

Intel Developer Forum
Tuesday, Sept. 26, 2006

JUSTIN RATTNER KEYNOTE

[Applause, music playing]

Announcer: Flight 236 for Anaheim is departing in 40 minutes.

Security Agent: Excuse me, sir. I'll have to be checking your notebook here -- security rules and all.

Justin Rattner: I'm trying to catch this flight. I hope it won't take too long.

Security Agent: Yeah. You're not carrying any liquids or gels with you, are you?

Justin Rattner: No, just my laptop.

Security Agent: Let me see. Is it a normal notebook? This is LCD, isn't it?

Justin Rattner: Yeah, I think that's an LCD display.

Security Agent: Hey, Charlie. We've got another one of those *Liquid* Crystal Displays.

Justin Rattner: Oh no! Is that going to be a problem?

Security Agent: No problem at all, but rules are rules. Let me just take care of this for you.

Justin Rattner: [Screams]

Security Agent: Yeah, okay, I'll have to leave this here; I'll meet you on the other side of security. Come on through, sir. Have a great flight.

[Laughter]

Announcer: Ladies and gentlemen, please welcome chief technology officer, Justin Rattner.

Justin Rattner: Well I understand they've relaxed the rules a little bit in the last couple days, so maybe it's okay to bring your liquid-crystal display. There's probably less than three ounces of liquids in that display.

Hey, it's great to be back to IDF for our fall session. I'd certainly like to add my welcome to Paul's and to Rob's. We're delighted to have everybody with us this week. I think it's just going to be an exciting week. I'm very excited to be here today because we've got some firsts for you this morning. It's always a lot of fun for me to be able to share those with the IDF audience. You know, we work on these things for years and years and years, and we finally get to that wonderful day when we can take the stuff out of the lab and literally give you the first viewing.

When I was here last spring, of course we were introducing the whole idea behind energy-efficient performance, EEP, and I had the opportunity to introduce the new Core microarchitecture, if you remember that. And of course that's been followed by more than 40 product announcements based on the Core microarchitecture and what we call the Core 2 Duo processor family. So that's been an exciting six months from the last IDF to this one.

This time I get to return to really more of the sorts of things that I worry about on a day-to-day basis, and that's pushing the frontiers of technology forward for Intel. And that's what I want to do this morning. I want to have you join me on this journey to the far horizons of computing technology. And in particular, today I want to explore the implications of having tremendous computing power literally at your fingertips. Teraops and teraflops of computing power, access to your data, access to applications from virtually any device -- desktop, notebook, handheld device, cell phone, what have you -- and really freedom from a lot of the things that makes using today's systems so difficult. You know, no more need to go through a complex installation, an application installation, freedom from viruses or any other forms of security breaches that might exist. We're really talking about a fundamental change in the way the whole computing infrastructure is built, and at the core of that infrastructure will be the future data center, what we refer to as the mega data center or simply the mega center.

We're also going to give you a glimpse of some of the technology, as I said, behind these developments in the mega center, and hopefully we'll give you a sense of how it relates to you as developers; what role, what opportunities these new directions will create for you and for your businesses. So if you'll bring your seatbacks to the upright position, we'll get started.

The demand on today's data centers is really staggering. It's already an enormous workload. How many of you are YouTube aficionados? How many like to download videos from YouTube? It's just amazing how that service has grown. There are over 100 million videos delivered every day by YouTube. And according to ComSource, Yahoo serves over 3 billion page views a day. Those are just staggering. When you think of a few years ago, we were just talking about delivering 1,000 page views per hour, and now we're talking about billions per day. The growth has really been astonishing. And of course that's driven the size of data centers up and up. And now, it's not unusual to talk about a data center with 10,000 processors or 50,000 processors and up. But I want to tell you that is just the tip of the iceberg, literally.

There is just this tremendous movement towards rich, Internet-based applications. You know, things like location-based services are part of that. People want to open their notebook or open their ultra-mobile device using location information and find out what services are available immediately around them. Of course that information is going to be supplied by high-bandwidth network connections like WiMAX. We'd like to have shared calendars. Much as we do in the enterprise today, we'd like to have shared calendars across the Internet, have a shared calendar with your family and friends. And of course we want to play a variety of online games. And believe me, as the complexity of those games increases, so will the need for sophisticated data centers and related services.

A lot of the growth in these applications is really being driven by new software frameworks -- frameworks like Ajax and Ruby on Rails and many others. It's really empowering developers to create these new applications that deliver this outstanding experience across the high-speed network.

Another thing that's emerging as part of this trend is actually the delivery of applications via the Web. So now I don't have to buy an application at the store, go to Fry's, go to CompUSA, or download it off the Web. Transparently, the key pieces of the applications are pushed directly into my client device of whatever sort, and all of the installation details are taken care of for me. The benefits of this, of course, will be reduced cost and improved reliability and dramatic simplification in maintenance because all of that will happen automatically.

And just to give you a sense of how much demand there is for that capability, IDC conducted a recent survey where they concluded that over half of the respondents wanted this notion of software applications on demand. So clearly we're at the start of a major inflection point, and those big changes are already underway.

