

neural.slang

A Standard Module for Inline Neural Networks in Shaders

Kai Zhang



What is Slang?

One modern shading language, compiled to every major GPU platform.



modern language features

cross-platform by design

automatic differentiation

Graphics compute keeps evolving

Graphics has always used better compute tools to build better approximations.



The latest step is letting neural networks learn some of the approximations for us.



Neural compute moves inside the pipeline

Start from the rendering pipeline we already know.

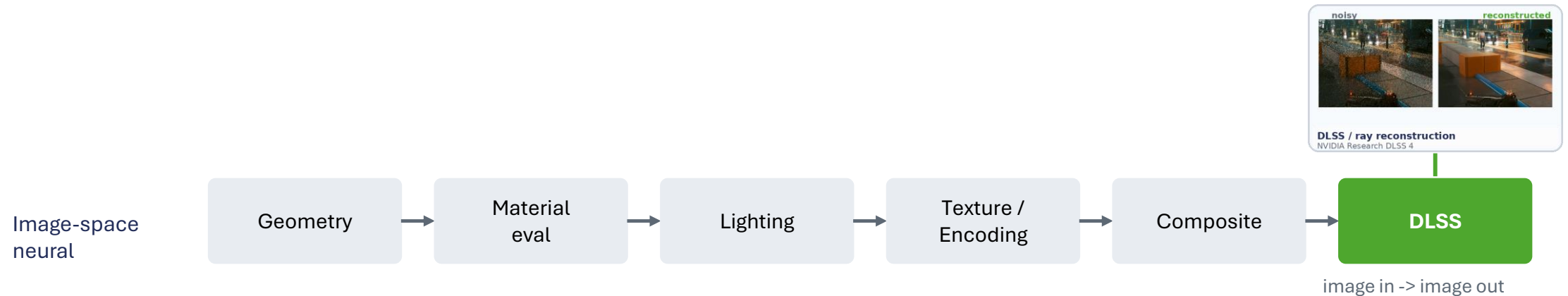


Geometry, material evaluation, lighting, texture/encoding, compositing, and final image processing.



Neural compute moves inside the pipeline

DLSS is the familiar case: neural compute after rendering.



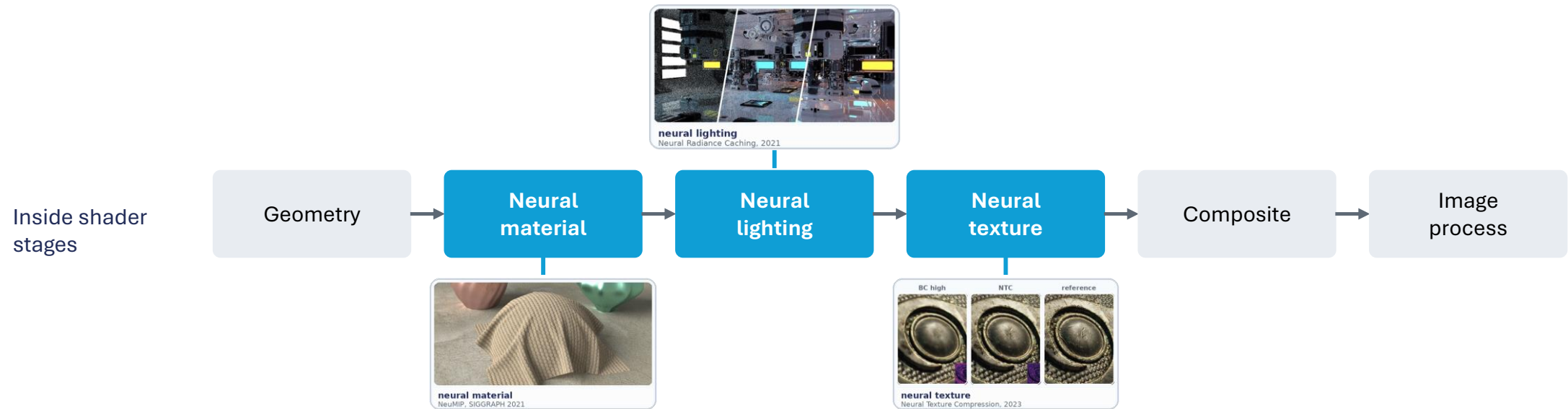
The neural network replaces the final image-processing box, while the middle of the pipeline remains traditional.

Image source: NVIDIA Research DLSS 4 / Ray Reconstruction



Neural compute moves inside the pipeline

Neural compute can live inside shader stages.



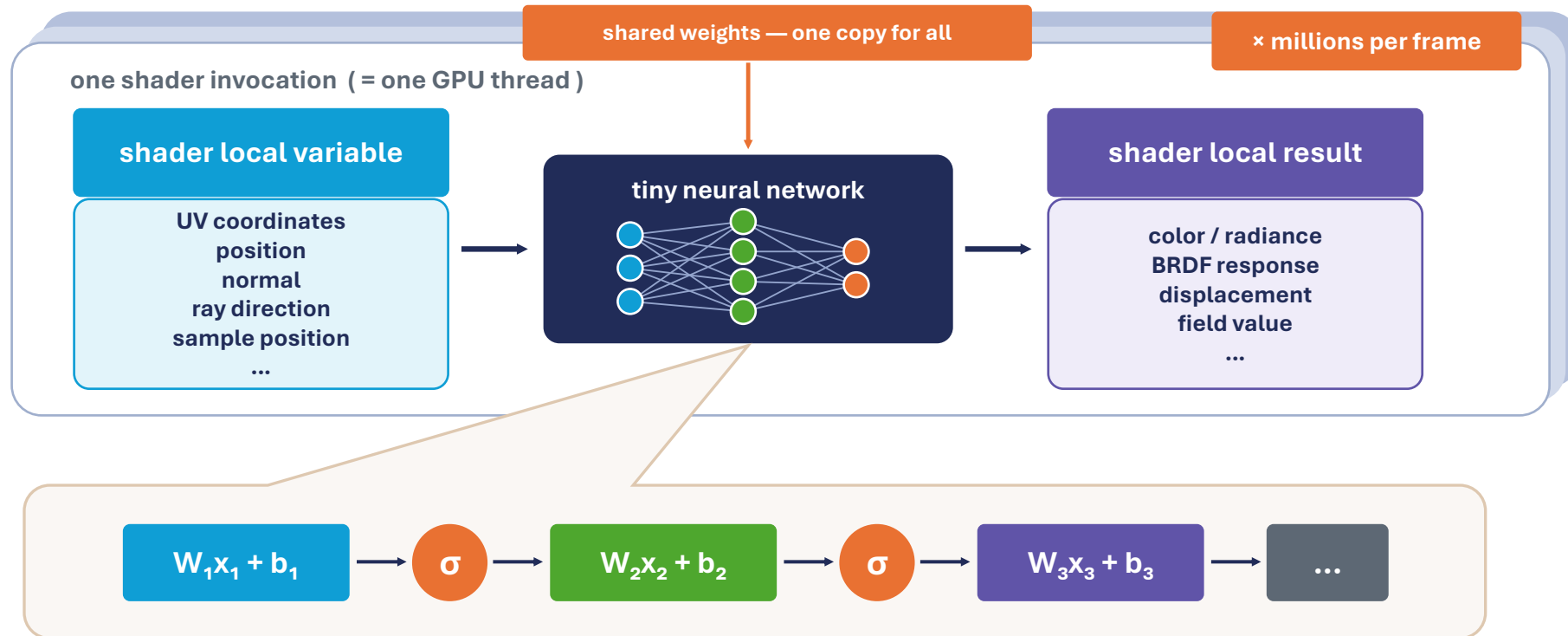
Same pipeline architecture; different compute inside selected boxes.

Image sources: NeuMIP (SIGGRAPH 2021), Neural Radiance Caching (2021), Random-Access Neural Compression of Material Textures (SIGGRAPH 2023)



The neural shading workload

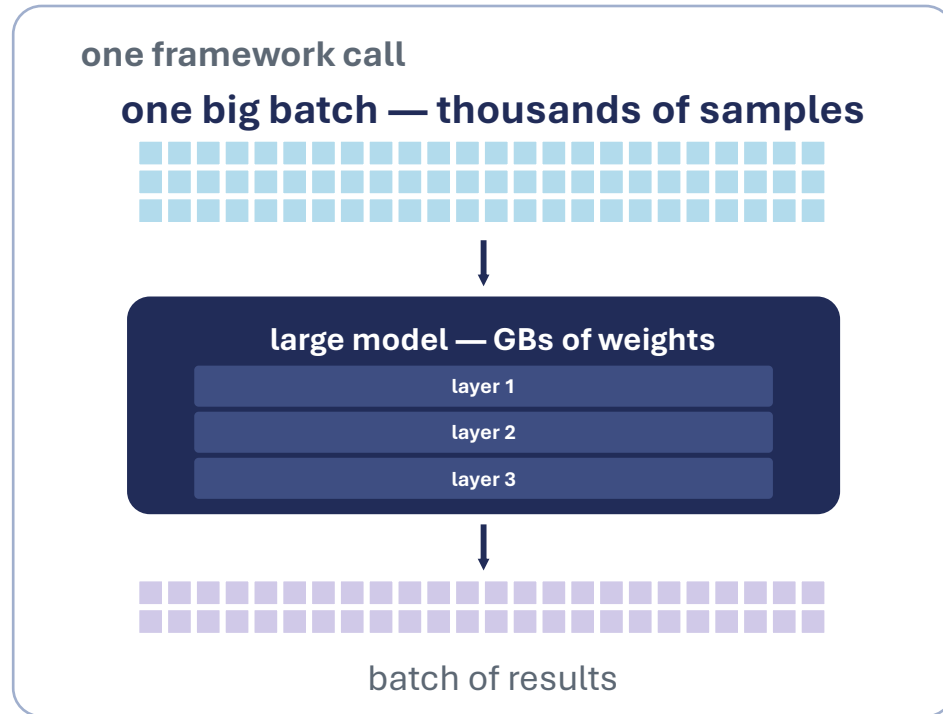
A neural building block is a tiny learned function inside the shader.



The workload = one tiny Mat-Vec chain × millions of GPU threads, every frame.

