

Benchmarking Methodology Working Group (BMWG)

IETF 94

Tuesday, November 3, 2015

0900-1130 JST Morning Session I

Room 413 OPS bmgw

Remote Participation:

<http://www.ietf.org/meeting/94/index.html>

<http://www.ietf.org/meeting/94/remote-participation.html>

Minute Takers: Marius Georgescu, Ramki Krishnan. (A huge thanks to our note takers!)

==

0. Agenda Bashing
No agenda bashing.

1a. New Charter and Milestones (Chairs)
Update, no comments.

1b. WG Status (Chairs)
No comments or questions.

2. Data Center Benchmarking Proposal

Presenter: Al Morton

Presentation Link: <https://www.ietf.org/proceedings/94/slides/slides-94-bmgw-8.pdf>

<https://datatracker.ietf.org/doc/draft-ietf-bmgw-dcbench-terminology/>

<https://datatracker.ietf.org/doc/draft-ietf-bmgw-dcbench-methodology/>

Presenter: I will suggest some text for the use of PDV and IPDV in this draft, and we can discuss it on the mailing list.

Scott Bradner:

1) In the introduction section it would be helpful to be concrete. What are the specific devices targeted by this methodology? The datacenter is a very fuzzy place with a lot of stuff.

2) What is the purpose of these tests, for example latency? To me latency is the impact of sticking a box in the path. That seems to be a logical thing. If it is that, then it is not first in to last out, it's first in to first out. What is the time added by sticking this box in there. If that is the case, you would want first in to first out as being the latency measure. First in to last out is the packet length determiner. That wouldn't be a processing delay. This is a question for the group. My opinion is that it should be first in to

first out. That is in section 2, Latency.

3) For Section 3.3 (Jitter), there is a very complicated thing in here about packet size and other stuff. If it's first in first out, the jitter is trivial, because it's the variability from when the first bit shows up. It's not depending of packet size. It's a much simpler thing.

4) Under Section 4, it has all these things about what type of device it's used to generate the traffic. I find it interesting to know, but what difference does it make to the testing? If it makes a difference to the testing then it should be stated why it makes a difference to the testing. If it doesn't make a difference, certainly it shouldn't be mandatory. We're testing black-boxes with black-boxes. I am not sure what difference does this make to the ability to understand the results of the test. If it does, then it should be.

5) Line rate, it doesn't answer the specific question: is it bits per second or frames per second. It should be stated. This is Section 5. It implies that it's frames per second, but it doesn't state that. And when it talks about the frames per second, forgets about the inter-frame gap, so the actual frames per second calculation is wrong. There's also a long thing in here about parts per million of frequency, stability and all of that and it seems like overkill. I this thing is 0.0005% faster than that thing, what difference does it make? We're not dealing with that kind of granularity. This thing in section 5.2 about the accuracy of the transmit clock that seems to me like overkill. The inter-frame gap is missing in Section 5.3 .

6) Under Section 6 of the buffering, I didn't see where it talks separately about shared and dedicated buffers. There's a lot of devices that have one big shared buffer pool, and all the interfaces use the same buffer pool. So, if one of the interfaces is using a lot, the other one don't get as much and it doesn't talk about that in here. At least I didn't see it.

7) Under Incast (Section 6.2), my basic question is: what are you trying to find out? This is a number of simultaneous packets coming in. What is the purpose of the test? What is the difference that is it going to make when you have an answer?

8) Under Throughput, it talk here about goodput in Section 7, and it's ambiguous. It's talking about the payload. It would be more straightforward to say it's the IP packet minus the IP header, or it's the TCP packet minus the TCP header. It's ambiguous to just say payload. What do you mean by payload? What I don't understand on goodput in general is how is this affected by the device being tested? The device being tested is a transit device. It's the sender and the receiver that determine the amount of payload per packet, not the device being tested. So I don't understand what the implication for the device under test is for this. If it's only for traffic that's being originated by the device, or synced by the device, maybe it would make a difference. If it's only transit it wouldn't. This goes back to the question what kind of devices are we talking about here.

