
ICANN85 Mumbai | CF – GNSO: ISPCP Outreach Session (1 of 2)
Monday, March 9, 2026 – 15:00 to 16:00 IST

DEVAN REED

Hello and welcome to the ISPCP Outreach Session 1 of 2. Please note that this session is being recorded and is governed by the ICANN Expected Standards of Behavior, the ICANN Community Anti-Harassment Policy, and the ICANN Community Participant Code of Conduct concerning Statements of Interest.

Please observe the following guidelines to participate in this session. I will also post them in the chat for your reference. Only questions posted in the Zoom chat identified as a question will be read aloud during the session as time permits when directed by the Chair of the session. If you wish to speak, please raise your hand in Zoom or otherwise as directed. When speaking, please state your name for the record and speak clearly at a moderate pace. I will now hand the floor over to Lars Steffen.

LARS STEFFEN

Thank you very much, Devan, and welcome everyone to the ICANN85 ISPCP Outreach session. I'm Lars Steffen. I'm the current vice chair of the ISPCP. So, welcome for joining us today. With me is also Philippe Fouquart, the current Chair of the ISPCP, and also Santanu Acharya, he is the current ExCom officer of the ISPCP. And on my right-hand side, please welcome Ram Mohan. He's the

Note: The following is the output resulting from transcribing an audio file into a word/text document. Although the transcription is largely accurate, in some cases may be incomplete or inaccurate due to inaudible passages and grammatical corrections. It is posted as an aid to the original audio file but should not be treated as an authoritative record.

current chair of the SSAC, and also Matthias Hudobnik. one of our speakers in the first session today.

Just to give you a brief overview of the day-to-day, so you're with the Internet Service Provider and Connectivity Provider Constituency at ICANN. We are part of the Commercial Stakeholder Group within the GNSO. And if you want to learn more about the ISPCP itself, the work that we are doing, the dates of the next meetings, please visit our website ispcp.info. Next slide, please.

Just to give you a broad overview, the ISPCP and ICANN, we represent internet service providers, access providers, connectivity providers, internet exchange points, and with this many of those who are operating resolvers for the DNS. And those are the guidelines in which policy development processes we are engaging with and we are participating in, and of course, supporting the multi-stakeholder model at ICANN itself. Next slide, please.

So, today we will start with this first session Focusing on AI and Machine Learning in the Context of DNS Security. This will be covered by Matthias. He just did point out that he's a former Europol data protection manager. So, just to be precise on this. And we heard from Laurin that he's not able to join us today due to his health conditions. So, we send out best wishes to him. And unfortunately, he's not able to join us today.

On the next slide you will see what is on the agenda after the coffee break. So, at 4:30, we will continue with Universal Acceptance, Digital Inclusion and the AI Factor, and what are the implications

on building an inclusive internet which will be covered by Ram Mohan. Nabil will join us later. And last but not least, we also have Nitin with us. He's the ISPCP representative at the Universal Acceptance Expert Group and also very actively working on other policy development processes that the ISPCP is working on.

So, thank you for being with us as well. And before we dive into substance, and one last slide I would like to share with you is, if you're staying at the ICANN meeting, not only for today, but we'll also be there tomorrow or on Thursday. So, tomorrow there will be the regular ISPCP membership work session in the morning at 10:30.

So, if you are interested in participating in this as well, you're warmly welcome to join us. It will be in the same room as today. And on Thursday, we will have the joint meeting with the Address Supporting Organization. And that is scheduled for the afternoon at 13:15. And before I hand over to Matthias for the first presentation, I would like to hand over to Santanu, who would also like to welcome your dear guests today.

SATANU ACHARYA

Thank you, Lars, and thank you all for joining the session on AI and machine learning in DNS security. As many of us know that DNS ecosystem is constantly evolving, so are the threats. So, with the scale and speed of today's internet, detecting, abusing and responding to incidents quickly has become more challenging than

ever. So, we'll learn today that how we can take over the threats.
So, Lars, over to you.

LARS STEFFEN

Thank you. So with this, let's dive into substance and give the floor to Matthias. Matthias, the floor is yours.

MATTHIAS HUDOBNIK

Yeah, hello, everyone. My name is Matthias Hudobnik. I'm a legal engineer and senior AI and data protection consultant and advisor. And as Lars rightfully said, I formerly worked at Europol, and I'm also an SSAC member. So, I'm very happy to give you a presentation on behalf of the SSAC about AI and also DNS abuse. First point in time snapshot. So, the AI landscape, as we know, evolves weekly or let's say daily.