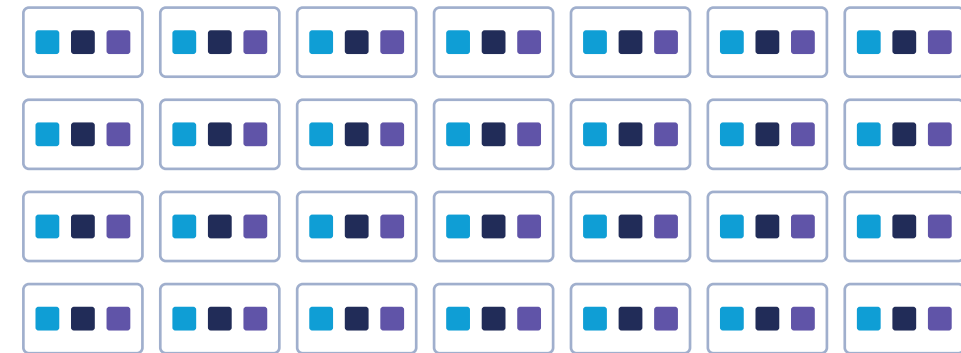
What an ML framework is built for

In a typical ML problem, the model is the main program.



VS

neural shading — the last slide



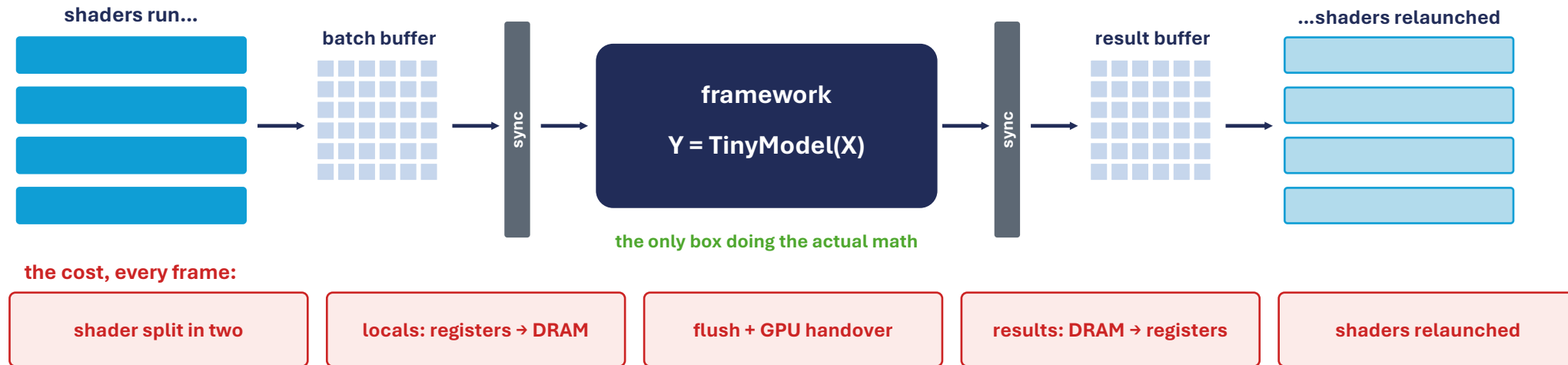
× millions per frame

each tiny card = one invocation:
locals → $Wx+b$ chain → result

a huge kernel call, driven by the framework — millions of parallel tiny calls, inside the renderer

Forcing it into the framework's shape

To use a framework mid-frame, the workload must be reshaped — twice.

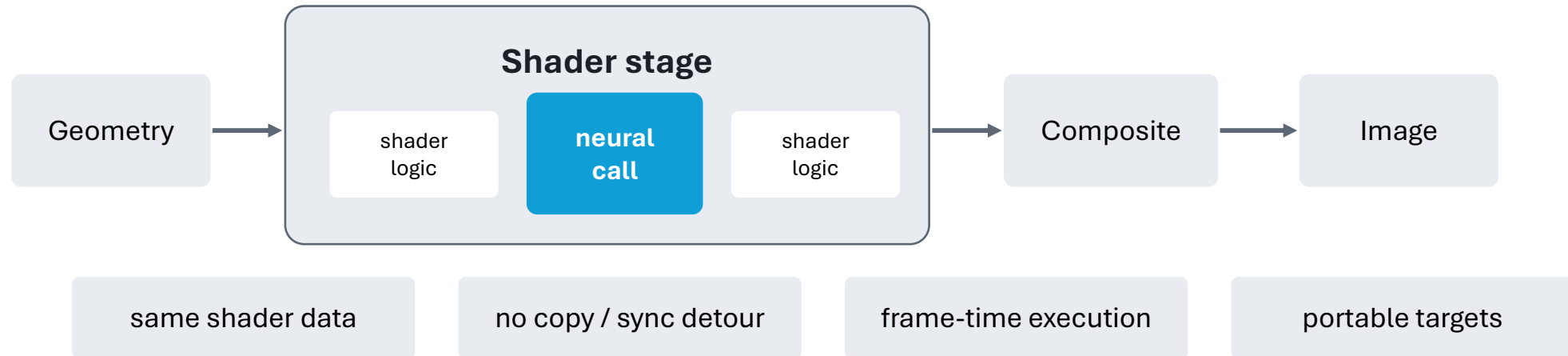


...repeated every frame, for every neural call site in the renderer

All of this plumbing — to deliver a few microseconds of $Wx + b$.

The right model: inline neural shaders

Neural code should be callable as a normal shader operation.



Goal: use neural networks without leaving the shader world.



What does a neural shader need?

If neural code is a normal shader operation, it still needs a few reusable primitives.

Feature encoding

map raw shader inputs
into a neural feature space

Differentiability

define once; support training
or fine tuning when needed

Neural evaluation

compute learned functions
inline with shader logic

Efficient execution

meet real-time budgets
across renderer targets

These are the building blocks we package as shader-language abstractions.



If we write the neural shader inline by hand

All four needs become code the shader author has to maintain.

forward	backward
<pre>Color evalTexture(Pos uv) { Feat f = Feat(0.0); for (int l = 0; l < 8; ++l) { float2 p = uv * scale(l) + offset(l); int2 base = int2(floor(p)); float2 t = frac(p); // ... bounds, masks, address math ... for (int cy = 0; cy < 2; ++cy) for (int cx = 0; cx < 2; ++cx) { uint s = hash(base + int2(cx,cy), l); float w = bilerp(float2(cx,cy), t); f[2*l+0] += w * enc[s+0]; f[2*l+1] += w * enc[s+1]; } } // ... vector packing and storage swizzle ... Hidden h; for (int o = 0; o < 128; ++o) { float a = bias(L0, o); for (int i = 0; i < 16; ++i) a += weight(L0, o, i) * f[i]; h[o] = relu(a); } // ... activation cache and target paths ... Hidden h2; for (int o = 0; o < 128; ++o) { float a = bias(L1, o); for (int i = 0; i < 128; ++i) a += weight(L1, o, i) * h[i]; h2[o] = relu(a); } // ... repeat for variants / layer shapes ... Color y; for (int o = 0; o < 3; ++o) { float a = bias(L2, o); for (int i = 0; i < 128; ++i) a += weight(L2, o, i) * h2[i]; y[o] = sigmoid(a); } // ... output transform and format conversion ... return y; }</pre>	<pre>void evalTexture_bwd(Color dy) { // replay or cache forward intermediates Feat f; Hidden h, h2; Hidden dh = Hidden(0.0); Hidden dh2 = Hidden(0.0); Feat df = Feat(0.0); ... for (int o = 0; o < 3; ++o) { float dz = dy[o] * sigmoidGrad(y[o]); atomicAdd(dBias(L2,o), dz); for (int i = 0; i < 128; ++i) { atomicAdd(dWeight(L2,o,i), dz*h2[i]); dh2[i] += weight(L2,o,i) * dz; } } // ... reductions for dMlp and dh2 ... for (int o = 0; o < 128; ++o) { float dz = dh2[o] * reluGrad(h2[o]); atomicAdd(dBias(L1,o), dz); for (int i = 0; i < 128; ++i) { atomicAdd(dWeight(L1,o,i), dz*h[i]); dh[i] += weight(L1,o,i) * dz; } } for (int o = 0; o < 128; ++o) { float dz = dh[o] * reluGrad(h[o]); atomicAdd(dBias(L0,o), dz); for (int i = 0; i < 16; ++i) { atomicAdd(dWeight(L0,o,i), dz*f[i]); df[i] += weight(L0,o,i) * dz; } } // ... vectorized and tiled variants ... for (int l = 0; l < 8; ++l) for (int cy = 0; cy < 2; ++cy) for (int cx = 0; cx < 2; ++cx) { uint s = replayHash(uv, l, cx, cy); float w = replayWeight(uv, l, cx, cy); atomicAdd(dEnc[s+0], w * df[2*l+0]); atomicAdd(dEnc[s+1], w * df[2*l+1]); } // ... dUv terms and hash collisions ... // ... target-specific atomics/reductions ... }</pre>

excerpts only: the full handwritten path is a few hundred lines of code

The code is the pain

forward	backward
<pre>Color evalTexture(Pos uv) { Feat f = Feat(0.0); for (int l = 0; l < 8; ++l) { float2 p = uv * scale(l) + offset(l); int2 base = int2(floor(p)); float2 f = frac(p); // ... bounds, masks, address math ... for (int cy = 0; cy < 2; ++cy) for (int cx = 0; cx < 2; ++cx) { uint s = h[base.x + cy * 2 + cx]; float w = f[2 * l + 0] * f[2 * l + 1]; f[2 * l + 1] = f[2 * l + 1] * w; } // ... vector ... Hidden h; for (int o = 0; o < 16; ++o) for (int i = 0; i < 16; ++i) h[o] = relu(a + weight(L1, o, i) * h[i]); // ... activation cache and target paths ... Hidden h2; for (int o = 0; o < 128; ++o) { float a = bias(L1, o); for (int i = 0; i < 128; ++i) a += weight(L1, o, i) * h[i]; h2[o] = relu(a); } // ... repeat for variants / layer shapes ... Color y; for (int o = 0; o < 3; ++o) { float a = bias(L2, o); for (int i = 0; i < 128; ++i) a += weight(L2, o, i) * h2[i]; y[o] = sigmoid(a); } // ... output transform and format conversion ... return y; } }</pre>	<pre>void evalTexture_bwd(Color dy) { // replay or cache forward intermediates Feat f; Hidden h; h2; Hidden dh = Hidden(0.0); Hidden dh2 = Hidden(0.0); Feat df = Feat(0.0); ... for (int o = 0; o < 3; ++o) { float dz = dy[o] * sigmoidGrad(y[o]); atomicAdd(dBias(L2, o), dz); ... for (int i = 0; i < 16; ++i) { float dz = dh[o] * reluGrad(h[o]); atomicAdd(dBias(L0, o), dz); for (int l = 0; l < 8; ++l) { float dz = dh[o] * reluGrad(h[o]); atomicAdd(dWeight(L0, o, l), dz * f[l]); df[l] += weight(L0, o, l) * dz; } } // ... vectorized and tiled variants ... for (int l = 0; l < 8; ++l) for (int cy = 0; cy < 2; ++cy) for (int cx = 0; cx < 2; ++cx) { uint s = replayHash(uv, l, cx, cy); float w = replayWeight(uv, l, cx, cy); atomicAdd(dEnc[s+0], w * df[2 * l + 0]); atomicAdd(dEnc[s+1], w * df[2 * l + 1]); // ... dUv terms and hash collisions ... } // ... target-specific atomics/reductions ... } }</pre>