If we look at the stack here, I think you'll get a better understanding of what I'm talking about in terms of extensive changes to the stack. We'll move from a model where we think of applications being installed, to applications being streamed, delivered via the web. We'll move from this notion of local data, my private data sitting on my hard drive, to local data really being a cache for massive amounts of data stored in the cloud and just bringing the necessary bits and pieces of that data to local storage as they're needed.

System management will also turn to a remote form using the capability of the net and these new software frameworks and new capabilities in the platforms, we'll be able to deliver very high-quality system management on a remote basis. And much of today's stateful software will transition to stateless software, making the migration of applications between platforms and between devices that much simpler. Literally, if a friendly TSA agent happens to destroy my notebook, I'll have no need to worry because I'll be able to get another one and the network will restore my state exactly where I was the moment the disaster took place. So tremendous capability, and a lot of those technologies are being put in place today.

All of those capabilities are really driving the demand on the data center and the growth of the data center. And we're beginning to see an exponential-like characteristic to that growth. In fact, one of the leading online service providers was kind enough to provide us with their projection for the growth over the next five years. And as you can see here from the graph, they're talking about a nine-fold increase in the number of servers, as well as the amount of power it will take to power those servers in the data center. And that nine-fold increase in their server capacity will actually produce a five-fold increase in the amount of outbound data traffic. So really astonishing gains over a relatively short period of time.

Now today, these mega data centers, or mega centers, really represent a relatively small, but still significant fraction of all servers that are delivered -- about 5 percent at the moment. But if you extend these trend lines out, not just for this one online service provider, but for all of the online service providers, you're really talking about an enormous potential market. Some estimate that by 2010, we'd be talking about 25 percent of all servers sold going into data centers with more than 100,000 servers. That's really astonishing, and yet it's possible within the next five years.

So with that kind of transformation underway, it's really hard to imagine, as a developer community, that we're not going to be impacted by this in a fundamental way. There are big problems to solve to enable the construction and the operation of these enormous data centers. And it doesn't matter where you are in the ecosystem, whether you're involved in provisioning the massive storage systems that will be part of the mega data centers, or you're part of the power community and working to deliver higher and higher energy efficiency within the data center, or you're working on information security and trying to figure out how these vast amounts of information can be protected from attacks. All of those people and everyone, literally, in the developer community will be impacted by the rise of the mega center.

But rather than tell this whole story myself, I thought it would be a good idea to bring out an expert, and we certainly have an expert here today, someone who's at the forefront of the kinds of challenges that I'm talking about. From Google, please welcome Luiz Barroso [Applause] Hey, great to see you, Luiz. Glad you could make it today.

Luiz Barroso: Glad to be here.

Justin Rattner: Google is a leader in innovative online services. I love Google Maps. I'm always showing people Google Earth and spinning around the globe. Can you explain, in relatively simple terms, the kinds of challenges that you face projecting those kinds of services across the net?

Luiz Barroso: Sure. I think probably the best way to do that is an example.

Luiz Barroso: So the example I chose is AdSense. It's a service we have that tries to place targeted ads that match the contents of a webpage. Now in this case, we have a page that's from the New York Times travel section on a city in Spain. And I guess I'd

ask you, Justin, which ads would you choose to show on a page like that?

Justin Rattner: Well, I guess if I'm looking at travel, a travel agent.

Luiz Barroso: That wouldn't be bad, but as it turns out, when you finish reading this article, you're not thinking about travel, you're really thinking about ham.

Justin Rattner: Ham?

Luiz Barroso: The page really is all about Spanish ham, jamón, and you really are craving that by the time you're done.

Justin Rattner: All right.

Luiz Barroso: And luckily these are the ads we happen to show on this page.

Justin Rattner: Okay. So this sounds like it has tremendous real-time demands on the system.

Luiz Barroso: It sure does. If you think about it, first, there are many more web pages out there than search results pages, so the traffic of something like this can be really tremendous. And then as you indicated, for every one of these pages, we have to crawl that content, we have to sort understand what they're about, and then go and try to match it against the huge database of ads we have to try to come up with the best ones.

Justin Rattner: Exactly.

Luiz Barroso: So in order to be able to do this well, really, and still afford it, we really need our computing infrastructure to be very, very efficient. And that's one of the things that I thought I'd bring up to talk to you and your audience here today.

Justin Rattner: Okay, so what do you see as the key challenge?

Luiz Barroso: Well, I thought about bringing up two of them that are important for us, today.

Justin Rattner: Okay.

Luiz Barroso: One of them is energy efficiency, which is clearly a theme in this forum today.

Justin Rattner: Right.

Luiz Barroso: The other one is programming efficiency, actually, something we care a lot about.

Justin Rattner: You mean programming productivity?

Luiz Barroso: Essentially that, yeah. What we like to do is, especially when you're programming large, complex services, we would like our engineers to be able to, first of all, program efficient services, and then try to do them as fast as possible. And one thing that should be said here today, that perhaps Intel and your developer community can help us out with, is what I call the storage gap.