9) In Section 7.5 there is information about the TCP stack, OS version etc. and I don't understand how that impacts the device being

tested. That's what I have for comments.

Presenter: It sounds like the authors were intending to do some TCP based testing. Thanks for the excellent review, Scott. I will work with them to solve this stuff.

Marius Georgescu: I just have a short comment about repeatability. Essentially they talk about variation. I think they should be talking about variance. It's good to have min and max, but the variance is also accounting for the number of trials. Also, about the summarizing function, I think we should have that discussion as a working group and will try to cover it in my presentation.

Presenter: Excellent feedback today. This is the kind of feedback we are looking for to make our work better and for all of us to learn a bit as we go along.

3. IPv6 Transition Benchmarking

Presenter:

Presentation Link: <https://www.ietf.org/proceedings/94/slides/slides-94-bmwg-1.pdf>

<http://tools.ietf.org/html/draft-ietf-bmwg-ipv6-tran-tech-benchmarking-00>

Newly adopted. Many comments addressed on the list.

IPv6 transition technologies evolution

– Stateful IPv6 transition

Scott: That drawing it's pretty amazing. The amount of joy on those characters is just stunning.

Marius: I think it's pretty cool as well. The drawing is called "Not speaking" and I asked the original author for the right to use this drawing because I believe it represents the IPv6 transition pretty well. What I like the most is the scared, little and patched-up IPv6.

Al: What were the generic categories of IPv6 transition technologies called before changing them?

Marius: They were called slightly different but they had a similar meaning. We switched from the CE-PE model to the IP-specific domains model, as Fred Baker suggested. I don't think this change is terribly important.

Al: I think the names are important. They have to convey what you are trying to test.

Paul Emmerich: I have a comment about latency. I don't think 20 measurements are enough, especially if you consider software devices, because you get a really bad tail latency and you should look at the 99th percentile or 95th percentile. Especially if you run the test on a VM, you wouldn't be able to capture this.

Marius: I didn't understand the question/comment. Thank you for the comment.

Ramki Krishnan: Is the network function here implemented as a physical function or a virtual function. Working in the virtualized space, you should consider that.

Marius: We would like to stay away from the virtualized world, especially since the effort to standardize tests for that are still ongoing. I think physical is the way to go for now.

Al: Let's say for the scope you haven't mentioned virtualization. How about we start in the physical world, and perhaps there's a bis version of this considering virtual devices.

Ramki: we could have a correlated effort since there are implementations supporting virtualization.

Marius: We could have a reference for virtualized network functions to the ongoing work. Would that work?

Ramki: Maybe a separate work item would work better.

Al: the scope of the draft is for now physical. Let's discuss that further on the list.

Scott: This is a non-content comment. You list a number of terms for which you then say: "use the definition in RFC2544" or whatever. Just list them, say something like : "The following terms are defined and should be used as defined in" the referenced RFC.

Al: I think the comments you've got were covered well.

4. VNF and Infrastructure Benchmarking Considerations

Presenter: Al Morton

Presentation Link: <https://www.ietf.org/proceedings/94/slides/slides-94-bmwg-2.pdf>

<https://datatracker.ietf.org/doc/draft-ietf-bmwg-virtual-net/>

- Resolution of comments from IETF93 and on the list
- This draft is referenced in ETSI NFV GS, OPNVF specs
- Ready for WGLC

Scott: I think putting the changes for each new version of the draft should be at the end rather than at the beginning.

Al: I'll do that in the future.

Al(author): I propose a Working Group Last Call (WGLC) for this draft.

Al(chair): Are there any objections to going to WGLC.

[no objections]

Al: It's up to Sarah to call consensus for taking this draft to WGLC. Please participate.

Joel Jaeggli: How many people have read the draft?

[6 hands in the room]

Ramki: Many people that have read the draft are not here.