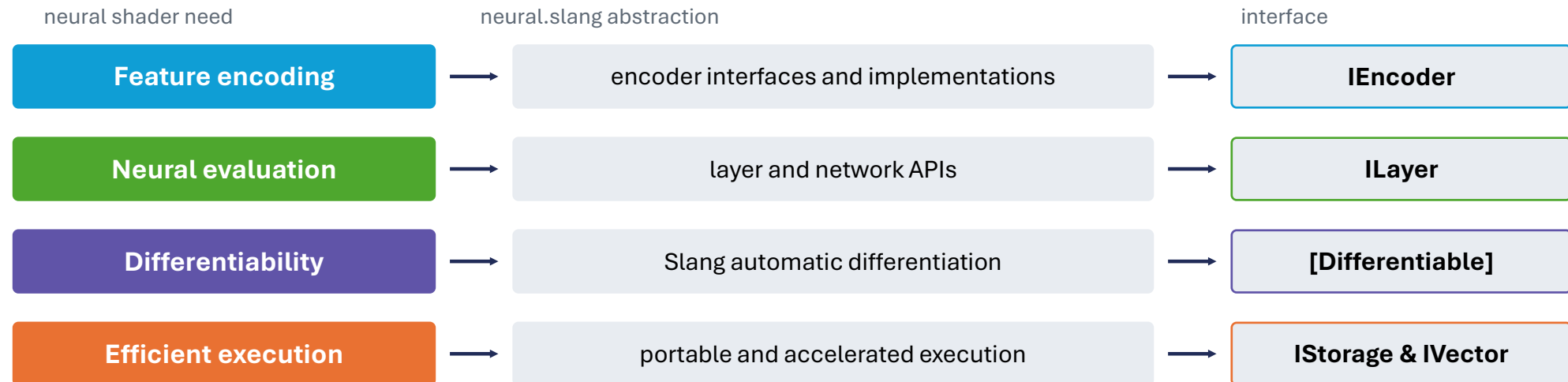
**painful
to maintain**

**wouldn't it be nice
if we could abstract
this away?**

Keep the neural shader inline; move the repetitive machinery behind reusable interfaces.

neural.slang abstracts this into interfaces

Reusable Slang abstractions, with explicit interfaces for each role.

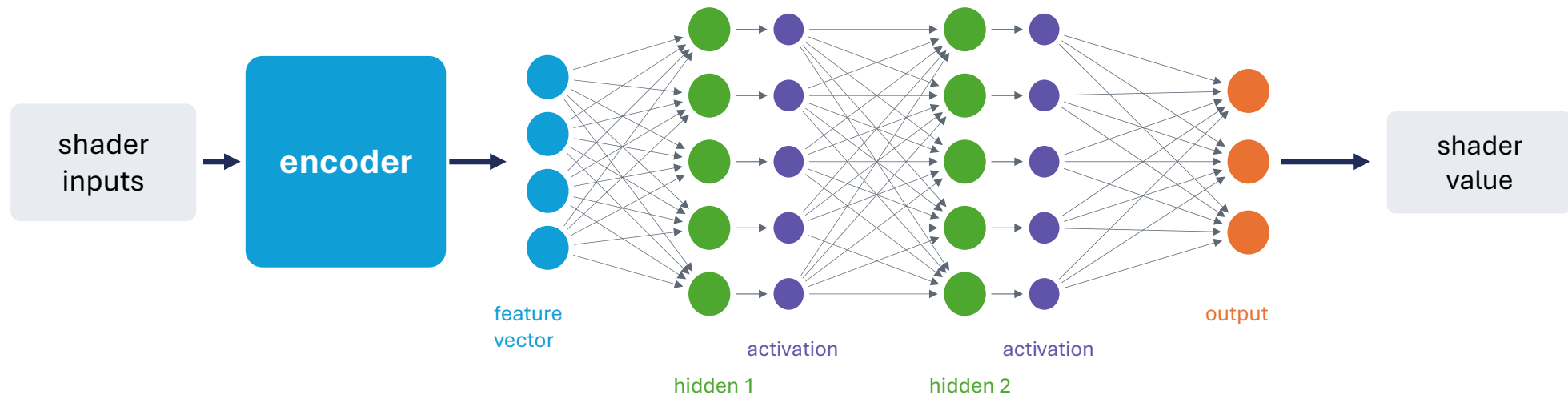


Same shader language; same compilation pipeline; same renderer data flow.

Slang targets CUDA, Vulkan/SPIR-V, Metal, Direct3D, WebGPU, and CPU backends.

How the interfaces work together

One neural shader function, with interfaces attached to the dataflow.

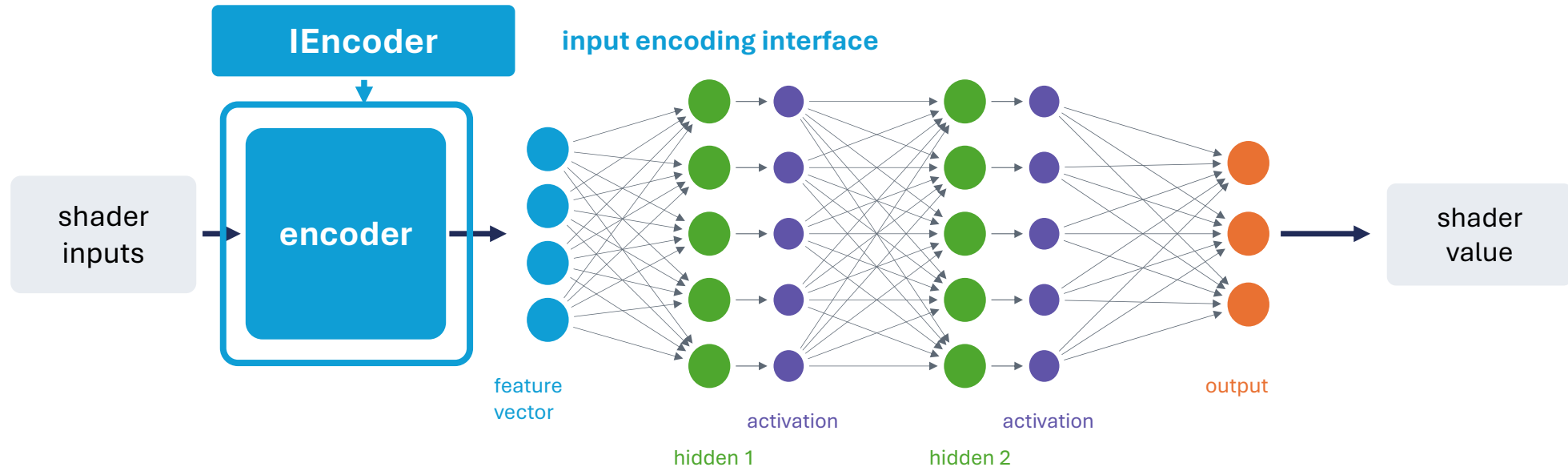


Click through to reveal the interface attached to each part of the network.



How the interfaces work together

One neural shader function, with interfaces attached to the dataflow.

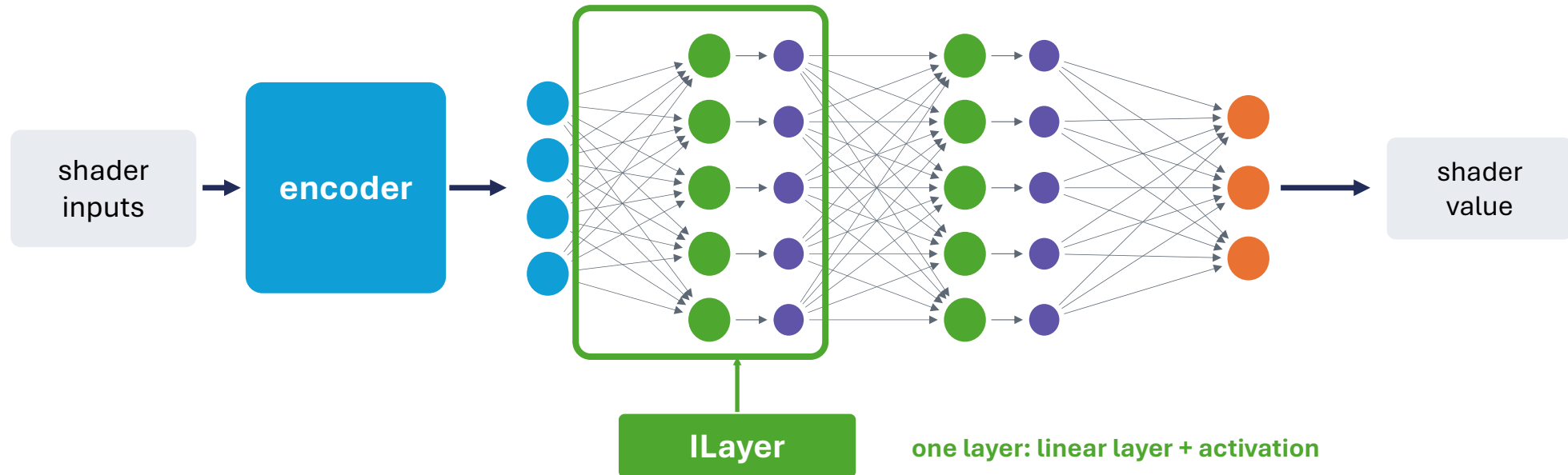


Same network structure; one interface boundary highlighted per build.



How the interfaces work together

One neural shader function, with interfaces attached to the dataflow.

