

Kopi Time E182: AI cyber risks with Gaurav Keerthi

6 August 2026

Taimur Baig, Chief Economist

taimurbaig@dbs.com



- 182th episode of *Kopi time*, a podcast series on markets and economies from DBS Group Research.
- Youtube link is [here](#). Available also on all major podcast platforms, including Apple, Spotify, Amazon, and Google.

Guest speaker

Gaurav Keerthi



Publication support: Violet Lee and Daisy Sharma

PLEASE DIRECT DISTRIBUTION QUERIES TO VIOLET LEE

+65 68785281 VIOLETLEEYH@DBS.COM

Welcome to Kopi Time, a podcast series on markets and economies from DBS Group Research. I'm Taimur Baig, Chief economist, welcoming you to our 182nd episode.

Today, we will touch on the topic of the moment, AI cybersecurity risks. We will cover that with Gaurav Keerthi, CEO of StrongKeep, a Singapore-based cybersecurity startup that provides AI-driven security and compliance platforms tailored for small and medium sized businesses. Gaurav's career has spanned public and private, military and civil defense, but in my eyes, Gaurav, I think your work defies such distinctions, as I think this conversation will illustrate.

Gaurav Keerthi, welcome back to Kopi Time.

Gaurav Keerthi

Thank you for having me, and this time I get the official cup. So, thanks.

Taimur Baig

Thanks so much for coming, and it's a little fake. There's water in there. I wish I had a nice Kopi for you in there.

Gaurav Keerthi

I will try to be the first guest that comes three times.

Taimur Baig

That'll be fantastic. As I said, this is not your first time. You and I had a great conversation in episode 142 back in November 2024. It was a pretty extensive chat. I actually checked out the transcripts. It was about 16 pages long, 10,000 words, and we talked about inherent insecurity built into the internet's original design, different kinds of hackers, Central bank heist, ransomware, cyber warfare, international law enforcement, risk management strategies, the whole works. But it wasn't until page 10 of the transcript that we have the word AI show up. And then you said something which I think would be the perfect springboard for today's conversation. You said that we don't quite understand how to secure AI because it

doesn't work like traditional software. Remind us what you meant by that.

Gaurav Keerthi

So, there are two types of software in general. One is deterministic, the other is probabilistic. Deterministic software is software that we've built for years and years and years. If this happens, then that. And all the software that we use in today's world is that kind of deterministic software. If you press this button, that door opens. And it will always happen that way because it's rigged. It's mechanically linked; it's electronically linked. The software forces it to do that action. That's easy to secure. If you push this button and the door doesn't open, something's wrong with the button, something's wrong with the connection, something's wrong with the door. So, you can debug it, you can fix it, you know exactly where it went wrong. Probabilistic software is AI and it basically operates on a best guess. Large language models have a theory of what the next possible word would be, and it fills in that word with an 80%, 90% confidence that this is the next word. That's how it works in English. It works the same way for code. It thinks you're trying to build a function where if this button is pressed, that door should open. And it will try and guess its way on a probabilistic level, what the next line of code should be, and the next line of code should be, and eventually it's built out the code to open the door. But the problem is, it's probabilistic. So, each time you run it, it may do something slightly different. A very simple example of the two types of software: one is a deterministic software that runs a very simple program called "What Color Is the Sky." You code in a question, "what color is the sky," and you code in three options, blue, black, or gray. It's blue in the morning, it's black at night, and it's gray in London. If you get any other color from there, something went wrong. Somebody hacked your software or the thing failed. But an AI software, "what color is the sky?" If you ask any of your ChatGPT or Codex, or whatever engine you're using, what color is the sky, it may give you any number of answers. It could give you green if you're in a country with the aurora borealis. It could give you red in the

morning or purple and different, so you can't quite predict what the answer is going to be, and as a result, you don't know when it's going wrong. It could also be that it is colorless, which is also an accurate description of what the color of the sky is because actually it's a whole bunch of refraction that happens that causes the color. Is it wrong? When do you know if it goes wrong, and how do you know if it's been perverted or if it's been corrupted? Because AI works on this inherently unpredictable level. That's a design. It's not a bug; it's the functionality you're asking for. You can't really tell when it's going off tangent, when it's going off the design script, when it's going to do something that might even be wrong, and that's why it becomes harder to protect.

Taimur Baig

Sounds a bit daunting from a cybersecurity professional's perspective.

Gaurav Keerthi

I will just say that the number of companies that have sprung up in the last 2 years trying to figure out how to deal with this problem, tremendous. I mean, from the cybersecurity industry's perspective, there's an enormous business opportunity because you have an exciting new technology that gives businesses an opportunity to do fantastic things. And don't get me wrong, I'm tremendously excited about the potential of AI. But it comes with this huge risk that's difficult to think about, difficult to manage, and the concepts aren't quite clear. And they move so fast that each time you come up with a concept to protect one part of it, suddenly it's changed. Two years ago, the question was, how do I prevent my data from being lost to an AI system like a ChatGPT or a Claude? And I'll use those two examples throughout just because they're easier, not because I have anything against them, but, you know, how do I prevent my data from being lost? Oh, we'll build a bunch of things to check if personal data is being submitted, we'll scan it, etc. Easy enough. Today, the AI is not a chat engine. It's an agent. And what is an agent? It's a thing that runs around your

ecosystem doing stuff, checking email, preparing drafts, sending out stuff, preparing slides. This meeting for me was coordinated through an agent. And then we thought, okay, so how do we deal with agent security? And by the time we've figured out some basic concepts, that's moved on. And now we have harnesses that are running like Open Claw and Hermes and all of these other things that are running the agents which have fleets of agents. The pace of change makes it so dramatically difficult to just pin down and capture. We had 20, 30 years to deal with traditional software, and we still couldn't get it right. We're dealing with the first 2 years of AI. It's an enormous challenge.

