, lines 2–7 describe the logic for making an sbuild wait for its consumer or producers.

Stream Sharing. `letstm` also translates to a sequential VHDL component. The queue is implemented as a circular buffer with one pointer per consumer. The VHDL entity has parameters for the queue capacity, number of consumers, etc. The SIROP VHDL generator simply instantiates this entity with appropriate parameters and connects its ports to the producer and consumers. Figure 9 is a block diagram of the hardware generated from a simple program with `letstm`.

Other Targets. Although this paper focuses on VHDL, SIROP is not limited to hardware. For example, one could compile SIROP to software where each sbuild is mapped to a thread and the handshake interfaces use queues and semaphores. The implementation is left for future work.

5 Lowering

SIROP can express all the functional parallel patterns that are typically used for hardware design: `StmMap`, `StmReduce`, `VecMap`, `VecReduce`, etc. The lowering transformation $\mathcal{L}[\cdot]$ converts them all to the core language.

`vbuild` appears in prior works [37, 42]. For example:

$$\mathcal{L}[\text{VecMap}(v : \text{Vec}[T, n], f)] = \text{vbuild}(n, i \mapsto f(v[i]))$$

Figure 10 shows a few examples of stream-related lowering rules. These are applied in a bottom-up manner, so producers are lowered before their consumers.

As mentioned in section 3, to simplify hardware generation, the core language does not support nested streams.

Algorithm 1 sbuild hardware generation

```

1 def emit(s, consumer):
2   # Handshake logic
3   send = s.valid_sig and consumer.ready_sig[s]
4   arpv = all( # all required producers valid
5     not s.ready_expr[p] or p.valid_sig
6     for p in s.producers)
7   can_step = arpv and (send or not s.valid_sig)
8   # Output registers
9   s.data_sig = create_register(name="data",
10    init=None, update=s.data,
11    update_if=can_step or send )
12   s.valid_sig = create_register(name="valid",
13    init=False, update=s.valid and arpv,
14    update_if=can_step or send )
15   # Accumulators
16   for a in s.accumulators:
17     create_register(name=a.name, init=a.init,
18      update=a.next, update_if=can_step )
19   # Producers
20   for p in s.producers:
21     s.ready_sig[p] =
22       s.ready_expr[p] and can_step
23   emit(p.stm, s)

```

Nevertheless, nested streams can be used in the higher-level language. Operations on nested streams are removed during lowering by flattening. Like the rest of the lowering pass, this is fully automatic; no manual refactoring is required.

Since the lowering pass flattens streams, `StmSplit` and `StmJoin` are trivial (figure 10). A useful consequence is that expressions like `StmSplit` $\circ$ f $\circ$ `StmJoin` will always lower to just f . If `StmSplit` and `StmJoin` were primitives, the optimizer would need yet another rewrite rule for this.

```

1 // /*RESET*/ stands for:
2 //   i != 2 && in_cnt == 2 && out_cnt == 1
3 //   || in_cnt == 2 && out_cnt == 0 && i == 2
4 sbuild(10){
5   if first then sdata(p) else sum + sdata(p),
6   i == 2
7 } {
8   (sum: u16) = {
9     init: 0,
10    next: if /*RESET*/ then 0
11          else if (first) then sdata(p)
12          else sum + sdata(p) },
13   (first: bool) = {
14     init: true,
15     next: /*RESET*/ },
16   (i: u4) = {
17     init: 0,
18     next: if /*RESET*/ then 0 else 1 + i },
19   (in_cnt: u32) = {
20     init: 0,
21     next: if /*RESET*/ then 0 else 1 + in_cnt },
22   (out_cnt: u32) = {
23     init: 0,
24     next: if /*RESET*/ then 0
25           else if (i == 2) then 1 + out_cnt
26           else out_cnt },
27 }{(p: Stm[u16, 30]) = { stm: s, ready: true } }

```

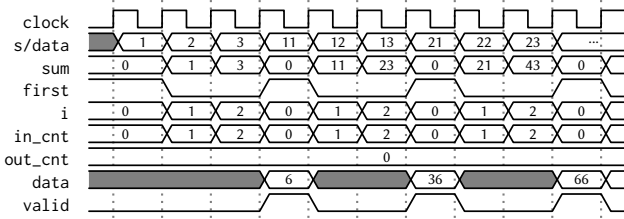


Figure 11. The lowered form of [expression 1](#) along with a timing diagram that illustrates its behaviour with input $[[1, 2, 3], [11, 12, 13], [21, 22, 23], \dots]$.

For an example of lowering `StmMap`, consider the following expression, which sums the rows of a 2D stream:

$s.\text{StmMap}((\text{row} : \text{Stm}[u_8, 3]) \mapsto \text{row}.\text{StmReduce}(+))$ (1)

The lowered expression for `StmReduce(row, +)` has already been shown in [figures 7](#) and [10](#). To implement `StmMap`, all we need is to reset the accumulators in `StmReduce` at the end of each row. [Figure 11](#) shows the lowered row sum expression. In general, lowering `StmMap` works by a “stream resetting” transformation $\mathcal{R}[\cdot]$. For each `sbuild` in the body of f :

1. Add a new accumulator to count valid outputs (e.g., `out_cnt` on lines 22–26 in [figure 11](#)).
2. Add accumulators to count inputs from each producer (e.g., `in_cnt` on lines 19–21 in [figure 11](#)).
3. Reset all accumulators, even the new counters, once the expected number of inputs have been read and the expected number of valid outputs have been produced.
4. Update the length of the stream.

```

1 sbuild(10){
2   if first then sdata(p) else sum + sdata(s),
3   i == 2
4 } {
5   (sum: u16) = {
6     init: 0,
7     next: if (i == 2) then 0
8           else if (first) then sdata(s)
9           else sum + sdata(s) },
10   (first: bool) = { init: true, next: i == 2 },
11   (i: u4) = {
12     init: 0,
13     next: if (i == 2) then 0 else 1 + i },
14 }{(p: Stm[u16, 30]) = { stm: s, ready: true } }

```

Figure 12. [Figure 11](#) after simplification ([section 6.2](#)). Notice that both `in_cnt` and `out_cnt` have been removed.

The technique described above works for `StmMap(s, f)` if $f : \text{Stm}[D_1, n_1] \rightarrow \text{Stm}[D_2, n_2]$. It can be generalized to functions of type $D_1 \rightarrow D_2$ and $D_1 \rightarrow \text{Stm}[D_2, n]$ by the “streamification” transformation $\mathcal{S}[\cdot]$, which transforms the function so its input and output are both streams. For example, if $f : D_1 \rightarrow D_2$, the streamified function $\mathcal{S}[\![f]\!]$ is

$$(x : \text{Stm}[D_1, 1]) \mapsto \text{sbuild}(1)(f(\text{sdata}(p)), \text{true}) \{ \} \{$$

$$(p : \text{Stm}[D_1, 1]) = \{ \text{stm} : x, \text{ready} : \text{true} \} \}$$

Functions of type $\text{Stm}[D_1, n] \rightarrow D_2$ are not allowed. In hardware terms, a function that has a sequential input must produce a sequential output. Functions like `StmReduce` return a stream of length one to preserve the notion of time.

