

GPU TECHNOLOGY CONFERENCE

Implementing CUDA Audio Networks

Giancarlo Del Sordo

✿ Acustica Audio

giancarlo@acusticaudio.com

Motivation

- Vintage gear processing in software domain
- High audio quality results
- Low cost and user-driven content approach



Outline

- Audio theory
- Basic engine solution
- Network implementation
- Results



Audio theory



Standard DSP approach



Example: simple 2 poles filter (source www.musicdsp.org)

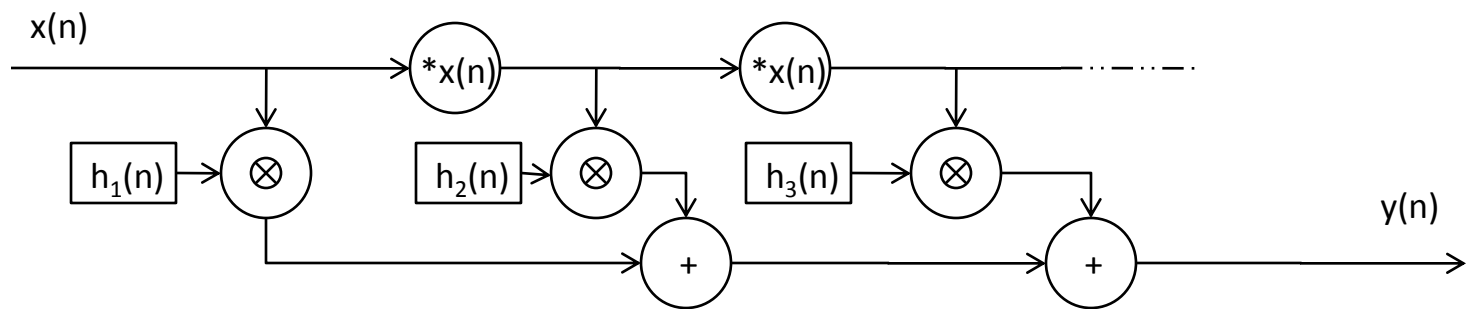
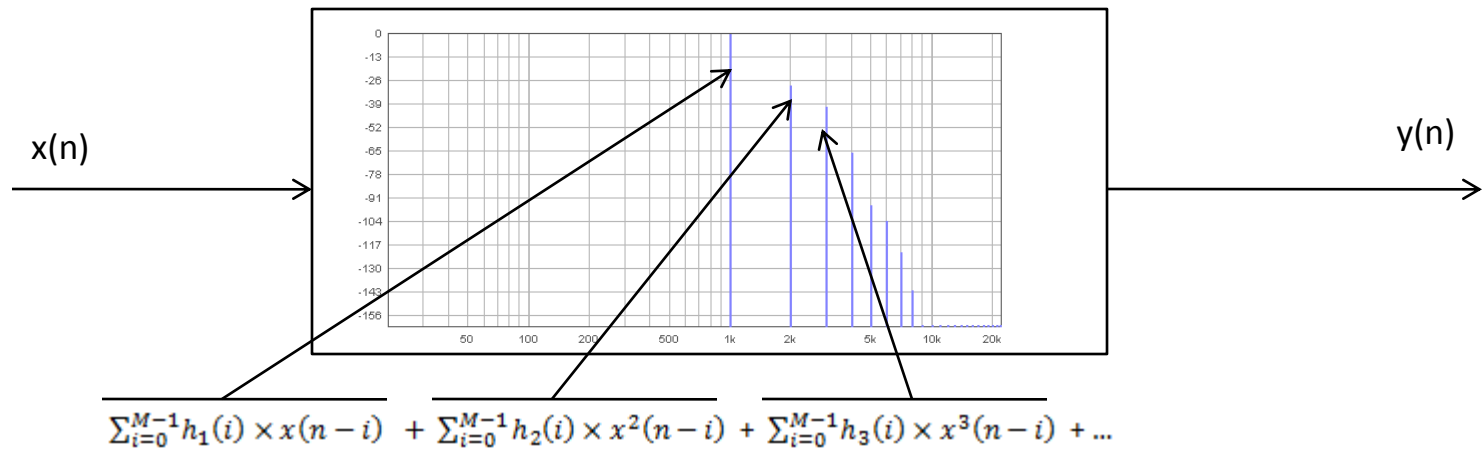
```
c = pow(0.5, (128-cutOff) / 16.0);  
r = pow(0.5, (resonance+24) / 16.0);  
  
//Loop:  
  
v0 = (1-r*c)*v0 - (c)*v1 + (c)*input;  
v1 = (1-r*c)*v1 + (c)*v0;  
  
output = v1;
```

Drawbacks:

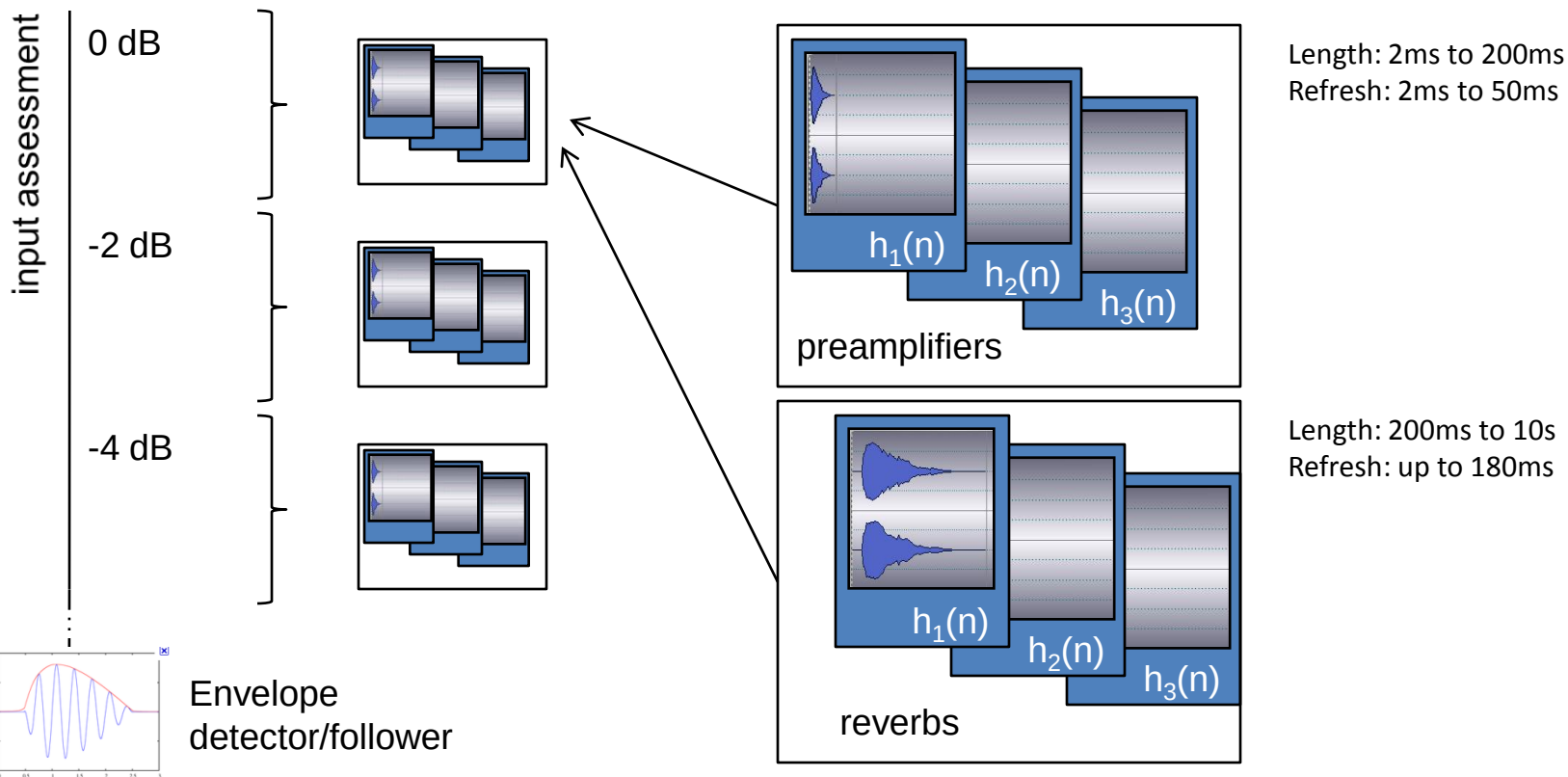
- time/costs required for achieving good quality algorithms
- complexity for taking advantage of GPU resources