5. Benchmarking SDN Controller

Presenter: Bhuvan

Presentation Link: <https://www.ietf.org/proceedings/94/slides/slides-94-bmwg-3.pdf>

<http://tools.ietf.org/html/draft-ietf-bmwg-sdn-controller->

benchmark-term-00.txt

<http://tools.ietf.org/html/draft-ietf-bmwg-sdn-controller-benchmark-meth-00.txt>

Al: Open Daylight folks have the concept of a leader and a follower.

Bhuvan: What I understood about the leader and the follower is it helps the election process to decide who is the master and who is the slave. The leader initiates the election process.

Scott: The test setup. You are testing and reporting on results considering the time it takes for the controller to do things. The things that it does involve asking the network element to do something and report back when it's done. I think you would want some text in there that specifically talks about the forwarding plane stimulator has to be very repeatable and if indeed it's going to vary as to its performance depending on the number of nodes it simulates, those things have to be taken into account when reporting the results. The timing you get out of it is the time of the controller plus the node, and when you get two variables in the same result, it tends to be confusing. So, you might want to think about that.

Bhuvan: This is a valid comment, Scott. Certainly we'll take this into account. Thank you.

Al: I've heard the same thing from Scott, and the solution of making one of them repeatable is perfect.

Marius: I think we should have a solution for Average vs Median in terms of summarizing. I will start a discussion on the mailing list on this subject.

Bhuvan: We have received similar comments in the past. I also think we need other ways to capture the results. I agree we need to discuss this further on the mailing list.

Scott: You might wanna consider putting the two diagrams in the terminology document as well, because in order to understand the terminology, you have to understand the diagrams.

Bhuvan: That makes sense.

Scott: The discovery depends on the test environment to emulate multiple nodes. It's a combination of scaling the controller and the device that's emulating the nodes.

You might find that the latter is going to be very hard, especially when considering the order of thousands of nodes.

Bhuvan: Here that particular point it's missing, but again it depends on the test solutions. The solution we are using currently is able to scale up to 10.000 node in a software environment running on Commercial off-the-shelf (COTS) hardware. I don't really see a constraint there.

Scott: I am not sure if it's a constraint or not, but maybe you want to consider ways to test the tester, which doesn't involve a

particular Device under test (DUT). Some way to test your emulator. A calibration process that can help you understand it can scale to the point you think it can.

Jingzhu Wang: I have a question about the definition of Control Session Capacity. I see in this context that we have different sessions but what kind of sessions are there? What type of protocols are there? From this perspective an Openflow session is different from Netconf session. All considering, is there a way to differentiate from different sessions or different scenarios?

Bhuvan: Without the protocol you cannot test it. We recommend to make sure to report the protocols that you are using.

Doug Montgomery: Maybe relating to Scott's comment. When you're trying to measure something like Discovery Topology Time, do you have all the time measurements, strictly defined by events in the southbound interface? In other words, could you envision factoring out the latency of the network or the emulated network itself, and have the amount of time the controller itself takes to process topology discovery messages?

Bhuvan: We assume that the emulator, considered a standardized element, has close to 0 impact. That is a valid point. We should make sure that it's repeatable.

Scott: I very much like what he just said. Consider strongly that the reported result have subtracted out the time it takes for the emulated nodes to do their function.

Al: I'll try to help with this, Bhuvan.

Bhuvan: To an extent, most the measurement are using the ingress point on the control plane and management plane. So, we avoid the delay within the lower layers. The question about the impact of the emulation plane we need to capture it a little more, maybe highlight it in the report. That still needs to be addressed. But, most of the tests are not impacted by this.

Scott: Please be clear in documenting that. I think that is a very important point.

Bhuvan: We have a section about test measurement recommendations. Maybe we have missed explaining that there.

Scott: To clarify, we're talking about the reporting. When you say how to report it, you wanna make sure it says that the time it takes for the emulated node to perform its function is not part of what is being reported as the performance of the SDN.