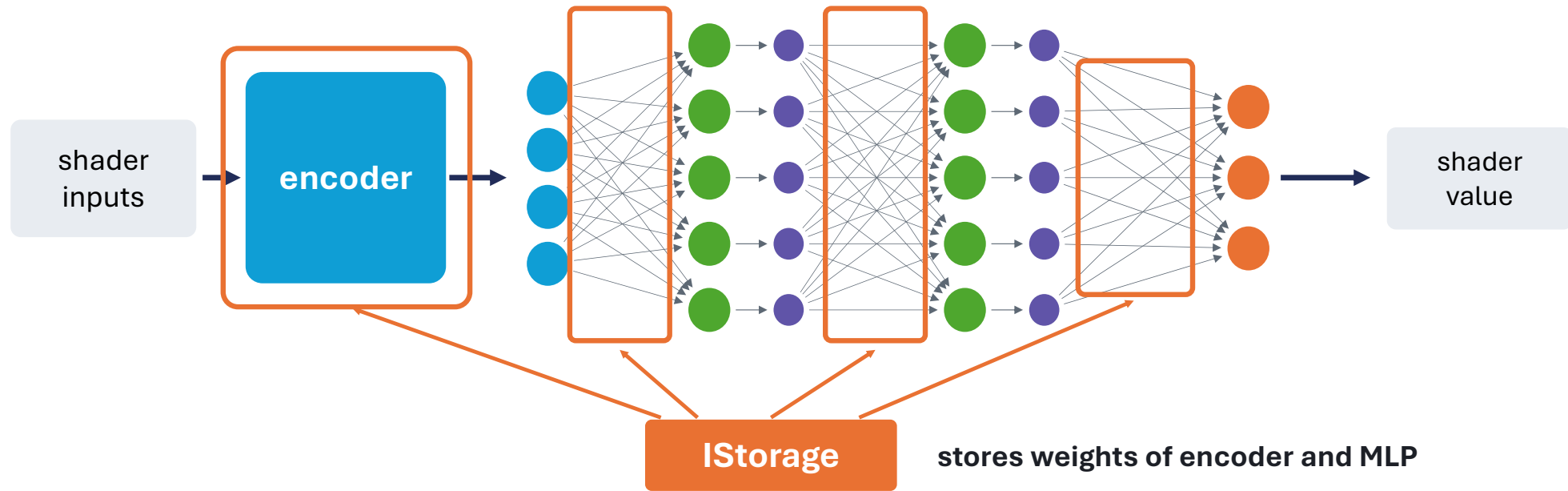


Same network structure; one interface boundary highlighted per build.



How the interfaces work together

One neural shader function, with interfaces attached to the dataflow.

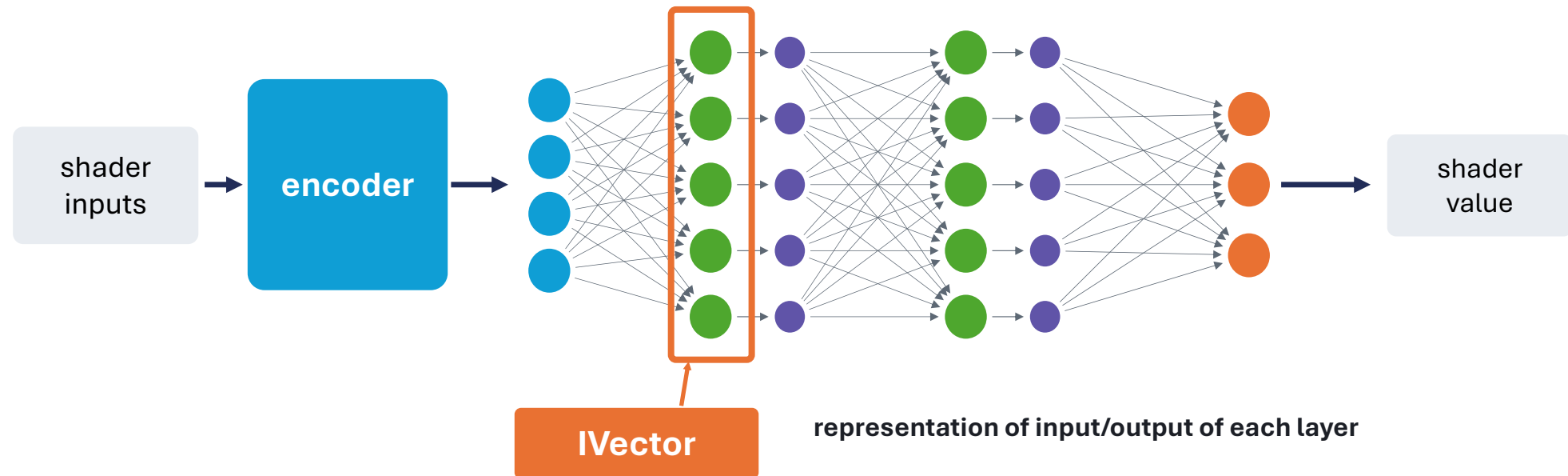


Same network structure; one interface boundary highlighted per build.



How the interfaces work together

One neural shader function, with interfaces attached to the dataflow.

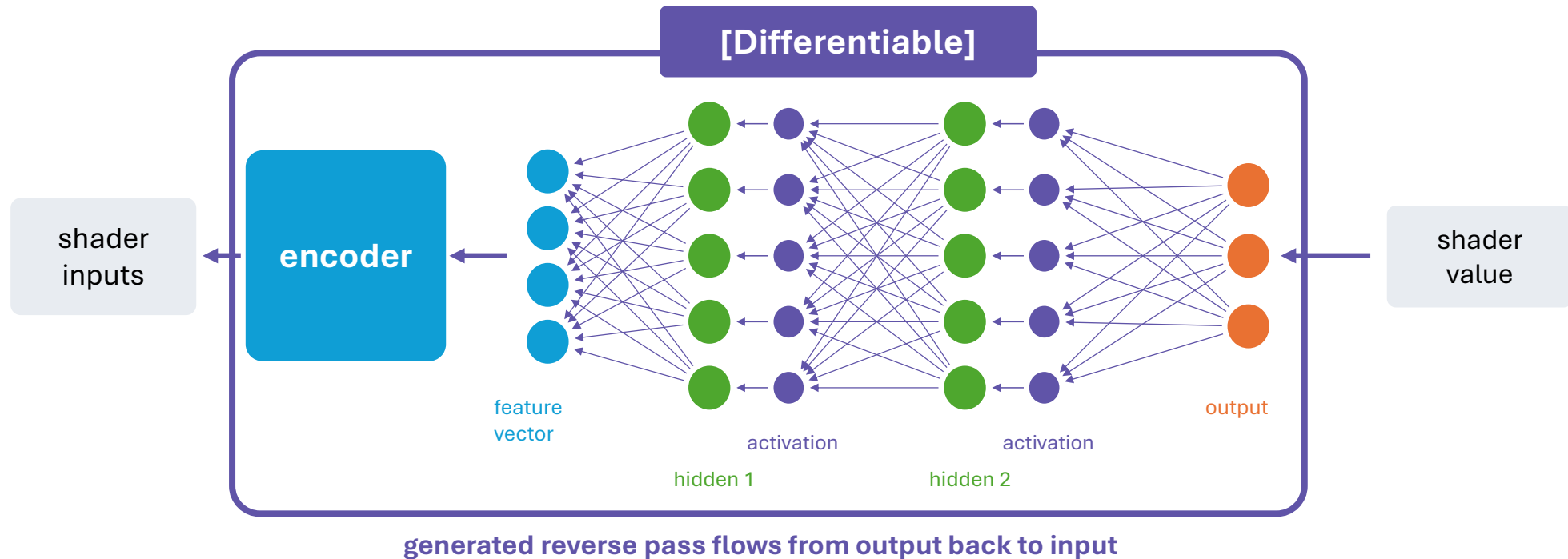


Same network structure; one interface boundary highlighted per build.



How the interfaces work together

One neural shader function, with interfaces attached to the dataflow.

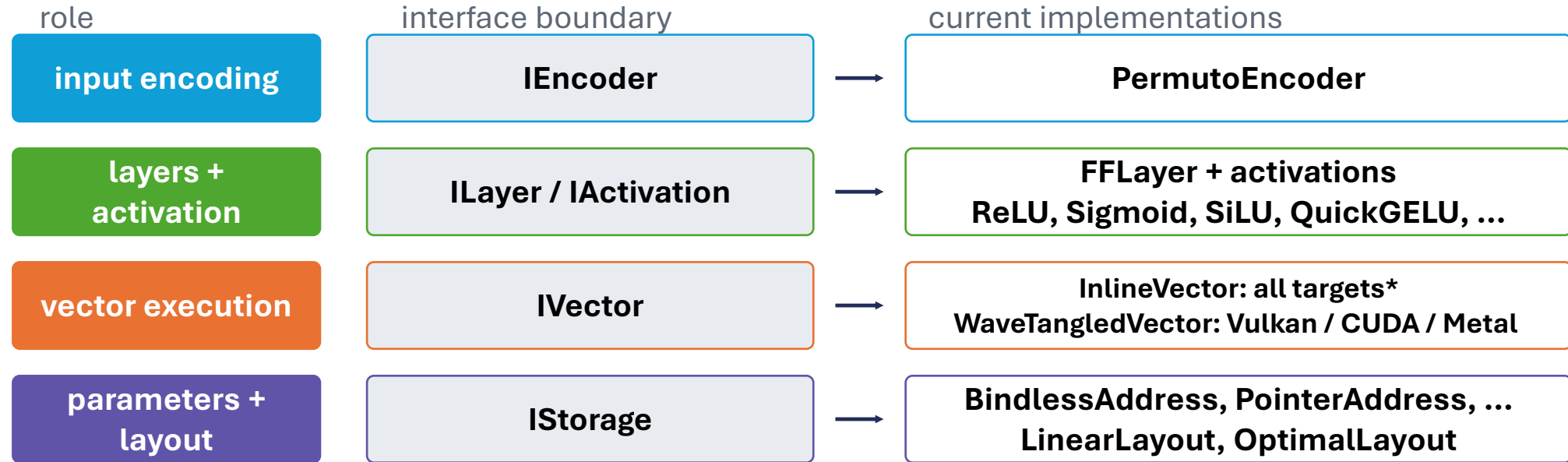


Same network structure; one interface boundary highlighted per build.



API structure

Key concrete types implemented in the neural.slang module.



***except WebGPU — parameter storage is bindless-only**

Users can add encoders, layers, activations, and vector representations; storage is low-level infrastructure only provided by neural.slang.