So here the situation is the following -- look at the service like AdSense we were talking about just now. A software engineer at Google has to develop a new service that has to serve a terabyte worth of data. Now today, that engineer has to choose between mapping that data on disk or on memory, essentially, and the gap, as I've shown in the slide between these two options, is enormous. And as you can imagine there might be many types of services that won't fit very well in either of those. So we'd love to have some solutions somewhere in the middle.

Justin Rattner: So you're really looking for the re-architecting of the memory hierarchy.

Luiz Barroso: Essentially, yes. What I'd like to see is what if we had storage solutions that would be much faster than disk drives and yet much cheaper per gigabyte than memory. I'm sure we could use that.

Justin Rattner: Okay, well I have some NAND technology. Maybe we should talk about it after the presentation this morning. But you also mentioned energy efficiency. Do you want to give us an idea of what's going on at Google in terms of energy efficiency?

Luiz Barroso: Sure. I'd like to talk about that, too. A while ago we actually plotted a trend like this. What we're seeing there is for a standard low-end server, when we project the energy efficiency trends we're going in a bad direction. That's sort of a year-old slide, in which power could essentially overtake the hardware costs if we don't do anything about it. Of course, technologies like chip multiprocessor and multi-core are a great part of that solution because they're inherently very energy efficient.

Justin Rattner: Well, your work certainly inspired a lot of our work in that area.

Luiz Barroso: Well, there were lots of baby steps way back then. Now we really need to look at the whole system, of course. Memory, I/O, cooling and power delivery and conversion as well are very important.

Justin Rattner: We've talked about energy efficiency. We've talked about developing new platform-level energy states, and we've also talked about our CMOS voltage regulator technology. Are you looking at any techniques like that?

Luiz Barroso: Oh, no. We're glad to have you guys take care of those problems. Those are hard ones. In that you're trying to make sure you save a lot of energy, but you can't sacrifice performance, and you can't sacrifice responsiveness.

Justin Rattner: I hope they're hard. I keep telling my boss, Mr. Otellini, that this is really tough stuff.

Luiz Barroso: That's certainly one of the hardest parts. We actually profiled our energy usage a while back and realized there were actually some simpler opportunities out there. One of them is in the area of power supplies. What you see today if you have a home PC or low-end server, the likelihood is that you have a power supply that's anywhere between 55 percent and 70 percent efficient. What that means is that, the power supply itself could be the component in your system that's consuming the most energy.

Justin Rattner: Really? Amazing.

Luiz Barroso: What we decided to do at that point is say, "Well, how can we make that a little bit better?" Typically it's a hard thing to do as well, because most of the mechanisms that increase efficiency will also increase cost. We happened to find one that increased efficiency and actually simplified the power supply design. I'm illustrating that on the two photos we see. On the left-hand side we see a more traditional power supply. That's one you would expect to provide 12 volts, 5 volts, 3.3 volts -- what you expect a low-end server, maybe, or PC to have. On the right-hand side you have a power supply that only provides one single 12-volt rail. So visually you can see that these power supplies are much simpler just by looking at a number of components and removing them. But the key ingredient there is that these power supplies actually can reach 90 percent efficiency.

Justin Rattner: Ninety percent efficiency? Now you're talking. You're getting to the area that I like to think about.

Luiz Barroso: So these are big gains, essentially a factor of four in terms of energy displaced. The tricky part is not the technology in this case again. The tricky part is really the ecosystem. If you want to make a change like this you need collaboration from folks like Intel and the developer community. So we've been engaging all of you guys in discussing whether we can come up with new open standards that could allow these types of efficiencies essentially to be available for everybody else.

Justin Rattner: That's a lot of what IDF is about, sort of educating the ecosystem and getting the ecosystem to address these major challenges. I think it's great that you've been here today and talked about your challenges. I think that will stimulate a number of people in the audience to maybe give you a hand. In fact, I'm going to give you a hand in another way. I've got some work in the power supply area that I think you'll find really interesting. If you'll take a seat and watch, you may find something that will find its way into a Google data center.

Luiz Barroso: Sounds great. Thank you very much.

Justin Rattner: Thank, Luiz. Luiz Barroso, everyone.

[Applause]

Justin Rattner: Luiz talked about power efficiency and how they took a very straightforward step of moving to a single-rail supply with very high efficiency. What we've been doing at Intel for some time is actually looking at the power problem, if you will, not simply at the board level, as we've done in the past, but really at the data center level; looking at the entire path, literally from the AC main all the way down to the individual chips and other components on the board.

What we've discovered in all of this is that there's sort of an inherent inefficiency going on today, which creates a big opportunity in the data center. If you look at the typical data center power system, it's basically delivering on a third of the measured power, measured on the power meter, to the actual compute load. Roughly two-thirds of the energy is going up in the form of waste heat. Even if we go to very high-efficiency power supplies, as Luiz just described, we're still likely to achieve only an improvement, let's say, up to 55 percent of total energy expended.

