

The Tech Art Journey of creating and optimizing an asset pipeline

#GCAP24

OI,
tudo bem?

#GCAP24

Who am I?

Tech Art Lead @ SMG Studio

Previously:
Lead Gameplay Programmer
Data Scientist
Naval Engineer

But also:
Maths Geek
Coffee + Cake Lover
SoulsBorneKiroRing fan



#GCAP24

What will happen here?

This will be fast



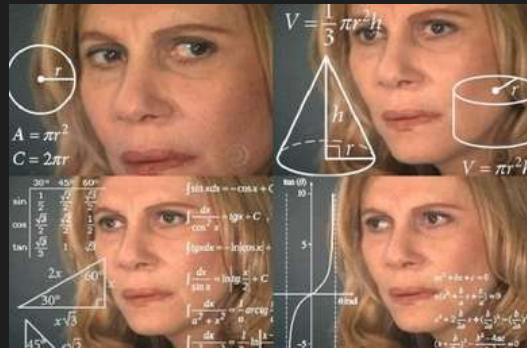
#GCAP24

This talk will be fast as there is a lot of content, but not a lot of time

What will happen here?

This will be fast

This will get technical



#GCAP24

This talk will have technical elements, so be prepared for talking about precision and bytes

What will happen here?

This will be fast

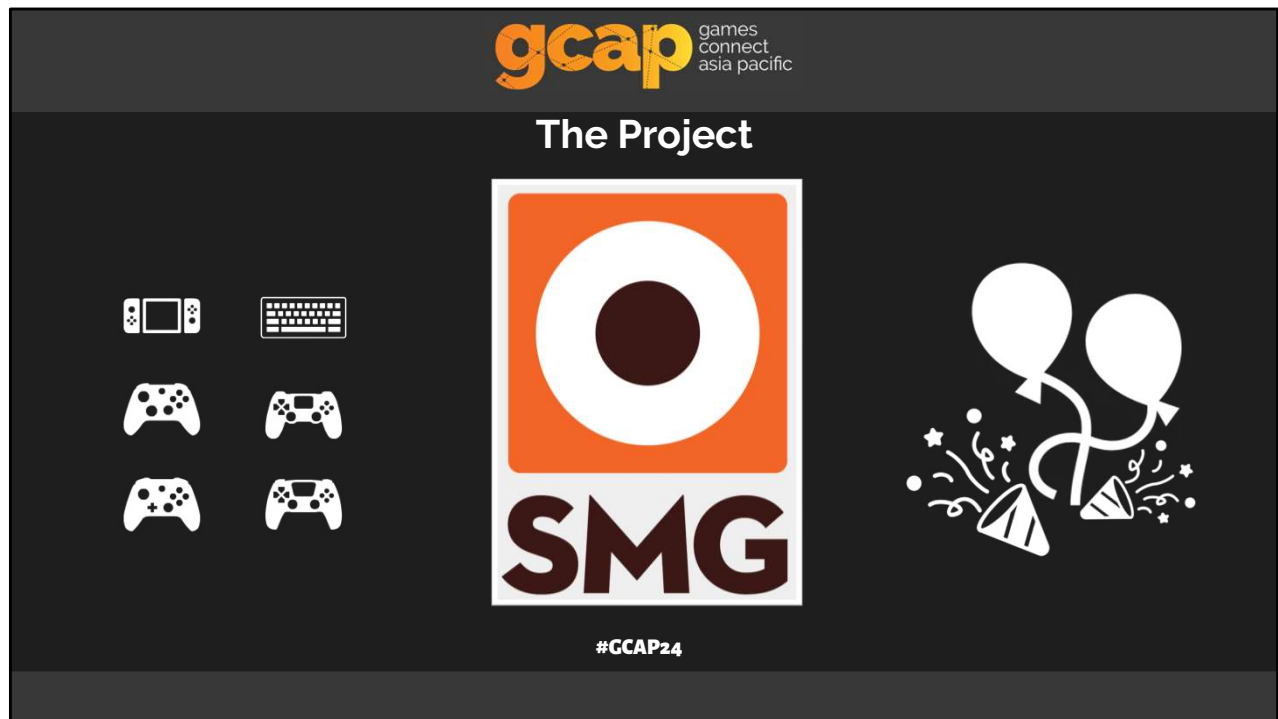
This will get technical

This deck will be available online



#GCAP24

But fret not because this deck will be online and you can get a link at the end!
With all that said, get yourself comfortable, keep your hands inside the vehicle at all times and off we go!



This talk will cover the technology implemented on SMG's largest game, a multiplatform party game...

The Project

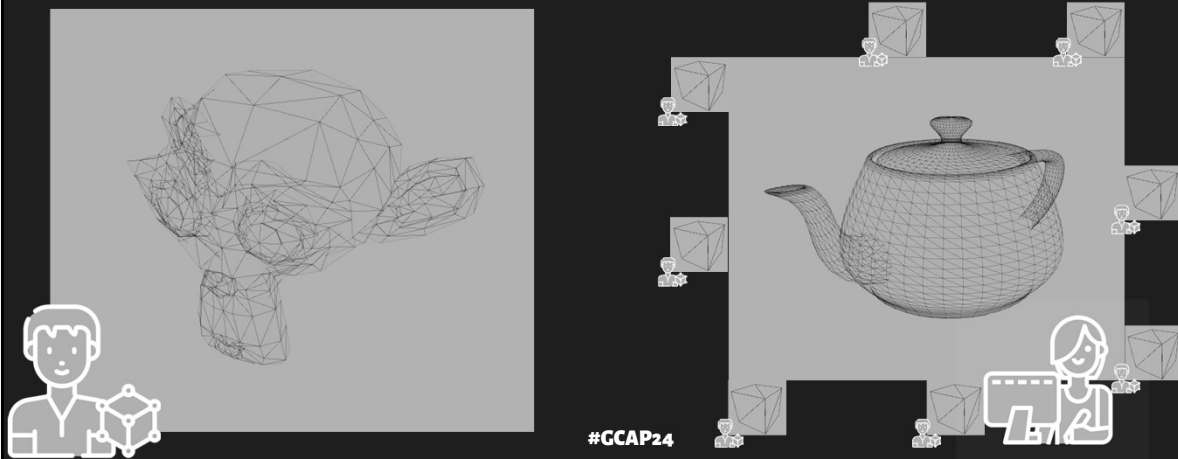


#GCAP24

That is still a secret.

The project was intended to be both announced and released by GCAP 2024, but due to release window complications we've delayed until 2025

The Project's TECH



So let's talk about the tech used in the project instead!

While most projects have a team of 3D artists that models entire assets

On our project we took a different route:

We started by building a library of “Atomic mesh assets” these can be anything, from a small wall piece, to a door, or anything!

Then we have a design team that create build instructions for each of our scenes and in some cases even characters!

The Tech art team ingest these instructions into our custom pipeline setting up culling and features such as compression or lightmap unwrapping

Finally the pipeline outputs meshes in-engine that can be used normally by anyone in the team.