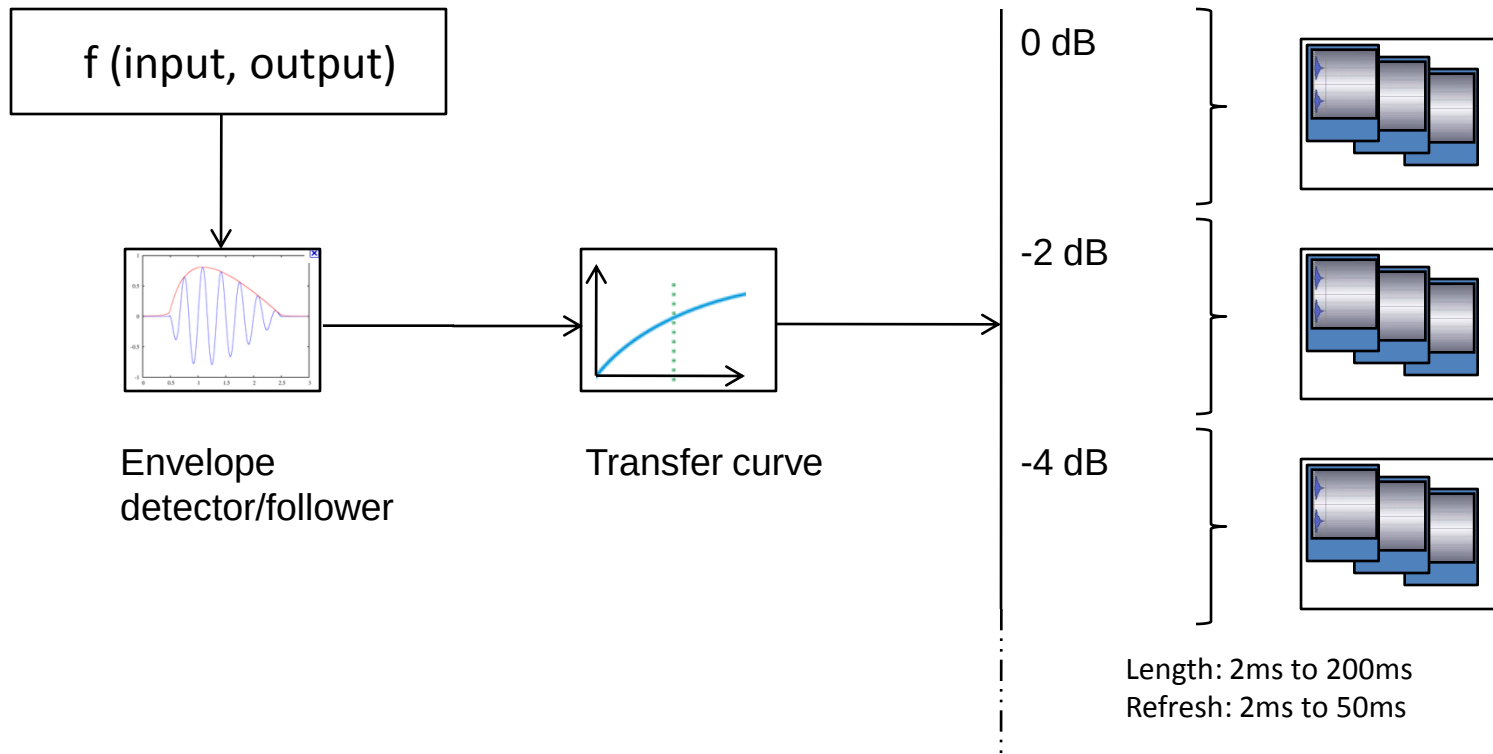
Volterra Series



Solution for dynamic processing (preamps, reverbs)

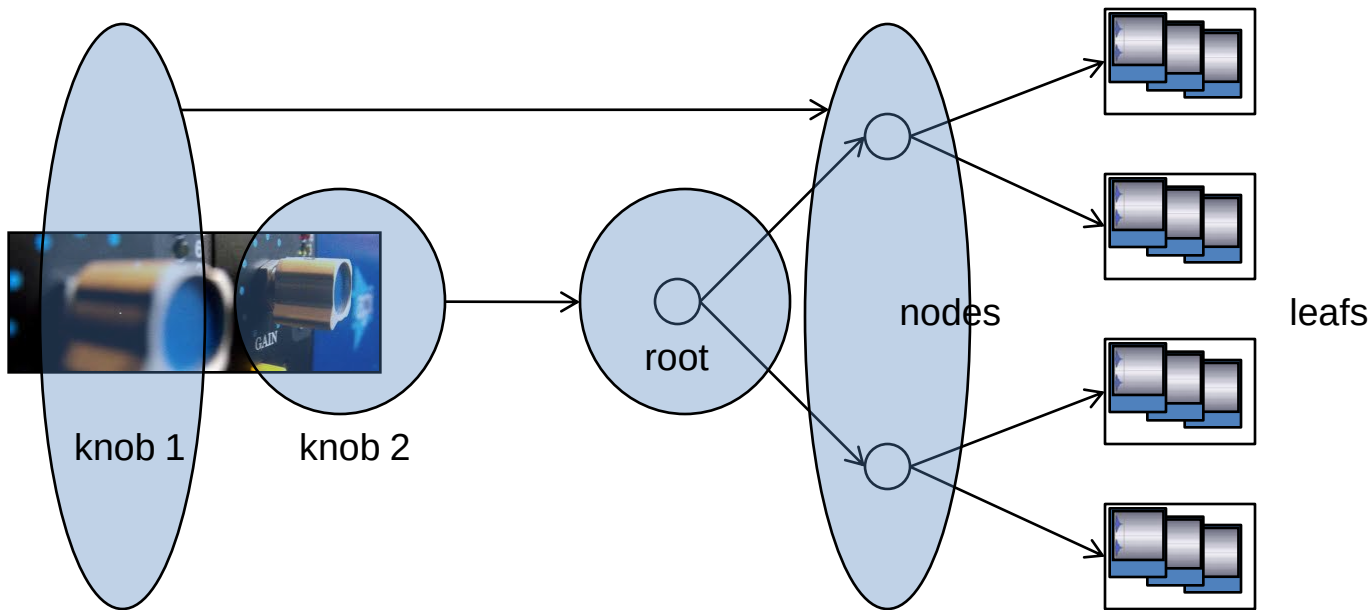


Solution for dynamic processing (comp/limiters)

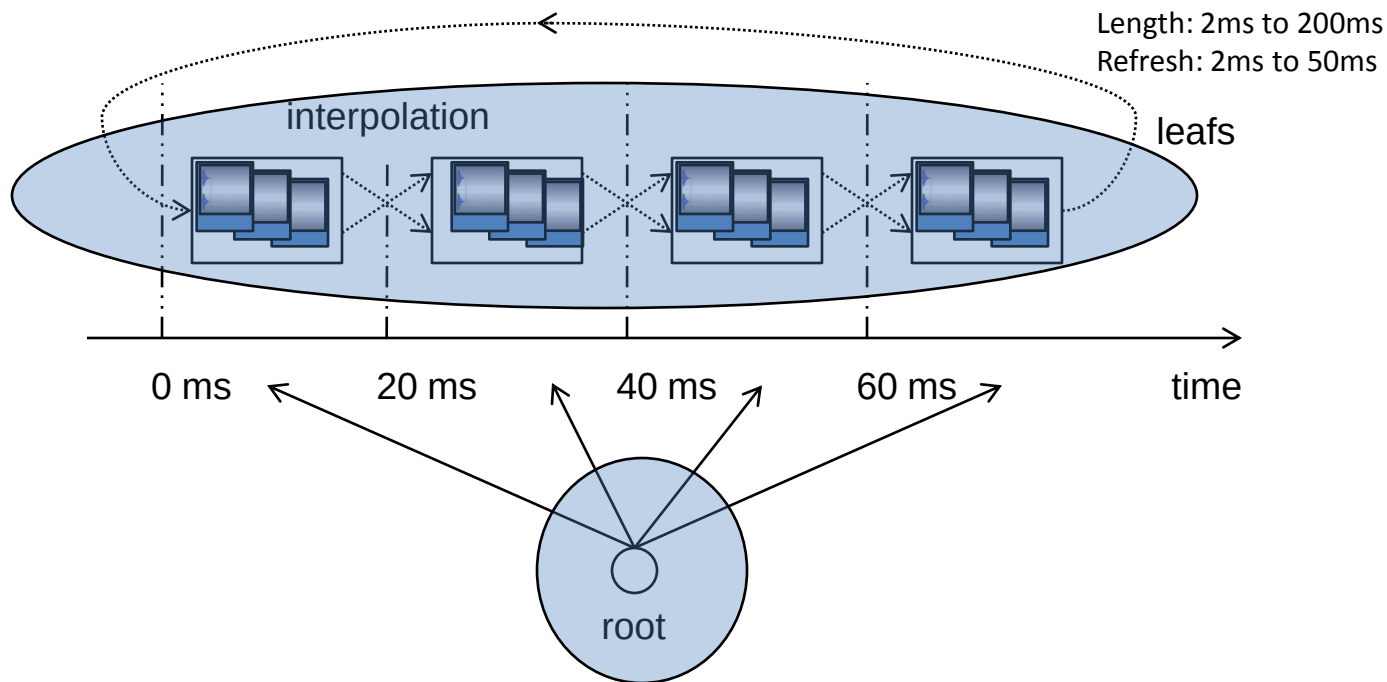


Solution for equalizers/filters

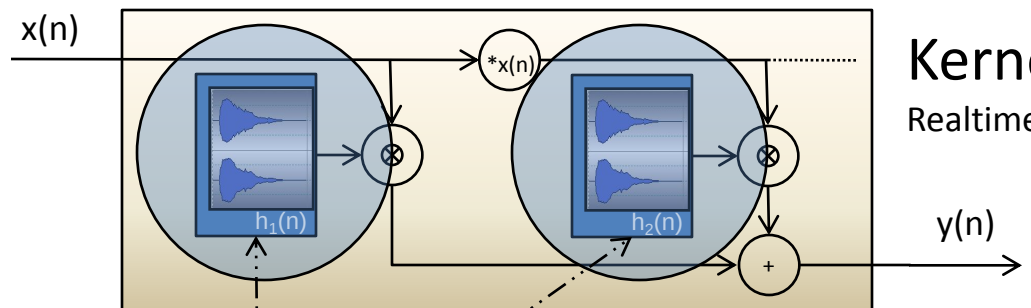
Length: 10ms to 500ms



Solution for time-varying fx (chorus, flanger, ...)