The problem, if you look at the following diagram, is really a result of the fact that there are so many conversions, basically AC-to-DC, AC-to-AC conversions along the entire path. Now, the diagram here shows you power starting at the UPS, and of course, if you're talking about a mega center, it's going to have uninterruptible power. It has to have uninterruptible power. Even if your data center is parked on the Columbia River right down from the Bonneville Dam, you want to have UPS. Then you have the power distribution complex, which generally involves another stage of AC-to-AC conversion. And ultimately, the final AC-to-DC conversion in the power supplies themselves as Luiz just described. As you go through this cumulative number of steps, you continue to lose efficiency. If there was some way to reduce the number of those conversions, there would be an inherent gain in energy efficiency.

Well, we've been looking at a solution just like that, and I'd like to invite Intel principal engineer and lead power system architect Tom Aldridge to come out and give you a little bit of a view into what he's been working on. Hey, Tom, good to see you today.

Tom Aldridge: Good to see you.

Justin Rattner: Thanks for coming. Tom Aldridge, everyone. Well, you're always up to something, Tom. I've worked with you long enough. There's always something in your lab that's going to blow people away. Tell me what you're doing at the moment.

Tom Aldridge: The graph that you showed on the screen up here, Justin, really gives a hint of what we're doing. What we're doing is taking very high-efficiency power supplies and AC distribution and we're starting to pull out redundant stages of conversion for data center applications. So, I've got a power supply that I've brought on stage here, Justin, that's an example of a modern high-efficiency power supply.

Justin Rattner: Is this one like Luiz's 90 percent efficient power supply?

Tom Aldridge: This is very similar to what Luiz was showing. This is a 90 percent efficient AC-to-12-volt DC power supply used in our latest servers, which are part of our demo unit over here.

Justin Rattner: All right. Don't tell me. You're going to dump this into the dustbin of obsolete technology.

Tom Aldridge: What gave you a clue, Justin?

Justin Rattner: Well, it's a tradition here for the technology keynotes at IDF. Here's our dustbin! Wait a minute! Looks like there's an Opteron in here. Must have been left over from another talk. Okay, dump away, Tom! There you go.

Tom Aldridge: So, the same power supply functionality, but now with DC input to this power supply, we've reduced the complexity and component cost quite a bit in this power supply. We've also increased the efficiency of the power supply and increased its reliability.

Justin Rattner: So I'm getting the sense that this could represent a massive change in the entire power system of the data center. Is that fair?

Tom Aldridge: It certainly could, Justin. This is not research. This is existing technology. Standard power supply technology. We had one of our leading vendors in the technology area of power supplies put these demo units together for us for what we're going to be showing you.

Justin Rattner: Okay. But I think the real question everybody is going to have is that a demo is fine, and we're going to take a look at it in a second, but what about a true full-scale test?

Tom Aldridge: We actually are moving in that direction in conjunction with Lawrence Berkeley Labs. The test is based on the diagram that was just up on the screen, showing that we've removed redundant stages of power conversion in the data center. So we built this prototype with Lawrence Berkeley Labs here in the Bay Area and showed significant power savings with it.

Justin Rattner: Looks like quite a data center here in the picture. All right. Well, can we take a look at the demonstration here?

Tom Aldridge: Yes, we certainly can, if we can bring up the live screen on the demo. What you see on the meter is the demo running on AC power right now. It's drawing about 3,800 watts with this high-efficiency implementation. With the assistance of Justin's team, we're going to switch over to DC power and you can see just how much we're saving.

Justin Rattner: Okay. Let me get my technical assistant out here. Igor, do you want to come out and throw the big switch?

Igor: Yes, Master.

Tom Aldridge: Where do you find these people?

Justin Rattner: I tell you, it's hard to get good help these days.

Tom Aldridge: Does he know what he's doing?

Justin Rattner: I don't know. Uh oh!

[Thunder]

Justin Rattner: No, Igor, not that switch, the other switch.

Igor: Yes, master.

Justin Rattner: Right, the DC switch.

Tom Aldridge: So Igor is making the transition here to DC, and soon you'll see the DC data.

Justin Rattner: Thank you, Igor.

Tom Aldridge: Okay. Now we're running purely on DC with these servers. And you can see we've dropped from 3,800 watts down to about 3,300 watts, a 14 percent savings in energy use of these servers. That's quite significant in data centers.

Justin Rattner: Yeah, that's really tremendous. What's the implication of that to a mega data center?

Tom Aldridge: Well, Justin, the implication is that we can get a 60 percent increase in servers in the data center compared to a conventional AC distribution implementation, so 60 percent per megawatt. We can get 600 more servers on the data center floor per megawatt for the same power.

Justin Rattner: Okay, so I could make that tradeoff.

Tom Aldridge: Yes you could.

Justin: I could either save the energy or take the increase in available energy and provision more servers.

Tom Aldridge: That's right.

Justin Rattner: Well that's really tremendous. All right, I think we'll look forward to further results from you in this area, and I want to thank you for joining us on stage today.

Tom Aldridge: Thanks, Justin.

Justin Rattner: Thanks so much. Tom Aldridge, everyone.