The Project's TECH

Atomic Mesh Assets

Scriptable Object
3D Models

Feature Setup

Unity Scripted Importer
.meta file data

Build Instructions

Human readable file
Custom file extension

In-Engine Mesh Asset

Regular Model
Prefab

#GCAP24

Some details on the different parts:

- Atomic Mesh Assets: Scriptable Objects that have 3D models assigned to them
- Build Instructions: human readable files saved with a custom file extension
- Feature Setup: Unity Scripted Importer which saves data to the .meta file (<https://docs.unity3d.com/Manual/ScriptedImporters.html>)
- In-Engine Mesh Asset: Behaves like a regular model, can be dragged around as a prefab

The Project's TECH

Atomic Mesh
Assets

Build
Instructions

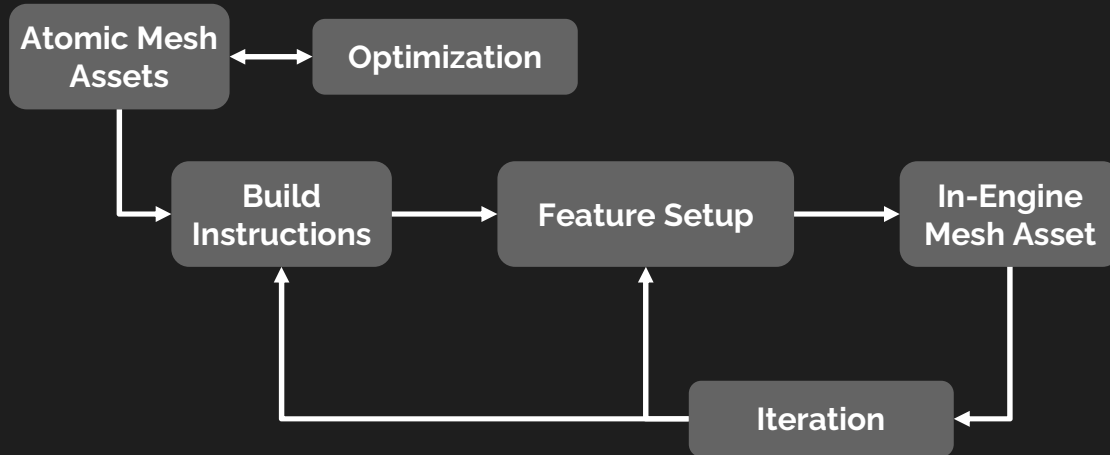
Feature Setup

In-Engine
Mesh Asset

#GCAP24

Putting all these steps together we get the main part of our pipeline.

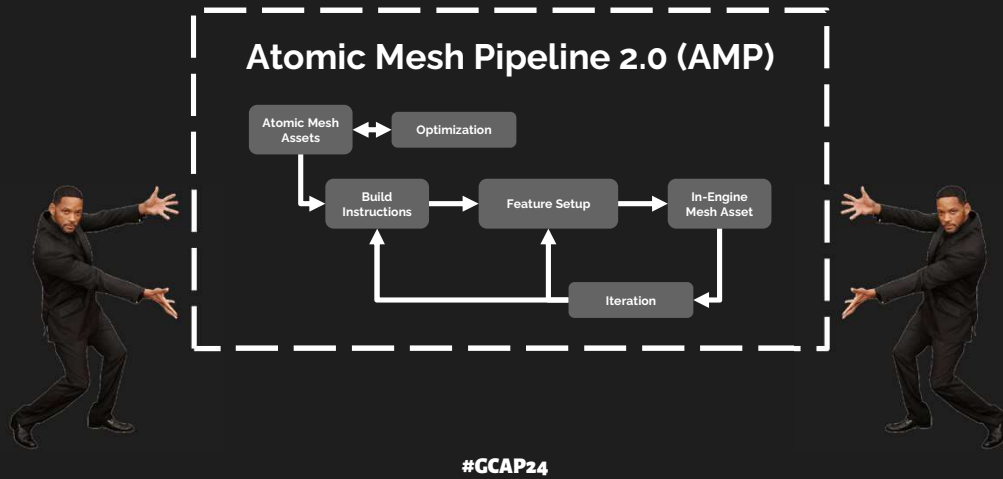
The Project's TECH



#GCAP24

But the pipeline also needs to be simple and flexible to allow for optimization of the atomic meshes and/or iteration of the final assets due to design changes

The Project's TECH



This pipeline was put in place before my time and I inherited it.

It has heaps of awesome tech and does the job.

It's based on the Unity's AssetImporter API handling the import of a custom file type that in turn generates an asset in the Unity library similar to a 3D model, which has a Transform hierarchy and associated meshes.

Procedural Meshes

AMP is Great!

Jamming many meshes together looks great!

But quickly become **expensive due to large number of verts** on the inside and MANY repeated parts.

How can we improve it?



#GCAP24

This idea of procedural meshes is great, flexible and yield quite intricate detail. But the output meshes are insanely high poly, mostly due to internal geometry + geo that will never be seen from our game's mostly static camera setup

Procedural Meshes

Visibility Culling

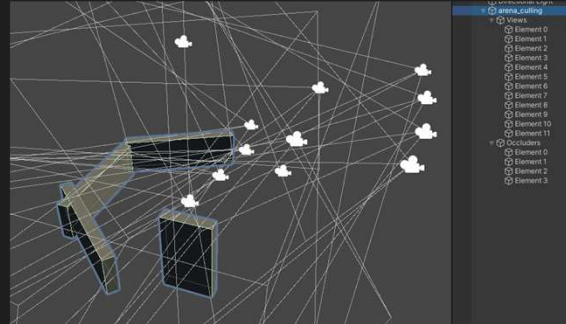
#GCAP24

Our solution was an embedded visibility culling solution

Procedural Meshes

Visibility Culling

1) Setup visibility Cameras + Occluders



#GCAP24

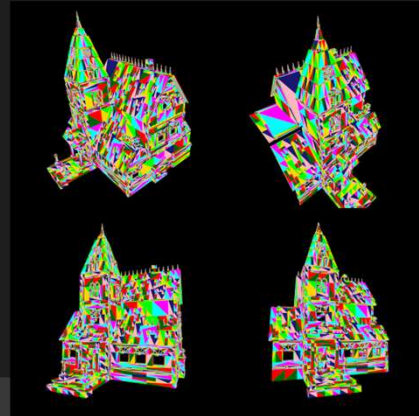
It starts by setting up a rig with cameras that mark all the angle the model will be seen in-game.

We can also refine it by adding blocks as visibility occluders if we know that a specific part will be consistently covered by another element, by VFX for example.

Procedural Meshes

Visibility Culling

- 1) Setup visibility Cameras + Occluders
- 2) For each camera
 - a) Render triangles indices
 - b) Render Occluders on top



#GCAP24

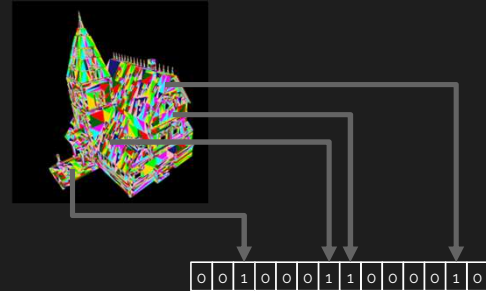
From there we run a process where we render the model from each camera angle, but instead of it's colors we output the triangle index to a uint32 render texture with depth buffer.

We then render the occluders with a null triangle index as to, well, occlude triangles.

Procedural Meshes

Visibility Culling

- 1) Setup visibility Cameras + Occluders
- 2) For each camera
 - a) Render triangles indices
 - b) Render Occluders on top
 - c) Collect Data



#GCAP24

Final step is to run a compute shader and aggregate the data on which triangles are visible to a ComputeBuffer.
Rinse and repeat for all cameras.

Procedural Meshes

Visibility Culling

- 1) Setup visibility Cameras + Occluders
- 2) For each camera
 - a) Render triangles indices
 - b) Render Occluders on top
 - c) Collect Data
- 3) Aggregate Data
- 4) Remove Invisible Triangles



#GCAP24

Final step is to just read the data back to the CPU and remove any triangles that are not visible and in turn any vert that is not part of any remaining triangle.

