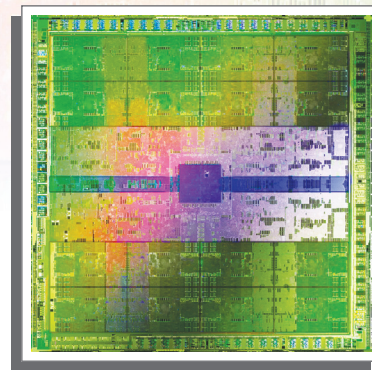


GPU Computing

از سیر تا پیاز



هومن سیاری
Sayyari@ComputerNews.ir

موتورهای پردازش موازی داخل چیپ‌های گرافیکی آنها را برای انجام طیف وسیعی از پردازش‌های متنوع آماده کرده‌اند ولی افسوس! کجایند برنامه‌هایی که بتوانند از این همه قدرت بالقوه استفاده کنند؟

ما می‌خواهیم به هر یک از سازندگان اصلی سخت‌افزار یا API‌ها نگاهی بیندازیم تا ببینیم آنها برای پشتیبانی از پردازش گرافیکی چه کرده‌اند؟

بررسی سخت‌افزارها

اگر بخواهیم به جستجوی سازندگان چیپ‌های گرافیکی که از GPGPU پشتیبانی می‌کنند بپردازیم خیلی زود به جواب می‌رسیم. پاسخ کاملاً واضح است فقط ۲ شرکت AMD و NVIDIA چنین چیپ‌های گرافیکی تولید می‌کنند. اجازه بدهید هر یک از آنها را به صورت مجزا مورد بررسی قرار دهیم:

CUDA و TESLA : nVIDIA

اولین تلاش برای دستیابی به GPGPU توسط nVIDIA انجام شد. شواهدی وجود دارد که نشان می‌دهد این شرکت با چیپ‌های گرافیکی اولیه سری GeForce 256 که حتی فاقد قابلیت سایه‌زنی قابل برنامه‌ریزی (Programmable shaders) بودند هم سعی در پیاده‌سازی GPGPU داشته است. البته این تلاش‌ها زمانی نتیجه داد که DirectX 9 با قدرت انعطاف‌پذیری به مراتب بالاتر در معماری سایه‌زنی قابل برنامه‌ریزی عرضه شد. مدل‌های DirectX Shader جزئیات بسیاری را جهت کنترل بر روی آنچه نمایش داده می‌شود برای افراد خواستار توسعه، ارائه می‌دهند و می‌توانند افکت‌های بسیاری نظیر سایه‌های پیچیده انعکاس نور، ایجاد مه و مانند آنها را بوجود آورند. مایکروسافت با ارائه Shader Model از ابتدا قدرت مانور زیادی را در اختیار توسعه دهندگان قرار داده است که نتیجه آن ایجاد پتانسیل بالا جهت واقعی‌تر و ملموس‌تر کردن تصاویر بوده است.

Shader Model II نسخه‌ای است که DirectX 9 در اختیار کاربران قرار می‌دهد. برای به تصویر کشیدن صحنه‌های واقعی و باور نکردنی در گیم‌ها و خلق سایه‌های درست و دقیق استفاده می‌شود. فعل و انفعالات پیچیده بین منابع نور، اشیاء و نشانه‌های گرافیکی شامل برنامه‌هایی می‌شود که جزئیات را با دقت بیشتری بررسی می‌کنند و در واقع در یک بازی، تمامی منابع نوری مرتبط به هر شیء باید مورد تجزیه و تحلیل قرار بگیرند. این بدین معناست که در سیستم‌هایی که از پردازنده‌های گرافیکی استفاده می‌کنند هر زمان که یک بازی یا یک برنامه کاربردی به بازسازی سایه‌ها بپردازد، Ultra Shadow II عملکرد نهایی