[Applause]

Justin Rattner: Well, turning from energy-efficiency now, let's take a look at another area that's going to be critical in the mega data center of the future and that's security. Now we all read about security every day. We hear about laptops getting stolen and Veterans' Administration records getting lost and so forth. It's been reported by some of the security watchdog groups, particularly the Security Privacy Group, that in the past 18 months alone, over 93 million records of U.S. residents alone have been compromised in some way. And I think what surprises me and I think will surprise you is the fact that increasingly those breaches are due to internal threats rather than external threats. In fact, as the number of servers in the data center increases, the internal threat begins to outweigh the significance of the external threat.

So what I'm talking about is, in most enterprises and in most datacenters, once you move inside the firewall, the system is relatively unprotected. And it's possible for a hacker of sorts to actually sniff the network and observe the packet traffic. And really the question is why aren't we protecting it when there are available technologies such as IPSec which were designed for this very purpose? Well IPSec has certain issues associated with it. In particular, it's fairly difficult to deploy, there are issues with key-management and so forth. It does have the advantage, though, that packets are encrypted in an end-to-end fashion. In fact, let's take a look at that.

In the diagram you can see here it's fairly straightforward. There are keys at each end of an IPSec path, so this is the end-to-end path, from source to destination, and keys are used at each end to decrypt packets. However, the packets moving across the network remain encrypted which makes it very difficult to do things like packet inspection as a data center manager might want to do in order to understand what sort of traffic is moving across the network.

An alternative to this is known as LinkSec where we actually encrypt at each link and decrypt at the end of that individual link. So we're encrypting and decrypting as we move from device to device, network device to network device across the network. And this gives us the opportunity, if we're equipped with the proper keys, to inspect the packet traffic, to look at denial-of-service attacks, which may be created as the result of internal malicious activity in the datacenter.

I'd like to give you a demo now of the LinkSec technology, so I'll ask for my assistants here to join me on stage. I think you guys moved my terminal here a little bit. Come over to this side. Larry here is my data center manager. Hi, Larry. And someone who looks remarkably like a TSA agent is playing the role of our hacker here. Let me bring up a picture of the setup we've got.

So obviously we couldn't bring an entire data center here today, so we're simulating a data center. We have a server as the data center. We have a switch and a hub and we have a couple of client devices, the ones you see in front of me and the ones, let's say, that, someone who's come into our building, into our business, and just routinely hooks up their laptop to the network either in a wired or a wireless fashion. But they represent a potential threat.

So the first thing I'll do is I'll send a message to Paul Otellini. Typing in front of thousands of people is never that much fun. "Great talk, Paul." All right, I better leave it at that. So this is a non-LinkSec machine. I should point that out, so anything I send here is going to be in the clear. And if I send it out, you'll notice that the message is transmitted, but it's also visible to the guest, if you will, the hacker, who is doing packet sniffing. So there's no encryption; the packet moves in the clear, because we're inside the firewall.

That's not a very good situation, and in a large data center, you know, you might have an employee that's got some other ideas about the information and its accessibility within the data center. The alternative is to deploy LinkSec technology, which is now going to encrypt and decrypt across every link in the data center. And this time if I send Paul that message, all right, so we'll turn on LinkSec here in the data center and send the message. All right. And now you see these encrypted packets going across the network. And while my email reaches Paul, of course, in the original form, having been encrypted and decrypted, the hacker can no longer see the message because he lacks the keys necessary to decrypt the packet. So this is really the capability of LinkSec, and we've been working with leading vendors in the industry to create this technology and make it available for deployment.

LinkSec has many advantages over available security technologies, one of which is the fact that it's incrementally deployable. So rather than having to convert the entire enterprise to the new security technology, it's possible to deploy LinkSec in an incremental fashion. LinkSec has other advantages. The encryption/decryption algorithms have been designed to operate at very high speeds, so I can use LinkSec across 10 gigabits per second and even higher speed links with good efficiency.

If you want to learn more about LinkSec, I'd encourage you to look up the class, which will be offered on Thursday, and you'll find out everything you need to know about the LinkSec technology, the state of the standards that are being defined around LinkSec, and how LinkSec can be deployed in your organization.

All right, let's turn to another issue, and that's really the one about building data

centers of the magnitude of a million servers. Here are a few quick facts to sort of calibrate you on what a server like that would require if it was built from today's technology. If we wanted to provision a million servers today, we'd have quite a challenge. We'd need an enormous amount of space, right? We're talking about 800,000 square feet of floor space. That's the equivalent of 18 football fields. We'd need an amazing amount of power; 500 megawatts of power. And to calibrate you there, that's more power than would be required for a quarter million homes. So an enormous amount of power necessary for the million-server data center using today's technology.

And literally, it wouldn't make a lot of sense to try to do the mega data center today. The only way we're going to make it feasible -- and given the growth in services, it's essential that we create the technology to make it feasible -- is to achieve dramatic improvement in server performance density as well as server energy efficiency.

Now, doing that kind of improvement in performance density is a big part of a program we've had underway for some time at Intel, what we call our Terascale Research Program. This program is based on our belief that there is a tremendous need for increased speed and density not just in the client, but in the server as well. And we've characterized this, for the Terascale Research Program, really in three ways. One is TeraFLOP of performance, trillions of operations per second. They could be ops, they could be flops, probably both. Terabytes per second of memory bandwidth -- if you're going to have that kind of computing power, you need to supply it with adequate memory bandwidth. And terabits per second of I/O performance; you need to get the data in and out of these high-performance processor chips and their memory.