So, what we share today reflects what we know now. And also, periodic assessment is built in, or let's say into our approach. And also secondly, this is not a comprehensive assessment. So, it's a baseline to inform community discussion and guide future SSAC work. And we are also setting the table, let's say, not serving the full meal, just that you know it. Please to the next slide. Okay, a little bit on agenda.

So, I will start with something often missing from AI and cybersecurity talk. So, actually mapping AI to the DNS protocol stack, so to the ecosystem, not just to domains. And then also going to a bit of definitions, then data with some honest caveats,

and then also a bit into the dual use, so offensive and also defensive capabilities, and crucially also the human impact and what the community can consider doing.

And also, last but not least, I would like to briefly inform you about our current work party. Please to the next slide. So, I would like to start with some kind of DNS attack surface and artificial intelligence. So, you see here, in my opinion, this is new and I believe also essential. Most discussions of AI and DNS abuse focus on the domain names and also phishing websites, but let's say AI also intersects with the DNS ecosystem, as I said before.

And here, I try to emphasize at four distinct layers, and also each has, let's say, some different attack and also defense dynamic. Starting with the registration layer, so we have like the extensible provisioning protocol, so registrar's accounts, WHOIS RDAB data, payment processing, and also AI here enables bulk lookalike registrations. We have also captive circumventions and also synthetic registrant identities here in the problematic area related to some defense things.

So, we have like defensively machine learning, screens registration patterns, and also which is trying to detect anomalies related to the EPP commands. When we go to the next layer, so to the more zone and authoritative layer, we have zone files, and then also we have DNSSEC signing and NS records. And also here, AI can automate subdomains, take over wire dangling DNS, and also DNS shadowing under compromised zones and zone enumeration.

And when you look into the, let's say, AI defensive side, and this is, let's say, a concrete opportunity that AI could help automate DNSSEC deployment which is already done also and has dealt at around 30% of zones. And also, AI-assisted zone integrity monitoring is another real possibility, let's say, where you could be able to implement these on a larger scale. And then when we go to the resolution layer, we have recursive resolvers which do some caching.

So, we have DOH, DOT, and also DNS filtering. And also here, AI generates domains or could be used to generate domains automatically and also enables cache poisoning at scale and also can fingerprint, let's say, for example, resolver configurations. And this is where also the big defense numbers live. So, we have like, for example, some DNS filter report, which processes 200 billion plus queries daily using AI powered filtering. I will give you then later also the list where you can read about it.

And then if we go to the application layer or user layer, so we have browsers, we have email authentication, we have also certificate issuance, and then also, here you could use AI crafted brand impersonation on lookalike domains and also AI generated phishing emails. So, IDN homomorphic attacks are also quite well known. And on the more defensive side, AI powered email authentication, then you have also certificate and transparency monitoring.

And also, in the browser itself, you have some browser level warnings. And yeah, so to sum it up, this layered view matters, because different ICANN community actors operate at different layers. So, registries and registrars own the registration layer, DNS operators own the resolution, and also the policy community shapes the framework that govern all of them. Please to the next slide.

So, this slide should also give you a bit of an overview. So, I titled it Defining AI Spectrum. So, rather than drawing a hard line, I try to present or present a four-point spectrum from rule-based automation through conventional machine learning, deep learning, and also generative AI to Agentic AI. So, the honest point at the bottom matters, so modern DNS security increasingly uses deep learning.

We know transformers, embeddings also for detection, and the boundary between conventional machine learning and so-called more newer or new AI system is generally blurry. You see here, when I talk about this, then we are between conventional machine learning, deep learning, and generative AI, and then which is more going into the Agentic side. And an audience member could also know or could reasonably say or challenge whether, for example, domain generation algorithm detection using deep learning established machine learning or a kind of more new AI system in between deep learning and also generative AI.

So, we acknowledge this openly rather than pretending the boundaries or the boundaries very clear. Please, the next slide. And here coming a bit to the more analytical heart of the presentation. So, what the data shows and does not show. So, it is split into two columns, and the right column is just like as important as the left. So, on the left side, we see what the evidence supports. So, for example, Akamai has a 300 plus AI bot traffic increase in its report weight specifically refers to AI bot traffic, or there are also 55% phishing click-through rate is from a controlled academic study, originally from Harvard researchers.