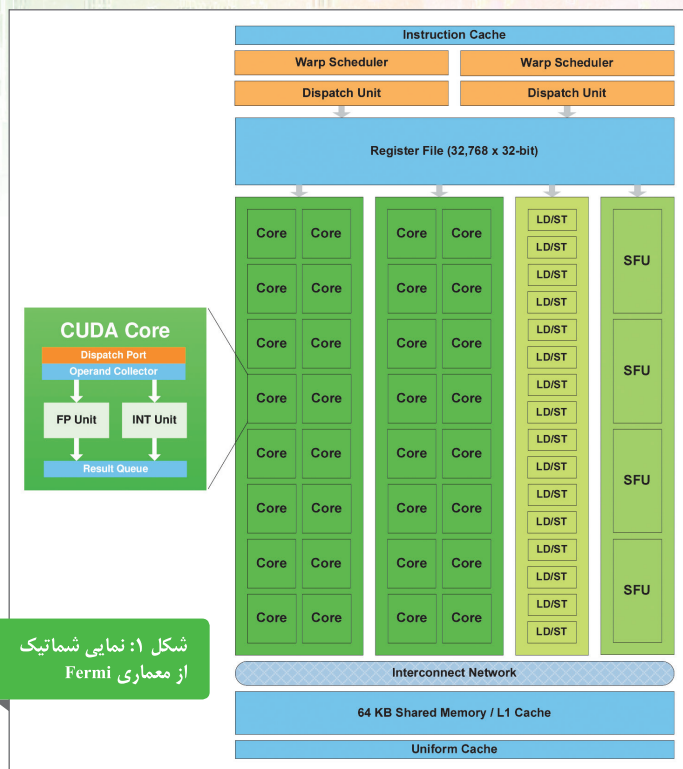
تا چند سال پیش که صنعت گیم‌های کامپیوتری چندان رشد نکرده بود CPU مجبور بود که بار سنگین گیم‌های آن زمان را بر دوش بکشد. چیپ‌های گرافیکی فقط کارهای ساده و اولیه‌ای مثل عملیات ساخت تصویر ارسالی به مانیتور، سوییچ بین حالت متنی و گرافیکی و یا سوییچ بین حالت نمایش ۱۶ رنگ یا ۲۵۶ رنگ را انجام می‌دادند. در آن زمان حتی سیستم‌عامل‌ها هم نیازی به گرافیک نداشتند چرا که متداول‌ترین سیستم‌عامل آن زمان یعنی DOS هم در یک محیط متنی و بدون نیاز به هیچ پردازش گرافیکی اجرا می‌شد.

بعد از ظهور سیستم‌عامل ویندوز، چیپ‌های گرافیکی کمی رشد کردند چرا که این سیستم‌عامل نیاز به پردازش گرافیکی بالاتری داشت و خودش هم در یک محیط گرافیکی نمایش داده می‌شد. در این زمان بود که بحث استفاده از Windows Accelerator برای رسیدن به راندمان بالاتر گرافیکی مطرح شد. سپس کارت‌های گرافیکی متنوعی که از قابلیت فوق پشتیبانی می‌کردند وارد بازار شدند از آن جمله می‌توان به 3dfx Voodoo و Rendition V1000 اشاره کرد. اما امروزه شرایط عوض شده است.

چیپ‌های گرافیکی می‌توانند پردازش‌های سنگینی را به صورت موازی اجرا کنند و با محاسبات میزبانی‌شده با دقت مضاعف را انجام دهند. پردازش گرافیکی به چیپ‌های ۳ بعدی جدید امروزی اجازه می‌دهد که فقط به کارهای گرافیکی فکر نکنند بلکه آماده انجام هر کار دیگری هم باشند!