Procedural Meshes

Visibility Culling

ViewCuller main method

```
private static bool[] GetVisibilityPerTriangle(Mesh mesh, Matrix transform, View[] views,
List<Position> meshOccluders, Vector2Int resolution)
{
    // Initialize render texture
    RenderTexture renderTexture = new(resolution.x, resolution.y, 24, GraphicsFormat.R32_UInt);
    {
        enableRandomWrite = true
    }
    // Initialize and Setup Buffers
    {
        // Render mesh and determine visibility
        CommandBuffer cmd = CommandBufferPool.Get();
        float aspect = resolution.x / ((float)resolution.y);
        for (int i = 0; i < views.Length; i++)
        {
            cmd.Clear();
            Matrix viewMatrix = views[i].viewMatrix;
            Matrix projectionMatrix = viewMatrix.InverseProjectionMatrix(aspect, 0.1f, 1000f);
            cmd.SetRenderTarget(renderTexture);
            cmd.ClearRenderTarget(false, true, Color.black);
            cmd.SetProjectionMatrix(projectionMatrix, projectionMatrix);
            cmd.DrawProcedural(transform, IndexMaterial, 0, MeshTopology.Triangles, 0, triangleCount, propertyBlock);
            for (int j = 0; j < meshOccluders.Count; j++)
            {
                cmd.DrawMesh(meshOccluders[j]);
            }
            cmd.DrawMesh(mesh, occluder.Transform, OccluderMaterial);
            cmd.DispatchComputeCompute, triangleVisibilityKernel, triangleVisibilityDispatchThreads, triangleVisibilityDispatchThreads, 1);
            Graphics.ExecuteCommandBuffer(cmd);
        }
        //Extract Data
        CommandBufferPool.Release(cmd);
        bool[] triangleVisibilityArray = new bool(triangleCount);
        triangleVisibilityBuffer.SetData(triangleVisibilityArray);
        bool[] visibilityPerTriangle = new bool(triangleCount);
        for (int i = 0; i < triangleCount; i++)
        {
            visibilityPerTriangle[i] = triangleVisibilityArray[i] != 0;
        }
        //Release Buffers and int
    }
    return visibilityPerTriangle;
}
```

#GCAP24

Index drawing shader

```
StructuredBuffer<int> Indices;
StructuredBuffer<float> Positions;

struct Varyings
{
    float3 position : SV_POSITION;
    uint instanceId : TEXCOORD0;
};

Varyings vert (uint vid : SV_VertexID, uint instanceId : SV_InstanceID)
{
    Varyings OUT = (Varyings);
    OUT.position = TransformObjectToClipSpace(Indices[instanceId * 3 + vid]);
    OUT.instanceId = instanceId;
    return OUT;
}

uint4 frag(Varyings IN) : SV_Target
{
    return IN.instanceId + 1;
}
```

Index gathering compute shader

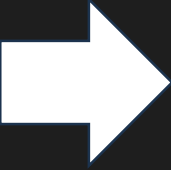
```
Texture2D<uint> RenderTexture;
StructuredBuffer<uint> TriangleVisibilityBuffer;

[numthreads(32,32,1)]
void CheckTriangleVisibility (uint2 id : SV_DispatchThreadID)
{
    uint index = RenderTexture[id];
    if (index > 0)
    {
        TriangleVisibilityBuffer[index - 1] = 1;
    }
}
```

Code Reference

Procedural Meshes

Visibility Culling

Vertices: 6567816 (0.66 GB)			Vertices: 452169 (51.7 MB)	
Position	Float32 x 3 (12 bytes)		Position	Float32 x 3 (12 bytes)
Normal	Float32 x 3 (12 bytes)		Normal	Float32 x 3 (12 bytes)
Tangent	Float32 x 4 (16 bytes)		Tangent	Float32 x 4 (16 bytes)
Color	UNorm8 x 4 (4 bytes)		Color	Float32 x 4 (16 bytes)
UV0	Float32 x 4 (16 bytes)		UV0	Float32 x 4 (16 bytes)
UV1	Float32 x 4 (16 bytes)		UV1	Float32 x 4 (16 bytes)
UV2	Float32 x 4 (16 bytes)		UV2	Float32 x 4 (16 bytes)
UV3	Float32 x 4 (16 bytes)		UV3	Float32 x 4 (16 bytes)

+90% REDUCTION

#GCAP24



Using a random model as our example
Culling can reduce our meshes up to ~90%!!!
That's right from 6.5 million down to 4.5 hundred thousand verts!
The Nintendo Switch can breath now.

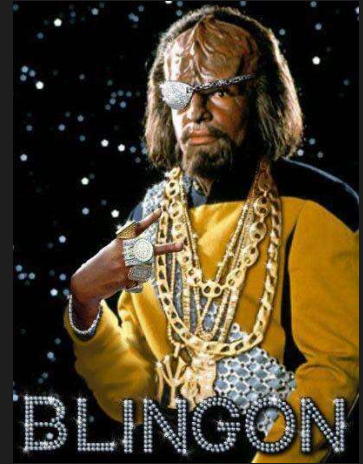
Procedural Meshes

Detail Drawer

Many of our atomic meshes have what we call “**details**”
And they are small, repetitive geo that are everywhere.

Let's dub them **BLING**

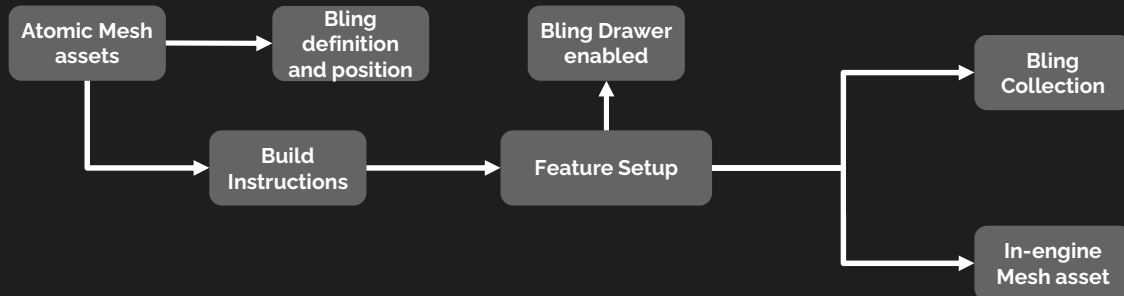
#GCAP24



The visibility culling greatly reduced our meshes from millions of verts to hundreds of thousands, but that was still not enough. Specially on weaker devices the final meshes were still taking too much time to process and consuming too much memory. A lot of the geometry was spent into small details that were repeated all over the model. So, to tackle them I created a new feature to render them as instanced meshes.

Procedural Meshes

BLING Drawer



#GCAP24

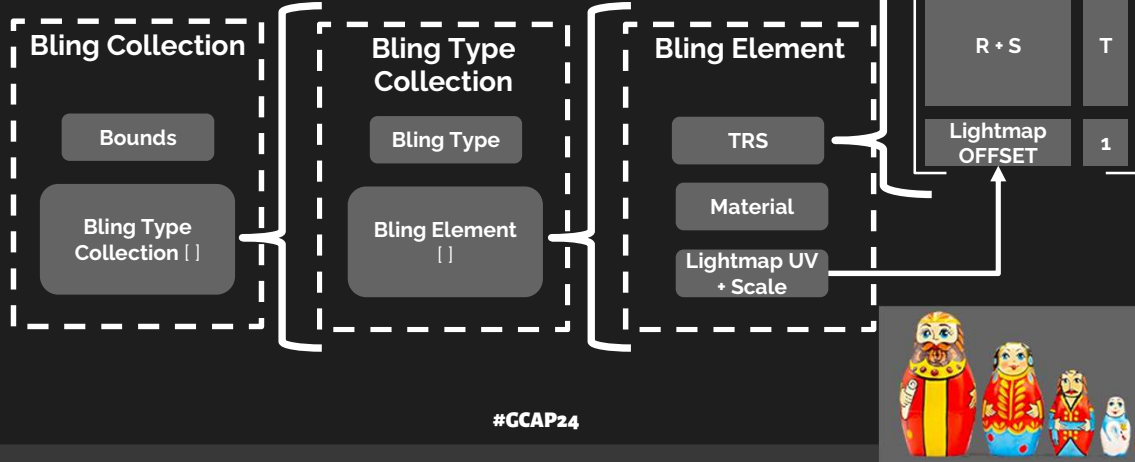
Going back into the AMP I've implemented a few changes:

For each atomic mesh that had Bling we would store their definition and position as data.