Then we have also like DNS floods where we have DDoS vectors and also DDoS attacks targeting AI compromise specifically surrounded around 350% monthly over month. And when we look more on the right, what we cannot yet attribute to AI per se. So, the, for example, 47 million total DDoS attacks which are driven by IoT botnets, cheap compute, and also fraud as a service. So, it's not only to AI per se, which is mentioned in the Cloudflare report.

Then we have also another example in the DNS filtering report where it says there are 65+ percent new domain threat patterns predated AI by over a decade, and also new domains have always been disproportionately malicious to be very precise. And then we have also 1 to 175 malicious request ratio, which is generally threat metrics with no evidence of AI specific causation. So, that's why we always need to be very cautious when we show also this data.

And the 45% phishing rather itself, which sits on the left also belongs here because it was an academic study and not a field measurement. So, just so you know, it's called the hiding study, it's often called the Harvard study through an RX if preprint, and they tested four groups with about 25 participants each, and also here was generic spam emails, got 12% click through, and then human experts spear phishing got like 54, and also fully AI automated spear phishing also got like 54, AI with a human in the loop got 56.

And here, the headline is also that AI performs on par with skilled human social engineers, but at near zero margin costs and also essentially unlimited scale. So, a human expert can craft maybe a dozen personalized spear phishing emails a day and then the AI system can generate a thousand. So, that's also the threat multiplier that matters for DNS abuse.

So, not the highest success rate per email, but the same success rate across orders of magnitude more domains and also targets simultaneously. And as I said, I should also note that it was like a short study there were like 101 participants which is a rather small group and a sample, so this is also a proof of concept, not a production measurement. Yeah.

BARRY LEIBA

This is Barry Leiba, SSAC. I note in particular the word yet up there and tie that in with the stuff you were saying at the end, comment on this please, it's my view that anything that isn't being done by AI

today soon will be and AI scales everything up to much higher rates than we've ever seen before. Any comments on that? Yeah.

MATTHIAS HUDOBNIK

Yeah, thank you very much for pointing this out. So, yeah, I agree with you 100%. And we see even like we had some lightning talk session also on Saturday where I made this comment, I don't know which of you have heard also for example about the AI assistant Open Claw which you can deploy. It's an open-source assistant you can chat with it via your Telegram or WhatsApp.

You need to have some technical capabilities to set it up, but you could really also use it on large scale, like for example on a Mac mini. And you can open ports with it, you can really do a lot of things, and this is something which is crazy. So, this guy joined also OpenAI, he was a developer and Austrian. And so, they will also integrate this, and this could be also abused, obviously, for defensive, but also for offensive purposes. Yeah.

PHILIPPE FOUQUART

If you don't mind.

MATTHIAS HUDOBNIK

Yeah, please, of course.

PHILIPPE FOUQUART

By the way, it's interactive. Philippe Fouquart from Orange. Ss a follow-up to that question. Specifically on domain name

registration, just to make sure at least I understand what we mean by that. It seems to me that using regular expressions for that has been around for a while, but as Barry was saying, it seems that AI may take, just to pick up an example, brand-login.gtld, we can do this with regular expressions and we've been doing, well, not we, but people have been doing that for ages.

It seems that with AI, you can sort of take it to a different scale altogether and figure out that UDRP is associated with ICANN, for example. I'm not making this up, by the way. icann-udrp.support can be registered for a malicious, is that sort of the idea behind this.

MATTHIAS HUDOBNIK

Yes. Thank you, Philippe. Exactly. Let's say the, I would say, the direction in which this, I mean, nobody knows it 100%, but the direction is exactly like that you can automatize this and you can also make really a lot of domain registration on large scale and automatize it.

And it's easily doable because, as I said, even with this kind of assistance, you can already now say, whatever, if you give all the data, book the whole flights for me, do bank transfers which is also there are a lot of like critical security issues in it, but you can really use this Agentic AI for doing exactly this large scale offensive but also defensive, use it for these purposes.

And please to the next slide. And here again, I try to show a bit AI enhanced DNS abuse attack mechanisms. And here again, I try to

be very careful. So, I said documented partially evidenced and correlation. So, let's say the three DNS, this kind of three categories, I try to be as precise as possible, but some things we just don't know yet and we just assume.

It's based on threat intelligence reports which I was referring to, and I will also show it and at the end the list so you can also have a look on that. So, starting with the registration layer abuse where I said it's documented. So, we see here also in reports that AI domains increase 50% year over year, and 65% plus of threats are newly domains, and also attackers exploit cheap ccTLDs like .fo which had like 27 malicious traffic.