Now, these three keys really form the foundation of our Terascale Research Program, and I'm really excited to be able to share with you some more details about the latest results from that research program. Paul gave you a bit of a preview of this. I'm going to take you into a little more depth.

We just got the silicon back earlier this week of our Terascale Research prototype. And as you can see in the accompanying diagram, each one of the 80 cores on this die consists of a simple processor. It has a simple instruction set -- not IA compatible; it's just has a simple instruction set that lets us do simple computations in floating point and basically push data across the on-die fabric that connects the 80-cores on this die together. Now, in aggregate, those 80 cores produce one teraflop of computing performance, so a single one of these Terascale Research prototypes is a one teraflop-class processor. It delivers an energy efficiency of 10 gigaflops per watt at its nominal operating frequency of 3.1 gigahertz. That's an order of magnitude better than anything available today.

Now, a critical thing when you pack that much computing power on a single chip is memory bandwidth. And to provide the adequate memory bandwidth for a teraflop of computing power, we developed what we think is a really novel solution. The novel solution involves stacking a memory chip directly underneath the processor chip. So as you can see here in the illustration, we bring the memory chip and the processor chip

together in a face-to-face manner. The chips have been bumped appropriately so when they meet, there are thousands of simultaneous electrical connections formed in that direct bonding step. That gives us a tremendous amount of memory bandwidth. For each core on that die, there's a 256-kilobyte static RAM, and the aggregate bandwidth between all of the cores and all of those SRAM arrays is one trillion bytes per second, truly an astonishing amount of memory bandwidth.

So now we've created two of the three Ts in experimental form here -- a teraflop of performance on a single chip and a terabyte per second of memory bandwidth. The last piece of the teraflop puzzle is a terabit per second of I/O bandwidth. Now, we've obviously done tremendous work in electrical signaling over the years and we've demonstrated a constant stream of improvements. We've improved the performance of USB, we've driven PCI Express as a new bus technology, our memory buses operate well in excess of a gigahertz. So we're no strangers to electrical signaling and we continue to make advancements in electrical signaling. But when we've taken a long-range look at electrical signaling, we see some danger signs. That caused us to look at a new approach to high-speed signaling, namely, to look at photonics. Optical signaling, if you will.

The program was launched about five years ago. I think you've heard about some of the developments. Some of them, in fact, I think astonished the scientific community. It was long thought, for example, that it would be impossible to create a laser out of silicon, but in fact, we demonstrated exactly such a design last year. We also addressed the other component needs for a complete photonic component family. That included things like modulators, multiplexers, deep multiplexers, and so forth. We started with a one gigabit per second modulator in 2004. We've recently demonstrated the 10 gigabits per second modulator. We have hopes to demonstrate even higher speed modulators in the near future.

Last week it was kind of hard to pick up a paper, certainly here in San Francisco since it was above the fold, we announced yet another breakthrough in the silicon photonics area. That was the development of the first electrically pumped, hybrid silicon laser. This is really one of the last remaining pieces of our photonics vision. To help explain its significance I'd like to bring out Professor John Bowers from UC Santa Barbara, one of our chief collaborators, who is really an expert in three-five compounds. John, you want to come on out?

[Applause]

Professor John Bowers: Hello, Justin. It's great to be here.

Justin Rattner: It's great to see you. Well, this is a tremendous achievement, this hybrid laser, but I thought it might be interesting to contrast last year's announcement of the Raman laser with the latest announcement of the Hybrid laser. Can you kind of share the wisdom here on how these differ?

Professor John Bowers: Sure. Last year Intel announced the first CW Raman lasers.

It was the first time people made a laser in silicon. That's what was so exciting. That was optically pumped. What we've done in this development is to work out a way to make electrically pumped lasers directly on silicon.

Justin Rattner: Do you employ some special materials in order to achieve this?

Professor John Bowers: That's correct, Justin. We deposit a thin layer of indium phosphide, which has high gain materials, and deposit it on silicon wave guides in a way that is CMOS compatible. That way we can make laser cavities and provide on-chip, directly-driven lasers on a CMOS chip.

Justin Rattner: So this mixture of silicon technology and three-five technology is really the basis. I guess that's why we call it "Hybrid Laser". It's sort of like a hybrid car, right?

Professor John Bowers: Very much. We combined the best aspects of three-five technology and CMOS technology.

Justin Rattner: Okay. Do you want to take us through the manufacturing process? I think that's of significant interest. I'll give you the controls here.

Professor John Bowers: Okay. I'll be glad to. Okay, what's shown here is a wafer of silicon. It's a SOI Wafer. What we've just shown there is etching a wave guide on that wafer of silicon. That wave guide is much like a wire for electrons, but that guides the light around the chip. In this case it's the laser cavities that we'll show. That would connect with the other devices you mentioned -- the modulators and detectors and other silicon photonic devices that one would use.