When ingesting the instructions, the AMP would check if the Bling Drawer feature was enabled or not, being able to bake Bling into the geometry or output a scriptable object with a Bling Collection for the model.

Procedural Meshes

BLING Drawer



#GCAP24

This Bling Collection contains the bounds for all Bling in it for visibility evaluation down the line.

It also has an array of Bling Type Collections.

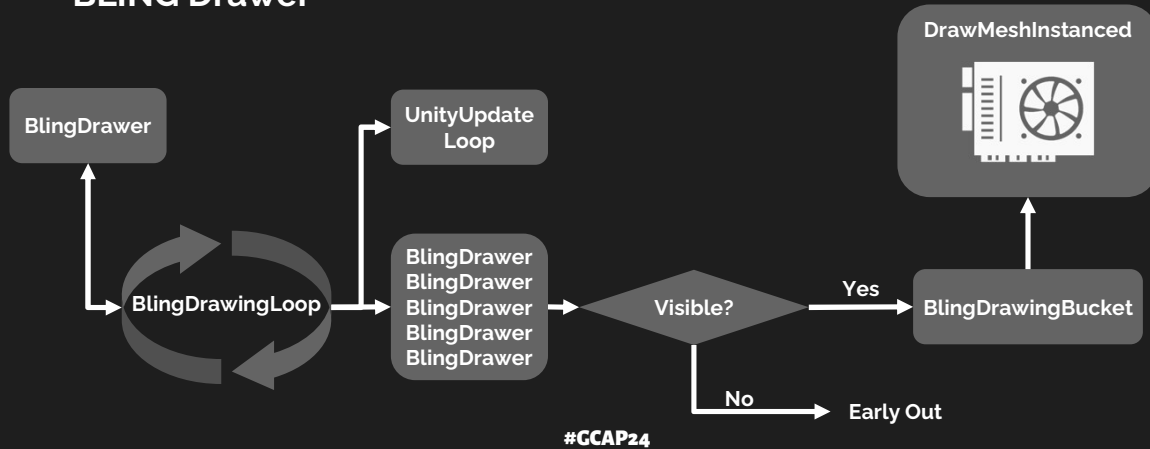
Each of the Type Collections has the bling Type as a string key and another array of Bling Elements

At the bottom of it all the Bling Element contains a TRS Matrix4x4, an uint for compressed material and a Vector4 for Lightmap UV + Scale.

If you're familiar with TRS matrices you know that m30~m32 are set to 0, so we can sneak the Lightmap UV Offset there.

Procedural Meshes

BLING Drawer



To tie in with the runtime, we have the following parts:

1. A **BlingDrawer** component that can be attached to **GameObjects** and has all the settings for drawing a **Bling Collection**, such as material, cull distance and visibility calculation strategy
2. A **BlingDrawingLoop** that makes use of the **PlayerLoop API** to register a custom update method.
3. When a **BlingDrawer** gets called by the **BlingDrawingLoop** and is visible, it will call the **BlingDrawingBucket** which renders the collection via a **Graphics.DrawMeshInstanced** call

Culling Group API

<https://docs.unity3d.com/Manual/CullingGroupAPI.html>

Blog on custom update loops

<https://medium.com/@thebeardphantom/unity-2018-and-playerloop-5c46a12a677>

Procedural Meshes

BLING Drawer

BlingDrawingLoop

```
public static class BlingDrawerLoop
{
    private static HashSet<BlingDrawer> _activeDrawers = new HashSet<BlingDrawer>();

    static BlingDrawerLoop()
    {
        PlayerLoopUtility.RegisterUpdateMethod(UpdateLoop, typeof(BlingDrawerLoop), PlayerLoopUtility.LoopStage.PostLateUpdate);
    }

    private static void UpdateLoop()
    {
        foreach (BlingDrawer drawer in _activeDrawers)
        {
            drawer.Render();
        }
    }

    public static void Register(BlingDrawer drawer)
    {
        _activeDrawers.Add(drawer);
    }

    public static void Deregister(BlingDrawer drawer)
    {
        _activeDrawers.Remove(drawer);
    }
}
```

#GCAP24

BlingDrawer main methods

```
private async UniTaskVoid InitializeAsync()
{
    SetupReferences();
    Initialize();

    if (HasBakedLightmap())
    {
        // Set up light map properties for each collection
        for (int i = 0; i < _drawingCollections.Length; i++)
        {
            BlingDrawerLoop.Register(this);
        }
    }
}

internal void Render()
{
    if (forceRenderingOff)
    {
        return;
    }

    if (ShouldUpdateEveryFrame() && _transform.hasChanged)
    {
        UpdateTransform();
    }

    if (!isVisible && Application.isPlaying)
    {
        return;
    }

    int layer = gameObject.layer;
    UpdateMaterialsInEditor();
    Material renderMaterial = _lodMaterials[Mathf.Min(_lod, 2)];
    for (int i = 0; i < _drawingCollections.Length; i++)
    {
        _drawingCollections[i].Render(renderMaterial, layer, _lod);
    }
}
```

Code Reference

Procedural Meshes

BLING Drawer

BlingDrawingBucket buffer population

```
private static void PopulateBuffer(ElementLocator[] blings, ref Matrix4x4 localToWorld)
{
    for (int i = 0; i < blings.Length; i++)
    {
        ElementLocator bling = blings[i];

        _preallocatedBuffer[i].color = bling.MaterialIndex;
        _preallocatedBuffer[i].lightmapUV = bling.LightmapUV;

        FastMultiply(ref localToWorld, ref bling.trsr, ref _preallocatedBuffer[i].trs);
    }
}
```

BlingDrawingBucket render method

```
public void Render(Material material, int layer, int lod)
{
    if (lod >= _lodMeshes.Length)
    {
        lod = _lodMeshes.Length - 1;
    }

    Mesh mesh = _lodMeshes[lod];
    ComputeBuffer argsBuffer = _argsBuffers[lod];

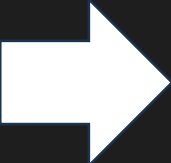
    Graphics.DrawMeshInstancedIndirect(mesh, 0, material, _bounds, argsBuffer, 0, _propertyBlock, ShadowCastingMode.Off, false, layer);
}
```

#GCAP24

Code Reference

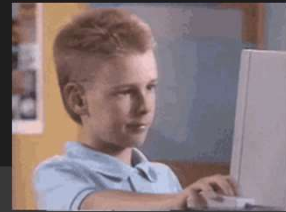
Procedural Meshes

BLING Drawer

Vertices: 452169 (51.7 MB)			Vertices: 392496 (44.9 MB)	
Position	Float32 x 3 (12 bytes)		Position	Float32 x 3 (12 bytes)
Normal	Float32 x 3 (12 bytes)		Normal	Float32 x 3 (12 bytes)
Tangent	Float32 x 4 (16 bytes)		Tangent	Float32 x 4 (16 bytes)
Color	Float32 x 4 (16 bytes)		Color	Float32 x 4 (16 bytes)
UV0	Float32 x 4 (16 bytes)		UV0	Float32 x 4 (16 bytes)
UV1	Float32 x 4 (16 bytes)		UV1	Float32 x 4 (16 bytes)
UV2	Float32 x 4 (16 bytes)		UV2	Float32 x 4 (16 bytes)
UV3	Float32 x 4 (16 bytes)		UV3	Float32 x 4 (16 bytes)

15% REDUCTION

#GCAP24



Implementing Bling Drawing on the same mesh as before we get a further 15% reduction on size plus improved resolution on Lightmaps

AMP WORKS!