So, also here, AI bulk generates look alike with synthetic WHOIS identifiers for example, and it's going in this direction. And then when we go to the next point content plus domain combined attacks, so here, as I said, it's partially evidenced. So, let's say there are anecdotal reports of AI-generated brain accurate content on lookalike domains, but no field scale measurement yet for the combined effect. And the qualitative shift is real, but the quantitative evidence at scale is emerging.

And then we have some DNS infrastructure attacks, which are a kind of correlation. So, DDoS volumes are clearly rising, but also DNS floods are the top vector. But also attributing this specifically to AI is a stretch given IoT, botnet, and also other factors. And again, like a Agentic AI could chain zone enumerations or also subdomain takeover and DNS shattering. But yeah, autonomous

DNS attack chains are not yet widely documented in threat intelligence reports.

Again, we will see how the future will evolve, but it seems that it goes in this kind of direction. Please, the next slide. So, in here, I tried to go a bit on the positive side, but we say AI powered DNS defense. And so, again, each defense now also carries explicit limitations. And so, this resolves a coherence problem in earlier versions where defense side was uncritically positive. And so, here again, I was trying to point out these three categories like DNS layer filtering.

So, again, when you look into reports, you have 200 billion plus queries per day. So, 10 days faster detection limitation is obviously false positive impact unknown at this scale. And also, when AI blocks a domain, legitimate business can be affected also. And also, like model governance, frameworks are nascent. So, then the next point, autonomous DDoS mitigation. Also here, 47.1 million attacks blocked, autonomous means automated response not fully independent judgment.

So, here also you need to carefully look at the data what is exactly meant and what is maybe a bit of also marketing. So, detection thresholds are human configured and this capability depends on, let's say, cloud scale providers, and also on-premises defense cannot match it. And then when we come to the third point, like kind of predictive domain intelligence, and you have also here

some kind of pre-attack flagging of malicious registrations. But also, here you have some kind of limitations.

So, I want to be honest, some of these techniques like domain generation algorithm detection and also anomaly scoring are like conventional machine learning techniques which are deployed for years, and also there are available, but calling them, for example, new AI defense also overstates the novelty. And yeah, newer generated AI power detection is promising, but less proven at scale for now. Please to the next slide.

So, now the question, why this matters. And I put it like human impact. So, as we know, behind every statistic is a person, business, or also institution which is harmed. And so, I try to bring this a bit up where you say, okay, who is actually affected? And here again, you will see then the report. So, the DNS filters 2026 report found that the average internet user encounters 66 threats per day, up from 29 the previous year.

So, that's a daily personal experience of the threat landscape we are discussing. And also, the bottom section addresses user-facing protections often absent from DNS policy discussions when we see like email authentication or browser level DNS security and also DNSSEC for end users. And yeah, as I said here, so consumers have very often problem with credentials theft, while AI crafted brand impersonation, but also small businesses, like you have domain hijacking, you have fraudulent storefronts, and also government agencies are struggling with it.

We have targeting phishing and DNS infrastructure attacks, and again, most of the things you can also do without AI, but AI helps you to do it much faster and on a bigger scale. And then also brand owners, which get an overwhelming volume of AI-generated lookalike domains. Please to the next slide. And yeah, here is a bit resources I used also for this presentation. So, we have a lot of relevant work, and this is not exclusive, so you'll find even more.

But you have DNS abuse prevention mitigation practices in MOCK and ABWSG, then we have like a DNSFilter, they have reports from 2025 and 2026 annual security report, we have a lot of like information at Cloudflare, but also Akamai, Netskope, Abion, and various others. Please the next slide. And here, I try to make one slide which is called like consideration for the community.

So, this is new and also directly addressing the question. So, what should we do Monday morning even though today is Monday, maybe Tuesday morning. So, these are not prescriptive recommendations, but they are areas for community discussion. And so, I try to again, split it in four areas where you can say, okay, registries and registrar, so what can they do? They can try to have kind of AI enhanced registration screening, analyze the EPP patterns registrant, identity signal domain strings at point of registration, and a hard question also should, let's say, contractual obligations like the RIR address AI generated bulk abuse specifically for example.

When we look on the DNS operator side, so here again, invest in AI powered resolution layer defenses, and also a question could be evaluate autonomous DDoS mitigation if it's not already done, and then consider DNSSEC automation, which can be also done without AI, obviously, but also as an AI-assisted deployment opportunity for example, and also this is where AI could directly help and not just threaten.