We then take a second wafer of indium phosphide. We deposit some layers on it. Then we expose that to an oxygen plasma. That's kind of the key thing we've developed. It's a low-temperature process, so it can come in at the end of the CMOS process in a CMOS-compatible way that doesn't cause contamination. That happens at about 300 degrees C. We expose the wafers in the plasma. Then we take this wafer and flip it over with a conventional silicon wafer bonder that is used in the Fabs today.

Justin Rattner: How thick is that layer, roughly?

Professor John Bowers: It's a wafer-scale approach. It's initially about half a millimeter thick, a standard wafer. We etch off that wafer, leaving just a thin film that's only about a micron thick. It's just a different way of depositing a layer on top of this.

Justin Rattner: The indium phosphide and the silicon actually stick to one another?

Professor John Bowers: That's right. The key thing is this oxygen plasma grows this very thin glue, basically -- it's about 25 atoms thick. It's this oxide-to-oxide bonding that's the same thing people use commercially in SOI wafers today.

Justin Rattner: I guess it's better than crazy glue.

Professor John Bowers: [Laughs] Absolutely, absolutely. The point is it's a wafer-scale process. The next slide here shows when we etch the indium phosphide material, we place contacts on this material that we use to inject current into here. That's the difference. If you inject current into a silicon device you typically get heat out. That's why last year's announcement was so surprising to people. Here we apply electrical current across this junction, and we get light out. You can see light building up in this cavity that carriers recombine together and produce photons. That laser light is what you see emanating out of that silicon wave guide.

The point is that the light is confined and entirely controlled by the silicon wave guide. So it's controlled by a CMOS-level process, not whatever one might do in a three-five device.

Justin Rattner: Okay, now all of this is using standard manufacturing technology?

Professor John Bowers: That's right. It's a standard wafer-bonder to combine the wafers, and it's a standard oxygen plasma system.

Justin Rattner: Okay. How many of these can you pack on a single die?

Professor John Bowers: That is the power of this announcement. What we've shown here is four lasers integrated together than emit light. You then combine in an external transmitter.

Justin Rattner: Did you bring any of these with you?

Professor John Bowers: Yes, I did, as a matter of fact. What I show here is a small die. It contains about 1,000 devices on it.

Justin Rattner: About the size of my thumbnail here.

Professor John Bowers: That's correct. And on the right-hand size is a bar of devices, and that's what we'll show today.

Justin Rattner: How many lasers on that one?

Professor John Bowers: That's about 26 lasers on that device.

Justin Rattner: Wow, so a single chip with 1,000 and a little bar here with 26. That's really impressive. Well, can we see this thing in operation? We've sort of teased everyone that they're going to see the first hybrid laser demo. And remember when you get home tonight to tell your family that you were here when Intel demonstrated the first hybrid silicon laser with our good friends from UC Santa Barbara.

Professor John Bowers: Okay. What you have here, you can see the shot here. The laser bar that I have in my hand is mounted here, it's connected up to bunch of wires to a cable over here that you'd use to drive externally. These are telecommunication lasers, so they're lased in the infrared -- telecommunication wavelengths. So we've got it imaged here onto an infrared camera, and you can see that here. In this case, there are four lasers lasing. They're spaced by a couple hundred microns apart. And that's four out of a total of about 26 lasers that are connected on here.

Justin Rattner: Okay, so if this is really operating here, I should be able to put my hotel room key into the beam and turn it off. All right.

Professor John Bowers: That's correct.

Justin Rattner: There you have it, everyone. Four hybrid silicon lasers first time publicly anywhere here at IDF. [Applause] All right. Well, you know, the press has obviously made a lot to do about the development of the hybrid lasers and, of course, we're happy about that, but you're an expert here, can you just share with us a bit about the significance of this development from your perspective?

Professor John Bowers: Sure. Well, historically the photonics community has been based on Indium Phosphide and it's fairly small-diameter wafers, and it's a fairly low-volume industry. So we all know that photonic has incredible capacity, right? You can put 10 terabits of data down a single optical fiber. So it really does address a lot of these issues that you're coming up with in the Terascale Program. But the problem is that the expense has always been extremely high. So a typical transmitter is at least \$100, so \$100 per pin is something that isn't very useful for your computer. But now this is a way to do this in a CMOS-compatible approach that could do high volume using facilities that you've already paid for with microprocessor products, as opposed to having dedicated lines and having to upgrade to the next level of technology, paid for by just single product lines.

Justin Rattner: Right. Well that was certain the vision we had when we embarked on the silicon photonics effort. And, we want to thank you for the tremendous collaboration, and we hope it continues. Thanks very much.

Professor John Bowers: My pleasure.

Justin Rattner: Professor John Bowers, everyone. [Applause] All right. You know, having the hybrid laser, of course, is pretty exciting, but what's more exciting is the potential to satisfy the third T that I mentioned, which is the terabit-per second I/O length. And what it would take to do this would be roughly 25 of the lasers that you saw John holding there, the tiny bar of lasers, connected by the silicon wave guides and feeding into a bank of modulators, as shown here. The lasers would be individually tuned, so we're going to use wavelength division multiplexing. If you want to think about it, we're using different colors to generate multiple streams of data. And in aggregate, taking all of these streams, deliver a terabit per second of optical I/O. Now the data would be impressed onto those

optical streams using silicon modulators that we talked about. In this case it would take a 40-gigabit-per-second modulator to achieve the aggregate performance of one terabit per second across the array.