Code example: neural.slang

17 lines define the neural function directly in Slang.

```
import slang.neural;
typealias Pos = InlineVector<float, 2>;
typealias Feat = InlineVector<float, 16>;
typealias Hidden = InlineVector<float, 128>;
typealias Color = InlineVector<float, 3>;
typealias Encoder = PermutoEncoder<float, 2, 2, 16, 8, Pos, Feat>;
typealias InLayer = FFLayer<float, Feat, Hidden, OptimalLayout, ReLU<float>, true>;
typealias HiddenLayer = FFLayer<float, Hidden, Hidden, OptimalLayout, ReLU<float>, true>;
typealias OutLayer = FFLayer<float, Hidden, Color, OptimalLayout, Sigmoid<float>, true>;

[Differentiable]
Color evalTexture(Pos uv, BindlessAddress<float> enc, BindlessAddress<float> mlp)
{
    var f = Encoder(buildHyperParameters()).encode<BindlessAddress<float>>(uv, enc);
    var h = InLayer().eval(f, mlp);
    h = HiddenLayer().eval(h, mlp.getOffset(InLayer.StorageSize));
    return OutLayer().eval(h, mlp.getOffset(InLayer.StorageSize + HiddenLayer.StorageSize));
}
```

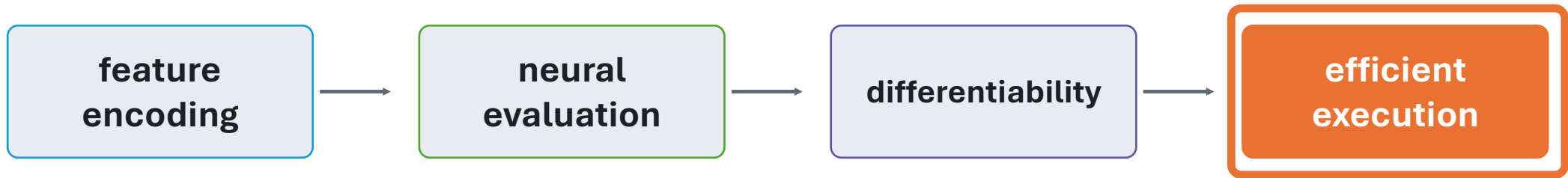
17 nonblank lines

2D coordinate -> encoder -> 16 -> 128 -> 128 -> 3



The missing piece: efficient execution

The API reduces authoring cost. The next question is whether the generated path can still run at real-time speed.

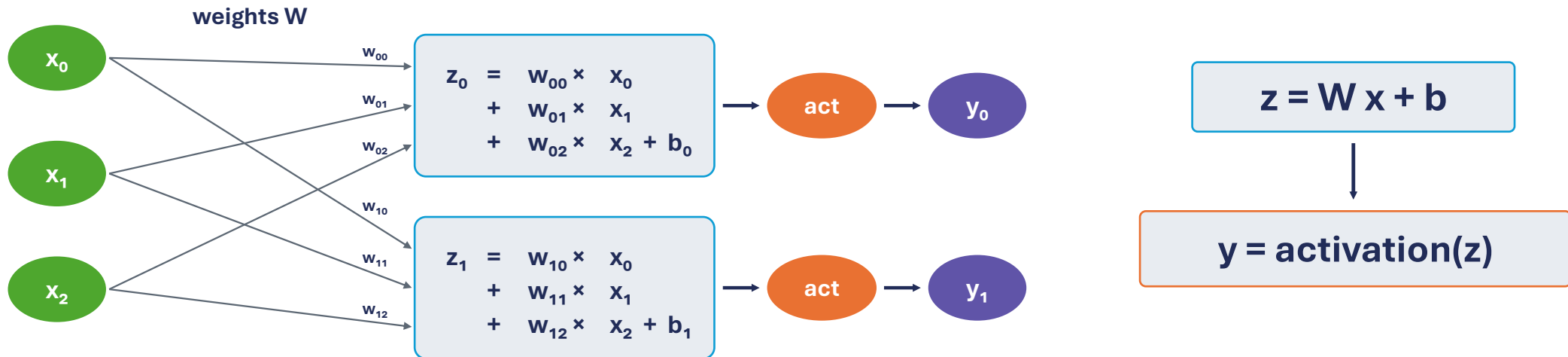


The goal is convenient code that still fits the frame budget.



Zoom in: one layer

A dense layer forms weighted sums, adds bias, then applies an activation.



Therefore: $y = \text{activation}(Wx + b)$



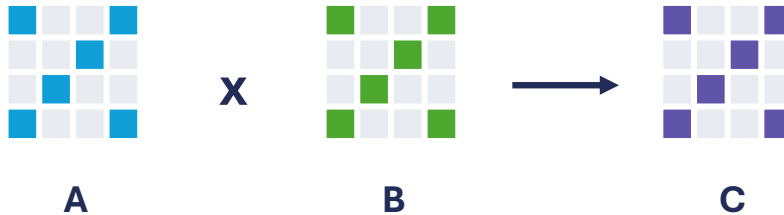
Modern GPUs expose matrix-matrix MMA

NVIDIA
Tensor Cores

AMD
Matrix Cores

Metal
SIMD Matrix

Other matrix
accelerators



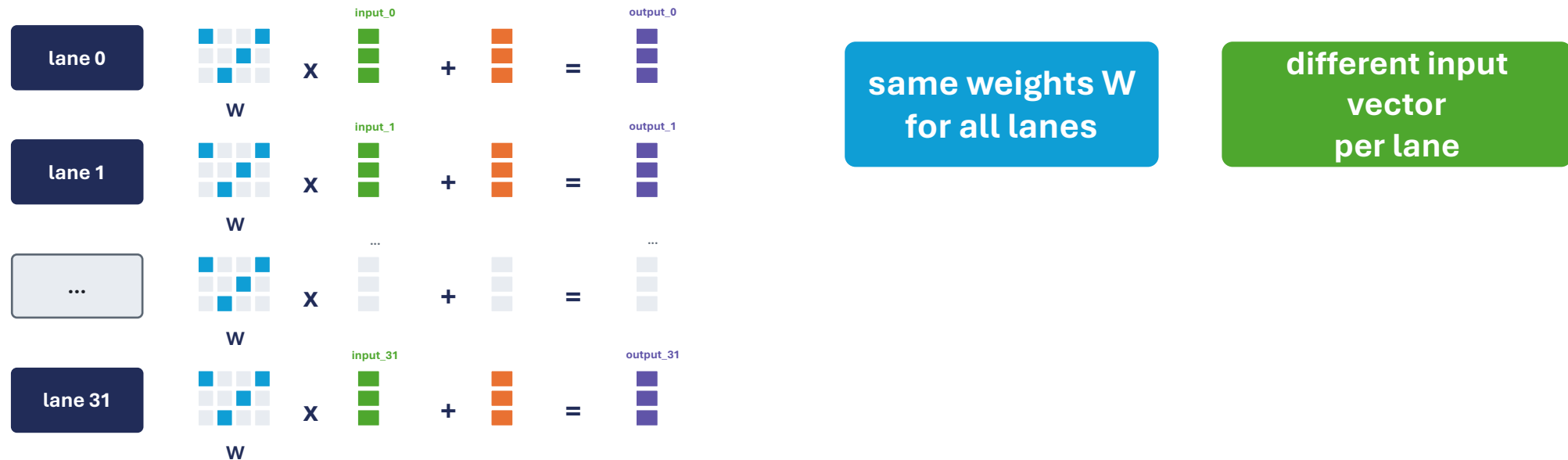
Hardware wants
tiled matrix x matrix

But each shader lane naturally has one vector for one layer evaluation.



Matvec to MMA: per-lane view

SIMT view: each lane evaluates the same layer with its own input vector.

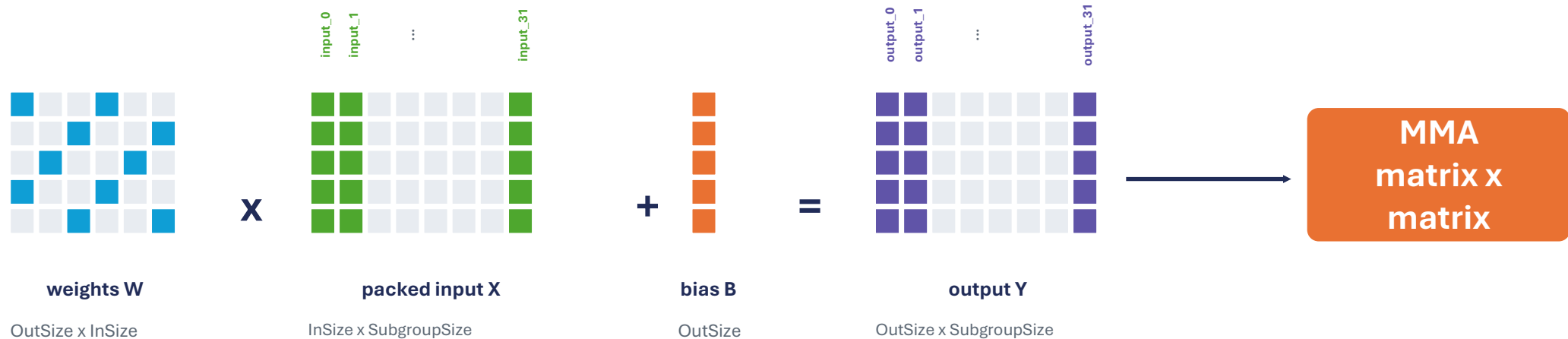


This is a batch of independent matrix-vector multiplies from the shader author's point of view.



Matvec to MMA: packed subgroup view

Cooperative execution packs lane-local vectors into matrix columns.