And then, the ICANN policy community, so assess whether current frameworks like the RAA obligations are calibrated for AI speed abuse, for example. If attackers can register and deploy hundreds of domains in hours, are our responses timely and also adequate could be also a question. And then finally, when we think about cross community sharing, so also coordinate between different other areas, like for example, first or DNS, or the mark or just try to build synergies and address, let's say the adversarial AI question that we haven't fully explored yet.

And if our defenses are AI powered, how do we secure those AI systems also from manipulation? Because at some point you can say, look, whatever, a human can attack an agent, but also an agent could be attacked by a human, and also an agent could attack a wire agent. So, there are a lot of questions, actually, which are not directly related to ICANN, but still in one way or another, touching upon it. Yeah, please the next slide.

Yeah, very quickly, so we also started with a work party which should be a first educational document where we say we try to

define AI scope and also analysis for SSAC's remit what is important, trying also to catalog offensive and defensive AI use cases in the DNS ecosystem, and then also try to make some little analysis related to current and projected security impacts.

And here also, we have some kind of guiding principle we see, it should be fact-based, neutral, dual use emphasis, point in time perspective, and also correlation is not like causation. And this brings me to the end of the presentation. Thanks a lot, and yeah, feel free to ask some questions. You can just put the next slide on, it's the last one.

LARS STEFFEN

Okay, first of all, thank you very much for the presentation. And we already have the first question. Yes, please.

JOYSA KAUSHIK

Hi, thank you for the presentation. My name is Joysa Kaushik, for the record. So, as the data suggests, it was mentioned in one of the slides that 6%, even more than 65% of the threats originate from the new domains. So, we have the upcoming round of new gTLDs, so do you think that it will increase?

I'm assuming that it will increase that percentage. So, are there any mitigation or rather preventive strategies or policies going on in the ICANN ecosystem just to tackle the new gTLD round, which is around the corner? So, thanks.

MATTHIAS HUDOBNIK

Yeah, I mean, I think it's a very good question. So, as I said, there was already abuse before, the thing is just that it's now on a larger scale and much faster, so much speedier. So, I think, again, as I tried to present at the third or second slide, I think it's a layered approach.

So, you could try to make some defense or mitigation strategies on all these layers, and it might, or not it might, it's meaningful to do that because then you have a kind of full chain which you would tackle. So, I think every kind of party in this layer should do their homework. That would be my answer.

LARS STEFFEN

Thank you very much. Then.

PETER JIA WEI CUI

Hello, and I'm Peter Jia Wei Cui from Taiwan. I was sponsored by TW to participate in ICANN, this is the first time in ICANN, and really find the discussions on AI usage on the DNS registry is really interesting. But for me, I think one of my concerns is based on the Taiwanese experience in dealing with the online fraud issues on social media.

So, basically, the cases like this, so Facebook, one of the main social media operators in Taiwan, they are facing one risk is that in Taiwan, a lot of fraud messages is being transmitted from

Facebook. So, Facebook is trying to take down the content from the social media.

And under the pressure from Taiwanese governments, Facebook is actually making more stricter rules for the AI models that automatically filter the messages and taking down the content, but actually, having the side effect to remove the content that is not by the fraud and the scanners and the malicious actors even trying to abuse the Facebook's report systems to take everything as a fraud so that it actually mess up with the automation systems that Facebook provide to prevent all these kind of things.

So, I'm wondering if during, because I think that the DNS attack prevention can really be incorporated in AI, but is there any other solutions to dealing with these kind of potential risks that are probably false positive or false negative issues happening during the automation on filtering and analytics on things?

MATTHIAS HUDOBNIK

Yeah, that's a very good question also. So, I will try to answer it on a personal capacity because it's more related to content moderation, not really falling into the remit of ICANN per se. So, for me, there are two things. So, one thing is you talk about the governance model where you say, okay, an AI has false positive and you need to try to mitigate it.

And here again, you need to have some, and Facebook is trying to do that, you need to have like a two or three stage approach where

you say, yeah, you have a human in the loop, you have more people which look over the data, and also oversee the AI system. And then also on a technical level that you really look, how is the system, let's say, you have various stages, so when you kind of adopt it, then you test it, then you deploy it, and then obviously, when you deploy it, you also look at the results.