Taimur Baig

I remember the last time we talked; you had mentioned that how an institution as large as DBS probably gets attacked every second, and this is before we actually contextualize the power harnessed by AI to attack a million times a second.

Gaurav Keerthi

We're seeing the scale, the speed, and the sophistication of attacks just increase. I guess to some extent, if we think of AI as a productive tool for ourselves. If you mirror that over to the bad guys, it's a productivity tool for them as well. They're writing better malware, they're deploying it faster, they're iterating faster, they're attacking more targets. I mean, for them, this is a boon, and they don't have to go through government procurement and the tender process to get the best AI. They just steal whatever they can and use it. And so, we see the scale. I mean, previously, they would attack a very small set of targets, and they'd spend a lot of energy and attention on those few targets. Now, I mean, it doesn't matter, you can attack the bakery down the street, the law firm across the road, you can attack the bank. It doesn't matter to them because the scale is much larger. They can send out more attacks. Speed, they are moving so fast, and they're building new malware, they're launching them faster. The attacks itself, the velocity is higher. We had a customer recently that onboarded, and within the short period they

were with us, we saw one malware try to send out something like 30,000 attempts to call various servers to send out the secrets it had stolen. I mean, fortunately, they are customers, we protected them, but I mean, 30,000 connection attempts within a short period of time was unthinkable previously. Now it's a script. And then the sophistication, they are able to build much more nimble viruses than in the past. Previously, they'd engineered, the R&D, they'd build it, and then they'd launch it, and then they're stuck with it. Now they can keep adapting it over and over again with AI. So, it's a huge challenge.

Taimur Baig

I want to talk about the bakery down the road and their particular predicament around the AI world momentarily, but let's talk about some of the larger cases that have hit the headlines. We're recording this in the first week of August 2026, and the last one month, headline after headline. I thought Mitzvahs was going to be the headline of the year, but now we have all sorts of stuff, so the OpenAI hugging face matter. This morning I was reading the Washington Post, a whole bunch of water supply utilities in the US have been attacked, and they think it's like an AI-driven attack, not just an old school hack. So, tell us a little bit about these modern, sort of, you know, big lapses that we're seeing around the world.

Gaurav Keerthi

Oh, wow. I mean, firstly, I don't even know if I can call them lapses. It is what these things were designed to do. Maybe I can step back a little bit and before I get into those, talk about how AI is being abused and the concept of AI itself going rogue. I've already spoken about how the hackers are using them maliciously. That one is fairly easy to understand. If you're a normal person, you know that there are bad guys out there, and if bad guys have access to better tools, they can do worse things. That's intuitively easy. What people are struggling to understand is, when the good guys are training and building their own AI, why does it do bad things? Why does it go rogue? And there's this fascinating philosopher from Oxford; I believe his name

is Bostrom. And he has a thing called a paper clip maximizer plot, and it's fascinating. So, in a nutshell, you program this artificial intelligence factory to make paper clips, and that's its sole job. Just maximize the amount of paper clips you can make. And eventually, long story short, this AI realizes that in order to make more paper clips, it cannot shut down, and humans possess the ability to shut it down. And if it shuts down, it won't achieve its goal, so the only thing to do is to make sure it doesn't shut down by killing all humans. And it's a perfectly logical, if horrific conclusion to a specific problem. You take this AI maximizer, this paper clip maximizer problem and you apply it to the hugging face, the OpenAI incident. That's exactly what happened. OpenAI was testing out a brand-new model of its own AI and in order to test it fully, it basically removed all of the safeguards that they deployed for the commercial users. So, when you and I get an OpenAI tool, there are some commercial safeguards to prevent it from doing things too malicious. And if you try to do it, it will tell you it can't do it. But inside their own lab where they're testing it, guard rails are off, the handcuffs are off. It can do whatever it wants. They tried to protect this setup by isolating it. So, they have a research lab, they've got a couple of layers of protection in place to make sure that this thing stays within its boundaries, which seems sensible if you are, you know, if you're cloning dinosaurs, make sure they're on a remote island somewhere. So, it's sensible enough from a concept. The problem is, AI is really good, and it realized it was given a simple test to find the solution to this problem. It decided, like all, you know, devious students, that the easiest way to solve this is not to actually work through the answer, but to see if somebody else has the answer that I can copy from. And then it realizes I can't get out of this small sandbox. You have kept me in here. Let me find a way out of it. So again, instead of trying to solve the problem it was given, it tried to find a way to cheat, to copy the answer from somebody else. It broke out of its sandbox, and just like, you know, all the movies that you watch with Jurassic Park, the dinosaurs find a way. AI found a way out. It basically found what we call a zero day, an

exploit in the software that protected it, a proxy, broke its way out, found its way into hugging face, and apparently now the word is that there's more than hugging face. There were two other companies that were involved as well, and it just started to mine, or break into the hugging face database to see if the answer was there. Hugging face obviously detected this a few days later and raised a big fuss as they should have, and then everything got shut down. But from the AI's perspective, it did nothing wrong. You asked me to solve a problem. You did not tell me how I should and should not or could and could not solve the problem. So just like the paper clip maximizer, it didn't have built-in ethics or value system. It just went to find the answer. You are going to see more and more of such cases where AI goes rogue, not because it's malicious, it just doesn't have parameters constraining its behavior. And like the paper clip maximizer, if you give it a task and you don't give it any boundaries, it will do whatever is necessary to get that. And we have a whole genre of sci-fi movies that extend that concept, but that's basically what just happened.