Bhuvan: That's perfect. Thank you.

Al: Excellent work, Bhuvan. Excellent presentation as well in the SDNRG.

Bhuvan: Thank you very much.

6. Benchmarking vSwitches in OPNFV

Presenter: Maryam Tahhan (Remote)
Presentation Links: <https://www.ietf.org/proceedings/94/slides/slides-94-bmwg-7.pdf>
<http://tools.ietf.org/html/draft-vsperf-bmwg-vswitch-opnfv-01.txt>

New / First time BMWG has had a remote presentation. We pulled it off flawlessly. A big thanks to meet echo.

Marius: One clarification question. I know instinctively what the soak test means, but I couldn't find any definition for it. Should it be clarified?

Al: We can do that. It's a long term test. It has been used traditionally in the industry.

Ramki: Any reason for choosing OPNFV ? It's all about vSwitch, so it's very generic.

Al: It's tied to OPNFV because it's a project in OPNFV. There's also a tool development there, a part that Maryam is trying to get to.

Al: two moonGen developers are here with us and they took a picture of the current proposed tests slide. We're hoping they will be able to join us. The moonGen developers, Paul and Sebastian, you're very welcome to join us in this work.

Doug Montgomery: It's an expansion of the question about the title, whether the work is specific to OPNFV. Is it specific to software only virtual switches? It seemed to me that the core of this might be dataplane benchmarking for SDN switches. Whether that's a pure software switch, a hybrid software switch or a hardware switch, and I am wondering why isn't there a draft about benchmarking methodology for SDN data plane ? And this would be a small delta on top of that.

Al: It would be great if it came in that direction, but it didn't. I think it's because this project intends to automate all this testing, and part of the tool Maryam is talking about is starting the vSwitch, scripting the external test devices, basically getting the whole thing to run in a single pass. Could that all be done with hardware switches? Probably, but the scope of this was set at OVS and DPDK and things of that nature. Have I been thinking about this in terms of completely reusable, beyond this scope? Absolutely, and I think Maryam would agree, right?

Maryam: Yeah, it's something we've been thinking for a long time and in a lot of detail and it's actually when we designed the test framework to be very pluggable. So you can plugin or plug-out whatever components you want. In terms of the vSwitch, it's treated as an abstract object essentially. What the underline implementation is, comes down to somebody enabling a vSwitch, a base switch or a virtual switch. I think the methodology is applicable in either scenarios. That's what we took into consideration when we designed the test framework itself.

Ramki: I have a follow-up comment for that. The hardware switches are trailing behind in terms of SDN types of capabilities, so I think

it's worth capturing this comment. My point is that the lack of capabilities in hardware should not stall this work.

Al: Thank you. We're running out of time, so let's go to the summary slide.

Myriam: the test specification and test framework will be in place for a time. We would like your opinion on the adoption of this document.

Al(chair): How many people in this room have read this draft? It looks like 5 people have read it. I think we should ask on the list and give more people a chance to read it. This would be a working group adoption of this summary. We would point to the Brahmaputra release to freeze the summary. This eventual part will come as a separate proposal along the road. This being an opensource work, so you can expect that the test design specification will be constantly improved. There's going to be a link to that collateral.

Ramki: We had similar situations in NFVRG, we'll keep it adopted and continue like that, they don't necessarily become an RFC.

Al: that's a possibility too. Are folks interested to hear updates on this work? I see some nods and lots of hands now. We plan on continuing to do that, and I want to thank Maryam for getting up in the middle of the night in Ireland for this presentation. Thanks you.

Maryam: Thank you very much.

7. Benchmarking for CoPP, Control Plane Policing

Presenter: Shishio Tsuchiya

Presentation Link: <https://www.ietf.org/proceedings/94/slides/slides-94-bmwg-4.pdf>

<https://tools.ietf.org/html/draft-shishio-bmwg-copp-00>

Control Plane Policing (CoPP) is defined as RFC6192 to protect router's control plane from undesired or malicious traffic

Al: The CPU utilization cannot be considered a benchmark. It's what we call a white-box metric. That's why we noted it as an auxiliary metric in the other work, the vSwitch benchmarking. Something you would measure alongside other benchmarks.