6 Optimizations

This section outlines optimizations in the `SIROP` compiler that improve resource usage or maximum clock frequency.

6.1 Partial Evaluation

Partial evaluation performs computations at compile time to improve runtime performance. The `SIROP` optimizer performs simplifications such as inlining calls to function literals and applying arithmetic identities (e.g., $e + 0 \rightsquigarrow e$).

A particularly interesting example is the fusion rule for `vbuild`, which has been used in prior works [\[37, 42\]](#) and is similar to shortcut deforestation [\[14\]](#):

$$\text{vbuild}(n, f)[i] \rightsquigarrow f(i)$$

Familiar cases like `map/map` fusion follow from this rule.

$$\begin{aligned} \mathcal{L}[\![\text{VecMap}(\text{VecMap}(v, f), g)]\!] \\ &= \text{vbuild}(n, i \mapsto g(\text{vbuild}(n, j \mapsto f(v[j]))[i])) \\ &= \text{vbuild}(n, i \mapsto g(f(v[i]))) \\ &= \mathcal{L}[\![\text{VecMap}(v, g \circ f)]\!] \end{aligned}$$

$$\begin{aligned}
& \text{Fuse}_x[[\text{sbuild}(n_c)(d_c, v_c) \{ a_1, \dots, a_\ell \} \{ p_1, \dots, p_m, (x : \text{Stm}[T, n_p]) = \{ \\
& \quad \text{stm} : \text{sbuild}(n_p)(d_p, v_p) \{ a_{\ell+1}, \dots, a_{\ell+j} \} \{ p_{m+1}, \dots, p_{m+k} \}, \text{ready} : r_c \} \}]] \\
& = \text{sbuild}(n_c)(d_f, v_f) \{ a_{f,1}, \dots, a_{f,\ell+j} \} \{ p_{f,1}, \dots, p_{f,m+k} \} \\
& \text{where } d_f = [d_p/\text{sdata}(x)] d_c \quad v_f = ([d_p/\text{sdata}(x)] v_c) \wedge (r_c \implies v_p) \\
& a_{f,i} = \begin{cases} \{ \text{init} : a_i.\text{init}, \text{next} : \text{if } (r_c \implies v_p) \text{ then } [d_p/\text{sdata}(x)] a_i.\text{next} \text{ else } a_i.x \} & i \leq \ell \\ \{ \text{init} : a_i.\text{init}, \text{next} : \text{if } (v_p \implies r_c) \text{ then } a_i.\text{next} \text{ else } a_i.x \} & \text{else} \end{cases} \\
& p_{f,i} = \begin{cases} \{ \text{stm} : p_i.\text{stm}, \text{ready} : \text{if } (r_c \implies v_p) \text{ then } [d_p/\text{sdata}(x)] p_i.\text{ready} \text{ else false} \} & i \leq m \\ \{ \text{stm} : p_i.\text{stm}, \text{ready} : \text{if } (v_p \implies r_c) \text{ then } p_i.\text{ready} \text{ else false} \} & \text{else} \end{cases}
\end{aligned}$$

Figure 13. The fusion rule for streams. Occurrences of $\text{sdata}(x)$ must be replaced by the producer data d_p . Furthermore, the next and ready expressions must be updated so that the consumer waits for the producer when needed and vice-versa.

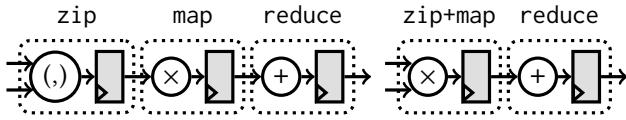


Figure 14. Dot product hardware before and after fusion.

In software, fusion is a well-known technique for eliminating the overhead of intermediate data structures [9, 14, 40, 45]. Vector fusion has similar benefits.

6.2 Basic sbuild and letstm Simplifications

There are several obvious ways to simplify the accumulators in `sbuild`. Unused accumulators are deleted; if an accumulator's value never changes then that value is inlined; and duplicate accumulators are merged. For example, in the timing diagram of figure 11, it is clear that `i` and `in_cnt` always have the same value and that `out_cnt` is always zero. Figure 12 shows the simplified expression.

For `letstm[n] x = e1 in e2`, the compiler can inline e_1 (i.e., rewrite to $[e_1/x] e_2$) if x appears at most once in e_2 . This completely eliminates the resource usage due to `letstm`.

6.3 Stream Fusion and Fission

As with vectors, it is sometimes helpful to fuse stream operators. This eliminates the overhead of the handshake interface. For instance, in the dot product example (figure 8), the optimizer will fuse the `zip` and `map` stages, as in figure 14.

Figure 13 shows the stream fusion rule, which essentially reproduces the handshake protocol within one `sbuild`. Only one rule is needed. By lowering an expression and then applying the one rule to rule them all, any combination of streaming operators can be fused: `StmMap` with `StmMap`, `StmMap` with `StmReduce`, etc. If each high-level operator was atomic, a rule would be needed for each combination. Furthermore, in prior works, the set of operators is not closed under fusion—some fusion rules are not expressible because the result is not in the list of built-in operators. For instance,

$$\begin{aligned}
& \text{Fission}[[\text{sbuild}(n)(f(d), v) \{ \vec{a} \} \{ \vec{p} \} \}] \\
& = \text{sbuild}(n)(f(\text{sdata}(x)), \text{true}) \{ \{ (x : \text{Stm}[D, n]) = \{ \\
& \quad \text{stm} : \text{sbuild}(n)(d, v) \{ \vec{a} \} \{ \vec{p} \}, \text{ready} : \text{true} \} \}
\end{aligned}$$

Figure 15. The fission rule for streams. f must not have any free variables.

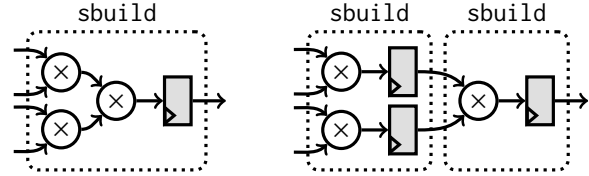


Figure 16. Parallel reduction before and after fission.

fusing `StmZip` with `StmZip` would yield an operator with three inputs. No such operator exists in `SHIR` or `Aetherling`.

To avoid generating long combinational paths, the optimizer uses a cost model for combinational delay. Fusion is only applied if the cost model estimates that the timing requirements will still be met. For example, if $C[[e]]$ is the cost of expression e , then

$$\begin{aligned}
C[[e_1 + e_2]] &= 33 + \max(C[[e_1]], C[[e_2]]) \\
C[[e_1 * e_2]] &= 100 + \max(C[[e_1]], C[[e_2]])
\end{aligned}$$

The maximum allowable delay is (somewhat arbitrarily) set at 100. Thus, up to three additions can be performed per cycle but multiplication cannot be chained with any other operations. This is highly simplistic; cost modelling is not the focus of this paper. Given a more sophisticated model, it would be straightforward to integrate it into the compiler.