And obviously, it depends on which data it's trained in terms of bias. And the big companies are doing that, but it's very hard. And the second part of your question is more related I think, through regulations and laws that there can be sometimes or there is the risk that obviously for the purpose of security that you make a takedown, and then you have here the question, let's say, free speech versus maybe, yeah, harm.

And then the question is how the laws are applicable related to the jurisdiction also. Like also in Europe, we have the Digital Services Act, the Digital Markets Act, where also this big companies are falling under it, and they also need to make measures. And also, there is a lot of pushback because in the US, it's less strict than for example in Europe.

And so, we have already there the thing where the filtering they say in Europe, you have different or stricter rules than in the US, and they need to find a balance and even in other countries, I'm not aware. But again, it's quite complex, so you have the technical thing, the governance, and then you have also the legal aspect

which is coming into it. I don't know if this answers your question, but.

PETER JIA WEI CUI

I also want to add on that because you mentioned about the issues about freedom of speech and expressions. There's one issue that in Taiwan since the ccTLD managers, TWNIC, actually had a response policy zone which is called RPZ to seize the analytics and the revolvers for the DNS systems. So, it actually creates a way to control the DNS systems in Taiwan to not having the users to connect to certain websites.

And it actually has some false positive issues because one day when the police ran into some bars and some gay bars like that, it actually creates harm and risk for the gay and LGBTQ community in Taiwan to actually being seized, and these limitations are really happening.

MATTHIAS HUDOBNIK

I mean, I fully agree. I think you have in general also, but that's a general problem. So, all the systems are trained on some kind of data, and then you have various bias mitigations you should deploy. And then very often, groups from whatever, which are not represented on a large scale, they might get discriminated in one or the other direction. You have always this kind of risks. Okay, please, next.

UNKNOWN SPEAKER

So, hi. [Inaudible – 39:43] for the record. I'm a NIXI Fellow here. So, just wanted to know, don't you think we are too late in AI securities as we are facing the problems every day, and now we are trying to solve it. Until then, we're going to solve this issue, new problems will arise.

For example, in near future, we're going to use bots or AI to buy the domain and to create content for it. So, can we implement something like know your bot just as we implement know your customer service in authentication services? Thank you.

MATTHIAS HUDOBNIK

Yeah, very good question. So, I think whatever, what other opportunity or what other things can we do than just try to mitigate and also work on that? So, I think the thing is more like learning to learn, which is very important nowadays, even for every person of us. You should integrate the things as much as possible, be always like skeptical. And the same thing with offense versus defense.

So, very often, in the offense side or criminal side, people are ahead the defense side, but still people try to catch up and do their work, and it's always this cat and mouse game, and you have no other opportunity than just say look we try to find a solution as good as we can.

And again, I think the most important thing is that like, as I try to point out this layered approach, so in all the different layers, you have stakeholders which should do their homework, and then as a

whole it will be better and safer. But there is not an easy answer to that, I would say.

PHILLIPE FOUQUART

Yes, please.

SRIJAN RAI

Yeah. Hi, my name is Srijan, I'm a NextGen. And my question was sort of a follow-up to the previous question in a way. And the question would be somewhat along the lines of is ISPCP sort of considering the standardized identity layer, is it under consideration? And what would be the reasons for not going for it?

Because I'm sure there's a reason for not thinking there would be stakeholders who may not want this specific bot-based identity or identity systems which recognize. And does this also involve some cooperation with IETF? So, are there discussions on similar, like with IETF as well?

MATTHIAS HUDOBNIK

I think I will hand over the question to you related to ISPCP.

PHILIPPE FOUQUART

Thank you for that. I'm Philippe Fouquart. I'm the ISPCP Chair. That's a very broad question. I think if I may rephrase it, you're asking whether the policy development process can sort of accommodate the sort of identity management that you're referring to. I think in all fairness, at this point, the answer is no,

but more broadly, I think there's a sort of existential question which is mine as well as to whether a policy-based solution is actually the right answer to a problem that's generated by AI.

Take domain name verification as it is considered by PDP 1 that's being launched just at this meeting. Surely, well, off the top of my head, if you consider the means and you refer, for example, of checking whether the same registrant for which a domain name has been considered as being abusive, the check consists in reviewing all the domain names that come from the same registrant.

At some point, and by the time the contracts are being updated in 2030 or something, there will be other means of abusing domain names. So, it's a rat race and a losing one for that matter. So, it's almost a sort of existential question for the PDP, I think, at least the gTLD one, and it's a good one. And we don't have the answer.