Block diagram of the engine (1)

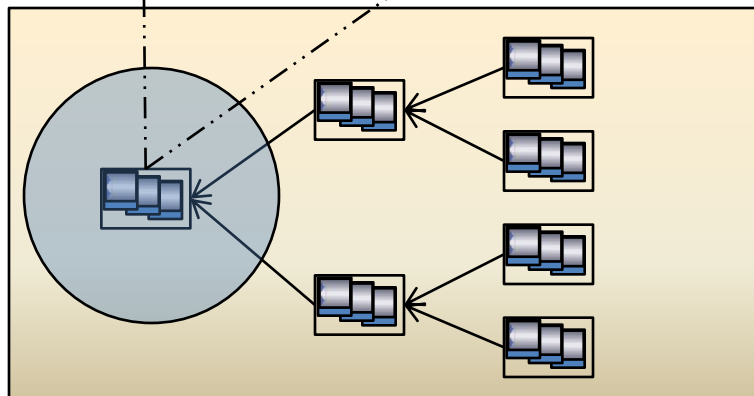


Kernel engine

Realtime

Advantages:

- low costs required for achieving accurate algorithms (sampling approach)
- easy CUDA implementation for Kernel engine

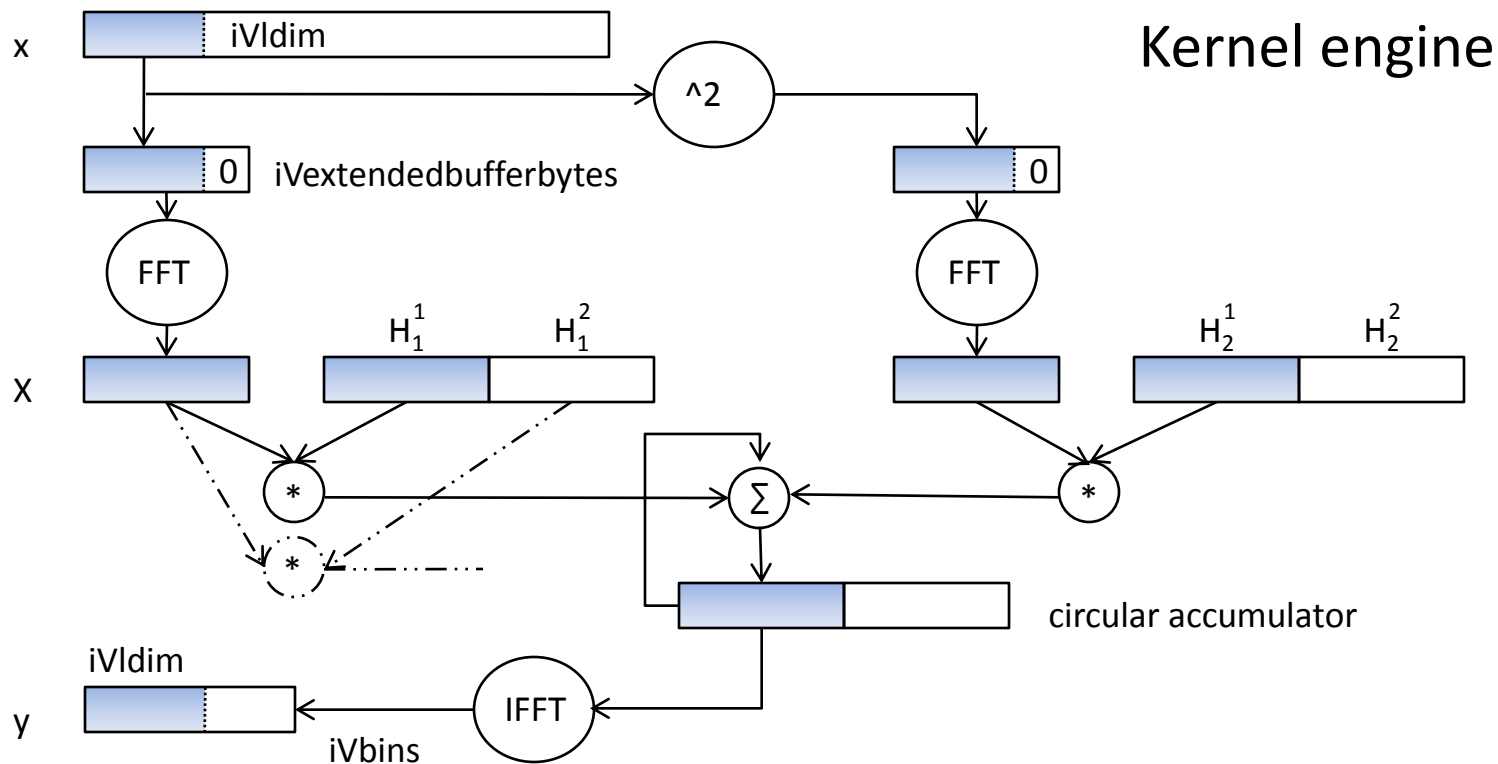


Vectorial engine

Refresh: 2ms to 180ms



Block diagram of the engine (2)



Basic engine solution



Inner block (1)

Input

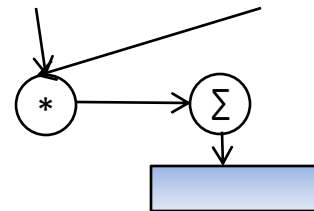
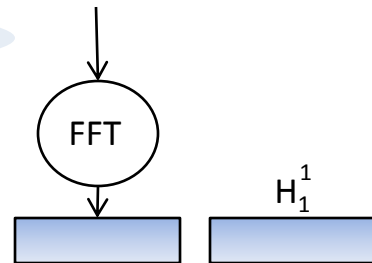
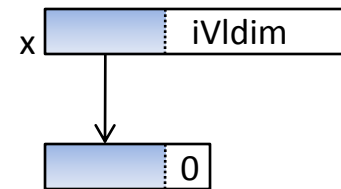
```
memset (cVtemp, 0, cOp32->iVextendedbufferbytes);
for (iV = 0; iV < cOp32->iVldim; iV++) cVtemp [iV].x = (float) dX_ [iV];
cudaMemcpyAsync ( cVtemp_, cVtemp, cOp32->iVextendedbufferbytes, cudaMemcpyHostToDevice,
                  cOp32->iVstream);
```

FFT (input and kernel)

```
if (!bUpdate) acfftExecC2C (cOp32->cVacplanx, cVtemp_, cVtemp_, CUFFT_FORWARD, cOp32->iVstream);
else {
    ...
    cudaMemcpyAsync ( cOp32->fVirfftpartition_ [iIrchannel], cOp32->fVirfftpartition
                     [iIrchannel], cOp32->iVcircularbufferbytes, cudaMemcpyHostToDevice,
                     cOp32->iVstream);
    acfftExecC2C ( cOp32->cVacplany [iIrchannel],
                  (cufftComplex*) cOp32->fVcoalescedirfftpartition_ [iIrchannel],
                  (cufftComplex*) cOp32->fVcoalescedirfftpartition_ [iIrchannel],
                  CUFFT_FORWARD, cOp32->iVstream);
}
```

Multiply and add

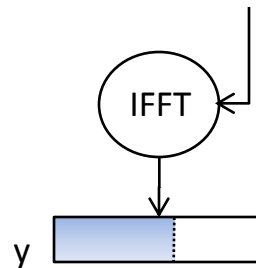
```
muacp<<<iCpar0, iCpar1, 0, cOp32->iVstream>>> (
    cVtemppointera_, (tCcomplex*) fVtemppointerb_, (tCcomplex*) fVtemppointerc_,
    cOp32->iVbins, cOp32->iVpartitions, cOp32->iVcircularpointer);
```



Inner block (2)