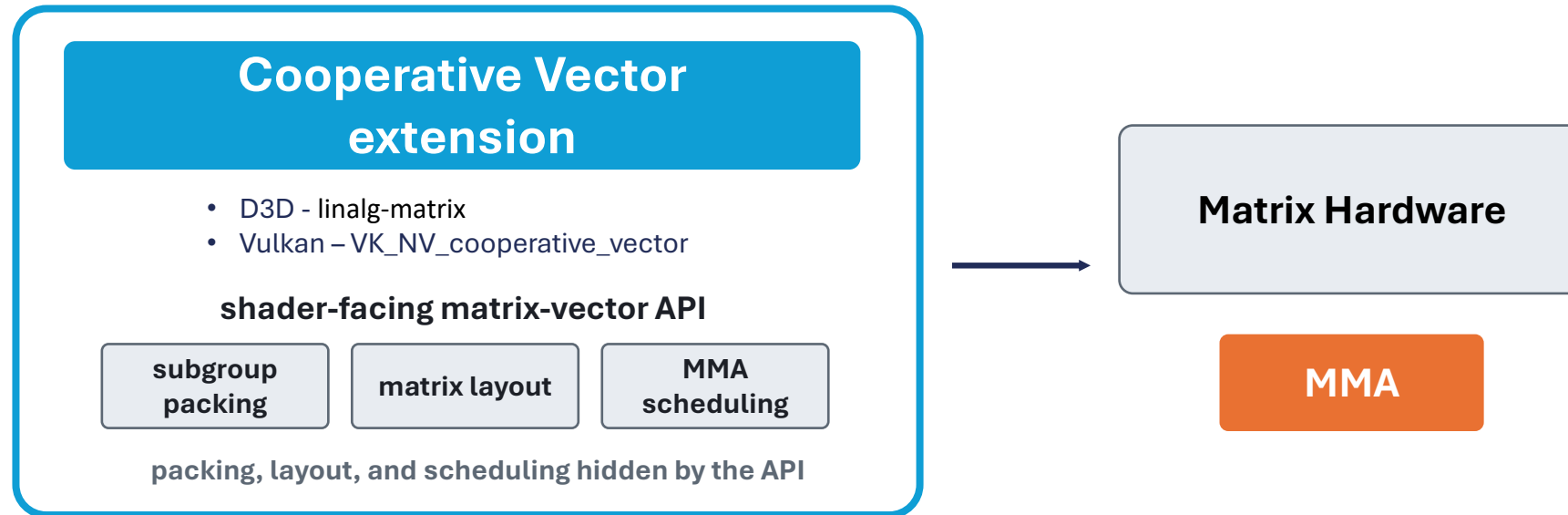


The hard part is constructing this packed representation without exposing it to the shader author.



Cooperative Vector hides the packing details

Vulkan/D3D path hides subgroup packing, layout, and MMA scheduling.

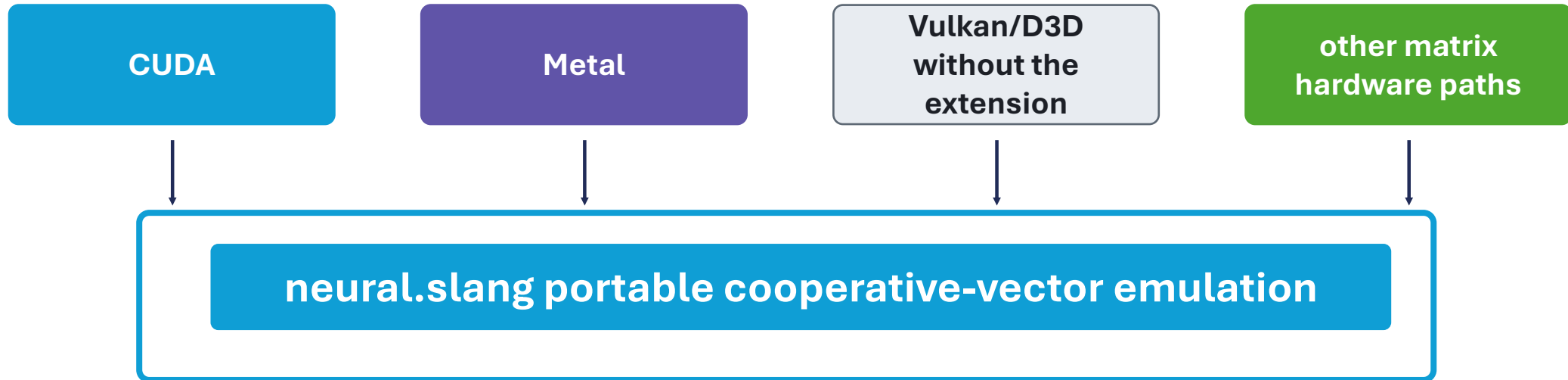


Great when this extension path exists.



What if Cooperative Vector is unavailable?

The same shader layer still needs packed-subgroup matvec.

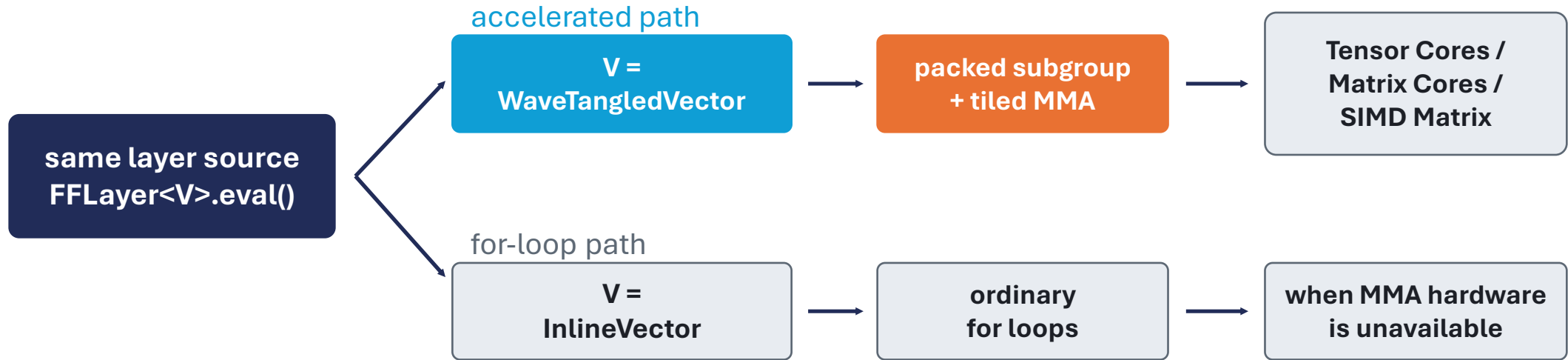


Same layer source; backend-specific code handles packing, tiling, and matrix instructions.



Vector type selects the execution path

In `FFLayer<V>`, `V` is the vector type that supplies matvec.



The API stays fixed; the chosen vector type controls how the layer matvec is evaluated.



WaveTangledVector key point: matrix tiling

Packed subgroup matvec becomes tiled matrix multiply.



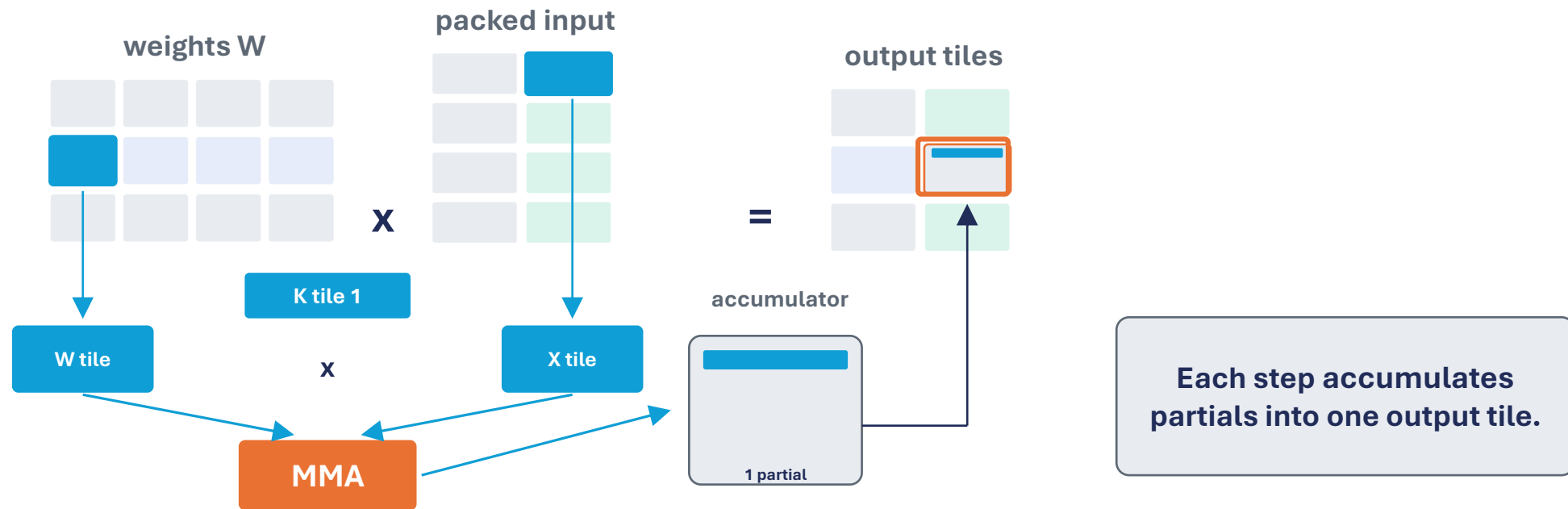
One output tile is built by walking across matching input tiles.

Each click advances one K tile.



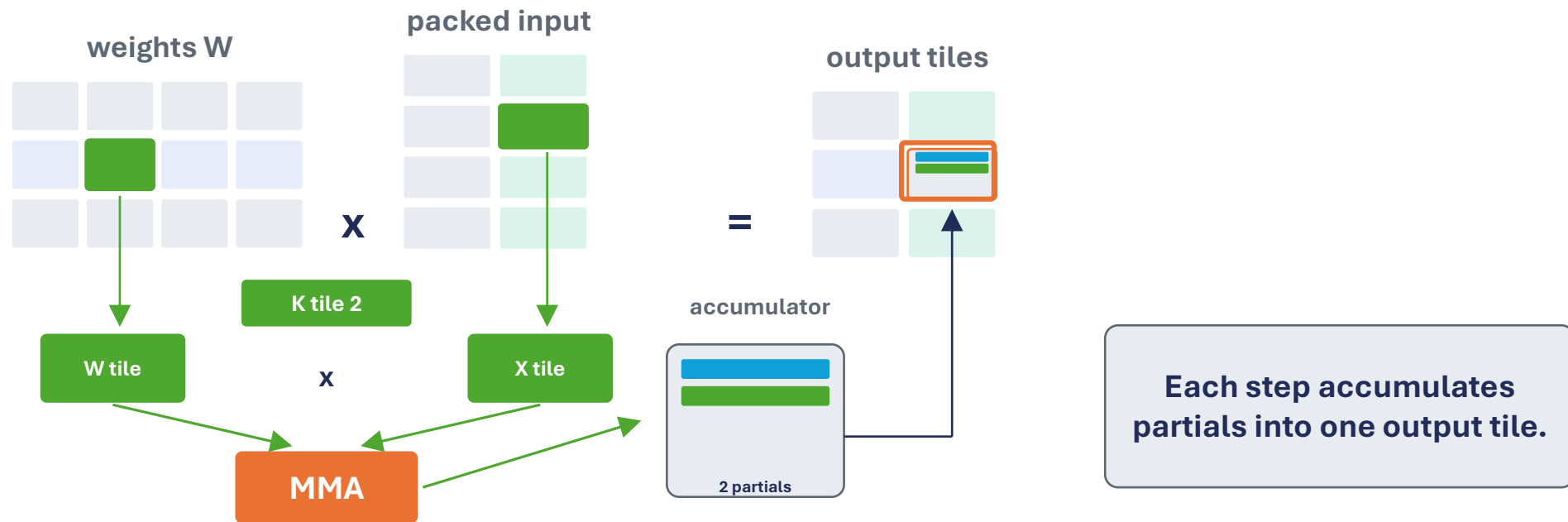
Tiled MMA flow: first tile pair

Load one W tile and one packed-input tile.