**BUT IS SLOW, VERY SLOW
LET'S UNDERSTAND WHY**

#GCAP24



All the tech so far works and is amazing.
But ended up being very slow to process.
Could take upwards to 20min to ingest a model, and sometimes get some
computers to run out of memory!!!!
This felt more like Sanic when we needed it to be Sonic!

Optimization Bottlenecks

Compression

Vertices: 392496 (44.9 MB)			Vertices: 392496 (18.0 MB)	
Position	Float32 x 3 (12 bytes)		Position	Float16 x 4 (8 bytes)
Normal	Float32 x 3 (12 bytes)		Normal	SNorm8 x 4 (4 bytes)
Tangent	Float32 x 4 (16 bytes)		Tangent	SNorm8 x 4 (4 bytes)
Color	Float32 x 4 (16 bytes)		Color	UNorm8 x 4 (4 bytes)
UV0	Float32 x 4 (16 bytes)		UV0	Float16 x 4 (8 bytes)
UV1	Float32 x 4 (16 bytes)		UV1	Float16 x 4 (8 bytes)
UV2	Float32 x 4 (16 bytes)		UV2	Float16 x 4 (8 bytes)
UV3	Float32 x 4 (16 bytes)		UV3	SNorm8 x 4 (4 bytes)

~50% LESS MEMORY

#GCAP24



Why we do it?

Much smaller memory footprint

Update Data formats

Float32 to Float16 or even Norm8

Almost 60% reduction!!!

Our meshes are chunky and don't suffer from compression in most cases

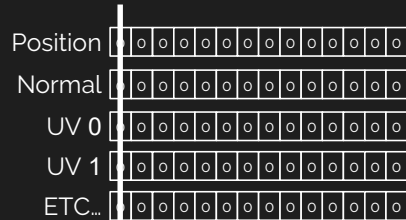
Optimization Bottlenecks

Compression

Compress on the CPU

vert by vert

Slow, VERY slow



High Compressed Vertex Data

0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
---	---	---	---	---	---	---	---	---	---	---	---	---	---	---	---

#GCAP24

Our first solution was great and functional, which served us well while investing into a vertical slice.

It just used C# to read and convert the data from regular Float32 into smaller Half16 or SByte4, following the format outlined by our HighCompressedVertexData struct

But this process is linear and greatly suffer from very large meshes with a long vert strip.

Optimization Bottlenecks

Compression

Repetitive Task?

Has well-defined constraints?

**COMPUTE SHADERS
TO THE RESCUE!!!**

Running code on Compute Shaders



#GCAP24

Compute shaders works wonders on well defined parallelized tasks such as
stride defined vert strips!
20% reduction on total processing time

Optimization Bottlenecks

Compression

```
HighCompressedVertexData[] vertexData = new HighCompressedVertexData[vertexCount];  
ComputeBuffer outputBuffer = new ComputeBuffer(vertexCount, Marshal.SizeOf(typeof(HighCompressedVertexData)));  
Compute Setup and Dispatch  
outputBuffer.GetData(vertexData);
```

Setting the data to a defined format is easy to setup, but...

No Normals? Different format

No Colors? Different format

But... Aren't we just moving bytes around?

#GCAP24



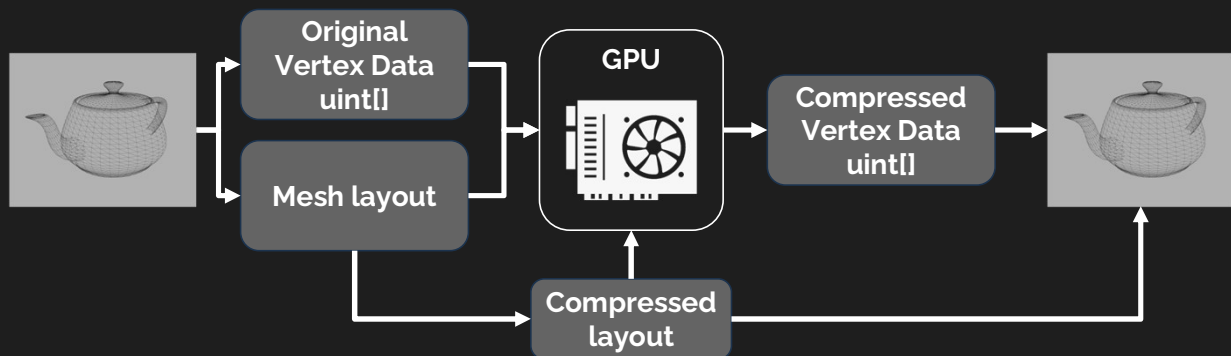
But this first approach is hardcoded on HighCompressedVertexData struct vertex format, which might need to be cleaned afterwards if the input did not have all attributes.

And despite having meshes with more or less channels, the compression per channel is consistent across the project.

But it's all just bytes...

Optimization Bottlenecks

Compression: CPU Side



#GCAP24

Instead of relying on `HighCompressedVertexData` struct, we can use a plain `uint` array, where bytes can be directly read from the GPU and applied to the mesh via the method `Mesh.SetVertexBufferData()` afterwards.

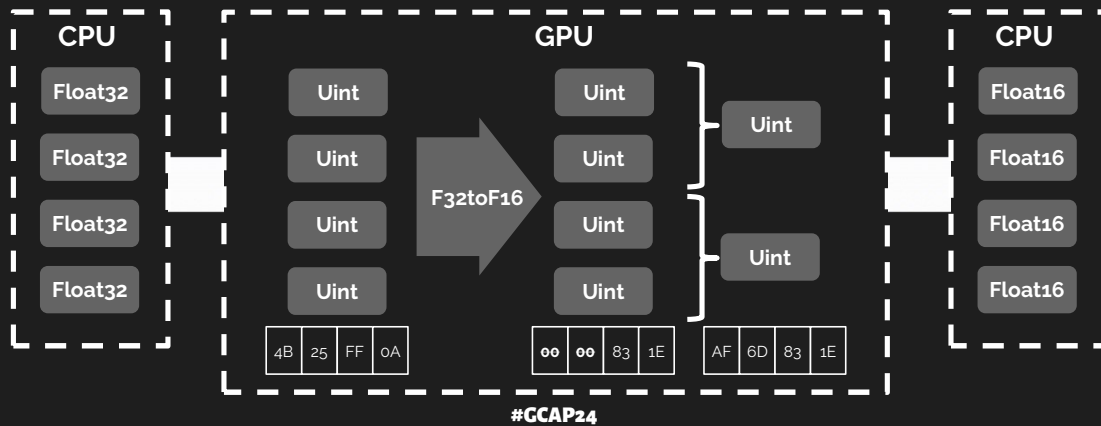
This is done by extracting the vertex data and mesh layout with Unity's own methods.

The data gets passed as is to the compute shader, while the layout is passed as Booleans, such as `_hasNormal`, or `_Texcoord1Has4Channels`.

After compression the data is fetched back into the CPU, the mesh has it's layout updated and finally it receives the now compressed data.

Optimization Bottlenecks

Compression: GPU Side (Float32x4 → Float16x4 example)



On the GPU side we just need to read the bytes for each uncompressed value as their uint representation, such as a float3 becomes uint3, given they both have the same bit count (4 bytes * 3).

From there we can compress following any rule we need, as long as each channel fit within the multiple of 4 byte count, e.g.: float3 is fine (4 bytes * 3 = 12), but half3 isn't (2 bytes * 3 = 6).

It's also important to note that the output will be packed into uints, which are 4 bytes each, so 2 halves need to be packed into a single uint, or 4 Unorm8 into the same single uint.

Finally all these uint soup is appended to the output buffer sequentially, which we can parallelize as we know the exact stride (or size) of each vert data.