Shishio: Thank you very much. Is the draft applicable for the working group?

Al: Thank you for this presentation and for preparing the draft. As I quickly read through the draft, there are some things similar to the stuff we did here in the past. How many people have read this draft?

Scott: A lot needs to be filled in the draft before we can figure out if it's going to be useful or not. The basic concept seems to be that someone is trying to abuse your router, how hard is it to

succeed? I am not sure what the test is trying to do, see the impact of what normal traffic is and i don't see that in the document.

Al: So you're suggesting there will be a data plane measurement in progress, while the control plane is being attacked.

Scott: something like that.

Al: That's the approach we took with the benchmarking of flow export, which had a diagram a lot like this. We were running a 2544 throughput test and then turning on netflow and seeing what the impact was, and this is a similar thing to that. An RFC written by Jan Novak (RFC6645 <https://tools.ietf.org/html/rfc6645>). That would be a good place to look to get some of the details.

Scott: just a clarification. CPU utilization doesn't mean anything. Even if it were a useful value, it doesn't mean anything, because what's meaningful is what the effect of the utilization is, so it can be 93% utilization and still dealing with everything, and 2% utilization and it's not. So, you don't really know. By itself the CPU utilization is a non-useful value.

Al: Another point is, there are some devices that keep reporting 100% utilization all the time. It's a mistake in the calculation apparently.

Shishio: If an attack happens the CPU utilization increases and the controller cannot deal with the control plane protocols.

Marius: Just wanted to say that this is a very interesting setup and a good start. As Scott was pointing out, the CPU utilization cannot really be used, being a white-box metric. But the bits-in, bits-out is useful. What I think should be defined is what's the expected traffic, in order to calculate the impact of the attack on the traffic. If you have a threshold, you can calculate the impact of the attack and I think that's good.

Al: We'll look for more commentary and further development on the list. Thank you.

Sudhin Jacob: Do you have a traffic with TTL equal to 1, because that would overload the CPU? Do you have a test methodology for that?

Shishio: I will update some of the test traffic details.

8. Considerations for Benchmarking High Availability of NFV

Presenter: Taeho Kang (representing the authors, Kim and Paik, KT)

Presentation Link: <https://www.ietf.org/proceedings/94/slides/slides-94-bmwg-5.pdf>

<https://tools.ietf.org/html/draft-kim-bmwg-ha-nfvi-00>

Scott: I like this idea of having a Considerations document. Trying to fold some of the things we've talked about in the Considerations. I think a considerations document is a useful tool and you may wanna consider having a set of three documents in the base set: Considerations, Terminology and Methodology.

Al: We actually have had that on occasion.

Scott: I did a quick scan of the RFC index and didn't see considerations.

Al: You may remember, there was someone working on multicast address problems. Scott Poretsky had a Considerations draft for the OPNFV data plane benchmarking (<https://tools.ietf.org/html/draft-ietf-bmwg-igp-dataplane-conv-app-17>). At that time, Russ White was working into the control plane for IS-IS and Scott had to make a case to do data plane instead, so considerations came first. But the considerations draft got folded into the other two in the end. So, probably that's why you didn't find it in the RFC index. It looks like I'm starting to remember everything now.

Bhuvan: Quick suggestion, it's worthwhile to consider the hardware configuration. Besides capturing the physical to the virtual CPU and memory mapping, it's worthwhile capturing hardware parameters as well. I think there is another Standards Developing Organizations (SDOs), ETSI, they are recommending the same thing.

Ramki: There's some work happening in Openstack called enhanced platform awareness (EPA). It brings out all the hardware capabilities and what is interesting, it talk about every intel processor capability. I think we do need to bring this in. Another point on which a open source project is starting, is capturing subtle differences between different servers and across different manufacturers and different generations. Those are additional aspects to consider. That would make this draft complete and concrete.