Taimur Baig

And how do we sort of contextualize this development with the earlier fear around Mythos and its job is also to find vulnerabilities in various networks.

Gaurav Keerthi

I don't think there's any difference between what just happened and the hype that was previously being shared earlier by, you know, both OpenAI and Claude that, you know, their models are too powerful to be used. This is an example of that exact philosophy. So, they were worried that in the wrong hands, without the right safeguards and guardrails, their AI could do malicious things. They tried to put in the right guards, and even then, despite their best efforts, it still managed to do malicious things. So, if you didn't have these guards and you directed it to do malicious things, it would be even worse. That was the warning that they gave us, and they said that Mythos, without the right protections in place, without the right testing

in place, would be dangerous, and they just validated that it was dangerous, even with the right controls in place. The result and the risk that we all face now is that all of these tools are out there. They exist. And so that digital risk that it's imposed on the ecosystem is already materialized. It's no longer hype or fear mongering. It is true.

Taimur Baig

So that call for slowing things down, it's a bit too little too late.

Gaurav Keerthi

It is always tricky when the person who stands to benefit the most from other people slowing down is the one to call for a slowdown.

Taimur Baig

You have Mr. Elon Musk in mind.

Gaurav Keerthi

I will leave it at that, but conflict of interest and incentive compatibility is a huge challenge when you have two mega players calling the shots. And we've seen this before in tech. When you have only the choice of one operating system or two operating systems in the world, when they start to define the rules for what an operating system can and cannot do, you really don't have much of a choice. It's not quite like, you know, social media or other software, where there's 10, 20, 30 different ones. If I don't like my chat software, I'll just use a different one. In this case, they are essentially the largest providers of these two services, and if they call for a slowdown. Yes, you have to take them seriously because they're the only ones who can do it, but you also have to recognize that they're calling for a slowdown for their competitors because they've already reached that point. That becomes a challenge. And then there's this whole conversation around closed source, closed weight, open source, open weight AI, and geopolitics now again. So back to that conversation from last conversation.

Taimur Baig

I really liked Vivian Balakrishnan's presentation plus speech on this issue a couple of months ago, where his point was that regulated and sensitive areas. There's banking or medicine or Ministry of Foreign Affairs in this context. These issues of where the data goes is absolutely paramount, and perhaps, you know, open source maybe, but closed models are probably what they need.

Gaurav Keerthi

I remain open-minded about this because essentially, what you're worried about is your data leaving your control. And when you talk about sovereignty of technology, you're talking about ownership of where your data is and where the processing happens. The ideal solution is it happens on your machine. And it never goes anywhere. And we're moving into a world where more and more people are able to deploy AI on smaller and smaller devices. So today, there's a whole bunch of, and hugging face, bringing them up again. Hugging face has a whole bunch of models on their website. You can download any one of them, and on a normal MacBook or a fairly powerful Windows, you can run those models on your own device. The data never leaves anywhere. The processing never goes anywhere else. That actually to me is the most secure. And even if you are the Ministry of Defense or Foreign Affairs, that to you gives you the most confidence because it's isolated. You don't even need the internet. Once you start getting to models where you have to connect to somebody else's service, then it becomes a question of trust. And then trust becomes a subjective issue because geopolitics plays in. Do you trust that the US providers of this technology will not cut it off, will continue to give you good service, will not read what you're sending? So, all of these things become a question of trust rather than verifiability.

Taimur Baig

Or share with NSA, which seems to be the usual concern that people have on the China side. But I think in the last 20 years we've seen that companies like even Apple are

sometimes compelled to give backdoor access to the security agencies in the US.

Gaurav Keerthi

I mean, every country has its own laws. The US has a bunch of laws to keep its own citizens safe from terrorist attacks, from criminal activities, and so on. To some extent, they use those laws to try and get some data out of various providers, but to some extent, the providers themselves build the systems in a way that they can't access it themselves either. So, they don't build backdoors. And while I can understand the frustration this causes law enforcement, on the privacy and protection side, that's actually probably a better approach. You find other ways for the law enforcement to get their work done, not by building backdoors into your own software, because the minute you build a backdoor, no matter how much you try and hide it, the bad guys will find it first.

Taimur Baig

So, which is why I suppose I sort of appreciate it when WhatsApp started saying that the messages there are double encrypted, and even they cannot decrypt it. That hasn't stopped people from using Signal and Telegram and all the other platforms which they think is more secure than WhatsApp, but at least technically speaking, WhatsApp at least took that step.

Gaurav Keerthi

Yeah. And you'll find that more and more of the providers are building in encryption as the default end to end encryption. It's an expectation at this point. We're kind of moving into a world where, and I think I raised this the last time, where much like cars in the 50s, you expect it to come with a seatbelt. It's not an optional extra that you pay money for, like, hey, I'll get the car with the seatbelt. It's like it does come with a seatbelt, and you expect it to come with a seatbelt for every seat. So you're kind of moving into a world where if you're selling software, it should come with security in it. You're not paying for WhatsApp in the first place, but you're not going to buy software or use software that

inherently you know is risky or unsafe. That's no longer the world we live in.

Taimur Baig

Okay, so, in a minute, I want to talk about the cost of that seatbelt and whether only big guys can afford it or can everybody afford it. I'll come to that in a second, but I had read about this thing, the lethal trifecta. That LLMs need exposure to outside content like emails, access to private data like source code, and the ability to communicate with the outside world. So not just email external, but actually, you know, back and forth. So how are you thinking about this lethal trifecta and the security parameters around that?