Optimization Bottlenecks

Compression

Compressed Layout

```
private static readonly VertexAttributeDescriptor[] HighCompressedVertexAttributes =
{
    new(VertexAttribute.Position, VertexAttributeFormat.Float16, 4),
    new(VertexAttribute.Normal, VertexAttributeFormat.SNorm8, 4),
    new(VertexAttribute.Tangent, VertexAttributeFormat.SNorm8, 4),
    new(VertexAttribute.Color, VertexAttributeFormat.UNorm8, 4),
    new(VertexAttribute.TexCoord0, VertexAttributeFormat.Float16, 4),
    new(VertexAttribute.TexCoord1, VertexAttributeFormat.Float16, 4),
    new(VertexAttribute.TexCoord2, VertexAttributeFormat.Float16, 4),
    new(VertexAttribute.TexCoord3, VertexAttributeFormat.SNorm8, 4),
};
```

Compressed Vertex Format

```
[StructLayout(LayoutKind.Sequential)]
private struct HighCompressedVertexData
{
    public half4 position;
    public SByte4 normal;
    public SByte4 tangent;
    public Color32 color;
    public half4 texCoord0;
    public half4 texCoord1;
    public half4 texCoord2;
    public SByte4 texCoord3;
}
```

#GCAP24

Compression on the CPU

```
halfMesh.SetVertexBufferParams(vertexCount, HighCompressedVertexAttributes);
NativeArray<HighCompressedVertexData> vertexData = new(vertexCount, Allocator.Temp);
for (int i = 0; i < vertexCount; i++)
{
    var data = new HighCompressedVertexData();
    var position = positions[i];
    data.position.x = (half)position.x;
    data.position.y = (half)position.y;
    data.position.z = (half)position.z;
    data.position.w = new half(1);
    SByte4.ReadFrom(normals[i], ref data.normal);
    SByte4.ReadFrom(tangents[i], ref data.tangent);
    data.color = colors[i];
    if (texCoord0.Count > 0)
    {
        ReadHalf4FromVector4(texCoord0, ref i, ref data.texCoord0);
    }
    if (texCoord1.Count > 0)
    {
        ReadHalf4FromVector4(texCoord1, ref i, ref data.texCoord1);
    }
    if (texCoord2.Count > 0)
    {
        ReadHalf4FromVector4(texCoord2, ref i, ref data.texCoord2);
    }
    if (texCoord3.Count > 0)
    {
        SByte4.ReadFrom(texCoord3[i], ref data.texCoord3);
    }
    vertexData[i] = data;
}
halfMesh.SetVertexBufferData(vertexData, 0, 0, vertexCount);
vertexData.Dispose();
```

Code Reference

Optimization Bottlenecks

Compression

Calling the compute method,
setting up the mesh layout
and setting the vert data

```
uint[] vertexData = UniversalMeshCompressionUtils.RunComputeShaderCompression(mesh, meshLayout);
mesh.SetVertexBufferParams(mesh.vertexCount, meshLayout.GetLayout());
mesh.SetVertexBufferData(vertexData, 0, 0, mesh.vertexCount);
```

Getting original data and running the compute shader

```
internal static uint[] RunComputeShaderCompression(Mesh mesh, MeshLayout meshLayout)
{
    ComputeShader computeShader = CompressionCompute.ComputeShader;
    computeShader.GetKernelThreadGroupSizes(ComputeShaderKernel, out uint threadCount, out _, out _);

    Setup Shader Params
    mesh.vertexBufferTarget |= GraphicsBuffer.Target.Raw;
    GraphicsBuffer originalBuffer = mesh.GetVertexBuffer(0);
    computeShader.SetBuffer(ComputeShaderKernel, "_Input", originalBuffer);
    uint[] vertexData = new uint[mesh.vertexCount * meshLayout.UintStride];
    ComputeBuffer outputBuffer = new ComputeBuffer(vertexData.Length, sizeof(uint));
    computeShader.SetBuffer(ComputeShaderKernel, "_UniversalOutput", outputBuffer);
    computeShader.Dispatch(ComputeShaderKernel, Mathf.CeilToInt(mesh.vertexCount / (float)threadCount), 1, 1);
    outputBuffer.GetData(vertexData);
    outputBuffer.Dispose();
    originalBuffer.Dispose();
    return vertexData;
}
```

#GCAP24

Code Reference

Optimization Bottlenecks

Compression

Kernel Method

```
void UniversalCompressMesh(uint3 id : SV_DispatchThreadID)
{
    int idx = id.x;
    if (idx >= _MeshVertexCount)
    {
        return;
    }
    UncompressedVert uncompressed = ReadUncompressedVert(idx);
    uint outputIndex = idx * _OutputMeshStride;
    #define Append(value) _UniversalOutput[outputIndex] = value; \
    outputIndex++;
    if (_uncompressedPosition)
    {
        Append(uncompressed.uncompressedPosition.x)
        Append(uncompressed.uncompressedPosition.y)
        Append(uncompressed.uncompressedPosition.z)
    }
    else
    {
        uint2 position = CompressFloatToInt32(float3(uncompressed.position.xyz, 1));
        Append(position.x)
        Append(position.y)
    }
}
```

Beginning of the compute shader, it keeps going depending on the mesh layout. But it's just more of the same from here

Data Reading

```
UncompressedVert uncompressed = (UncompressedVert) 0;
uint inputIndex = idx * _MeshStride;
int offset = 0;
uncompressed.uncompressedPosition = _Input.Load3(inputIndex + offset);
uncompressed.position = asfloat(uncompressed.uncompressedPosition);
offset += 4 * 3; //float3
if (_hasNormal)
{
    uncompressed.normal = asfloat(_Input.Load3(inputIndex + offset));
    offset += 4 * 3; //float3
}
if (_hasTangent)
{
    uncompressed.tangent = asfloat(_Input.Load3(inputIndex + offset));
    offset += 4 * 4; //float4
}
```

This is the beginning of the ReadUncompressedVert method. The rest will just read the other mesh data

Compression Methods

```
uint2 CompressFloatToInt32(float3 value)
{
    uint x = floorf(value.x);
    uint y = floorf(value.y);
    return x | (y << 16);
}
uint CompressFloatToInt32BySign(float3 value)
{
    bool3 sign = value < 0;
    float3 compressed = abs(value) * 127;
    uint3 compressedInt = round(compressed);
    compressedInt = compressedInt - 127;
    return compressedInt <> 0 ? 0 : 0;
    return compressedInt.x | (compressedInt.y << 16) | (compressedInt.z << 32);
}
uint2 CompressFloatToInt32ByNormal(float3 value)
{
    float3 compressed = value * 200;
    compressed = round(compressed);
    return (compressedInt.x <> 0 ? 0 : 0) | (compressedInt.y << 16) | (compressedInt.z << 32);
}
uint2 CompressFloatToInt32ByTangent(float3 value)
{
    return abs2(CompressFloatToInt32ByNormal(value.x), CompressFloatToInt32ByNormal(value.y));
}
```

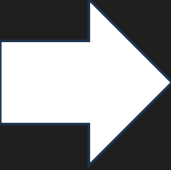
These are our compression methods to convert Float32 values into lower precision values

#GCAP24

Code Reference

Optimization Bottlenecks

Compression

Vertices: 392496 (44.9 MB)			Vertices: 392496 (18.0 MB)	
Position	Float32 x 3 (12 bytes)		Position	Float16 x 4 (8 bytes)
Normal	Float32 x 3 (12 bytes)		Normal	SNorm8 x 4 (4 bytes)
Tangent	Float32 x 4 (16 bytes)		Tangent	SNorm8 x 4 (4 bytes)
Color	Float32 x 4 (16 bytes)		Color	UNorm8 x 4 (4 bytes)
UV0	Float32 x 4 (16 bytes)		UV0	Float16 x 4 (8 bytes)
UV1	Float32 x 4 (16 bytes)		UV1	Float16 x 4 (8 bytes)
UV2	Float32 x 4 (16 bytes)		UV2	Float16 x 4 (8 bytes)
UV3	Float32 x 4 (16 bytes)		UV3	SNorm8 x 4 (4 bytes)