SIROP also supports the opposite transformation: splitting one `sbuild` into two. For example, in figure 16, the optimizer splits up the `sbuild` to avoid performing two multiplications in one cycle. Fission is guided by the same cost model as fusion. As with fusion, only one rule is needed; see figure 15.

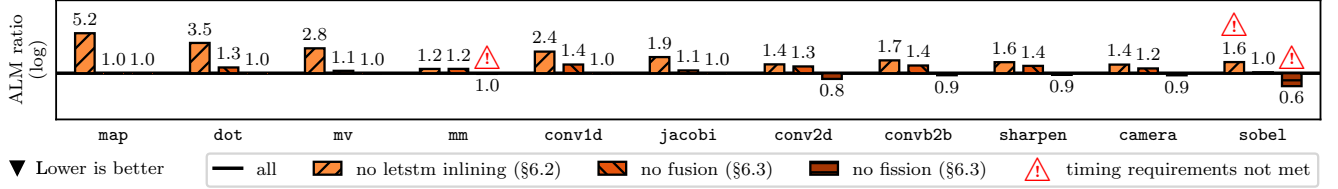


Figure 17. ALM usage ratio (other / baseline) when *turning off* various optimization passes.

7 Evaluation

This section evaluates SIROP’s hardware generation and optimization capabilities. First, the effectiveness of the optimizations described in section 6 is demonstrated. Next, SIROP is compared with Aetherling [11], SHIR [36], and the Intel HLS compiler to show that SIROP is expressive and can be used to generate fast and resource-efficient hardware.

7.1 Methodology and Benchmarks

SIROP is evaluated on a superset of the benchmarks from Aetherling [11] and SHIR [36], as shown in table 1. Designs are synthesized using Quartus Prime 19.2 Pro Edition for an Arria 10 GX FPGA. The target clock frequency is 175 MHz, like in the Aetherling paper. Resource usage is measured and timing closure is verified post-place and route. Quartus reports the numbers of ALMs (Adaptive Logic Modules), which include lookup tables and registers; BRAMs (Block RAMs), which are on-chip memory units; and DSPs (Digital Signal Processing units), which are used for multiplication.

For each program, a testbench checks correctness and reports the latency, in clock cycles, from the first input to the last output. To shorten simulation time, testbenches for large programs use inputs with fewer rows (e.g., conv2d is tested with a 16×1920 input rather than 1080×1920).

Generating VHDL via SIROP for all 124 programs in sections 7.2, 7.3, and 7.4 takes about eight minutes on a laptop with an Intel Core i7 CPU and 16 GB of RAM running Ubuntu 22.04. The slowest program takes less than 30 seconds.

7.2 Effects of Optimizations

Setup. Figure 17 shows the effects of the optimizations from section 6 on ALM usage and clock frequency. BRAM usage follows similar trends to ALM usage. DSP usage and latency do not change much because the SIROP compiler does not change the degree of spatial parallelism. The degree of parallelism can instead be chosen while lowering from the “algorithmic level” (e.g., map, reduce) to the “architectural level” (e.g., StmMap, VecReduce), as will be discussed in section 7.4.

Results. Disabling letstm inlining and fusion tends to increase ALM usage—i.e., when enabled, these passes save resources. letstm inlining is particularly effective because

Table 1. Benchmarks used to evaluate SIROP.

Name	Input Streams	Description
map	$200 \times u_8$	Add five to each element
dot	$840 \times u_{16}$ (both)	Dot product
mv	$256 \times 256 \times u_{16}$	Matrix-vector product
mm	$256 \times 256 \times u_{16}$ $256 \times 256 \times u_{16}$	Matrix-matrix product
conv1d	$16 \times i_8$	1-D convolution
jacobi	$1024 \times 128 \times u_8$	Four-point stencil $\frac{1}{4}(x_{i-1,j} + x_{i,j-1} + x_{i,j+1} + x_{i+1,j})$
conv2d	$1080 \times 1920 \times u_{32}$	3×3 convolution
convb2b	$1080 \times 1920 \times u_{32}$	3×3 , then 2×2 conv.
sharpen	$1080 \times 1920 \times u_{32}$	Image sharpening
camera	$1080 \times 1920 \times u_{32}$	Demosaic and sharpen
sobel	$1080 \times 1920 \times u_{32}$	Sobel-Feldman operator

most variables are used only once. sbuild accumulator simplification has little effect here, but disabling it does sometimes hinder other optimizations in the experiment in section 7.4. Disabling partial evaluation often yields designs so poor the simulator crashes, so it is not included in the figure.

Disabling stream fission sometimes decreases resource usage, especially for the sobel benchmark. However, this may also lead to failure to meet the timing requirements.

In short, letstm inlining and fusion reduce resource usage while fission is important for achieving timing closure.

7.3 Comparison with Prior Works

Setup. Figure 18 compares SIROP with prior works: the Intel HLS compiler, SHIR, and Aetherling. For Intel HLS, each benchmark is implemented in C++ using the `ihc::pipe` streaming interface. The `#pragma unroll` directive is added wherever spatial parallelism is desired. The shortest program, map, takes 5 lines of C++. The longest, camera, is about 140 lines. In Aetherling, each benchmark is written in the higher-level language L^{seq} using operators like map and reduce. Whenever possible, benchmark code written by the Aetherling authors themselves is used. In SHIR, each benchmark is implemented at the “architectural level” using operators like StmMap and StmReduce. The SIROP programs use the same operators as the SHIR programs. For illustration, appendix A shows the source code for the mv benchmark.

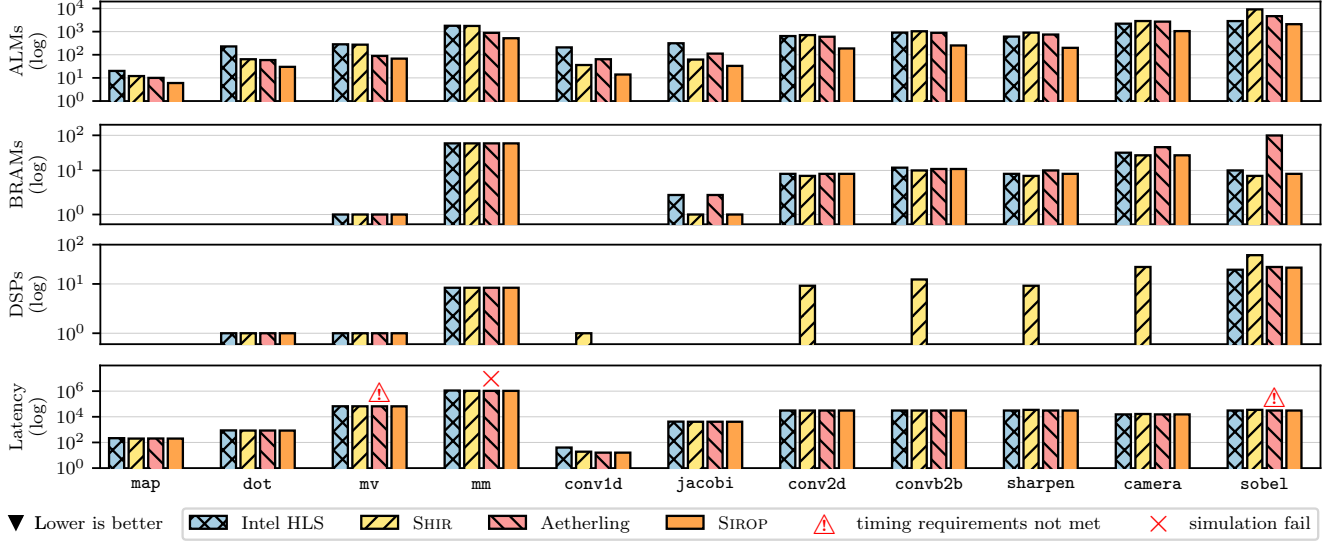


Figure 18. Resource usage and latency of designs generated by SIROP compared to prior works.