Output

```
switch (cOp32->iVchainedposition) {  
    case iCchainedfirst_:  
        cudaMemcpyAsync (&cOp32->fVchainedbuffer_ [iVtemp2],  
                        fVtemppointerc_, cOp32->iVextendedbufferbytes,  
                        cudaMemcpyDeviceToDevice, cOp32->iVstream);  
  
        ...  
        break;  
    case iCchainedmiddle:  
        addf1<<<iCpar0, iCpar1, 0, cOp32->iVstream>>> (fVtemppointerc_,  
                                                    &cOp32->fVchainedbuffer_ [iVtemp2],  
                                                    cOp32->iVextendedsection);  
  
        ...  
        break;  
    case iCchainedlast_:  
        addf1<<<iCpar0, iCpar1, 0, cOp32->iVstream>>> (&cOp32->fVchainedbuffer_ [iVtemp2],  
                                                    fVtemppointerc_,  
                                                    cOp32->iVextendedsection);  
  
    case iCchainednone_:  
        acfftExecC2C (cOp32->cVacplanxinv, (cufftComplex*) fVtemppointerc_, cVtemp_,  
                     CUFFT_INVERSE, cOp32->iVstream);  
  
        ...  
        cudaMemcpyAsync (cOp32->fVextendedbuffer, cVtemp_, cOp32->iVextendedbufferbytes,  
                        cudaMemcpyDeviceToHost, cOp32->iVstream);  
}
```



Inner block (3)

In-depth analysis: multiply and add

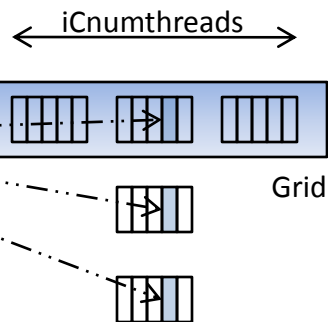
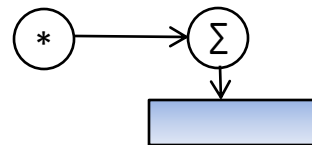
```
muacp<<<...>>> (...);
```

```
static __global__ void muacp ( tComplex* cA, tComplex* cB, tComplex* cC, int iSize,
                             int iPartitions, int iCircularpointer) {
    const int iNumthreads = blockDim.x * gridDim.x;
    const int iThreadid = blockIdx.x * blockDim.x + threadIdx.x;
    int iVtemp;

    for (int iV = iThreadid; iV < iSize; iV += iNumthreads)
        for (int iVinner = 0; iVinner < iPartitions; iVinner++) {
            iVtemp = iSize * ((iCircularpointer + iVinner) % iPartitions) + iV;
            cC [iVtemp] = masce (cA [iV], cB [iSize * iVinner + iV], cC [iVtemp]);
        }
}
```

```
static __device__ __host__ inline tComplex masce (tComplex cA, tComplex cB, tComplex cC) {
    tComplex cVreturn;

    cVreturn.x = cC.x + cA.x*cB.x - cA.y*cB.y;
    cVreturn.y = cC.y + cA.x*cB.y + cA.y*cB.x;
    return cVreturn;
}
```



Inner block (4)

In-depth analysis: FFT

acfftExecC2C (...)

```
void acfftExecC2C ( ACFFFTLOCHANDLE cPlan, float2* cInput, float2* cOutput,
                    int iDirection, int iStream = 0) {

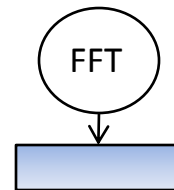
    ((ACFFFTLOC*) cPlan)->cVexecuteFun [iDirection <= 0 ? 0 : 1] (cInput, cOutput, cPlan, iStream);
}
```

```
(*cPlan)->cVexecuteFun [0] = cVacfftFunfw [...];
```

```
int iVacfftSize [] = {8, 16, 32, 64, 128, 256, 512, 1024, 2048, 4096, 8192, 16384};
ACFFFTFUN cVacfftFunfw [] =
{acfnv, acfnv, acfnv, acfnv, acfnv, acfnv, acf01, acf02, acf03, acf04, acf05, acf06};
```

```
void acf06 ( float2* cInput, float2* cOutput, ACFFFTLOCHANDLE cAcfft, int iStream = 0) {
    ACFFFTLOC* cVloc = (ACFFFTLOC*) cAcfft;

    v2_16<<< grid2D(cVloc->iVbatches*(16384/16)/64), 64, 0, iStream >>>(cInput, cOutput);
    v1024<<< grid2D(cVloc->iVbatches*16), 64, 0, iStream >>>(cOutput, cOutput);
}
```