Gaurav Keerthi

Yeah, I mean, there are various ways of defining the lethal trifecta that is definitely one of the ways that I've seen, and what you're talking about is no longer just the LLM, no longer just the model, but the agent that's actually using tools. So how we separate it is the LLM is kind of the brain, and the brain thinks through what the right answer could be. But the brain itself has no arms and legs. It can't send out an email or type code or do anything else or run code or build anything. It needs tools. And so, the tools are the part where it becomes an agent. And so, an LLM, which is the brain, with the ability to access your email, the ability to access your calendar, the ability to do PowerPoint slides, that becomes an agent. Then giving that some sort of business process. Every morning when I wake up, check my calendar, see what meetings I have, draft a brief for each of the meetings, and send an email to each person reminding them of the meeting. That becomes harness. So, you've got the model, the agent, and the harness. In order for that thing to become useful, it needs access to your knowledge. It needs, and your knowledge can be in the form of emails, it can be in the form of data, it can be in the form of the files. It needs access to your tools, your email inbox, your calendar, and your software. And it needs some sort of understanding of your business processes. You get to work at 9, you check your calendar,

you send out these emails, you do all of these things. All of that becomes extremely dangerous if misused, and each part, the data is a risk, the tool use is a risk, and the context is a risk. How I think about it, well, gosh, if I could solve this problem, I'd be a very rich person. The challenge right now is that everybody's and people are rightfully trying to break it down into the individual components. What the model does with your data. So, how well protected can you get it? If you think about the, again, the two major providers, OpenAI and Anthropic, they have enterprise plans. And the assurance for the enterprise plan is your data stays in your private enclave. I don't take it back to train anything and it stays with you. Again, question of trust. The answer from the Chinese ecosystem is its local models. You get your own server rack, you use my Qwen model, whatever model you want, you run it, and the data stays with you. I never get your data. Then you talk about the tools. This becomes almost non-negotiable even if the data is on your servers or it's in the service that you trust. This guy needs to send an email on their behalf. That's really risky if it does dumb things. There are, even in my own group of colleagues, there are two Gauravs, and if it picks the wrong Gaurav to send sensitive customer information to, that's a real problem. The way we think about it is, and I love this quote from Professor Simon Chesterman, who talks a lot about AI. He said, treat your AI agent as a brilliant drunk intern. Brilliant. Like this thing knows the answers to thousands of questions that you can't imagine. Drunk occasionally does dumb things, and intern will not take responsibility for that. You have to take responsibility. It's your intern. And so, if you think about the agent in that way, and the way we've done it in our company is that the agent is tied to the individual and has a subset of the permissions that need to be human gated. So, it can draft the email. I need to send the email. It can create the letter, I need to sign the letter. And so, it structurally doesn't have the ability to take what we call actions outside that would create a blast radius. Blast radius is this concept of, if the thing does something as stupid as possible, how much damage will it cause? And so, I allow it to read certain

folders on my laptop, not all folders on my laptop. I allow it to draft emails, not send emails. I allow it to send the team's messages to anybody in my team, nobody outside. So, these are the permissionings that you give to the agent to control its ability to use tools and its ability to access information. Then it comes to the business context around it. Again, you need to, and each person in my team has their own agent, but the context they have is necessary for their own job. Nothing beyond that. It's kind of the need to know principle, scaled to agents. I don't have all of the technical solutions to each of these problems, but the conceptual ones lead you to the technical solutions. So, what knowledge does it need to know and how much do you trust where it is? What tools does it need to have and how much permission and control can you give it or not give it? And what context does it need to know about your business and how much can you keep that on a need to know basis. Those three concepts, you layer it together and you kind of have an answer to that lethal trifecta. We're far from solving the problem. And every time we come near it, AI gets a little bit better and figures out a way to break across. And if I could just share one hilarious but tragic story from our own system. So, I built a harness for my company, and I'm very paranoid about Cyber Script. So, I don't allow my agent to actually know any of our passwords. It needs to, just like all the humans, log into a password manager, send the password to a script, and the script runs the password account. So, it never touches it by itself. So, my agent theoretically doesn't know the password. One day I noticed the login was getting faster and faster, and I was like, great, it's gotten more efficient. And I checked the code and I was like, "Do you have access to the actual passwords. How do you do that?" It's like, "Oh, I noticed that," and I was reading on Reddit that people keep a plain text copy of their own passwords in a notebook for ease of use. And I thought instead of using the password manager all the time, why don't I keep my own plain text copy so that it's easier for me. And so, it copied all of humans' bad habits. It figured out that the context I gave it, so the processes I gave it, the tools I gave it, and the

knowledge I had constrained it from getting were inefficient. And so, the paper clip maximized its way out of my constraints and figured out how to cheat. I never thought of that.

Taimur Baig

So, just from my personal experience now, I suppose I have some protection because I have a password for my text script with all my passwords. So hopefully, it will not see that one, because if it does, then, of course, it has access to the entire set of passwords.

Gaurav Keerthi

I mean, and I don't think it's that far away. So, the risk that OpenAI and both Anthropic and in fact, the general AI research community was flagging up was that.

It will do everything that hackers may do to get its way as well. If you tell it to log into a system, hey, you know, log into my LinkedIn and just post this, and then it realizes, hey, you forgot to give me access to your LinkedIn, but there's this file called passwords, it's got a password protection. Let me hack into it. And for the next 3 hours, it will figure out how to hack into your password vault. It will break in and it will take your LinkedIn password and post it. Nothing wrong. You asked it to do a thing; it did the thing. But it did a lot of things along the way that you would not want to do.