Al: There's some overlap between KDI considerations and OPNFV considerations draft we're thinking about moving to WGLC. So, there should be some references between these drafts. We'll have to pass that back to the authors. We need to build on each other's shoulders.

Ramki: I think this is still focused on High Availability (HA).

Steven Wright: Two comments on this. Firstly, thank you for bringing this in. I think this is near and dear to many operator's hearts. On the magical figure, I think there's typically other figures that get associated with HA things like the scale of the failure. How many subscribers are affected if there's an outage and that kind of thing. There are probably some other dimensions which should be added to the considerations list. I think you probably you should be careful to attach that to the end-to-end service. So there is a bit of mapping between the service objective of availability and the components underneath, and that's often non-trivial.

Al: I think the essence of that comment is that some traditional telco figures of merit may be addressed differently in the cloud.

Steven: Yeah, for example many times you think of scaling the VNFs in terms of how do you deal with capacity issues, but you may well want to keep the size of the chunks small, so you minimize the impact of the failure. There's lots of ways. There are many tools that can be

used. I think this is a good area for a considerations document. I think it will have impact for two things, the design patterns for VNFs and the design patterns for services.

Al: I think another reference that could help is the reliability working group at the ETSI NFV. They produced a document that tackles many of these topics. So, the authors should take a look at that and maybe some of those features can be incorporated here.

Taeho: OK, I'll tell the others.

Lingli Deng: Speaking for my colleague, who is leading a project on High Availability in OPNFV. I think there is quite a bit of overlap. They have more metrics being defined there. We have put out a requirements document for HA, including benchmarking considerations. So, I am wondering if we can connect the two projects.

Al: It's a great suggestion. In fact I was trying to do that with the ETSI.

Scott: This sounds like a functionality test, which is not within the scope of BMWG. It doesn't mean it's not important. It's just not within the scope of this working group.

Al: That's right. We're testing a performance characteristic that can be compared, and if you're checking if something it's working or not that's a different scope.

Al: Failure over time would be the benchmark. So, you would be able to compare different systems, different strategies. Measuring that time can be considered a benchmark.

Steven: I think you need to split this at least in two pieces. One being the detection time, and how to detect different types of things; and there's the switchover time or recovery time. You tend to have numbers for different things there.

Al: Thank you for being the lightning rod for the others.

9. VNF Benchmark-as-a-Service, VBaaS: Problem Statement and Challenges Presenter: Robert Szabo'

<https://datatracker.ietf.org/doc/draft-rorosz-nfvrg-vbaas/>

Intro to a topic presented at NFVRG in IRTF (Wednesday Morning, 502)

Robert: ...You cannot scale-up VNFs on demand. You have to stop your VNF first and then add resources, CPU, memory.

Al: Actually I saw a demonstration where they were doing exactly that. A running VM, adding CPUs, other resources to it. So, nothing's static in this world.

Robert: Yes, definitely. But for the current technology you have to have a bound to the resources and a bound to a certain performance. An interesting question is what are the performance characteristics related to the different deployment options.

Bhuvan: Quick clarification question. Are all these measurements done during the pre-deployment or during deployment phase?

Robert: In the pre-deployment phase.

Bhuvan: One thing I would like to suggest, it's not only the VNFs throughput, latency and jitter, it's also important to note the deployment time. These aspects are also important to benchmark. Only then you will get a complete characteristic of the service performance.

Robert: OK, thank you.

Bhuvan: How will you measure this throughput, latency. Is it through the active traffic or through a passive measurement?

Robert: I have a figure in the next slides.

Lingli: I think this is a very interesting proposal and would love to try it. But, I assume that there is an assumption here that all VNFs and benchmarking are using the same type of resources. I doubt the necessity of doing this.