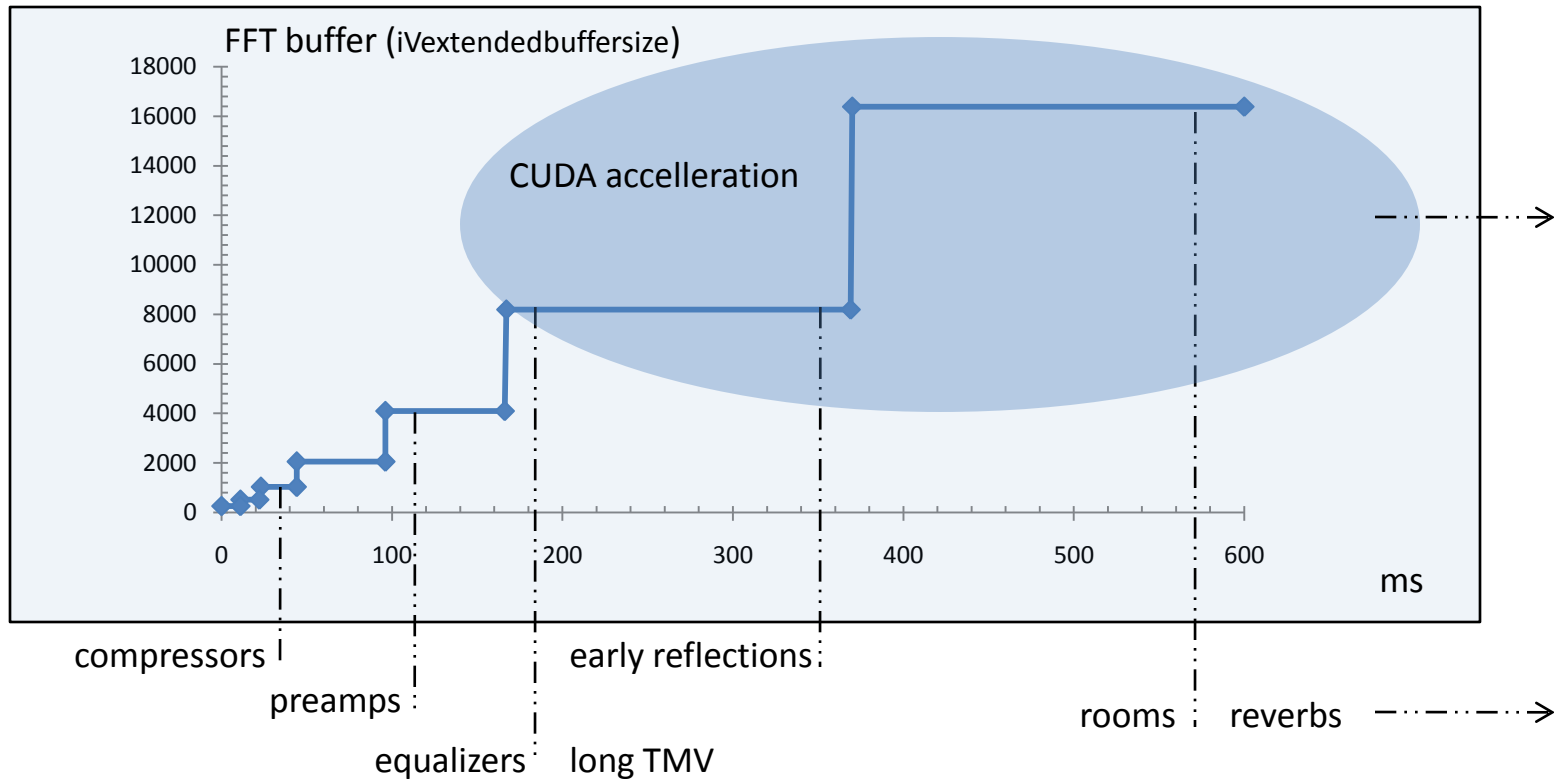


custom FFT:
acfnv = basic c2c radix2

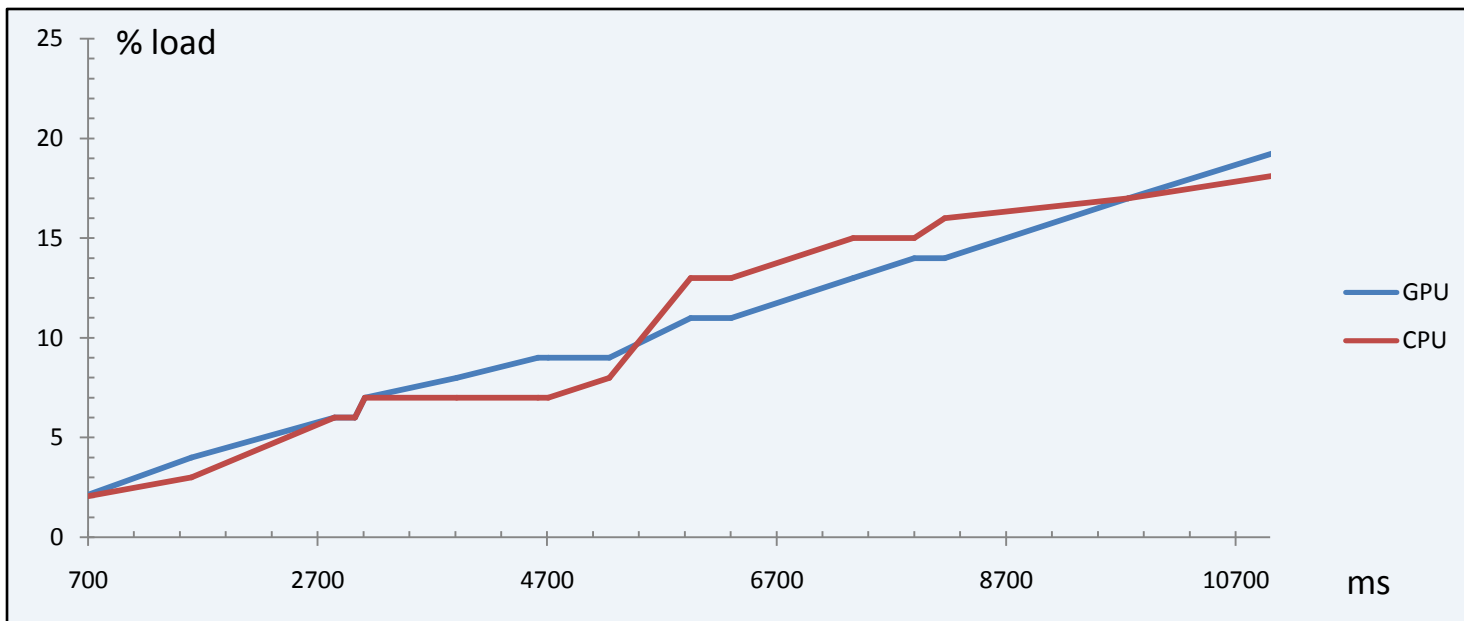
Vasily Volkov inner blocks:
1024 = 16*4*16
16384 = 16*1024



Buffer tuning and latency (1)



Buffer tuning and latency (2)



CPU: Core2 Duo T5800 2.0 Ghz

GPU: GeForce 9600M GT

Results:

- real-time resources are doubled adding GPU-powered instances



Buffer tuning and latency (3)

CPU: Core2 Duo T5800 2.0 Ghz

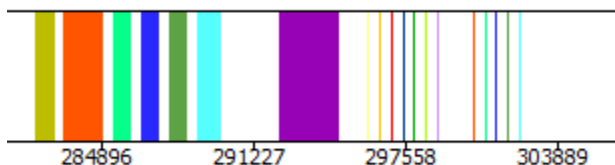
3 instances, 1 kernel, 5.5 seconds, FFT buf = 16384, 30 partitions

GPU: GeForce 9600M GT

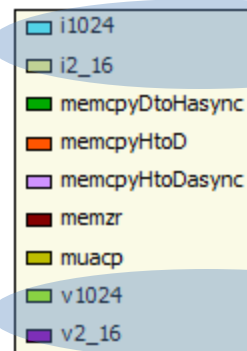
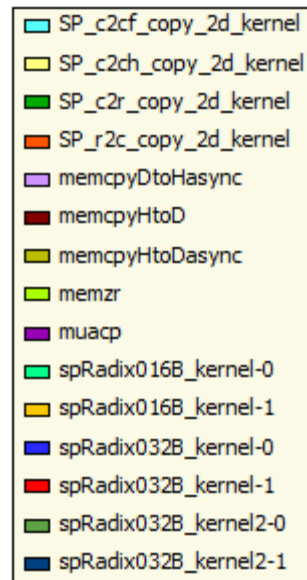
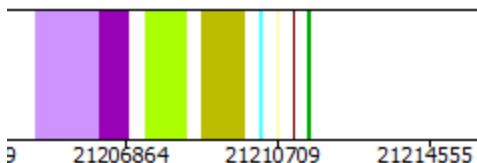
IPP (Intel Performance Primitives): CPU 40%