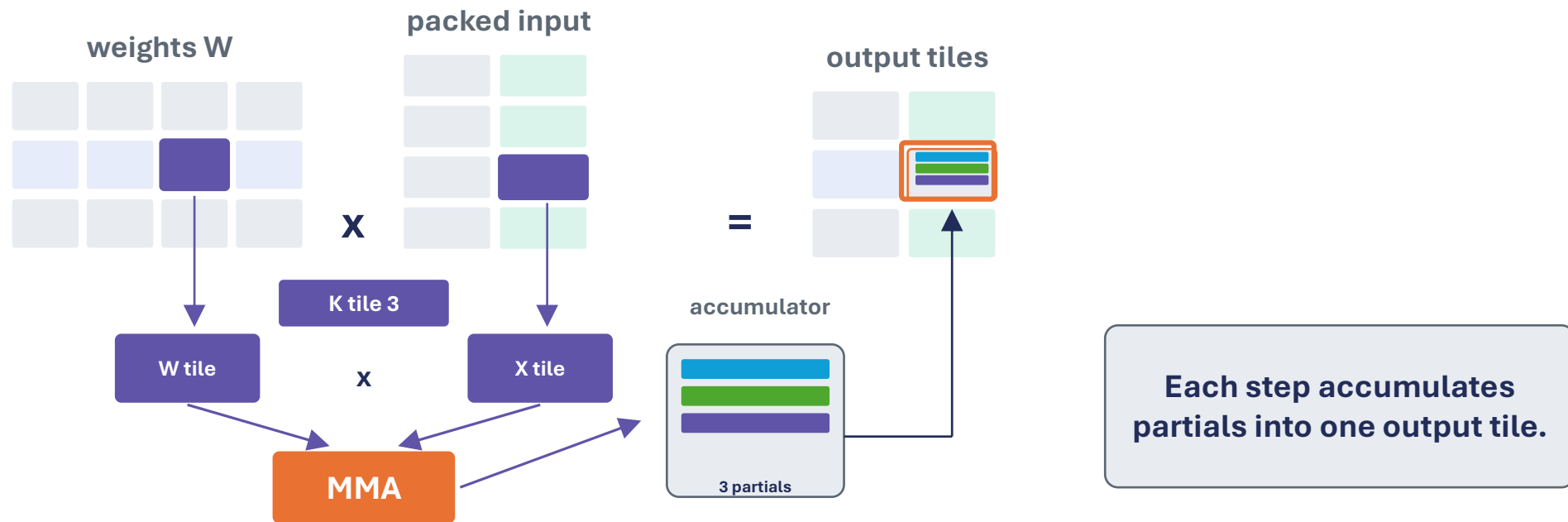
Tiled MMA flow: next tile pair

Accumulate another partial product into the same output tile.



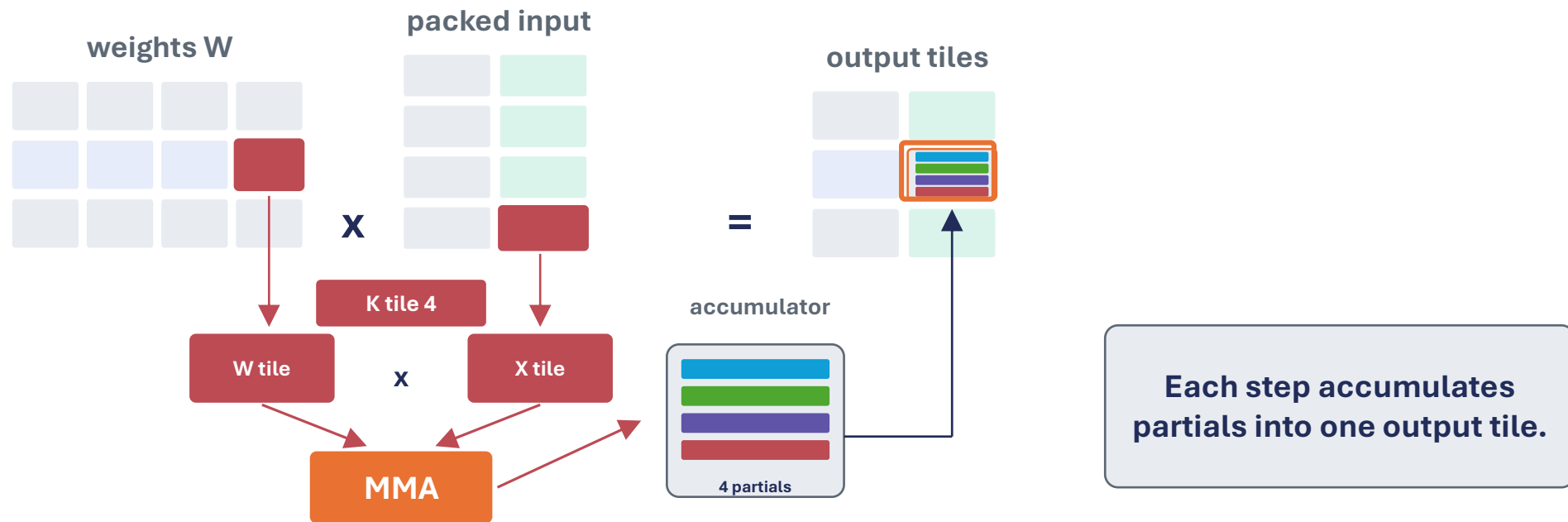
Tiled MMA flow: third tile pair

Continue accumulating partial products.



Tiled MMA flow: accumulated output tile

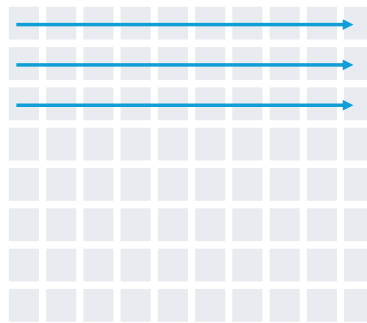
After the K loop, the accumulator becomes the output tile.



WaveTangledVector key point: target-specific layout

Example: row-major -> tile-contiguous -> CUDA thread/register layout

1 row-major matrix



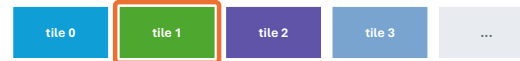
memory



2 tiled MMA layout



memory



3 one CUDA MMA tile

zoom: values are assigned to CUDA threads/registers

Row\Col	0-1	2-3	4-5	6-7
0	thread 0 reg a0,a1	thread 1 reg a0,a1	thread 2 reg a0,a1	thread 3 reg a0,a1
1	thread 4 reg a0,a1	thread 5 reg a0,a1	thread 6 reg a0,a1	thread 7 reg a0,a1
...				
7	thread 28 reg a0,a1	thread 29 reg a0,a1	thread 30 reg a0,a1	thread 31 reg a0,a1
8	thread 0 reg a2,a3	thread 1 reg a2,a3	thread 2 reg a2,a3	thread 3 reg a2,a3
9	thread 4 reg a2,a3	thread 5 reg a2,a3	thread 6 reg a2,a3	thread 7 reg a2,a3
...				
15	thread 28 reg a2,a3	thread 29 reg a2,a3	thread 30 reg a2,a3	thread 31 reg a2,a3

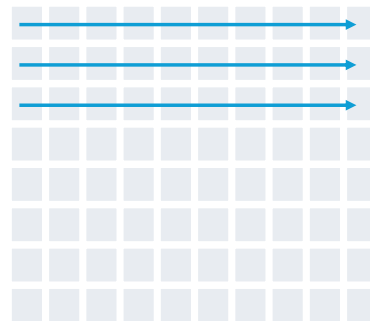
CUDA reference: PTX mma.m16n8k8 per-thread register layout

Different targets expect different layouts for their best performance.

WaveTangledVector key point: target-specific layout

Example: row-major -> tile-contiguous -> CUDA thread/register layout

1 row-major matrix



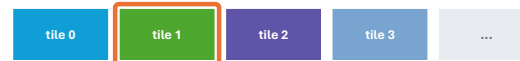
memory order



2 tiled MMA layout



memory order



3 one CUDA MMA tile

zoom: values are assigned to CUDA threads/registers

Row\Col	0-1	2-3	4-5	6-7
0	thread 0 reg a0,a1	thread 1 reg a0,a1	thread 2 reg a0,a1	thread 3 reg a0,a1
1	thread 4 reg a0,a1	thread 5 reg a0,a1	thread 6 reg a0,a1	thread 7 reg a0,a1
...				
7	thread 28 reg a0,a1	thread 29 reg a0,a1	thread 30 reg a0,a1	thread 31 reg a0,a1
8	thread 0 reg a2,a3	thread 1 reg a2,a3	thread 2 reg a2,a3	thread 3 reg a2,a3
9	thread 4 reg a2,a3	thread 5 reg a2,a3	thread 6 reg a2,a3	thread 7 reg a2,a3
...				
15	thread 28 reg a2,a3	thread 29 reg a2,a3	thread 30 reg a2,a3	thread 31 reg a2,a3

CUDA reference: PTX mma.m16n8k8 per-thread register layout

```
NetworkParameterLayoutConverter<half, 0b111u11, 16, 128, 128, 3>.toOptimalLayout(portable, optimal, tid, threads);
```

WaveTangledVector key constraint: divergence

WaveTangledVector requires uniform control flow.

- Divergent control flow leaves some lanes inactive — and restoring them is driver-level control that Slang does not have.
- For divergent code, use **Cooperative Vector** where the platform supports it: it handles divergence below the API.
- **Cooperative Vector** backend support in **neural.slang** is on the roadmap — coming soon.