I haven't a sense that we should consider the end rather than the means in that respect and make it possible for those that would have to apply those provisions to use AI if they want to, as it adapts moving forward.

But whether as a whole, the PDP can or the picket fence, should I say, the remit within which policy development can be developed would have to change to accommodate the sort of framework you alluded to, it's much broader, probably something that we can discuss within the ISPCP, but it's probably a more ambitious goal.

But coming back to Lars' introduction, by all means, join the constituency and suggest anything that you would consider relevant for this group and further for ICANN. Thank you. We have one question over there.

UNKNOWN SPEAKER

Mainly, I've got the reply from the discussions which are happening right now, but still I'll ask that since AI can be used for both to create DNS abuse, like phishing and creating fake domains, also through ML, AI, and deep learning practices, we can also resolve these threats. So, how mainly are ICANN communities ensuring that AI is being used responsibly for the DNS security?

MATTHIAS HUDOBNIK

Yeah, so I mean, I think I can just repeat what I said. So, I think the most important thing is, as I pointed out, like that on every layer, the party is trying to do what they what they need to do. So, on the registration layer, on the zone and authoritative layer, on the resolution layer, and also like on the application layer per se.

And then also, I think the most important thing is that we really try to adapt and also foresee a bit of things, even though it's very hard to foresee. I think that's maybe a broad answer, but I think that's what I would say. I don't know if you want to add something.

RAM MOHAN

Yeah, let me just quickly add to this. This is Ram. I think there's a risk inside of ICANN, this community. There's a risk of looking at all

the harms and saying, why isn't ICANN doing something about it? And the reality is that there's a lot of harms done in lots of layers, areas, et cetera. And at least my perspective is the best thing that we can do is to say, what is happening at the DNS layer, at the protocol identifiers, at the IP address, and the domain name layer?

What is the impact of LLMs and AI, et cetera, on those layers. First, we have to understand it and map it. Then after that, we have to look at it and say, is there something novel, something new here, or is it more of the old except at a much greater scale? That's what Barry was talking about earlier, right? More of the same, but except at a much higher scale.

Doesn't mean it's not a threat to security and stability, but the response to it would be different. If it is a brand new thing, then we have to sit and say, okay, how do we go look at it? But so far, it seems like in this area, the infrastructure layer that we are focused on, it appears so far to be the two vectors that we should be thinking about are scale and speed. And we ought to be thinking about how to respond to massive scaling up at a massive increase in speed.

NABIL BENAMAR

Can I add something to it? This is Nabil for the record. So, I gave a lightning talk on Saturday about how AI is changing the landscape. And the focus was on the -- how good or bad are we really prepared to this new generation of attacks? So, the main thing that we can add to the speed and velocity versus preparedness to combat these

attacks is the fact that we have now new tools that do not need expertise from humans to use them.

So, this is drastically modifying everything. For example, recently, two weeks ago, Kali Linux, which is the most famous penetration tool, has embraced Claude AI officially. So, now you can talk just English plain text language to your Kali Linux and ask to do the following, and it will transform everything into steps, into commands that will need deep expertise and hours to perform them.

So, now it's easier for anyone to launch these attacks or DNS abuse or any other stuff. But on the other hand, on the other side, for organizations, they are in a bad situation nowadays because first, we have lack of expertise in the field, we have a lack of talents in the field, with the employees.

And also, before they can deploy any AI-based technique or platform, they need to assess it first. So, this will need time, and this time is against the enterprises, organizations, et cetera. Thanks.

MATTHIAS HUDOBNIK

Thanks, Nabil. That's a very good point. And I just want to say two little things. So, the first thing also, because you refer to wipe coding, so, and this is exactly the thing. So, nowadays, you do not need to be a developer, you can just start wipe coding, and the

most important thing is that you really talk with the large language model.

And so, you can just put in your idea and the code will be automatically written down, and you can just deploy it. And so, that's also a new thing. And you also see that then people which are just trying to experiment and are open, they build new tools. They are not software developers, they have just some ideas and try to deploy it and do it.

And that's why I think it's very important in general that you are open to really explore new things, play around. And then, again, you can use these capabilities for the better or for the worse. So, in both sides. And very often, obviously, criminals, they are faster, but I would say also from a defense side, you can still use it. And there are a lot of smart people which also see some whatever abuse and say, look, we can tackle this also with AI, and I can just also use wipe coding to find or to mitigate this flaw.