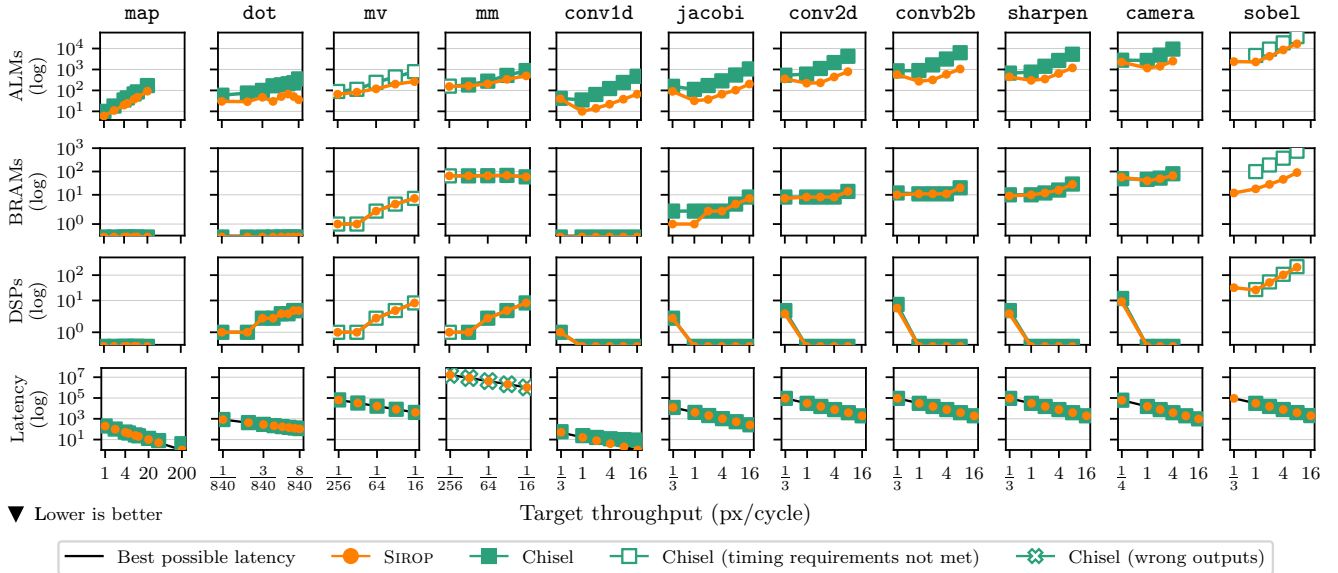


Figure 19. Resource usage and latency of designs generated using SIROP compared to Aetherling’s Chisel backend. The x-axis is the target throughput given to the Aetherling compiler, which determines the degree of spatial parallelism. The “best possible latency” is $\max(\text{in_len}, \text{out_len}) / \text{throughput}$. The plots are log-log because the target throughputs generally increase exponentially. An empty shape means the design does not meet the timing requirements or requires too many ports.

Results. SIROP consistently generates designs with fewer ALMs than all prior works: on average, 76% fewer than Intel HLS, 68% fewer than SHIR, and 61% fewer than Aetherling. BRAM and DSP usage is usually similar across tools. No DSPs are needed for conv1d, jacobi, conv2d, convb2b, sharpen, and camera because the kernels only multiply by powers of two and can therefore use bitwise shifts.

One factor in SIROP’s better ALM efficiency is its ability to fuse any two stream operators, as discussed in section 7.2.

Another reason is that, in SIROP, it is straightforward to omit or remove unnecessary accumulators. In the map benchmark, for instance, SHIR uses a hardware template that is general enough to handle nested streams. The general template has two counters. In SIROP, by contrast, the final design does not have any counters. The stream resetting transformation (section 5) does add input and output counters; however, for a 1D stream, these counters will be unused and it is trivial for the optimizer to remove them (section 6.2).

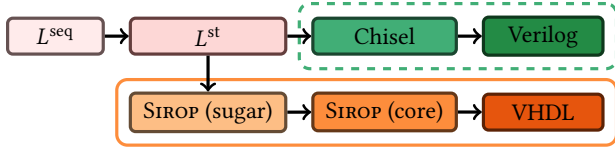


Figure 20. Chisel and SIROP compilation flows.

Notes. The Aetherling mm design produces incorrect results. One stream reaches the multiplier with a delay of one cycle, whereas the other has a delay of two cycles. SIROP uses dynamic scheduling, which rules out such issues.

For the sharpen, sobel, and camera benchmarks, SHIR requires the programmer to manually insert registers (using a Registered primitive in the IR). Without manual intervention, latency increases by a factor of 4, 4.8, and 2.8, respectively, and the sobel design does not meet the timing requirements. Some Aetherling designs also do not meet the timing requirements. Intel HLS and SIROP achieve timing closure and reach the expected throughput for every benchmark without manual intervention.

7.4 Design Space Exploration with Aetherling

Setup. SIROP is also evaluated by using it as a backend to Aetherling, as shown in figure 20. Benchmarks are written in Aetherling’s high-level language, L^{seq} . Aetherling lowers these programs to its “space-time” language with streams and vectors, L^{st} . The lowering process aims to minimize resource usage while achieving a given steady-state throughput. This step determines the degree of spatial parallelism. The resulting L^{st} programs are parsed and compiled using SIROP. Aetherling’s existing backend emits Chisel [3], which is then translated to Verilog. With the Chisel backend, Aetherling has been shown to produce designs with lower resource utilization than Spatial [23], a high-level HDL, and Halide-HLS [33], an extension of the Halide image processing language.

Results. Figure 19 shows the resource usage for different target throughputs. As expected, higher throughputs require more spatial parallelism and thus more resources. Like in section 7.3, the designs produced by SIROP always use fewer ALMs than those produced by Aetherling’s Chisel backend. Figure 19 also shows that, for both Chisel and SIROP, the measured latency is close to the ideal latency for the given target throughput. This shows that SIROP can accurately express all the same design points as Aetherling.