در این مقاله، منظور از **پردازش گرافیکی** انجام هر گونه پردازش اعم از گرافیکی و غیر گرافیکی توسط چیپ گرافیکی (GPU) و نه پردازنده اصلی (CPU) است. ریشه پردازش گرافیکی به یک تحقیق آکادمیک به نام **GPGPU** که مخفف **General Purpose computing on Graphics Processing Units** است برمی‌گردد. در واقع GPGPU به مفهوم انجام هر نوع پردازش (و نه فقط پردازش‌های گرافیکی) توسط چیپ‌های گرافیکی است. اوایل، GPGPU با مشکلات زیادی دست و پنجه نرم می‌کرد چرا که چیپ‌های گرافیکی برای انجام عملیات میزبانی‌شده با DirectX پایین‌تر از نسخه ۹، مشکل داشتند اما در دوران DirectX 11 این فناوری جان تازه‌ای گرفته است و آماده هر گونه هنرنمایی است!

نسل جدید پردازش گرافیکی خیلی گسترده‌تر شده است. ویژگی DirectCompute تعبیه‌شده در DirectX 11 تقریباً تمامی کارت‌های مبتنی بر DirectX 11 را پشتیبانی می‌کند. OpenCL هم از پلتفرم چند سیستم‌عاملی پشتیبانی می‌نماید.



کار را ارتقا می‌بخشد. با استفاده از این تکنولوژی هر چه نیاز بیشتری به نورپردازی و محاسبه سایه‌ها در یک صفحه داشته باشیم بهبود قابل توجهی در نوع عملکرد حاصل خواهد شد و در حقیقت هر چه قدر ترکیبات پیچیده‌تر باشد، نتایج قابل ملاحظه و برجسته‌تری به دست خواهد آمد. در مجموع این تکنولوژی نسبت به نسخه قبلی می‌تواند میزان کارایی سایه‌زنی را تا چهار برابر افزایش دهد. یکی دیگر از مزایای این تکنولوژی این است که به نویسندگان نرم‌افزارهای گرافیکی این امکان را می‌دهد که سایه‌ها را با سرعت بسیار بالاتری از طریق حذف نواحی غیر ضروری بازسازی نمایند. در واقع با مشخص کردن پهنای باند یک صفحه و تمرکز بر روی نواحی که بیشتر تحت تاثیر منابع نوری هستند، نرم‌افزارنویسان می‌توانند فرآیند پردازش سایه‌ها را به طور عمده‌ای تسریع نمایند. با به کارگیری سایه‌هایی که امکان تنظیمات پارامتری را برای افزایش کارایی دارند، نویسندگان نرم‌افزارهای گرافیکی می‌توانند تصاویر گرافیکی باور نکردنی که به واقعیت بسیار نزدیک می‌باشند را خلق نمایند.

ان‌ویدیا از همان اول متوجه شده بود که GPUها پتانسیل انجام پردازش‌های سنگین را دارند. بنابراین مهندسان آن به این فکر افتادند که چگونه می‌توانند از GPU برای پردازش‌های غیر گرافیکی استفاده کنند. تا آن موقع کارت‌های گرافیک برای انجام امور گرافیکی مناسب بودند اما ساختن برنامه‌های غیر گرافیکی که بخواهند از قدرت کارت گرافیک برای پردازش مورد نیاز خود استفاده کنند مشکل بود. در برنامه‌نویسی کارت‌های گرافیکی امکان استفاده از برخی از قابلیت‌های رایج برنامه‌نویسی مثل حلقه‌ها و بازگشت‌ها میسر نبود.