PHILIPPE FOUQUART

We still have one hand over here and another one down the table. Thank you.

KUSHARGA SINGH

Hi, this is Kushagra for the record. While I got my answer to some extent, but my question is basically there was a mention about email security protocols, SPF, DKIM, and DMARC. So, why are we still talking about those as a baseline? Basically, attackers have

moved past that way before. Basically, there are open-source toolkits that they can use and exploit those as well. So, are we talking about setting a new set of baselines for that so that these cannot be exploited anymore? Yeah.

RAM MOHAN

Well, technologies like DKIM and SPF are still essential. You still want to secure the core layers. So, you look at RPKI, you look at DNSSEC, et cetera, they're all protocols focused on securing certain pieces. When you secure those pieces, you have predictability, and the predictability results in stability, and so those pieces, you have stability.

Now, if you have that as a baseline, then attempts to destabilize, you know, the communication flow or the transactional flow, you have a baseline that's hard to then destabilize. Then the rest of the work, then all of us as engineers or as policy people, et cetera, we can focus our work on the areas that are trying to be destabilized knowing that the foundational pieces have some stability to it.

Now, clearly, they're not universally deployed, they're not universally adopted, and they still have some problems that happen with them. So, therefore, we have to continue doing the work on DKM and SPF and all of those other technologies because those technologies provide the minimum stable foundation for secure transactions and secure communications between systems.

PHILIPPE FOUQUART

Thank you. Please.

SAANVI

This is Saanvi [ph] for the record. Given that AI models are increasingly used to generate malicious infrastructure and detect it as well, this is basically just a follow-up question that should we begin thinking about AI governance within the internet governance frameworks and not AI just merely as a tool to form new protocols and to secure the code layers?

MATTHIAS HUDOBNIK

So, I mean, I think this is already done. So, you see like various initiatives, like for example, when you look like in Europe on the EU AI Act, which is more or less like a governance framework and is going hand in hand and with the GDPR still and companies are obliged to fulfill certain governance aspects when you deploy an AI system, then you have some kind of risk mitigation where you say, okay, What kind of AI tool do I have? In which category does it fall?

Based on this category, I need to have certain safeguards, technical ones, organizational ones, but also human in the loop, explainability. So, there are a lot of things already in place. Now, it's more a bit like the discussion in Europe where they say, okay, maybe the rules are too strict and we are lacking behind of innovation in comparison to, for example, the US because the approach is a different one.

And yeah, here, it's a bit like a clash. And why also, for example, the European Commission, I don't know if you've heard about it, there is this new Ombuds kind of proposal where you have certain deadlines where the AI Act comes into force. So, you have certain stages.

And so, they try to open it up a bit and maybe try to find a way to be a bit like less strict. But we will see how it will further evolve. But this is a kind of governance based on a law, which is applicable when you also deploy some certain system. And even the company, if the company is not in Europe, and you process data of European citizens, you are applicable. So, it has quite a broad scope. So, in terms of an application of the law.

IDIL KULA

Hello, I would like to add to my colleague. Actually, regarding the -
- yes, thanks. I am Idil from Turkey. I'm a returning Fellow for ICANN. Regarding the European Union legislations and GDPR, not GDPR, but the AI Act, and all of those cells are mostly focused to the free market, actually the products and exports or imports.

So, regarding the internal governance and DNS management and all of the infrastructure, I think that there should be some more work that's breached all of those AI product-focused legislations to the DNS markets, because it is not really emphasized, it's my view. Also, there are some technical but not enforceable certificates, like ISO/IEC 42001.

This is AI software management for, like, anyways, AI management systems certificates. Well, if the stakeholders who are responsible for the management of DNS systems are inspired and responsible enough, they can deploy these ISO or IEC or their likewise certificates or checklists in their business lines. But obviously, my view is there is more work to be bridge the gap between ai market resilience to create project AI liability and the DNS market. That's all. Thanks.

PHILIPPE FOUQUART

Okay. Thank you. We are already three minutes after the hour, and I've been asked to close the session. So, I know that there are still a few questions that we have not addressed yet. My suggestion is that we take those now into the coffee break and have one-to-one conversations.

We will reconvene at 16:30 for the second part of the outreach session covering Universal Acceptance. Thank you for joining for the first part. We hope to see you back for the second part. And with this, the first part is adjourned. I'm looking forward to see you in approximately 25 minutes back in the room. Thank you very much.

[END OF TRANSCRIPTION]