Note. Some designs require more ports than the target device provides. For example, moving 16 values of type u_{32} in and out per cycle requires at least $2 \cdot 16 \cdot 32$ ports, which is too many; thus, conv2d cannot be synthesized at a throughput of 16 with the bitwidth set by the Aetherling authors. This is the case for both SIROP and Aetherling.

Table 2. Lines of code for the hardware generation step.

	Scala	Haskell	VHDL	Verilog
Aetherling [11]	1,411	771	–	168
+ Chisel [3]	+ 8,577			
SHIR [36]	3,931	–	8,009	–
SIROP	2,657	–	289	–

Due to a minor bug in Aetherling’s L^{st} to Chisel conversion, no Verilog is generated for sobel with throughput $1/3$. Furthermore, as mentioned in section 7.3, Aetherling’s mm designs produce wrong outputs.

7.5 Engineering Effort

We now show how SIROP simplifies compiler development.

Lines of Code. Table 2 shows the size of each functional HLS compiler’s hardware generator. The SIROP VHDL generator is mostly written in Scala, with some VHDL for the letstm hardware template. It uses no external dependencies except for logging and file I/O. Hardware generation in Aetherling itself is more concise, but this is because it targets Chisel, a higher-level HDL. The Chisel project is far larger than SIROP’s hardware generator. SHIR has over 100 hardware templates for various operators, which add up to over 8,000 lines of VHDL.

Extensibility. It is difficult to express a 2D convolution in SHIR without using host memory. To obtain the results in section 7.3, a new streaming operator, StmSlide2D, is added to both SHIR and SIROP. In SHIR, this requires providing a VHDL implementation for the operator. In total, the change took 95 lines of VHDL and 109 lines of Scala code. In SIROP, the operator is instead lowered to the core language, which took only 88 lines of Scala code. The SIROP optimizer and hardware generator can handle the new operator with no changes. The SHIR optimizer, by contrast, would need new rewrite rules to be able to transform the new operator.

In addition to SIROP being a higher-level target, the compiler includes features that simplify testing and debugging. For example, it has a simple interpreter that can be used to test the correctness of the new operator directly in Scala and inspect the value of each sbuild accumulator at each time step. In SHIR, one must write and run a VHDL testbench.

Summary. The simplicity of the core language enables more concise hardware generation than prior works. Furthermore, it is easy to extend the higher-level language with new syntax sugar. All this comes without sacrificing hardware quality. In short, this new, minimal set of streaming primitives is a strong foundation for functional HLS.

8 Related Work

Functional HLS. SHIR [36] and LIFTHLS [24] compile functional parallel patterns to dynamically-scheduled streaming accelerators. A key feature of SHIR is its ability to represent and optimize host memory accesses. There has also been recent work exploring resource sharing in SHIR for large designs [21]. Coarse-grained resource sharing in SIROP is left for future work, as are host memory accesses and memory-related transformations such as banking. We expect many of the techniques used in SHIR to be transferable to SIROP.

Aetherling [11] compiles similar functional programs to streaming accelerators with static schedules. All Aetherling operators have fixed latency and communicate at fixed rates. This enables static scheduling with no overhead—not even a control unit—but makes it impossible to support variable-latency operations. For example, consider a program that decides whether a stream contains a given element. In SIROP, the design can produce an affirmative result immediately if the element is found, even if it occurs near the beginning of the stream. In Aetherling, by contrast, one would have no choice but to consume the entire stream before answering.

SHIR, LIFTHLS, and Aetherling all have many stream and vector operators, which complicates hardware generation and optimization. SIROP has a far more compact IR and generates more resource-efficient accelerators.

Imperative HLS. Many HLS tools use imperative languages, especially C and C++. This is convenient for C and C++ programmers, but imperative languages are a poor fit for hardware design. On one hand, many features (e.g., dynamic memory allocation) do not map well to hardware. On the other hand, it can be difficult to infer pipeline and spatial parallelism from loop-based code. C-based HLS frameworks such as the Intel HLS compiler, Xilinx HLS, and LegUp [5] rely on pragmas to guide synthesis. Other tools, such as Halide-HLS [33], HeteroCL [26], Allo [6], and POM [48] use a separate scheduling language to clarify intent. In practice, the generated hardware tends to have unpredictable performance and source code with pragmas is non-portable [19, 29].

Some authors have tried addressing the portability and reliability issues. Dahlia [29] has a type system to enforce unwritten rules of HLS. AnyHLS [32] generates more portable C-HLS code starting from a yet higher-level language. One can also augment C-HLS with functional parallel skeletons [19]. However, forgoing imperative programming for HLS altogether might be more reliable long-term.

There has been work [1, 6, 13, 30, 46, 48, 49] on using MLIR [27] to optimize hardware designs at multiple levels of abstraction. Although the SIROP compiler is currently implemented as a standalone project, the SIROP IR could just as well be implemented as an MLIR dialect. SIROP is particularly well-suited to implementing the parallel patterns used in prior works [11, 19, 24, 25, 36] while being easy to translate to hardware.

Other Hardware Languages. Some languages have been designed to map well to hardware while providing a higher level of abstraction than traditional HDLs. Notable examples include Esterel [12], Bluespec [31], C λ SH [2], Spatial [23], Fleet [41], and Reticle [43]. There are also embedded HDLs and libraries (e.g., Chisel [3], Lava [4], PyRTL [8], MyHDL) that allow metaprogramming in a popular host language. All these are still best suited for users with hardware design experience. Languages like SHIR aim to be even higher-level, letting users focus only on the algorithm. This paper has the same use case in mind, although it focuses on the IR rather than the highest-level language.

Domain-Specific Languages. There are many DSLs for HLS in machine learning [10, 38, 44, 47], image processing [7, 16, 17, 28, 33, 34, 39], etc. SIROP aims to be general enough to generate streaming accelerators in any domain.

Minimal Array Languages. `vbuild` is used in prior IRs for software, including for destination-passing style [37] and idiom recognition [42]. Other works on deforestation [9, 14, 40] use similar ideas. This paper applies the principle of using a minimal IR, but in a new context (HLS) and with a new IR.

Dataflow Programming. SIROP is similar to KPNs (Kahn Process Networks) [22] in that programs consist of many “computing stations” (`sbuild`) running in parallel. However, to minimize resource usage, each `sbuild` node communicates with its consumer directly, without a queue. Sending data must therefore be blocking. FIFO queues are only added for `letstm`, and these have a fixed capacity. Also, unlike in KPNs, `sbuild` can wait for *many* producers at once; this allows `StmZip` to be implemented without any memory.

SIROP, like SHIR and Aetherling, currently only supports acyclic pipelines. This limitation could perhaps be lifted by adding a fourth stream primitive `letrec  $x = e_1$  in  $e_2$` , where e_1 can use x to access its own output.

9 Conclusion

This paper has demonstrated the benefits of using a minimal functional language for high-level synthesis. With fewer primitives, there are fewer hardware templates to implement and test. Furthermore, key optimizations like stream fusion and fission can be implemented once and for all. All the familiar functional primitives can be expressed as syntax sugar; this new language is *more* expressive than Aetherling and at least as expressive as the subset of SHIR that does not use memory. Moreover, SIROP generates fast and resource-efficient hardware—there is no loss of hardware quality.