مشکلات دیگر آن بود که DirectX 9 بخشی از معماری چیپ گرافیکی را قفل کرده بود این بدان معنا است که پیاده‌سازی ویژگی‌هایی که برای انجام کارهای غیر گرافیکی لازم بود خارج از استاندارد DirectX شناخته می‌شد و پشتیبانی نمی‌شد. با وجود این‌ها به نظر می‌رسید که چیپ‌های گرافیکی کاندیدای مناسبی برای انجام پردازش‌هایی نظیر SIMD (اجرای یک دستورالعمل بر روی چند داده) و عملیات ممیز شناور محض باشند. ان‌ویدیا برای پیشرفت هر چه بیشتر GPGPU یک فریم‌ورک نرم‌افزاری جدید و کاربرپسند به نام CUDA 1.0 را طراحی کرد. **CUDA** مخفف عبارت Compute Unified Device Architecture است. حالا برنامه‌نویسان قادر بودند که از زبان برنامه‌نویسی C در کنار امکاناتی که ان‌ویدیا برایشان بر اساس این زبان برنامه‌نویسی فراهم کرده بود استفاده کنند و برنامه‌های خود را بسازند. این روش خیلی راحت‌تر از روش برنامه‌نویسی بر اساس روش قبلی یعنی برنامه‌نویسی سایه‌ها است. به عبارت دیگر در این روش جدید برای نوشتن برنامه‌های عمومی نیازی به استفاده از کدهای گرافیکی نبود. CUDA با چیپ‌های GeForce 8800 GTX و به بعد کار می‌کند. اما اولین نسل از چیپ‌های گرافیکی که اختصاصاً بر اساس پردازش گرافیکی ساخته شدند Tesla 870 بودند.

ان‌ویدیا از زمان چیپ‌های ۸۸۰۰ سعی کرد که چیپ‌هایی بسازد که قابلیت GPGPU آنها کاربردی‌تر و بهتر باشد، هر چند به دنبال آن نبود که GPU را جایگزین CPU کند. CPUها هنوز در پردازش خطی یا چند رشته‌ای برتری محسوسی داشتند هر چند GPUها این پتانسیل را دارند که در پردازش‌های حجیم که به صورت موازی بر روی صدها رشته عملیاتی (Thread) که بر روی داده‌های جداگانه ولی هم نوع انجام می‌شوند بازدهی بسیار بالایی را از خود نشان دهند. این روش برنامه‌نویسی بسیار مناسب کاربردهایی مثل محاسبات علمی و یا تحلیل‌های مالی است.

شاید این موضوع خیلی معنی‌دار باشد که ان‌ویدیا آخرین معماری خود معروف به Fermi را به عنوان یک پلتفرم پردازش توسط گرافیک (GPU compute platform) معرفی کرده است و نه یک پردازنده گرافیکی (Graphics processor)!

معماری Fermi، پیشرفت‌های سخت‌افزاری قابل توجهی را به ارمان آورد که در نتیجه این چیپ را برای پردازش‌های عمومی و نه صرفاً گرافیکی بسیار ایده‌آل کرده است. از جمله این پیشرفت‌ها می‌توان به Unified Memory Architecture، Atomic Memory Operations و Context Switching بهتر و موارد دیگر اشاره کرد. از زمانی که Fermi معرفی شد ان‌ویدیا چندین بار پلتفرم نرم‌افزاری CUDA را آپدیت کرده است.

ان‌ویدیا به پردازش گرافیکی فقط به دید ابزاری برای انجام محاسبات مثلاً استخراج نفت و یا پردازش‌های آکادمیک نگاه نمی‌کرد بلکه به دنبال آن بود که این قابلیت را کاربردی‌تر و عمومی‌تر نماید. مثلاً ان‌ویدیا قابلیت PhysX را چندین سال پیش عرضه کرده بود هر چند از یک سخت‌افزار اختصاصی برای پیاده‌سازی آن استفاده نمی‌کرد ولی تلاش داشت که با APIهای طراحی شده بتواند آن را توسط GPU پیاده‌سازی کند. ان‌ویدیا همواره با شرکت‌های سازنده گیم‌های کامپیوتری کار می‌کرد و آنها را به استفاده از قابلیت‌های پردازش گرافیکی در گیم‌های جدیدشان و در کاربردهایی مثل شبیه‌سازی آب، افکت‌های لنزهای نوری، انفجار و ... تشویق می‌کرد. این شرکت همچنین با شرکت‌های مهمی مثل Adobe، ArcSoft، CyberLink و ... همکاری داشته است تا آنها قابلیت تبدیل فرمت‌ها را با استفاده از پردازش گرافیکی در محصولاتشان بگنجانند.