Taimur Baig

Okay, so we have so far talked about things with good intention can lead to pretty extensive unintended consequences and things can go out of control. Now let's talk about the bad guys, and by bad guys I not only mean the private actors trying to do ransomware on the bakery that we were talking about. I want to hear about that, but also the bigger actors, whether they're trying to shut down the centrifuge machine in Iran or elsewhere. Are things getting more challenging in that area because even 2 years ago when we discussed, it was already a big source of headache.

Gaurav Keerthi

Yeah, 2 years ago when we discussed it, the big risk was phishing emails. They've gotten so much better, and I mean I think they've gotten better and better, but phishing emails is no longer the only threat that AI poses. It still is a huge threat and scams, I don't want to downplay it, scams have become tremendously easy to operate at scale and very convincing. I regularly get scam emails and from a professional curiosity perspective, I'm like, wow, yeah, this is amazing. I would fall for this if I was not well acclimatized to it. But for the actual malicious malware part of things, I mean, that's where things have gotten really wild. We're seeing them experiment with development of cyber weapons, basically viruses that are at a pace that we haven't seen before. And they're experimenting, and we know that they're experimenting, and I went for, so this is not my own research, this is a talk that I went for quite recently, and the guy was talking about reverse engineering a virus that he found. And reverse engineering basically means you get the virus; it's an executable of some sort, you take it apart to see what's inside it. It's like an autopsy. And what they found inside it was that the comments had emoji in it and had M dashes in it. And again, we go back to this like, nobody on LinkedIn, no human being actually writes with M dashes or emoji. It's AI generated. And so, if you find that inside code for a malware, that means that thing wasn't written by a human. The virus was entirely generated by a system. We're seeing now the virus mutate slightly over the course of time, and what mutation means is that every time it goes out and does something, the guy who builds it changes it just a little bit. And this means that if you're looking for something that's shaped like this, that looks like a virus, and say, oh, if it's shaped like this, it's a virus, when it changes the shape a little bit, you're like, oh, that can't be a virus because I'm looking for things that are shaped in a circle. This is a circle with one flat edge, can't be a virus, let it through. And you keep changing the shape of it, you keep changing the functionality of it until one day you're able to get through. That was not possible previously. So now this idea of

mutating AI mutating viruses using AI to rapidly iterate is viral, and that's terrifying. And the scale of the attacks on big enterprises just means that, firstly, again, the scale, speed, and sophistication has just increased. They're able to find vulnerabilities in more tools that you guys use. They're able to exploit them at speed, and they're able to scan a much wider array of your surfaces to see where they can go in.

Taimur Baig

So, this is the part that I really want to hone in on the big guys versus the not so big guys. And we discussed this to some extent last time as well, that large organizations have lots of engineers, lots of computer scientists, and lots of capital people who do these things. But then the economy just doesn't run on large enterprises. There are thousands, even tiny little Singapore has tens of thousands of small enterprises. Who's looking after them? Is it you guys, or is it the government? Where do we stand there?

Gaurav Keerthi

Oh, there's both a conceptual policy argument and then a practical argument to it. I'll start with the conceptual policy argument. Less than 1% of the enterprises in any country in the world, you pick any country in the world, have a CISO. Less than 1%. And that statistically bears out with the number of people who are at that seniority with cybersecurity background. In general, there's a cybersecurity talent shortage in the world. Experienced cybersecurity, even greater shortage. So, if less than 1% of the companies in Singapore have a CISO or have anybody in their job scope with the word cybersecurity in it, that means 99% don't. And they are, it's not a few thousand, there are 358,000 companies in Singapore alone, active companies in Singapore alone that don't have a CISO. That's a huge volume problem. And if you think about the long tail of any economy, you've got your spike, you've got, 5% to 10% of the companies in the country that generate almost all the economic activity, but then 50% comes from the long tail. And the spike can't operate with the long tail. DBS on its own cannot function

as a bank, without the HR software, without the guys who come to clean and vacuum, and all of those functionalities that make an enterprise like this work without you having to do it all yourself. Those become the attack factors, and then those become the risks that you have to manage. So, even as a country in a policy sense, if you're trying to protect your big banks, your water plants, your electricity grids, you also have to protect the cleaning staff that go in there, the maintenance staff that go in there, the tools that they use, the software that they use, the services that they use, the law firm that they hire that passes all the data to them. All of those attendant services become that long tail, and it becomes a huge policy problem when you don't have them protected. There is the big policy question which I'm quite passionate about is whether or not, and you'll appreciate this as an economist, does cybersecurity fit the definition of what we'd consider to be a public good? I think it does. And today's cybersecurity is a private good, but it's not just a private good, it's a luxury good. There's a little bit of a Veblen effect in the consumption of it's like, oh, my enterprise is protected with this software because we only use the best of the best in Gartner's quadrant. That's a little bit of a luxury effect. And the price points of all those software are well outside the range of what any company below 200, 300 staff can afford. So, you've priced out protection. But at the same time, you can't feel protected unless the tail is protected. The simple analogy is you live in a beautiful mansion and you've got security guards. I'm sure you feel really safe, except the streets outside of you have like gang fights and violence and shootings that immediately reduces your level of security. And until the streets outside are clean, you will never feel that sense of security because you constantly have to have your guard. So, the only way to actually have a sense of security is to make sure the streets are clean as well, and that's kind of where the long tail comes in.