Acknowledgments

We acknowledge the support of the Natural Sciences and Engineering Research Council of Canada (NSERC), grants RGPIN-2020-05889 and RGPIN-2026-06336, as well as the CIFAR AI Chairs Program.

Data Availability Statement

The artifact for this paper is available on Zenodo [18]. It includes SIROP, SHIR, Aetherling, and Intel HLS source code for all benchmarks. One example is provided below for easy reference. The artifact also includes a parser that translates any L^{st} program to an equivalent high-level SIROP program. This shows that SIROP is at least as expressive as Aetherling.

A Benchmark Source Code

The following listings contain the source code for the mv benchmark, which performs matrix-vector multiplication.

SIROP. The SIROP implementation of matrix-vector multiplication is very similar to the dot product code in figure 4. Here, the matrix `mat` is a 2D stream and `StmMap` is used to compute the dot product row-by-row.

```
1 (mat : Stm[Stm[u16], 256], 256]) =>
2 (vec : Stm[u16], 256] ) =>
3   mat.StmMap( (row: Stm[u16], 256]) =>
4     /* dot(row, vec) */
5     StmZip(row, vec)
6     .StmMap( (x) => x.0 *% x.1)
7     .StmReduce( (x) => x.0 +% x.1 ) )
```

SHIR. The SHIR source code looks superficially different from the SIROP source code because SHIR is embedded in Scala rather than having a custom parser. However, the actual operators used are identical.

```
1 val u16 = IntType(16)
2 val mat = ParamDef(OrderedStreamType(
3   OrderedStreamType(u16, 256), 256 ))
4 val vec = ParamDef(OrderedStreamType(u16, 256))
5 val row = ParamDef(OrderedStreamType(u16, 256))
6 val xy = ParamDef(TupleType(u16, u16))
7 val acc = ParamDef()
8 val x = ParamDef()
9 ArchLambdas(Seq(mat, vec),
10  MapOrderedStream(
11    ArchLambda(row,
12      // StmReduce
13      FoldOrderedStream(
14        // (acc, x) => acc + x
15        ArchLambdas(Seq(acc, x),
16          AddInt(Tuple2(
17            ParamUse(acc),
18            ParamUse(x) ) ) ),
19        // StmMap
20        MapOrderedStream(
21          // (xy) => xy.0 * xy.1
22          ArchLambda(xy,
23            TruncInteger(
24              MulInt(ParamUse(xy)),
25              16 ) ),
26          // StmZip(row, vec)
27          Zip2OrderedStream(Tuple2(
28            ParamUse(row),
29            ParamUse(vec) ) ) ) ),
30    ParamUse(mat) ) )
```

Aetherling. The Aetherling source code differs from the SIROP and SHIR code in that one must explicitly repeat the input `vec`. SIROP and SHIR automatically insert a streaming operator to perform this repetition.

```
1 big_mvm =
2   let mat =
3     com_input_seq "I0"
4     (Proxy::Proxy (Seq 65536 Atom_UInt16))
5   in
6   let vec =
7     com_input_seq "I1"
8     (Proxy::Proxy (Seq 256 Atom_UInt16))
9   in
10  -- Doesn't work :( -----
11  -- mapC (\row -> dot row vec) $
12  --   partitionC (Proxy @256) (Proxy @256) mat
13  -----
14  let repeated_vec = unpartitionC $
15    up_1dC (Proxy @256) $
16    partitionC (Proxy @1)
17    (Proxy @256)
18    vec in
19  -- Element-wise product of vec with each row
20  let products2d = partitionC (Proxy @256)
21    (Proxy @256) $
22    map2C (\x -> \y -> mulC $
23      atom_tupleC x y)
24    mat repeated_vec in
25  -- Sum each row
26  unpartitionC $ mapC (reduceC addC) products2d
```

Intel HLS. The Intel HLS code is written in C++. Inputs and outputs are provided using the `ihc::pipe` class, which exposes a ready/valid interface like the one in figure 3. The input `vec` is saved in an array so it can be read repeatedly.

```
1 #include "HLS/ac_int.h"
2 #include "HLS/hls.h"
3
4 constexpr int N = 256;
5 template<unsigned SystemID> class PipeID {};
6 ihc::pipe<class PipeID<0>, uint16> mat;
7 ihc::pipe<class PipeID<1>, uint16> vec;
8 ihc::pipe<class PipeID<2>, uint16> out;
9
10 component void matvec() {
11   uint16 vec_arr[N];
12   for (int i = 0; i < N; i++) {
13     uint16 sum = 0;
14     for (int j = 0; j < N; j++) {
15       uint16 v;
16       if (i == 0) {
17         v = vec.read();
18         vec_arr[j] = v;
19       } else {
20         v = vec_arr[j];
21       }
22       sum += mat.read() * v;
23     }
24     out.write(sum);
25   }
26 }
```