ان‌ویدیا توانست با Tesla که بر پایه معماری Fermi طراحی شده حدود ۱۰۰ میلیون دلار در سال ۲۰۱۱ فروش کند.

AMD با قدرت در حرکت

AMD کار بر روی پردازش گرافیکی را با تاخیر شروع کرد ولی با شتاب بالا ادامه می‌دهد. ATI Stream نامی بود که این شرکت در مقابل CUDA انتخاب کرده بود. اولین کارت گرافیکی که اختصاصاً برای پردازش گرافیکی طراحی شد FireStream 580s بود. البته برای اولین بار در سری Radeon کارت Radeon X1900 هم از GPGPU پشتیبانی می‌کرد. البته واقعیت آن است که تا زمانی که کارت‌های سری Radeon HD 4000 عرضه شدند عملاً AMD حرف چندان در زمینه

شبیه سازی کند که ساختار این برنامه هیچ ارتباطی به گرافیک نداشت. در اینجا بود که یک نیاز اساسی به شدت احساس می شد و آن استانداردسازی رابط برنامه نویسی برای کارت های گرافیکی بود. البته این اتفاق در تاریخ کارت های گرافیکی تکرار شده بود یعنی ابتدا یک تکنولوژی معرفی می شد و سپس استانداردهای آن ارایه می شدند.

CUDA

CUDA از شرکت انویدیا یکی از اولین تلاش ها برای استانداردسازی برنامه نویسی برای پردازش گرافیکی بود. انویدیا همیشه تلاش می کرد که این موضوع را جا بیندازد که CUDA یک استاندارد مختص این شرکت نیست بلکه همه شرکت ها می توانند از آن استفاده کنند ولی نه AMD و نه اینتل حاضر به استفاده از آن نشدند!

CUDA با یک کامپایلر به زبان C و یک کتابخانه از توابع مربوطه کار برنامه نویسی برای پردازش موازی توسط کارت گرافیک را آغاز کرد. بعد از چند سال چندین کامپایلر و دیباگر از شرکت های دیگر در این زمینه معرفی شد و حتی برنامه نویسی برای CUDA در مجموعه برنامه نویسی قدرتمند Microsoft Visual Studio گنجانده شد.

CUDA بیشترین موفقیتش را در بازار کامپیوترهای خانگی و سوپر کامپیوترهای آکادمیک داشته است اما CUDA اهداف بزرگتری را دنبال می کند. شرکت معروف Adobe از CUDA در برنامه Adobe Premiere Pro استفاده می کند و از آن برای تبدیل فرمت های HD و یا برای برخی از Transition ها استفاده می نماید. نرم افزار MotionDSP از CUDA برای کاهش اثر Shaky-cam در ویدیوهای خانگی استفاده می کند. بسیاری دیگر از برنامه های سطح بالا و حرفه ای هم از CUDA استفاده می کنند تا زمان پردازش های خود را به میزان چشمگیری کاهش دهند. این همان چیزی است که انویدیا به شدت به دنبال آن است تا سازندگان معتبر نرم افزار بر اساس CUDA نرم افزارهای خود را طراحی کنند تا کاربران این نرم افزارها مجبور به خرید کارت های انویدیا شوند.

ATI Stream

همانگونه که قبلا اشاره شد سابقه AMD در پردازش گرافیکی به مراتب کمتر از انویدیا است و این شرکت بیشتر بر روی OpenCL و DirectCompute متمرکز شده است. در واقع Stream نام روشی بود که AMD برای رقابت با CUDA انتخاب کرد اما بعد از مدتی AMD به این نتیجه رسید که استفاده از پلتفرم های استاندارد بسیار جذاب تر از روش اختصاصی خودش است!