FFTW: CPU 44%

CUDA – cufft library: CPU 5%, GPU 90%



CUDA – custom fft implementation: CPU 5%, GPU 36%

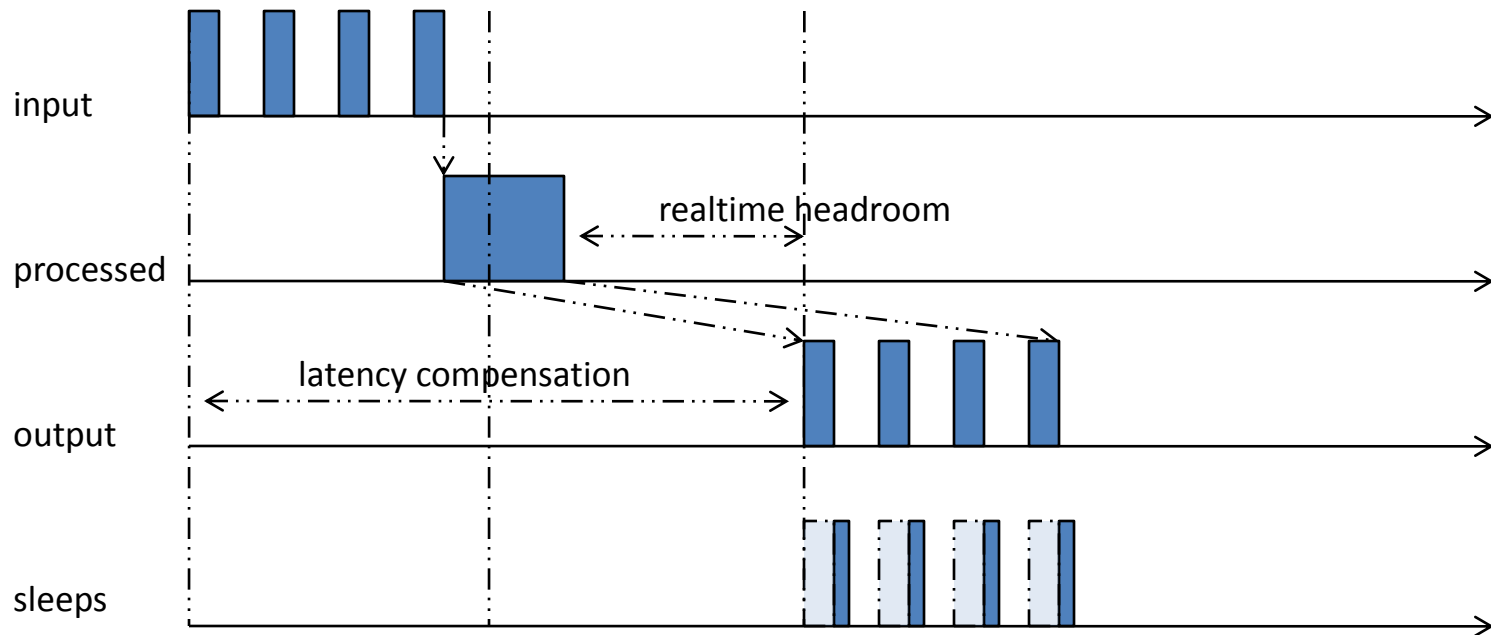


IFFT

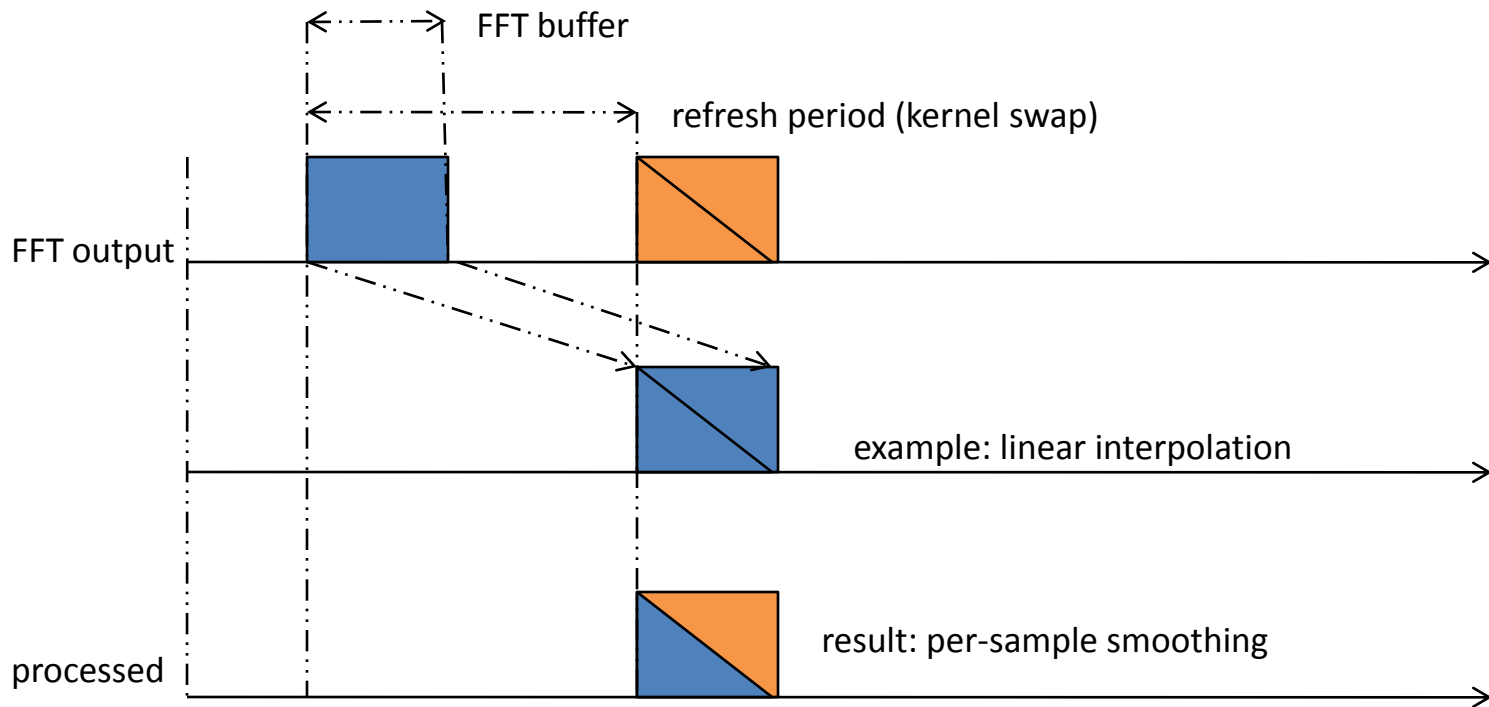
FFT



Audio processing buffers (1)

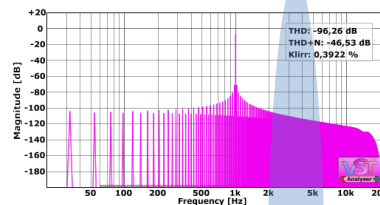
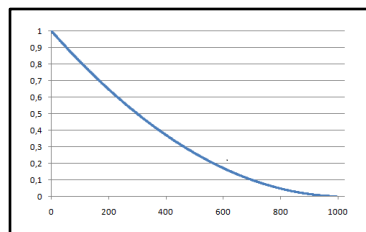
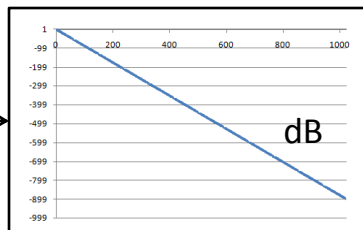
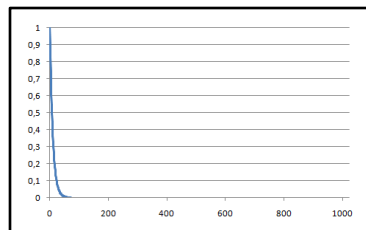


Audio processing buffers (2)

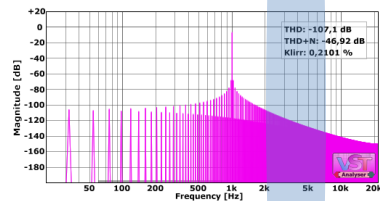


Audio processing buffers (3)

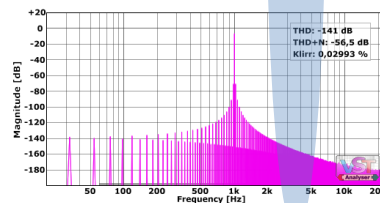
Example: preamplifier featuring 50 ms kernels