References

- [1] Nicolas Bohm Agostini, Serena Curzel, Jeff Jun Zhang, Ankur Limaye, Cheng Tan, Vinay Amatya, Marco Minutoli, Vito Giovanni Castellana, Joseph Manzano, David Brooks, Gu-Yeon Wei, and Antonino Tumeo. 2022. Bridging Python to Silicon: The SODA Toolchain. *IEEE Micro* 42, 5 (2022). doi:10.1109/MM.2022.3178580
- [2] C.P.R. Baaij. 2015. *Digital circuit in CλaSH: functional specifications and type-directed synthesis*. PhD Thesis - Research UT, graduation UT. University of Twente, Netherlands. doi:10.3990/1.9789036538039
- [3] Jonathan Bachrach, Huy Vo, Brian Richards, Yunsup Lee, Andrew Waterman, Rimas Avizienis, John Wawrzyniak, and Krste Asanović. 2012. Chisel: Constructing hardware in a Scala embedded language. In *Proceedings of the 49th Annual Design Automation Conference (DAC '12)*. doi:10.1145/2228360.2228584
- [4] Per Bjesse, Koen Claessen, Mary Sheeran, and Satnam Singh. 1998. Lava: hardware design in Haskell. *SIGPLAN Not.* 34, 1 (1998). doi:10.1145/291251.289440
- [5] Andrew Canis, Jongsok Choi, Mark Aldham, Victor Zhang, Ahmed Kammoona, Tomasz Czajkowski, Stephen D. Brown, and Jason H. Anderson. 2013. LegUp: An open-source high-level synthesis tool for FPGA-based processor/accelerator systems. *ACM Trans. Embed. Comput. Syst.* 13, 2, Article 24 (Sept. 2013). doi:10.1145/2514740
- [6] Hongzheng Chen, Niansong Zhang, Shaojie Xiang, Zhichen Zeng, Mengjia Dai, and Zhiru Zhang. 2024. Allo: A Programming Model for Composable Accelerator Design. *Proc. ACM Program. Lang.* 8, PLDI (June 2024). doi:10.1145/3656401
- [7] Nitin Chugh, Vinay Vasista, Suresh Purini, and Uday Bondhugula. 2016. A DSL Compiler for Accelerating Image Processing Pipelines on FPGAs. In *Proceedings of the 2016 International Conference on Parallel Architectures and Compilation (PACT '16)*. doi:10.1145/2967938.2967969
- [8] John Clow, Georgios Tzimpragos, Deeksha Dangwal, Sammy Guo, Joseph McMahan, and Timothy Sherwood. 2017. A pythonic approach for rapid hardware prototyping and instrumentation. In *2017 27th International Conference on Field Programmable Logic and Applications (FPL '17)*. doi:10.23919/FPL.2017.8056860
- [9] Duncan Coutts, Roman Leshchinskiy, and Don Stewart. 2007. Stream fusion: From lists to streams to nothing at all. In *Proceedings of the 12th ACM SIGPLAN International Conference on Functional Programming (ICFP '07)*. doi:10.1145/1291151.1291199
- [10] Javier Duarte, Song Han, Philip Harris, Sergo Jindariani, Edward Kreinar, Benjamin Kreis, Vladimir Loncar, Jennifer Ngadiuba, Maurizio Pierini, Dylan Rankin, Ryan Rivera, Sioni Summers, Nhan Tran, and Zhenbin Wu. 2019. Fast Inference of Deep Neural Networks for Real-time Particle Physics Applications. In *Proceedings of the 2019 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays (FPGA '19)*. doi:10.1145/3289602.3293986
- [11] David Durst, Matthew Feldman, Dillon Huff, David Akeley, Ross Daly, Gilbert Louis Bernstein, Marco Patrignani, Kayvon Fatahalian, and Pat Hanrahan. 2020. Type-directed scheduling of streaming accelerators. In *Proceedings of the 41st ACM SIGPLAN Conference on Programming Language Design and Implementation (PLDI 2020)*. doi:10.1145/3385412.3385983
- [12] Stephen A. Edwards. 2002. High-Level Synthesis from the Synchronous Language Esterel. In *11th IEEE/ACM International Workshop on Logic & Synthesis (IWLS 2002)*. <https://www.cs.columbia.edu/~sedwards/papers/edwards2002high-level.pdf>
- [13] Schuyler Eldridge, Prithayan Barua, Aliaksei Chapyzhenka, Adam Izraelevitz, Jack Koenig, Chris Lattner, Andrew Lenharth, George Leon-tiev, Fabian Schuiki, Ram Sunder, Andrew Young, and Richard Xia. 2021. MLIR as Hardware Compiler Infrastructure. In *Workshop on Open-Source EDA Technology (WOSET)*.
- [14] Andrew Gill, John Launchbury, and Simon Peyton Jones. 1993. A short cut to deforestation. In *Proceedings of the Conference on Functional Programming Languages and Computer Architecture* (Copenhagen, Denmark) (FPCA '93). doi:10.1145/165180.165214
- [15] Sergei Gorlatch and Murray Cole. 2011. *Parallel Skeletons*. doi:10.1007/978-0-387-09766-4_24
- [16] James Hegarty, John Brunhaver, Zachary DeVito, Jonathan Ragan-Kelley, Noy Cohen, Steven Bell, Artem Vasilyev, Mark Horowitz, and Pat Hanrahan. 2014. Darkroom: compiling high-level image processing code into hardware pipelines. *ACM Trans. Graph.* 33, 4, Article 144 (July 2014). doi:10.1145/2601097.2601174
- [17] James Hegarty, Ross Daly, Zachary DeVito, Jonathan Ragan-Kelley, Mark Horowitz, and Pat Hanrahan. 2016. Rigel: flexible multi-rate image processing hardware. *ACM Trans. Graph.* 35, 4, Article 85 (July 2016). doi:10.1145/2897824.2925892
- [18] Louis Hildebrand and Christophe Dubach. 2025. *Sirop: A Small IR for HLS with Parallel Patterns (LCTES '26 Artifact)*. doi:10.5281/zenodo.20090976
- [19] Lana Josipović, Nithin George, and Paolo Ienne. 2016. Enriching C-based high-level synthesis with parallel pattern templates. In *International Conference on Field-Programmable Technology (FPT '16)*. doi:10.1109/FPT.2016.7929527
- [20] Lana Josipović, Radhika Ghosal, and Paolo Ienne. 2018. Dynamically Scheduled High-level Synthesis. In *Proceedings of the 2018 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays (FPGA '18)*. doi:10.1145/3174243.3174264
- [21] Tzung-Han Juang, Christof Schlaak, and Christophe Dubach. 2023. Let Coarse-Grained Resources Be Shared: Mapping Entire Neural Networks on FPGAs. *ACM Trans. Embed. Comput. Syst.* 22, 5s (Sept. 2023). doi:10.1145/3609109
- [22] Gilles Kahn. 1974. The Semantics of a Simple Language for Parallel Programming. In *IFIP Congress*. <https://api.semanticscholar.org/CorpusID:18030506>
- [23] David Koeplinger, Matthew Feldman, Raghu Prabhakar, Yaqi Zhang, Stefan Hadjis, Ruben Fisel, Tian Zhao, Luigi Nardi, Ardavan Pedram, Christos Kozyrakis, and Kunle Olukotun. 2018. Spatial: a language and compiler for application accelerators. In *Proceedings of the 39th ACM SIGPLAN Conference on Programming Language Design and Implementation (PLDI 2018)*. doi:10.1145/3192366.3192379
- [24] Martin Kristien, Bruno Bodin, Michel Steuwer, and Christophe Dubach. 2019. High-level synthesis of functional patterns with Lift. In *Proceedings of the 6th ACM SIGPLAN International Workshop on Libraries, Languages and Compilers for Array Programming (ARRAY 2019)*. doi:10.1145/3315454.3329957
- [25] Morihiro Kuga, Kansuke Fukuda, Motoki Amagasaki, Masahiro Iida, and Toshinori Sueyoshi. 2017. High-level Synthesis based on Parallel Design Patterns using a Functional Language. In *Proceedings of the 8th International Symposium on Highly Efficient Accelerators and Reconfigurable Technologies (HEART '17)*. doi:10.1145/3120895.3120918
- [26] Yi-Hsiang Lai, Yuze Chi, Yuwei Hu, Jie Wang, Cody Hao Yu, Yuan Zhou, Jason Cong, and Zhiru Zhang. 2019. HeteroCL: A Multi-Paradigm Programming Infrastructure for Software-Defined Reconfigurable Computing. In *Proceedings of the 2019 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays (FPGA '19)*. doi:10.1145/3289602.3293910
- [27] Chris Lattner, Mehdi Amini, Uday Bondhugula, Albert Cohen, Andy Davis, Jacques Pienaar, River Riddle, Tatiana Shpeisman, Nicolas Vasilache, and Oleksandr Zinenko. 2021. MLIR: Scaling compiler infrastructure for domain specific computation. In *Proceedings of the 2021 IEEE/ACM International Symposium on Code Generation and Optimization (CGO '21)*. doi:10.1109/CGO51591.2021.9370308
- [28] Jiajie Li, Yuze Chi, and Jason Cong. 2020. HeteroHalide: From Image Processing DSL to Efficient FPGA Acceleration. In *Proceedings of the 2020 ACM/SIGDA International Symposium on Field-Programmable Gate Arrays (FPGA '20)*. doi:10.1145/3373087.3375320
- [29] Rachit Nigam, Sachille Atapattu, Samuel Thomas, Zhijing Li, Theodore Bauer, Yuwei Ye, Apurva Koti, Adrian Sampson, and Zhiru Zhang. 2020. Predictable accelerator design with time-sensitive affine types.