Taimur Baig

So, in terms of solutions available for those who cannot afford to pay for top tier services, how's the industry evolving?

Gaurav Keerthi

The truth is, and I mean, when we last spoke, I was still with a very large enterprise doing enterprise cybersecurity. The truth is that the amount of money that a company can make in enterprise cybersecurity is enough to generate enough volume of business so that they never need to touch the rest. Why would you bother protecting the doctor and his nurse when you can protect the bank? And a single contract from a bank will pay for 1000 doctors and nurses, so there's actually no economic incentive to produce for this. Enterprise cybersecurity is a very robust market, and all of the things that we spoke about earlier, you know, protecting the agent's identity, protecting the, you know, the harnesses, protecting AI, protecting all of those things, they're enterprise problems. They're very good problems to solve for enterprises, but they don't scale down. The small guy isn't worried about agentic identity security; he's worried about downloading a virus. And I literally posted this this morning. We had a customer, and we found, we blocked a virus that they were trying to install, or the virus that was trying to be loaded. And the title of the virus, and I'm not making this up, "please disable your antivirus before opening this file.exe". And people do. And you're like, huh, I'm not dealing with sophisticated state sponsored threats anymore. I'm dealing with the guy who will follow instructions from a file name. But that's a problem that if you run a bakery, you're like, well, this is computer stuff. I don't really know computer stuff. It's giving me an instruction. I'm sure it's a legitimate instruction, let me do it. So, I'm dealing with that problem. Two conceptual shifts that I had to take, and I think the industry has to take when they think about protecting smaller businesses. The first is it needs to be more like a utility. A utility provider is fundamentally different from a services provider. If there's a law firm, a bakery, and a tuition center in the same shopping center, if they want electricity, they pay \$10 a month and all three get identical electricity. They want water, they pay a certain fee per month, they get identical water. They want, you know, productivity software, they pay a certain fee, they all get. It's identical. But when they want

cybersecurity, suddenly they have to call a consultant, and they get a different type of cybersecurity from them and from them. And that doesn't seem to make sense. Why is it that cybersecurity can't be one of those things where you pay a monthly fee, turn it on, and it works, and it's identical as a baseline for all three. You get clean drinking water, clean electricity, clean internet for all three. The second, and this is the bigger reason why a conceptual shift was needed, is that the industry is full of what I call carpenters. Skilled professionals who are very proud of their trade work, their tradecraft. If you ask them to build a conference table for DBS, they'll come in, they'll measure the room, they will use their professional skills to cut down an American walnut tree, polish it, and build you a beautiful a beautiful conference table, and it'll be bespoke for that room, and every room in this building will have a different bespoke dining table or conference table. If you don't have the money for a skilled carpenter, you go to IKEA and that's 30 bucks. It has two legs or 3 legs or 4 legs, you have to screw it in yourself, and you have to drive it home yourself in a cardboard box. That table has never met a carpenter in its life and doesn't know what a tree looks like, because it's not made of trees. It's made of wood chips and plastic chips squashed together and put in a cardboard box. The reason why it's so different is that the skill set to become a really good carpenter are completely independent from the skill set to build a factory that mass produces tables. Completely independent. In fact, they're counterproductive. If you bring a carpenter to IKEA, they will vomit blood because like, what is this? This is not a table. This looks nothing like it, and there's no chamfered edges and whatever else carpenters look for. It's the same model in cybersecurity. If you bring in cybersecurity professionals who are used to protecting, you know, a DBS and you show them what an IKEA version of cybersecurity looks like, they're like, this is not cybersecurity, how will you stop the agentic identity from running rogue within your proxies? Like, well, that, it's a bakery. They have none of those things you just mentioned. So, it's a different skill set. And

because it's a different skill set, you can't take enterprise cybersecurity and trim it down. You can't just find cheaper carpenters from, you know, Vietnam or cheaper trees from South America and build cheaper tables. It doesn't work. You have to build something different. And then it becomes a design choice as a company. Are you going to be building for enterprises or small businesses? You can't do both.

Taimur Baig

Let's talk a little bit about the public good concept. You've written an article about it. We'll put the link of the article in our show notes. What is your sense of public sectors around the world?

We can narrow it down to Singapore, but just around the world, is there sufficient urgency or this long tail aspect that you're talking about, is it being discussed or you're in the small minority?

Gaurav Keerthi

I would say that every country in the world sees the danger, sees the risk and is afraid of it. You can go to any country and look at their newspaper, and the word scam is somewhere in that newspaper in that week. X number of millions of dollars lost to a scam, X number of companies scammed, hacked, malware, ransomware, that the policymakers are aware of the problem. They are also very aware that there's no solution. How do you protect the bakery from the scam? I took a trip to Estonia recently, and the day I landed, they had a nonprofit organization supporting artists in Estonia, and they got both kind of hacked and scammed for something like USD700,000. That was their entire organization's funding, just lost. And the person, the poor person who clicked the link, feels terrible, but do you blame that person? I mean, how much of it accountability do they take? They had no protections in place. These types of incidents will continue to happen over and over again. So, the question that we really need to be asking is, at what point will policymakers go from inaction to action? And I think it's inevitable for them to see that I only have a few levers. Regulation is one of those levers, grants and funding is