DirectCompute

DirectCompute با فریم ورک DirectX 11 عرضه شد و بنابراین فقط در ویندوز ویستا و ۷ و ۸ قابل استفاده است. بنابراین بر روی ویندوزهای قدیمی تر مثل ویندوز XP و سایر سیستم عامل های غیر مایکروسافتی قابل اجرا نیست. همچنین بر روی ویندوز فون ۷ و ۸ و Xbox 360 نیز کار نمی کند. DirectCompute بر روی تمام GPU هایی که می توانند از DirectX 11 پشتیبانی کنند اجرا می شود. این GPU ها شامل سری NVidia GTX 400 و مدل های جدیدتر از آن و نیز سری AMD Radeon HD 5000 و مدل های جدیدتر از آن می گردند. اینتل هم تنها در پردازنده های Ivy Bridge از DirectX 11 پشتیبانی می کند.

مزیت کلیدی DirectCompute آن است که تنها از یک زبان سایه زنی پیشرفته به نام HLSL برای برنامه نویسی پردازش گرافیکی استفاده می کند. این موضوع باعث می شود که کار برنامه نویسان به مراتب ساده تر شود. مزیت دیگر هم این

GPGPU نداشت. سری Radeon HD 5000 پیشرفت قابل توجهی در این زمینه کرد و سپس در سری Radeon HD 6000 قابلیت پردازش گرافیک به اوج خود رسید.

اگر کمی دقیق تر نگاه کنیم متوجه می شویم که در ابتدا AMD سعی داشت که فقط از انویدیا تقلید کند بدون آنکه نوآوری داشته باشد اما بعد از مدتی مسیر خود را عوض کرد و به آغوش استانداردهای باز مثل OpenCL و DirectCompute پیوست!

اخیرا AMD مسیر خود را به سمت کاربردهای حرفه ای سوق داده است. این شرکت کارتهای اختصاصی خود را که برای پردازش گرافیکی طراحی می کند را تحت عنوان FireStream عرضه می نماید. حالا AMD قصد دارد که توانایی کارتهای FireStream را در چیپ های گرافیکی Radeon که در پردازنده های جدید این شرکت با نام Fusion گنجانده شده اند قرار دهد. Fusion نام پردازنده های جدید شرکت AMD است که یک SoC محسوب می شوند یعنی هم CPU و هم GPU در کنار هم در داخل یک چیپ قرار دارند.

پردازنده های Fusion هم در کامپیوترهای دسکتاپ و هم موبایل کاربرد دارند. AMD به این روش mainstream GPU compute App Acceleration می گوید.

تعداد برنامه هایی که از این قابلیت AMD استفاده می کنند چندان بالا نیست. این برنامه های کم می توانند بر روی GPU سریعتر اجرا شوند اما راندمان متوسط CPU منجر می شود که نتوان مقایسه درستی بین آن و پردازنده های معادل اینتل انجام داد. البته AMD معتقد است که در آینده تعداد برنامه هایی که از این قابلیت استفاده خواهند کرد به مراتب بیشتر خواهد شد.

فارق بین اینتل و AMD در زمینه پردازش گرافیکی آن است که AMD این قابلیت را به پردازنده های جدید خود که همان Fusion باشند هم انتقال داده است چرا که یک پردازنده Fusion از یک چیپ گرافیکی Radeon با پشتیبانی از قابلیت GPGPU و یک CPU ساخته شده است اما اینتل در پردازنده های جدید خود از جمله Ivy-Bridge از یک چیپ گرافیکی به نام HD 4000 بدون پشتیبانی از قابلیت GPGPU و یک CPU استفاده کرده است.

اینتل: پلی به پردازش گرافیکی

اینتل پیشرفت پردازش گرافیکی را با دلواپسی مشاهده می کرد. اینتل تلاش کرد که کارت گرافیکی خودش را بسازد اما این پروژه با شکست مواجه شد. ایده ای که اینتل برای این کارت گرافیکی شکست خورده داشت امروزه به صورت محدود در برخی از کاربردهای پردازش گرافیکی با راندمان بالا دیده می شود.