- In *Proceedings of the 41st ACM SIGPLAN Conference on Programming Language Design and Implementation (PLDI 2020)*. doi:10.1145/3385412.3385974
- [30] Rachit Nigam, Samuel Thomas, Zhijiang Li, and Adrian Sampson. 2021. A compiler infrastructure for accelerator generators. In *Proceedings of the 26th ACM International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS '21)*. doi:10.1145/3445814.3446712
- [31] Rishiyur S. Nikhil. 2008. *Bluespec: A General-Purpose Approach to High-Level Synthesis Based on Parallel Atomic Transactions*. doi:10.1007/978-1-4020-8588-8_8
- [32] M. Akif Özkan, Arsène Pérard-Gayot, Richard Membarth, Philipp Slusallek, Roland Leißa, Sebastian Hack, Jürgen Teich, and Frank Hannig. 2020. AnyHLS: High-Level Synthesis With Partial Evaluation. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems* 39, 11 (2020). doi:10.1109/TCAD.2020.3012172
- [33] Jing Pu, Steven Bell, Xuan Yang, Jeff Setter, Stephen Richardson, Jonathan Ragan-Kelley, and Mark Horowitz. 2017. Programming Heterogeneous Systems from an Image Processing DSL. *ACM Trans. Archit. Code Optim.* 14, 3, Article 26 (Aug. 2017). doi:10.1145/3107953
- [34] Oliver Reiche, M. Akif Özkan, Richard Membarth, Jürgen Teich, and Frank Hannig. 2017. Generating FPGA-based image processing accelerators with Hipacc: (Invited paper). In *2017 IEEE/ACM International Conference on Computer-Aided Design (ICCAD)*. doi:10.1109/ICCAD.2017.8203894
- [35] Kyle Rupnow, Yun Liang, Yanan Li, and Deming Chen. 2011. A study of high-level synthesis: Promises and challenges. In *9th IEEE International Conference on ASIC*. doi:10.1109/ASICON.2011.6157401
- [36] Christof Schlaak, Tzung-Han Juang, and Christophe Dubach. 2022. Memory-Aware Functional IR for Higher-Level Synthesis of Accelerators. *ACM Trans. Archit. Code Optim.* 19, 2 (Jan. 2022). doi:10.1145/3501768
- [37] Amir Shaikhha, Andrew Fitzgibbon, Simon Peyton Jones, and Dimitrios Vytiniotis. 2017. Destination-passing style for efficient memory management. In *Proceedings of the 6th ACM SIGPLAN International Workshop on Functional High-Performance Computing (FHPC 2017)*. doi:10.1145/3122948.3122949
- [38] Hardik Sharma, Jongse Park, Divya Mahajan, Emmanuel Amaro, Joon Kyung Kim, Chenkai Shao, Asit Mishra, and Hadi Esmaeilzadeh. 2016. From high-level deep neural models to FPGAs. In *The 49th Annual IEEE/ACM International Symposium on Microarchitecture (MICRO-49)*. doi:10.1109/MICRO.2016.7783720
- [39] Robert Stewart, Kirsty Duncan, Greg Michaelson, Paulo Garcia, Deepayan Bhowmik, and Andrew Wallace. 2018. RIPL: A Parallel Image Processing Language for FPGAs. *ACM Trans. Reconfigurable Technol. Syst.* 11, 1, Article 7 (March 2018). doi:10.1145/3180481
- [40] Josef Svenningsson. 2002. Shortcut fusion for accumulating parameters & zip-like functions. In *Proceedings of the Seventh ACM SIGPLAN International Conference on Functional Programming (ICFP '02)*. doi:10.1145/581478.581491
- [41] James Thomas, Pat Hanrahan, and Matei Zaharia. 2020. Fleet: A Framework for Massively Parallel Streaming on FPGAs. In *Proceedings of the Twenty-Fifth International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS '20)*. doi:10.1145/3373376.3378495
- [42] Jonathan Van der Cruysse and Christophe Dubach. 2024. Latent Idiom Recognition for a Minimalist Functional Array Language Using Equality Saturation. In *Proceedings of the 2024 IEEE/ACM International Symposium on Code Generation and Optimization (CGO '24)*. doi:10.1109/CGO57630.2024.10444879
- [43] Luis Vega, Joseph McMahan, Adrian Sampson, Dan Grossman, and Luis Ceze. 2021. Reticle: a virtual machine for programming modern FPGAs. In *Proceedings of the 42nd ACM SIGPLAN International Conference on Programming Language Design and Implementation (PLDI 2021)*. doi:10.1145/3453483.3454075
- [44] Stylianos I. Venieris and Christos-Savvas Bouganis. 2016. fpgaConvNet: A Framework for Mapping Convolutional Neural Networks on FPGAs. In *IEEE 24th Annual International Symposium on Field-Programmable Custom Computing Machines (FCCM '16)*. doi:10.1109/FCCM.2016.22
- [45] Philip Wadler. 1990. Deforestation: Transforming programs to eliminate trees. *Theoretical Computer Science* 73, 2 (1990). doi:10.1016/0304-3975(90)90147-A
- [46] Hanchen Ye, HyeGang Jun, Hyunmin Jeong, Stephen Neuendorffer, and Deming Chen. 2022. ScaleHLS: a scalable high-level synthesis framework with multi-level transformations and optimizations: invited. In *Proceedings of the 59th ACM/IEEE Design Automation Conference (DAC '22)*. doi:10.1145/3489517.3530631
- [47] Chen Zhang, Guangyu Sun, Zhenman Fang, Peipei Zhou, Peichen Pan, and Jason Cong. 2019. Caffeine: Toward Uniformed Representation and Acceleration for Deep Convolutional Neural Networks. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems* 38, 11 (2019). doi:10.1109/TCAD.2017.2785257
- [48] Weichuang Zhang, Jieru Zhao, Guan Shen, Quan Chen, Chen Chen, and Minyi Guo. 2024. An Optimizing Framework on MLIR for Efficient FPGA-based Accelerator Generation. In *2024 IEEE International Symposium on High-Performance Computer Architecture (HPCA)*. doi:10.1109/HPCA57654.2024.00017
- [49] Ruizhe Zhao, Jianyi Cheng, Wayne Luk, and George A. Constantinides. 2022. POLSCA: Polyhedral High-Level Synthesis with Compiler Transformations. In *2022 32nd International Conference on Field-Programmable Logic and Applications (FPL)*. doi:10.1109/FPL57034.2022.00044

Received 2026-03-20; accepted 2026-05-01