Without interpolation



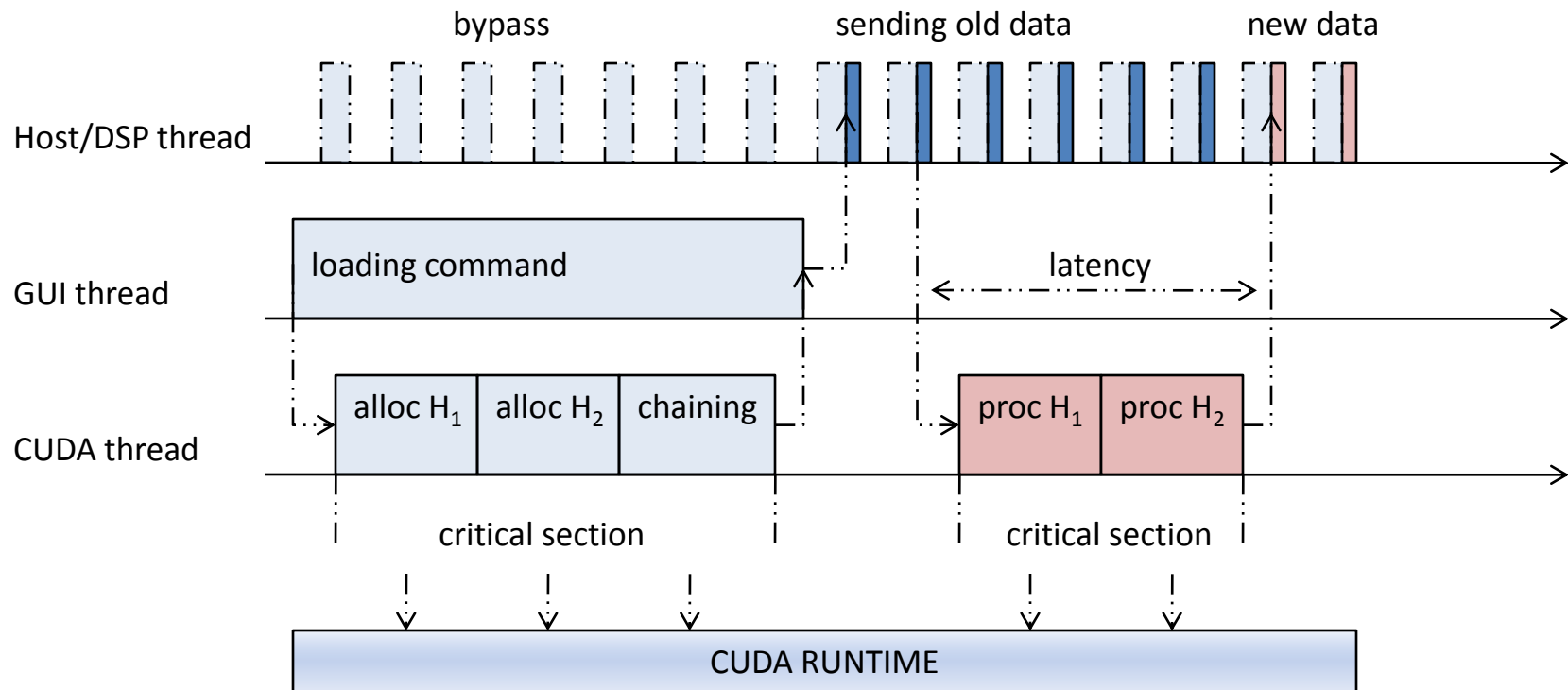
Log interpolation



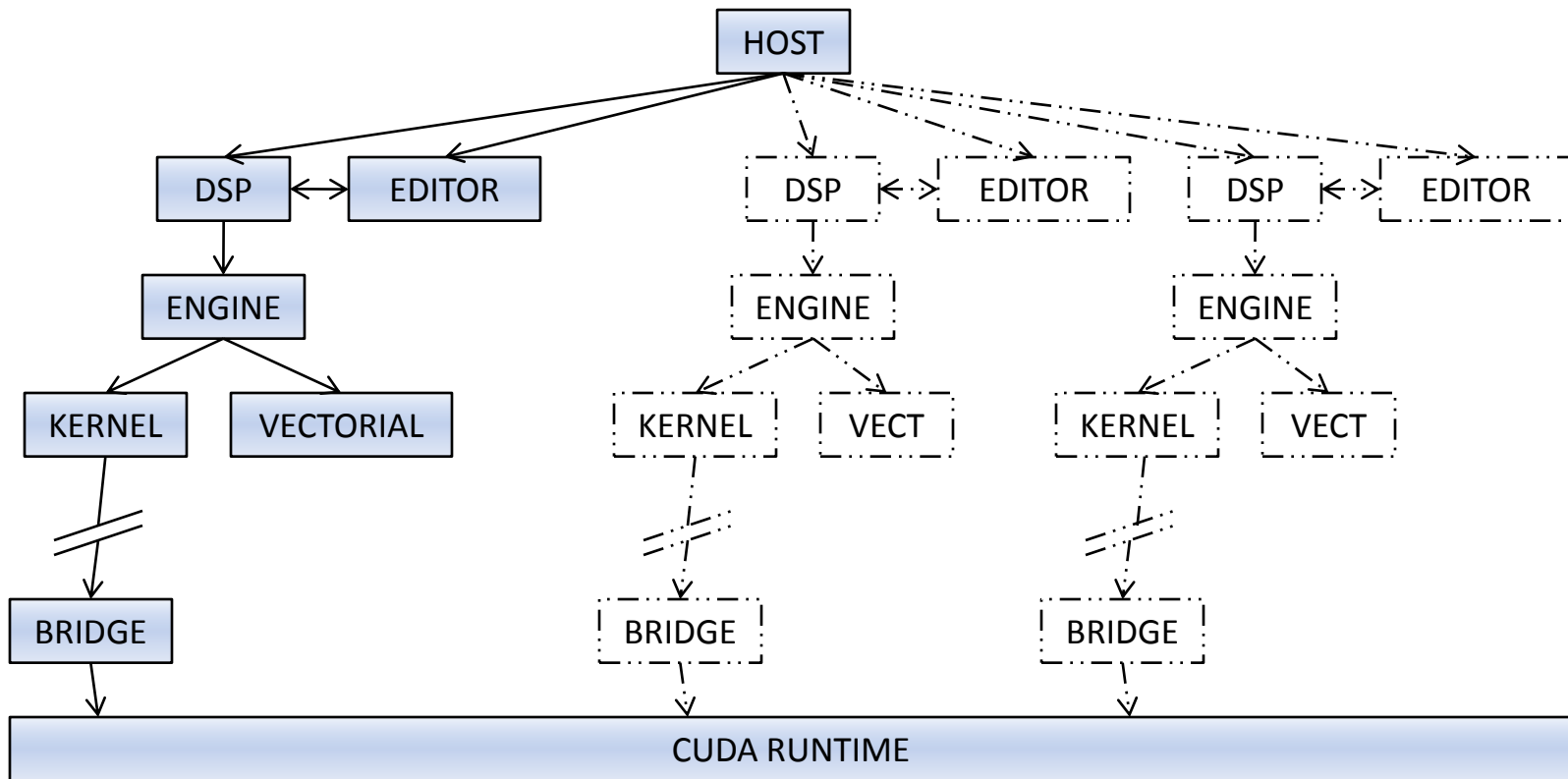
Custom interpolation



Threading model



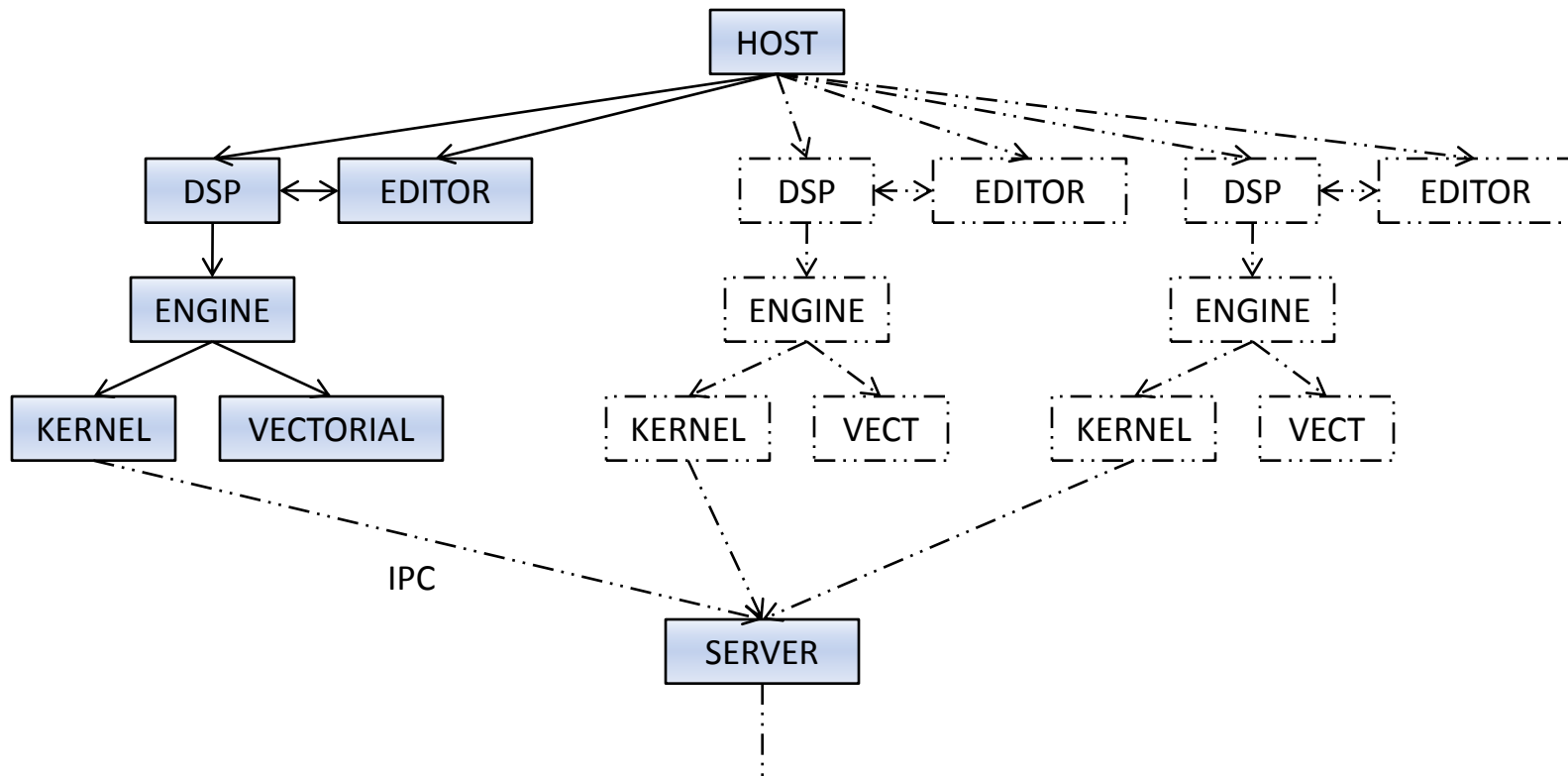
Block diagram of a plug-in implementation



Network implementation



Network implementation

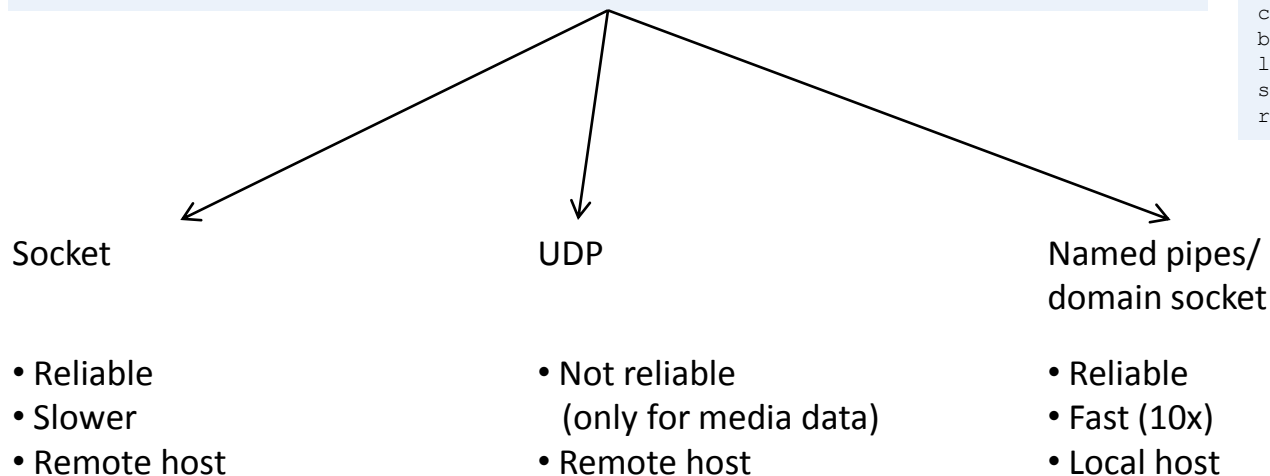


Inter-process communication (1)

Custom socket implementation

```
cVsocket = new CoreSckt (type);
```

```
accept  
connect  
bind  
listen  
send  
recv
```



Inter-process communication (2)

Example: set a 2 state variable bean

```
cVprocess->setbb(CorePrCs::iCbeanbypass, true);
```

Client

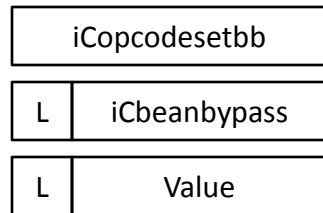
```
cVtransportopcode->setvl (0, iCopcodesetbb);
cVtransportopcode->send_ (cVlocalsocket);
cVbeanstructb->setvl (0, iObject);
cVbeanstructb->setvl (1, bValue);
cVbeanstructb->send_ (cVlocalsocket);
cVlocalsocket->lflsh ();
cVtransportack->recv_ (cVlocalsocket);
```