از طرف دیگر اینتل مشکلاتی با گرافیک گنجانده شده در پردازنده های جدیدش دارد. این گرافیک ها در رده متوسط محسوب می شوند هر چند کارهایی مثل کدینگ فایل های ویدیویی را به خوبی انجام می دهند. البته کدینگ ویدیو یک کار استاندارد محسوب می شود و چنانچه یک کار غیر استاندارد به این چیپ گرافیکی واگذار شود عملکرد ضعیفی خواهد داشت. با این وجود با توجه به اینکه بسیاری از کاربران از پردازش گرافیکی فقط به دنبال کدینگ ویدیویی آن هستند لذا اینتل می تواند نیازهای آنها را پاسخ دهد.

پردازنده جدید اینتل یعنی Ivy-Bridge دارای راندمان بالایی در بخش CPU است و چیپ گرافیکی آن هم از DirectX 11 به طور کامل پشتیبانی می کند. این مجموعه در واقع سیگنالی است که نشان دهنده غروب کارت های گرافیکی رده متوسط است.

تلاش های اولیه برای استفاده از کارت گرافیک به عنوان یک پردازنده با مشکلات زیادی روبرو شد چرا که برنامه نویسان باید برنامه هایی را بر اساس کارت گرافیک می نوشتند که اصولا مربوط به کارت گرافیک نبود! یعنی آنها می خواستند با برنامه های خود کاری را انجام دهند که در حوزه وظایف آن کارت تعریف نشده بود. مثلا باید برنامه هایی می نوشتند که بتواند با استفاده از کارت گرافیک سایه زنی را

شکل ۳: مقایسه API های رایج

CUDA	OpenCL	DirectCompute
Nvidia proprietary	Developed by a consortium led by Khronos Group	Developed by Microsoft
Oldest GPU compute platform software	Strong implementation on Mac OS X	Part of DirectX 11
Lots of support in high-performance computing	Available on most OS platforms	Graphics-like API
Available on multiple OS platforms	Runs on Nvidia, AMD, Intel (CPU). Intel GPU support in Ivy Bridge	Strong performance on Windows Vista, 7, 8
Mostly Nvidia H/W with CPU fallback		Runs on GPUs; no CPU fallback



شکل ۲: کارت گرافیک Tesla

جنگ API ها

امروزه وضعیت API های پردازش گرافیکی مشابه شرایط API های گرافیک ۳ بعدی اواخر دهه ۱۹۹۰ میلادی است. پلتفرم پیشرو همان CUDA است. با وجود ادعای NVIDIA هنوز هم CUDA به صورت اختصاصی توسط خودش استفاده می شود. CUDA دارای سهم بزرگی از توسعه دهندگان و برنامه های کاربردی است اما واقعیت این است که هیچگاه تاریخ با استانداردهای اختصاصی مهربان نبوده است! شاید همین باعث می شود که NVIDIA تلاش می کند که سایر رقبا را هم به استفاده از CUDA تشویق کند هر چند تا کنون که موفق نبوده است! البته می توان ادعا کرد که DirectCompute هم اختصاصی است چرا که فقط بر روی سیستم عامل ویندوز اجرا می شود و حتی ویندوزهای ماقبل از ویستا را هم در بر نمی گیرد هر چند ویندوز حاکم مطلق کامپیوترهای دنیاست و از طرف دیگر تمامی کارت های گرافیکی که از DirectX 11 پشتیبانی می کنند هم می توانند از DirectCompute بهره ببرند.

OpenCL با پشتیبانی از چندین سیستم عامل و پشتیبانی از طیف وسیعی از سخت افزارها آینده درخشان تری را ترسیم می کند. OpenCL توسط سیستم عامل MAC OS X به صورت خودکار پشتیبانی می شود و استفاده از آن مخصوصا در MAC Book ها راندمان بالاتری را نسبت به نوت بوک های ویندوزی به ارمغان می آورد. اما واقعیت آن است که OpenCL هنوز اول راه است و هیچ توسعه دهنده بزرگی به صورت رسمی برنامه نویسی برای آن را در محصولات خود نگنجانده است.