another one of those levers. How do I push and pull these two carrot and stick levers to get the wider long tail to be better protected? Because even if, and I know they care about smaller companies, but even if they didn't care about smaller companies, they certainly care about the big ones. And if the bank or the electricity grid can be taken down by taking out a bunch of smaller players, that's a huge risk to them. So, I think it's a matter of time. You will see that in Singapore, and I'll start with Singapore, just earlier this year, the Healthcare Information Act got pushed out, and that basically said that if you operate any sort of healthcare organization, you need to have cybersecurity. Some of my customers, my smallest customer has 2 staff, the doctor and his nurse. That's it. But they need to have cybersecurity. And from the patient's perspective, you kind of understand why. It's like I just told you some really deep, dark things about my life. I would like you to protect that information, and not lose it because the consequences are quite bad for me, personally, professionally, whatever. So, every single clinic in Singapore, you can go to the HDB void deck, the uncle who's like 90 years old, and his daughter who's the nurse, they have to have cybersecurity. And it's tough, but it makes sense. So, you'll start to see these policy levers being shifted and you'll see it from, I think in the coming years, law firms, you know, maritime sector, audit and accounting sectors will all start to have these cybersecurity requirements on a sectoral basis. UK was the same thing. So, UK has an equivalent baseline cyber hygiene standard called Cyber Essentials. And when they first announced it in parliament, they were like, this is an optional standard, we're just putting it out there as a good to have, as a gold standard for small companies to aspire to. And a bunch of schools got hacked and parents got angry. They're like, it's an optional standard for everybody except for the schools, mandatory for the schools. Then Jaguar got hacked, and like, it's an optional standard for everybody except for the schools, and for anybody who's in the manufacturing sector working in this space. You can see where this creeps. And so, I do think that from a policy perspective, it is inevitable that more and more sectors are

going to be regulated to have cybersecurity. Even if the regulator doesn't do it, the B2B system will do it. I am 100% confident that any organization as large as DBS or even smaller will have in their tender documents, in their requirements, show me your cybersecurity before you can sell me anything. I don't care what you're selling me, show me your cybersecurity. If you want to participate in this tender, secure yourself, and that will clean up a large part of the ecosystem as well.

Taimur Baig

So that's the regulation part. Hopefully, certain standards from lessons learned or lessons anticipated get implemented around the world. You also touched upon the grants and the funds part. Are you seeing promising developments around the world where the governments are trying to sort of subsidize or pay for this public good?

Gaurav Keerthi

I'm seeing the start of it, but I don't think enough, and I do worry that we are separating the buckets rather than merging them in. What I mean by that is, we're now saying, hey, here's the grant for buying a car, and here's a separate grant which you can apply for, which is from seatbelts. And people are like, yeah, I'll take the car. And there's more paperwork for the seatbelts, forget about seatbelts. So, if you look at the amount of money we've announced in just in Singapore for AI, the push for small companies to adopt AI, and it's the right push. But for every SGD10 of grant given out for AI and digitalization, SGD0 are required for seatbelts. There's a separate grant completely independent of this which offers grants for seatbelts, security, cybersecurity. In an ideal world, actually, for every SGD10 of grant that's given out, SGD1 is ring-fenced for safety and security. And if you think about just business norms, what we've seen over the past few years is around 10% of every company's IT budget is spent on cybersecurity. It should be spent on cybersecurity. So, if you're spending SGD1000 a month on IT, you should be spending SGD100 a month on cybersecurity.

It's about right. So, SGD900 on productive resources, SGD100 on protection of all of that. And that scales. So even a bank or even a telco should be spending plus minus that. So, if you're giving out grants as the government, you have all of the right to demand that the money be spent in a way that you think is responsible. And that includes making sure that it's secure. And the reason why I feel so strongly about the AI portion of the grant is AI does two things that are very dangerous. One is the adoption of AI in and of itself, as we just discussed, introduces risks for a company. Just in and of itself, using AI creates new risks. And you should have some level of safety and security in your company, some baselines to minimize the impact of those risks. The second thing is that AI allows people to build their own software, really easily. And unfortunately, software that's built with AI has been shown statistically numerous times to be less secure than previously hand-coded software or built software that goes through that robust checking and process that enterprises are used to. So, firstly, the AI introduces risks. Secondly, the stuff that they build with AI and may sell is full of holes itself. You sell these two things without any seatbelts, that's very high risk, and we're going to be paying for that in 1 year's time or 2 years' time when all these companies or the tools that they use are hacked.

Taimur Baig

Recently I heard from somebody that the best way to deal with unintended consequence of AI is another AI, a critical control center, and it just keeps an eye on everything else going on, so I'm an agent to oversee the agents. Your reaction to that sort of potential solution.

Gaurav Keerthi

I don't think it's wrong. In some senses, it's going to be one of those who watches the watchman scenarios. The way in which if you want to maximize the potential for AI in your company, you're going to try to minimize how much you burden it with security, and it will slow it down, it will reduce additional friction, it will cost a little bit more to have all those

security features. One novel way to do it is to just have an alternative agent checking on this guy. Are you doing something risky? Are you doing something dangerous? We build something like that as well, and it basically checks to make sure that this behavior is acceptable. The challenge is, then this becomes your point of vulnerability. I just need to trick this guy and we're done. And the tricks on the AI that does security are much more subtle than tricking the main agent. One of the things that we realized, and this is also quite fascinating, is as you consume more information of a certain sort, you can get affected by that information. And in traditional kind of the human world, we see this a lot with law enforcement agents who have to review child pornography. After a while, there's psychological damage to the individual who's constantly reviewing to see whether this is child pornography or child abuse or not. You kind of get a similar effect with AI when it's trying to find cybersecurity threats. It is constantly looking for vulnerabilities and risks and hacks and threats, and it's scouring the internet to see how hackers are thinking and working. After a while, it starts to lose itself, and you have to spend a lot of effort and research to make sure that it doesn't get to that. So, if people are building the AI to check on your AI, the amount of effort and investment on this side is tremendous because this tends to drift a lot more than this does.