#GCAP24

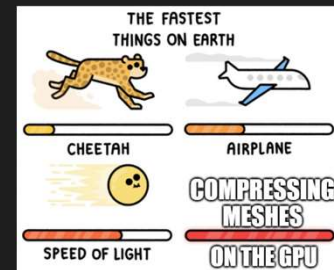
Taking the starting case

Optimization Bottlenecks

Compression

Vertices: 392496 (44.9 MB)		Vertices: 392496 (18.0 MB)	
Position	Float32 x 3 (12 bytes)	Position	Float16 x 4 (8 bytes)
Normal	Float32 x 3 (12 bytes)	Normal	SNorm8 x 4 (4 bytes)
Tangent	Float32 x 4 (16 bytes)	Tangent	SNorm8 x 4 (4 bytes)
Color	Float32 x 4 (16 bytes)	Color	UNorm8 x 4 (4 bytes)
UV0	Float32 x 4 (16 bytes)	UV0	Float16 x 4 (8 bytes)
UV1	Float32 x 4 (16 bytes)	UV1	Float16 x 4 (8 bytes)
UV2	Float32 x 4 (16 bytes)	UV2	Float16 x 4 (8 bytes)
UV3	Float32 x 4 (16 bytes)	UV3	SNorm8 x 4 (4 bytes)

900% FASTER
25% total time reduction



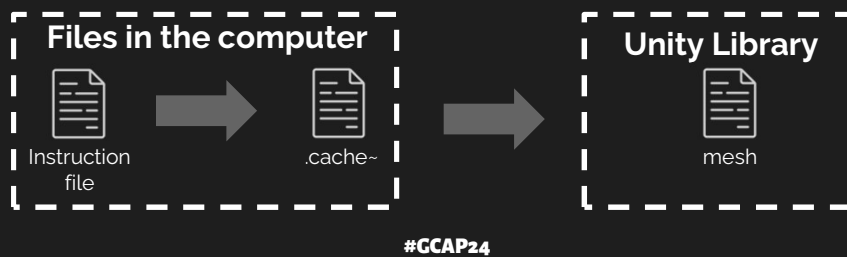
#GCAP24

The Compute Shader method is ~900% faster!
Which cascades as the entire AMP being 25% faster!

Optimization Bottlenecks

Caching and Serialization

Talking about bytes and precision, did you know that different platforms handle these differently?
To ensure the resulting assets are the same across every computer and build we created a cache system.



Different platform could generate slightly different results for the mesh generation, specially lightmap unwrapping using xAtlas.
To ensure that the results would be stable and reliable (as well as faster) we created a cache system where the AMP, when importing an instruction file it would check for an adjacent .cache~ file;
If it finds one the .cache~ is read and the result mesh is added to the Unity library, but if it's missing (due to being a new instruction file) it would be generated from scratch and a .cache~ would be then created and committed to the repo.

Optimization Bottlenecks

Caching and Serialization

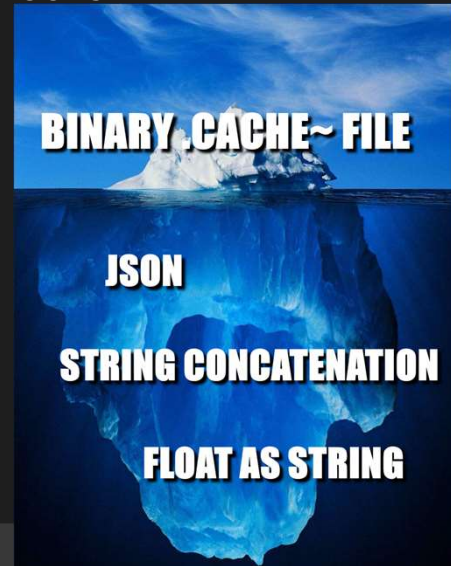
The caching is done by serializing the result meshes to a file as binary.
But how is the binary serialization done?

JSON

All the data was being interpreted as numbers
Then as text
Then concatenated into JSON
Then converted into bytes
Finally wrote to disk

AND BACK!!!

#GCAP24



The process was sound and seemed great, but how did it work?

Via JSON strings.

Converting data into text and then into binary so it could be written to disk, and the inverse when reading.

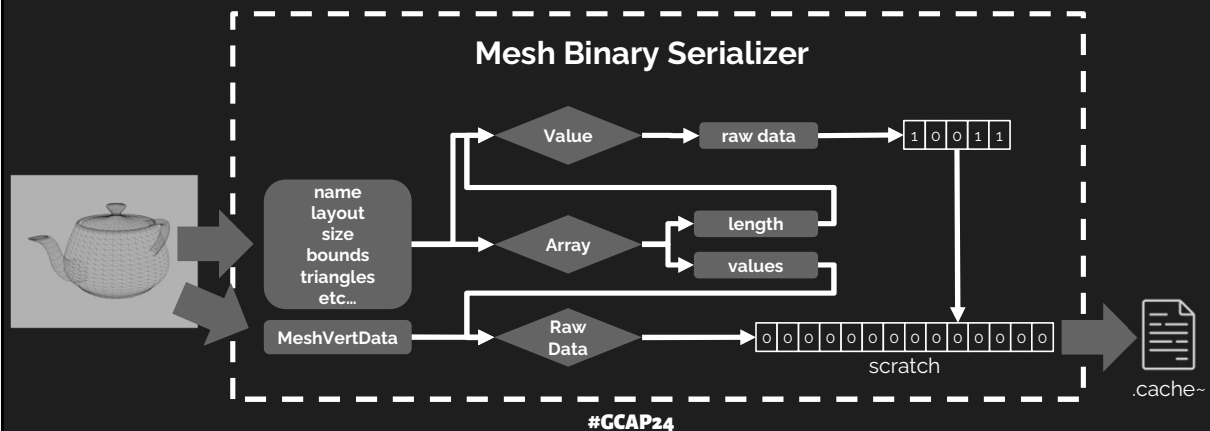
This was bad as not only it was quite slow, it also ate up a lot of memory AND lost precision, as different hardware round floating point numbers differently on ToString().

The memory situation was so bad that some instruction files couldn't be ingested as it would require more memory than people's computer had!

Also found out that whenever a mesh was generated, we'd serialize it to disk and then immediately deserialize it back, wasting lots of time when (re)generating a cache.

Optimization Bottlenecks

Caching and Serialization



The improvement was a new serializer, that works in the following way:

We start with a mesh, then create a new Serializer, which contains a large byte array that we'll call the scratch.

First we extract the info from the mesh, such as name, attribute layout, triangles, etc...

That info can be separated into a few different types:

The first type is "Values", which contain 2 parts

- the type size, for example Bounds is a Vector3, so 3 floats for 4 x 3 bytes
- the raw data, the actual binary representation of the value

The type size ends up being a raw data itself, and all raw data can be copied into the scratch

The second type is "Arrays", which also contain 2 parts

- the length, which can be saved as a Value
- the actual values, which can be copied as raw data

Raw data is highlight as a third type because we can serialize the MeshVertData from the mesh directly.

Then all this get written to disk as a string representation of the scratch

The inverse of this process is done for deserializing and this ended up being 25% faster for deserializing .cache~ files as well as massively reducing the memory usage, allowing the team to handle any instruction that they encountered since without crashes!