پردازش موازی در آینده

CPU ها هرگز از گردونه خارج نخواهند شد چرا که همیشه نیاز به پردازش خطی وجود خواهد داشت و برخی از کاربردها اصولا قابل اجرا به صورت موازی نیستند. آینده اینترنت و کامپیوترها به سوی ویژوال شدن پیش می رود. ویدیوهای دیجیتال، عکسبرداری و بازی ها گام های اولیه به سوی ویژوال شدن هستند و اینترنت و ویژوال از طریق استانداردهایی مثل WebCL و HTML5 Canvas تجربه ۳ بعدی را بر روی وب عملی خواهد کرد. برنامه نویسان بیشتری از کدهای پردازش موازی استفاده خواهند کرد. چپ های گرافیکی چه از نوع مجزا (Discrete) و چه از نوع داخلی (Integrated) هر دو از پردازش موازی با قدرت بیشتری پشتیبانی خواهند کرد.

یادتان باشد که پردازش گرافیکی هنوز یک طفل نوباوست! ■

است که این کدهای نوشته شده می توانند به راحتی بر روی طیف وسیعی از کارت های گرافیکی اعم از AMD و NVIDIA اجرا شوند بدون اینکه نیاز به تغییری در کد نوشته شده وجود داشته باشد.

البته باید به این نکته هم توجه کرد که کدی که بر اساس DirectCompute نوشته شده است اگر بر روی سیستمی اجرا شود که گرافیک آن از DirectX 11 پشتیبانی نمی کند دچار خطا شده و اجرا نمی گردد. بنابراین برنامه نویسان باید در برنامه های خود این نکته را هم مد نظر داشته باشند و کد دیگری هم بنویسند تا در شرایط فوق جایگزین کد مبتنی بر DirectCompute شود و البته توسط CPU اجرا گردد!

OpenCL

در ابتدا OpenCL توسط اپل توسعه داده شد و سپس اپل آن را در قالب یک استاندارد باز به یک کمیته به نام Khronos Group واگذار کرد. اپل نام آن را برای خود ثبت کرده است ولی استفاده از آن را آزاد گذاشته است.

OpenCL تقریبا بر روی همه پلتفرم های موجود از CPU های کامپیوترهای قدیمی تا GPU های تجهیزات موبایل مثل گوشی های هوشمند و تبلت ها اجرا می شود. فقط توجه داشته باشید که کدی که بر اساس OpenCL نوشته می شود ممکن است تفاوت اجرای قابل توجهی داشته باشد زمانیکه بر روی یک GPU گوشی موبایل اجرا می شود نسبت به زمانیکه همان کد بر روی یک گرافیک قدرتمند مثل NVIDIA GTX 680 اجرا می گردد!

پشتیبانی از OpenCL به طور کامل انجام شده است. AMD آنقدر شایسته OpenCL بود که حاضر شد از ATI Stream SDK خود دست بکشد و به یک SDK جدید بر مبنای پردازش موازی که اختصاصا بر اساس OpenCL طراحی شده است روی بیاورد. OpenCL حتی در وب هم خودنمایی می کند. یک نسخه از OpenCL به نام WebCL وجود دارد که توسط مرورگرها مورد استفاده قرار می گیرد. این نسخه به جاوا اسکریپت ها اجازه می دهد که کدهای OpenCL را فراخوانی کنند. این بدان معناست که شما می توانید کدهای پردازش گرافیکی را در داخل مرورگر خود اجرا کنید.

از طرف دیگر OpenCL هنوز در دوران طفولیت خود به سر می برد. ابزارها و میان افزارهای آن تازه معرفی شده اند و هنوز توسعه دهندگان نیاز به زمان بیشتری دارند تا کتابخانه ها و توابع مورد نیاز را عرضه کنند. به طور مثال در حال حاضر توسعه دهندگان بزرگ مثل مایکروسافت از OpenCL در Microsoft Visual Studio پشتیبانی نمی کنند.