Server

```
case iCopcodesetbb:
    cVbeanstructb->recv_ (cVlocalsocket);
    cVbeanstructb->getvl (0, iVsubopcode);
    cVbeanstructb->getvl (1, bVtemp);
    cVbeanstructb->getvl (2, cVbeanstruct.iVobject);
    cVbeanstructb->getvl (3, cVbeanstruct.iVsubobject);
    cVprocess->setbb
        (iVsubopcode, bVtemp, &cVbeanstruct);
    cVtransportack->send_ (cVlocalsocket);
```

class is wrapped

message



opcode

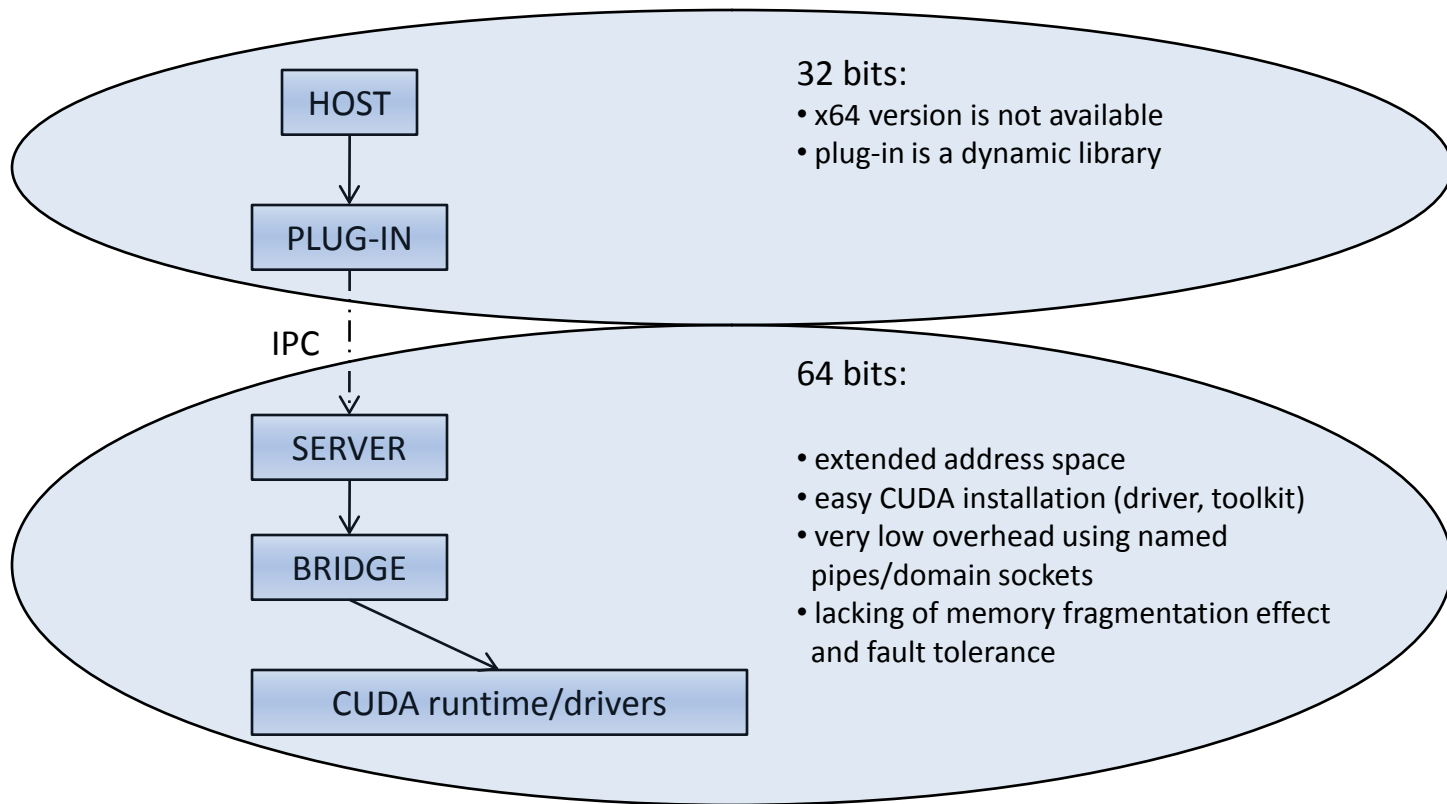
subopcode/bean item

Different threads on server side:

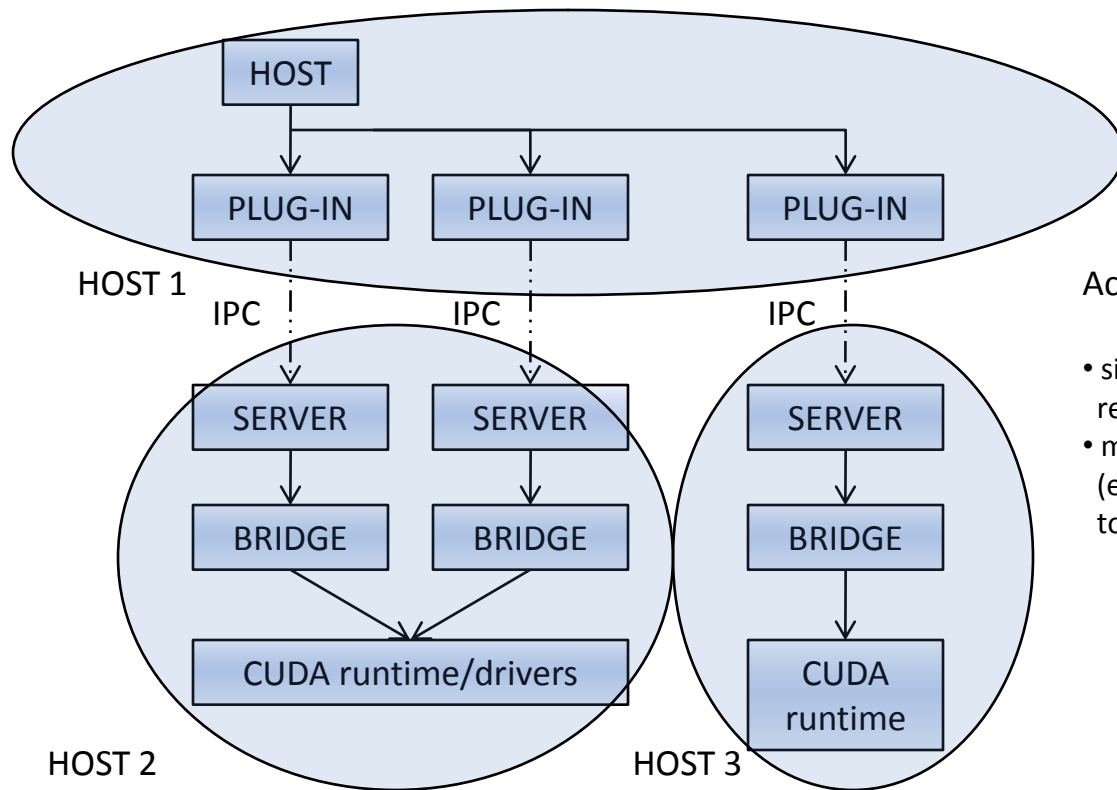
- Audio processing
- Editor: low priority beans



32 bits legacy hosts in 64 bits OS



Multi-servers/Multi-clients block diagram



Achievements:

- single control point, shared resources, easy maintenance
- more OS are available at the same time (example: MAC OSX host connected to cheap windows servers)



Results

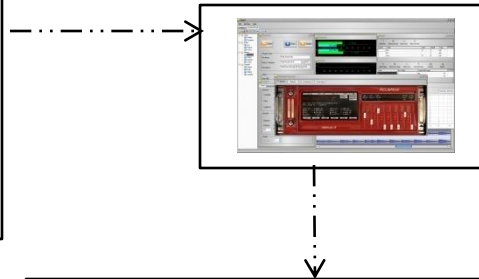


The library

- MAC OSX Tiger..Snow Leopard, Windows XP... 7
- Support for Microsoft, Borland/CodeGear, Apple compilers
- Support for IPP library
- Optional support for FFTW library
- Direct convolution implemented in ASM ia32 and x64
- Automatic application aimed to vintage gear sampling
- Huge number of emulations available and used daily in professional audio productions

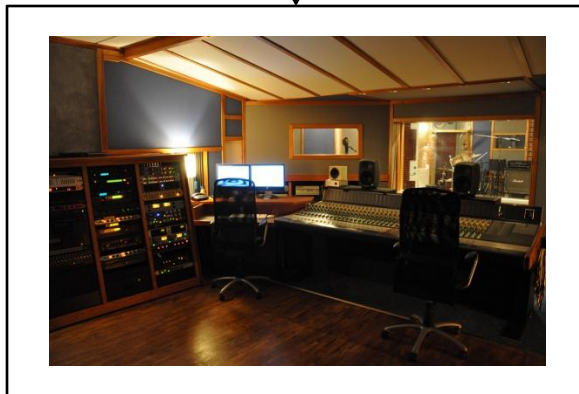


On the road



sampling

- Big recording facility
- Client/Server as replacement of missing outboard



Next steps

- CUDA implementation of the vectorial engine
- Resource optimization (clients connected to the same server instance could share resources)
- New features based on client-to-client communication (ie cross-talk modeling)