Training adds more matrix work

For $y = W x + B$, the chain rule gives a transposed matvec, an outer product, and a reduction.

$$dx = W^T \times dy$$

$$dW = dy \times x^T$$

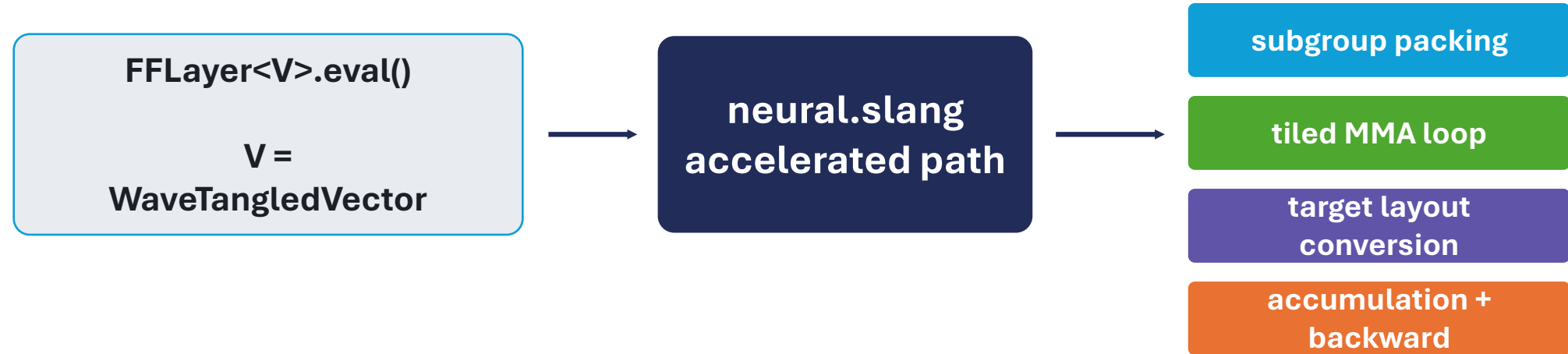
$$dB = \text{reduce}(dy)$$

Full backward is also an MMA problem.



The API chooses the performance path

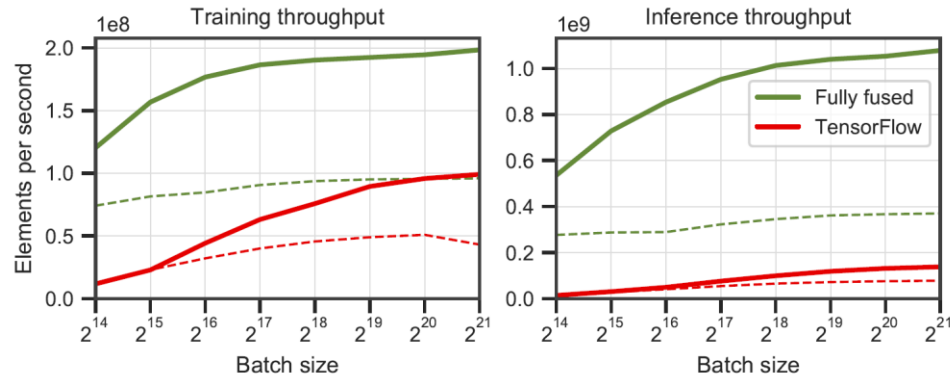
Developers still write one layer call.



The vector type selects the execution strategy; the layer API stays the same.



Performance baseline

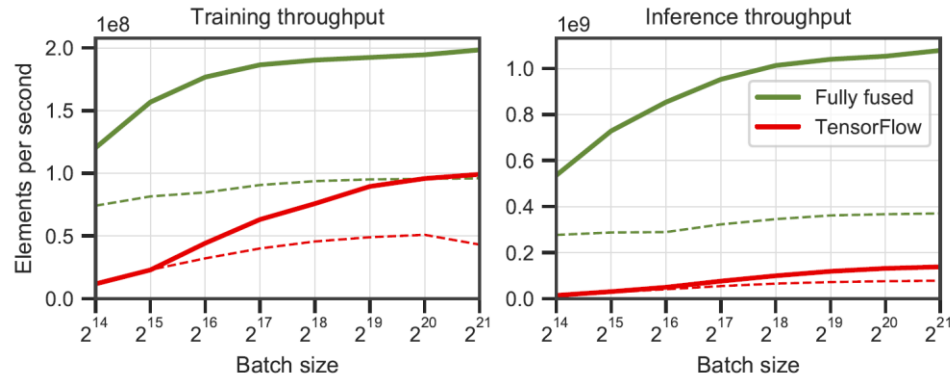


Source: Müller et al., Real-time Neural Radiance Caching for Path Tracing, SIGGRAPH 2021.

Fully fused MLPs are already much faster than general ML frameworks.

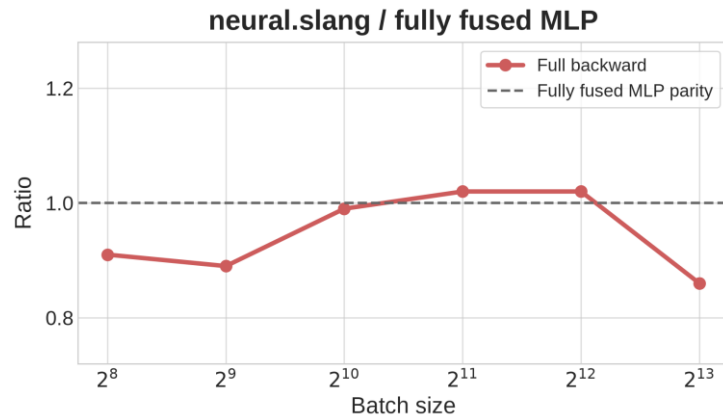


Performance baseline



Source: Müller et al., Real-time Neural Radiance Caching for Path Tracing, SIGGRAPH 2021.

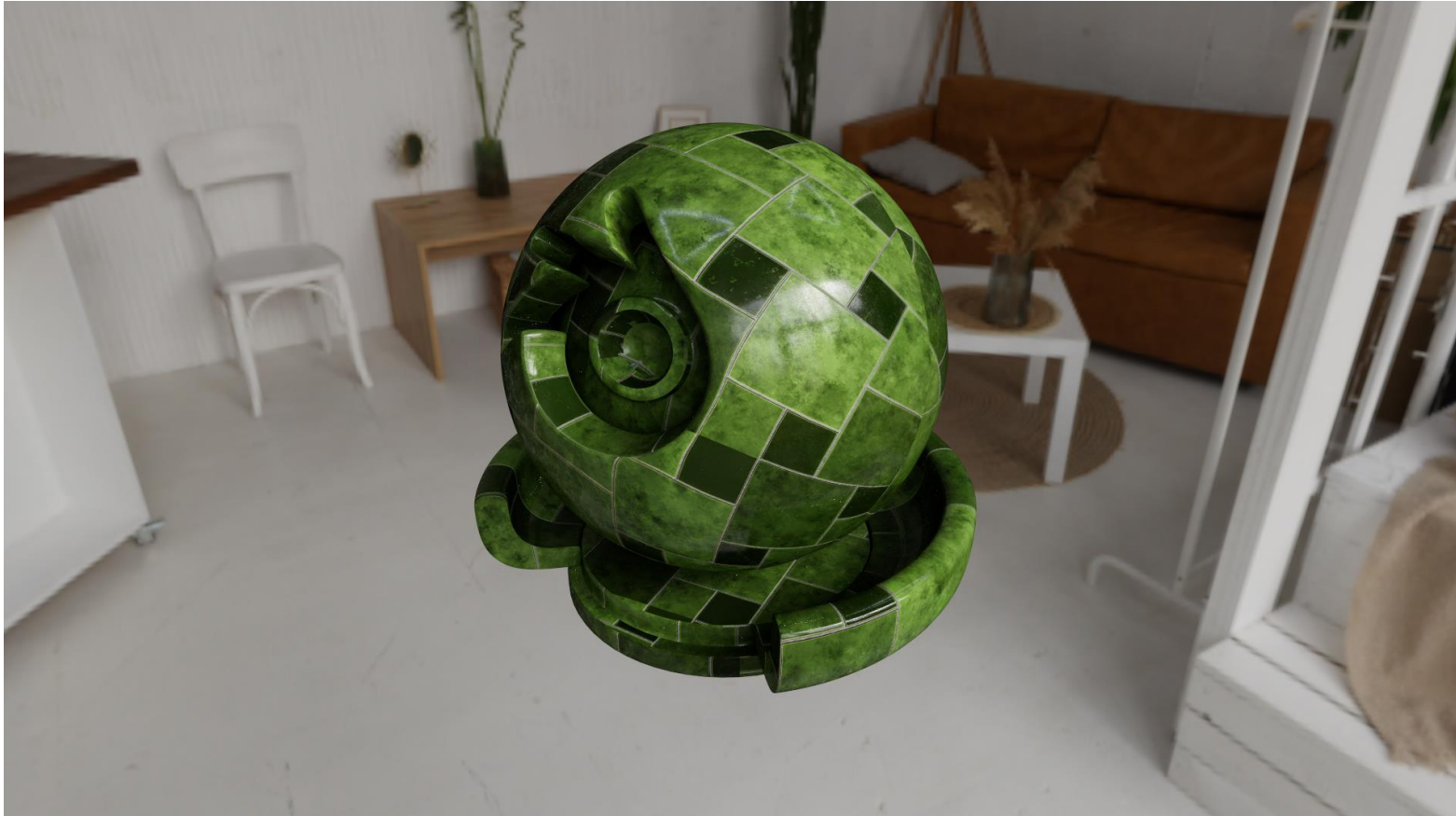
Fully fused MLPs are already much faster than general ML frameworks.



Ratio near 1.0 means fully fused MLP parity.

**Shader-native API
same performance class**

Example: neural material in a path tracer



A path tracer with a neural material — and the result is beautiful.



What developers get

shader-native
MLP authoring

17 lines
vs 100+

portable
across targets

for-loop +
accelerated
paths

**neural.slang makes small neural networks
practical as shader building blocks.**



HPG
2026

More resources

- Siggraph 2026 course:
An Introduction to Neural Shading

- 9:00 am – 12:15 am, July 19
- Room 515A



- Sample code for neural.slang:

[shader-slang/slangpy-samples/experiments/neural_slang](https://shader-slang.com/slangpy-samples/experiments/neural_slang)