Those different modulated signals would then go through a passive multiplexer, again built from silicon, and feed into a single -- and this is important -- a single optical fiber. So one optical fiber would deliver a terabit per second of I/O performance. At the other end of that single optical fiber length would be the receiver. The receiver would contain a demultiplexer to separate the different optical frequencies, and silicon-based photo detectors to recover the data from the signal stream and then the various electrical interfaces to transform that information into the form that can be used by our computer system.

Now in addition to delivering an astonishing amount of performance, single-fiber, small-chip, terabit-per-second of I/O capacity, the link is very efficient. And, in fact, if you can compare trying to do this over a copper link, you see that the photonic link represents a 50x improvement. So rather than have literally hundreds of wires if you were going to do this electrically, you reduce it all to one single optical fiber.

So there you have it. That's our look at the future technologies that will be at the heart of tomorrow's data centers, driven by the popularity of online services, which continue to grow as consumers and enterprises begin to take advantage of these new software models. As a developer community, we need to embrace this change and stand up to solve the kinds of challenges that you've heard talked about this morning. And we believe that working together we can make this vision of the mega data center a reality for everyone. We hope you'll join us on this journey. I promise you it will be an exciting ride. Thanks for being here today.

[Applause]

Female Voice: Ladies and gentlemen, please welcome David Dixon.

David Dixon: Good morning. We ask that the members of the press and analyst community stay in their seats for just a few minutes while we set up the stage. There will be time for the paparazzi to have their fun after the Q&A. We'll begin in about three minutes. Thank you.

Female Voice: Ladies and gentlemen, sessions will begin promptly at 11:00 on the second floor. Members of the press and analysts, please remain seated while we prepare the stage for the Q&A session. Once again, the sessions will be starting promptly at 11:00 on the second floor. We'll see you there. Thank you and have a great morning.

This presentation contains forward-looking statements that involve a number of risks and uncertainties. These statements do not reflect the potential impact of any mergers, acquisitions, divestitures, investments or other similar transactions that may be completed in the future. The information presented is accurate only as of today's date and will not be updated. In addition to any factors discussed in the presentation, the

important factors that could cause actual results to differ materially include the following: Intel operates in intensely competitive industries that are characterized by a high percentage of costs that are fixed or difficult to reduce in the short term, significant pricing pressures, and product demand that is highly variable and difficult to forecast. Additionally, Intel is in the midst of a crossover to a new microarchitecture on 65nm process technology in all major product segments, and there could be execution issues associated with these changes, including product defects and errata along with lower than anticipated manufacturing yields. Product roadmap projections are affected by the timing and execution of the manufacturing ramp, product defects and errata (deviations from published specifications), changes in the timing of qualifying products for sale and changes in customer demand. For a variety of reasons, we may terminate product development before completion or delay or decide not to manufacture and sell a developed product. For example, we may decide that the product might not be sufficiently competitive in the relevant market segment, or for technological or marketing reasons, we may decide to offer a different product instead. Revenue and the gross margin percentage are affected by the timing of new Intel product introductions and the demand for and market acceptance of Intel's products; actions taken by Intel's competitors, including product offerings, marketing programs and pricing pressures and Intel's response to such actions; Intel's ability to respond quickly to technological developments and to incorporate new features into its products; and the availability of sufficient inventory of Intel products and related components from other suppliers to meet demand. Factors that could cause demand to be different from Intel's expectations include customer acceptance of Intel and competitors' products; changes in customer order patterns, including order cancellations; changes in the level of inventory at customers; and changes in business and economic conditions. The gross margin percentage could vary significantly from expectations based on changes in revenue levels; product mix and pricing; variations in inventory valuation, including variations related to the timing of qualifying products for sale; excess or obsolete inventory; manufacturing yields; changes in unit costs; capacity utilization; impairments of long-lived assets, including manufacturing, assembly/test and intangible assets; and the timing and execution of the manufacturing ramp and associated costs, including start-up costs. Expenses, particularly certain marketing and compensation expenses, vary depending on the level of demand for Intel's products and the level of revenue and profits. Intel is in the midst of a structure and efficiency review which may result in several actions that could have an impact on expense levels. Intel's results could be affected by the amount, type, and valuation of share-based awards granted as well as the amount of awards cancelled due to employee turnover and the timing of award exercises by employees. Intel's results could be impacted by unexpected economic, social, political and physical/infrastructure conditions in the countries in which Intel, its customers or its suppliers operate, including military conflict and other security risks, natural disasters, infrastructure disruptions, health concerns and fluctuations in currency exchange rates. Intel's results could be affected by adverse effects associated with product defects and errata (deviations from published specifications), and by litigation or regulatory matters involving intellectual property, stockholder, consumer, antitrust and other issues, such as the litigation and regulatory matters described in Intel's SEC reports. Please refer to Intel's most recent Earnings Release and most recent Form 10-K or 10-Q filing for more information on the risk factors that could cause actual results to differ.