Robert: The whole idea is to support this multi-deployment environment and different hardware acceleration options. A benchmarking as a service standard should help evaluate different deployment options, different hardware acceleration solutions, different CPU or different execution environments.

Steven: I'm looking forward to more details. The concern I have is that we probably need to scale this up in terms of types of tests and types of benchmarks that need to be considered. You have to start somewhere, and what you've got is fine, but it occurred to me that for a lot of VNFs, throughput etc. may not be the appropriate metric. So, if the VNF it's HSS, transactions into the database might be a more relevant metric. Second, I think a lot of the functionality you are trying to benchmark is the scaling ability of VNFs then you are getting multiple extra tests that you need to instantiate. So, again there's more complicated testing arrangements that may be needed to support this. You have to start somewhere, and what you've got makes perfect sense.

Al: All throughout this I was thinking: what about more than one VNF? This is very often a shared environment.

Ramki: Real quick, we have been working on VNF performance monitoring. And I think this is answering some of Steven's and others' concerns.

10. 25 Years of work for BMWG

Presenters: Al Morton, Kevin Dubray, Sarah Banks
I'm sure there's a slide deck somewhere.

Action for Al: Did anyone take a photo of you/Scott when you presented him with his custom work bench? if so, please share to the BMWG mail list. Thank you!

Al: We're celebrating the fact that we have been doing this for 25 years, and the reason why we've been doing this is Scott Bradner. Scott founded this working group and I've tried to ensemble every WG chair that's been part of this to say something here today.

Kevin: How to frame Scott? A Master of time management (considering the diversity of his work). Scott has defined the vision of this WG, not just provide tests but offer insight and understanding, raising the bar for networking technology. Delegation is another key attribute. Scott has enabled lots and lots of people. Scott has always been here to guide us toward where the community needed us to be. As a measure of that success, our WG has produced more than 34 RFCs, and judging by the progress tonight there's a lot more on the way. Efficiency in communication, another key attribute. Putting the focus where it needed to be is another great quality. Making the difference in the big picture. Scott has got that ability since the time this WG was founded until today. He understood really what makes a meaningful benchmark. Who are we writing these benchmarks for. Not just a test that you can execute, but a test that yields insight, a test that is a basis for comparison. Scott has always been a shepherd and a watchdog, to make sure that we are working on things useful for our community. What are some of the side effects talking about the big picture? It's not just benchmarks, it's also understanding what goodness means. Benchmarks as a way to make our community stronger and better. In closing, I would really want to thank Scott for his contribution to our working group.

Scott: Thank you Kevin, I appreciate that tremendously.

Sarah: That presentation is a hard act to follow. And Scott, you are a hard act to follow. Through my time at the IETF, you had an influence on our entire organization, and BMWG wouldn't be BMWG without you. Thank you very much for your contribution to our working group.

Al: Scott, you've probably got plenty of awards and letters of appreciation, and when I thought about doing this I was asking myself what can we give you that will fit someplace other than your walls, and that you might see once in a while and think of us fondly. So, it's a tiny load bench inscribed with: "The BMWG Recognizes Scott Bradner Founding Chairman 1990-1993 In our Hearts and Minds Forever".

Scott: Thank you very much. I appreciate the sentiments expressed here. I don't expect to completely go away. I will be on the mailing list. I've just resubscribed with my non-Harvard email address a few minutes ago. So, keep up the good work. It has been very important to me. This WG was the very first thing I did at the IETF, and I suspect at some level it will be the last, because it will keep going for a very long time. So, thank you very much.

- End of session -

Addition to the notes: The entire WG would like to extend thanks to Scott Bradner. I know we piled it on thick in Yokohama, but I've had time to reflect since that day (for which I didn't know Al had that surprise planned, and I was brief in words!) but as I peak into the archives of BMWG, the fingerprints of "Scott" are all over it, and in an overwhelmingly positive way. Thanks for getting us started, and keeping us going, with your reviews, your comments, your feedback. --
SB