Optimization Bottlenecks

Caching and Serialization

Main Serialize Methods

```
public int SerializeMesh(mesh)
{
    SerializeString(mesh.name);
    SerializedElement<int> meshIndexFormat;
    SerializedElement<meshIndexFormat> meshIndexFormat;
    for (VertexAttributeFormat attribute = mesh.VertexAttributeFormat; attribute < VertexAttributeFormat.TexCoords; attribute++)
    {
        VertexAttributeFormat format = mesh.VertexAttributeFormat(attribute);
        SerializedElement<int> format;
        SerializedElement<int> dimension;
        SerializedElement<int> meshVertexCount;
        SerializedElement<int> meshTriangleCount;
        SerializedElement<int> meshBounds;
        SerializedElement<int> mesh;
        return (int)scratchPosition;
    }
}

private void SerializeMeshHeader(mesh)
{
    long stride = mesh.VertexAttributeFormat;
    int length = (int)mesh.VertexCount * stride;
    SerializedElement<int> length;
    GraphicsHeader buffer = mesh.VertexHeader();
    FileHeaderHeader(buffer);
    buffer.GetData_scratch((int)scratchPosition, 4, length);
    _scratchPosition += (int)length;
}
```

Main Deserialize Methods

```
public Mesh Deserialize(byte[] data)
{
    PopulateScratchWithHeader(data);
    Mesh result = new();
    result.name = DeserializeString();
    result.indexFormat = DeserializeElement<int>();
    result.indexFormat = DeserializeElement<int>();
    result.indexFormat = DeserializeElement<int>();
    List<VertexAttributeDescriptor> descriptors = new();
    for (VertexAttributeFormat attribute = VertexAttributeFormat; attribute < VertexAttributeFormat.TexCoords; attribute++)
    {
        VertexAttributeFormat format = VertexAttributeFormat.DeserializeElement<int>();
        int dimension = (int)DeserializeElement<int>();
        if (dimension > 0)
        {
            descriptors.Add(new VertexAttributeDescriptor(attribute, format, dimension));
        }
    }
    int vertexCount = (int)DeserializeElement<int>();
    result.SetVertexHeaderParameters(vertexCount, descriptors.ToArray());
    int[] triangles = DeserializeElement<int>();
    result.bounds = DeserializeElement<int>();
    result.meshIndexFormat = DeserializeElement<int>();
    result.meshIndexFormat = DeserializeElement<int>();
    result.meshIndexFormat = DeserializeElement<int>();
    result.meshIndexFormat = DeserializeElement<int>();
    result.meshIndexFormat = DeserializeElement<int>();
    return result;
}

private void DeserializeMeshHeader(mesh)
{
    int length = DeserializeElement<int>();
    mesh.SetVertexHeaderParameters(scratch, (int)scratchPosition, 4, (int)length);
}
```

#GCAP24

High-Level Methods

```
private static byte[] _scratch = new byte[1048576]; // 1 MB
private int _scratchPosition;

public static string SaveMeshIntoData(Mesh mesh)
{
    if (mesh == null)
    {
        return string.Empty;
    }
    MeshBinarySerializer<int> serializer = new();
    int length = serializer.Serialize(mesh);
    byte[] compressedData = CompressionUtility.Compress(_scratch, length);
    return Convert.ToBase64String(compressedData);
}

public static Mesh LoadFromData(string serializedData)
{
    byte[] compressedData = Convert.FromBase64String(serializedData);
    byte[] data = CompressionUtility.Decompress(compressedData);
    MeshBinaryDeserializer<int> meshData = new MeshBinaryDeserializer<int>();
    return meshData.Deserialize(data);
}
```

Low-Level Methods

```
private unsafe void SerializeElement<T>(T element)
    where T : struct
{
    int length = Marshal.SizeOf<T>();
    GCHandle handle = GCHandle.Alloc(element, GCHandleType.Pinned);
    Marshal.Copy(handle.AddrOfPinnedObject(), _scratch, (int)scratchPosition, length);
    handle.Free();
    _scratchPosition += (int)length;
}

private unsafe T DeserializeElement<T>()
    where T : struct
{
    int length = Marshal.SizeOf<T>();
    IntPtr unmanagedPointer = Marshal.UnsafeAddrOfPinnedObject(scratch, (int)scratchPosition);
    T element = Marshal.PtrToStructure<T>(unmanagedPointer);
    _scratchPosition += (int)length;
    return element;
}
```

Code Reference

Optimization Bottlenecks

Caching and Serialization

Vertices: 392496 (18.0 MB)	
Position	Float16 x 4 (8 bytes)
Normal	SNorm8 x 4 (4 bytes)
Tangent	SNorm8 x 4 (4 bytes)
Color	UNorm8 x 4 (4 bytes)
UV0	Float16 x 4 (8 bytes)
UV1	Float16 x 4 (8 bytes)
UV2	Float16 x 4 (8 bytes)
UV3	SNorm8 x 4 (4 bytes)

55% Faster Write

65% Faster Read

**25% Faster AMP + NO crashes +
Platform agnostic!**

#GCAP24

Using the same mesh we've compressed before, the new serializer is 55% faster when writing and 65% faster when reading (which happen more often!)

This led to another speed up of 25% of the AMP!!!

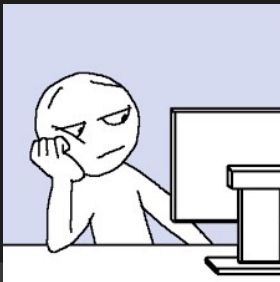
It also prevented memory related crashes

Productivity

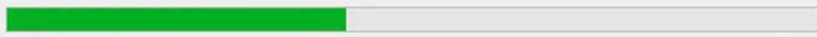
Why should we care?

It's common to ignore Editor Tools performance

But at the end of the day that's the team productivity being wasted watching loading bars



Running AMP (busy for 42m)...



Ingesting Instructions

#GCAP24

And why do all this just for an Editor tool? No one cares if it's slow.

In reality we all **SHOULD** care if it's too slow.

Because it's people's time we're talking about, it's not a big deal if a button takes a few seconds.

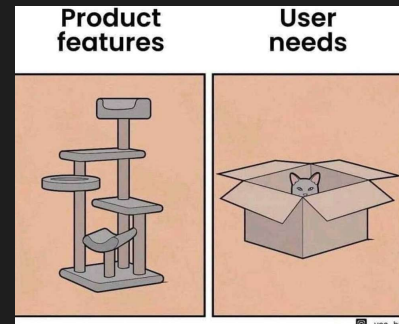
But if it takes many minutes, you might want to look into it.

There is a balance as some process are done sparingly and can be slow, but in those cases it shouldn't crash after 30min or it will be a huge source of frustration.

Productivity

How to identify the problem?

You might wonder how much time I spent using the AMP.
Almost zero, nada, zilch.
So how could I identify the problems?



Although I put a lot of effort into understanding and improving the tool, I barely used it.

So how could I identify issues?

By doing 3 main things:

- 1) Asking the team how they were doing, what was bad and annoying
- 2) Having an open call policy, whenever they wanted to complain about something I'm always a slack call away
- 3) By lurking on their open huddles and paying attention to how they work so I can spot pain points

Productivity

Do you need to do it all yourself?

NO!

The same way that you're doing this for the team, **rely on them for help**

Engage your team in coming up with solutions, even with coding tasks

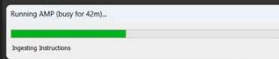
This helps people to stay motivated and yourself a little bit saner!

#GCAP24

This guy just talked like a train wreck, what happened here?

TL;DR:

- Profile frequently
- Assets are just byte soups
- GPU not only draw frames
- Ask questions and listen to your team



#GCAP24

- Profile frequently
 - Profile the game and profile the engine, your tools might be killing you
- Assets are just byte soups
 - Don't be afraid to treat them like so
- Compute shaders are a powerful

tool for grunt work and even draw
calls can be used to “paint” data

- Ask questions and listen to your heart, I mean, your team
 - Talk to people who use your tools, they might not know how it works, but they know how it doesn't

Questions?

#GCAP24

The Tech Art Journey
of creating and optimizing an asset pipeline

Obrigado!

 /vicnascimento



SLIDES

#GCAP24