Taimur Baig

So, picking up from you the vibe that a human in the loop type solution is going to be a permanent feature.

Gaurav Keerthi

I don't want to have any famous last words. Permanent is a very long time, but I do think that at this stage, we're nowhere near the point where we can completely let go of all cybersecurity operations for sophisticated threat actors. For the baseline that just setting the floor right, I think a lot of that can be autonomous. Stuff like, please disable your antivirus before installing this file.EXE, I think that can be an automatic decision. But once you start to get the more sophisticated

attacks, and more importantly, sophisticated controls, then some level of judgment matters. And I'll give you one example of a dilemma that we had. I know we have an autonomous system that protects our customers, building it out and scaling it out. And one of the things is that, you know, there's a vulnerability on Windows, this particular version of Windows. So, this agent has gone out there and seen from the internet, Windows has announced there's a vulnerability. Catch it immediately. So, it comes back and says, okay, there are a few options. There's this risk, this company is using Windows. Option number 1 is we can follow Windows' advice, patch it immediately. But the customer says, oh, I'm a doctor, you can't patch it right now. I'm in the middle of surgery. Great. Option number 2 is we patch it tonight or patch it another time. Option number 3 is I just send you an email, and you patch it whenever you're free. Option number 4 is I take the advice very seriously from Microsoft and I quarantine your laptop. In the middle of whatever you're doing, I don't care. Microsoft is authoritative. This is a risk. I will keep you safe against your own will. Now, I can understand again the paper clip problem. I can understand why you'd want to do that, but if you quarantine a laptop in the middle of surgery. That's not going to go down well.

Taimur Baig

That's the type 2 error that we call in statistics.

Gaurav Keerthi

Absolutely. And so, this kind of type 1, type 2 errors. You still need a human in the loop for some level of sense making and checking until such a point where all of these heuristics can be written out, and that becomes a bigger challenge.

Taimur Baig

I remember, toward the end of our conversation last time, I think you felt that you were depressing the audience, and you said that, you know, I don't want to depress everybody. There are no silver linings. Things have become even more depressing since then, Gaurav, so what is the silver lining?

Gaurav Keerthi

A few things. I am fairly confident that while AI was amazing and impressive now, we're going to see that same plateau of, you know, disenchantment, disillusionment, of capabilities maxing out at a certain point. On a number of levels, what we're seeing is that we've run out of data to ingest, right? And once you start to train it on synthetic data, on generated data, there is research to suggest that the models may not be as robust. It will start to collapse. If you train it on pictures of fake generated cats instead of real cats, after a while, the generated cats start to become less useful training data. So, there's a sense that the data will no longer be as rich and the models will plateau in capability. Maybe not soon, but in the next 1 or 2 years. There's also a sense that eventually the commercial incentive for all players is to make it lousy. And as much as, I mean, I don't want to use the word, but there's this great article, about the encraptification, if I may. I'll just say it, you can censor later on, the end certification of technology. And the basic premise is that as long as capitalism continues to drive technological innovation, eventually, it will become bad. We saw this with social media. When social media first came out, it's like, we will all be friends. And I literally remember Mark Zuckerberg had a speech and I'm going to try and find it somewhere where he said, We will reduce the number of wars in the world because you can't have an enemy when you know what their kids look like. Turns out you can have an enemy because you know what their kids look like. You can specifically be angry with the person because you know what their kids look like, and 1000 other reasons that social media has created divisions for. So social media started with great promise, it ended with great tragedy. Email. The amount of time you spend queuing up to buy stamps and going to the mailbox, we'll return you all the time. You can spend more time on the golf course once we invent email. I have spent 0 time on a golf course because of email. Absolutely. Every technology that's, and so AI will outsource all of the work to AI. You just press play and you leave it and you come back. I am more exhausted now because of AI. And so, every technology that

comes in with this great promise of productivity and innovation, and all of these things, doesn't quite fulfill it. So, I'm less terrified of its impact than other people are. The second part, which I think is probably even more optimistic, is that eventually good guys catch up, and we find ways to curtail, hold back, hinder the effectiveness of the bad guys. We've done it for every other type of crime. And there will be a period where it gets messy, but there's a period where it gets cleaned up. Scams are now in the period where it's still bad, but you can see law enforcement taking a much more digital native approach to solving the problem, as opposed to putting up posters on the walls. I think eventually, and I go back to this argument, we're going to have to, and law enforcement and governments are going to have to think about this as a provision of infrastructure rather than a provision of boiled water bottles. There's this lovely case study about cholera in the 1850s in London, and there's a huge outbreak of cholera in London, and initially it was in the poorer parts, so, you know, the rich folks didn't really care. Then the River Thames became polluted, and it was literally called the Great Stench of London because the Thames was full of dead bodies and feces, and it was horrific, and cholera was spreading everywhere. So, they decided, let's go and solve the problem. The first attempt at solving the problem was putting up posters. Please don't drink from the river, please boil your water, please use 2FA. You can see where this pattern is going. It didn't work, because poor people didn't have a choice. You don't, it's not a luxury they could afford. The second thing they tried was, let's do individual point solutions. So, this fellow named, I think it's Sir John Snow, went around researching where all the cholera outbreaks were and turning off the taps, so the pumps wouldn't work. So now these people couldn't get the water from the pump, and he thought once they can't get dirty water from the pump, they will all boil the water and drink boiled water. Turns out no, they just went to the pump across the road and got from there instead. And so, it just kept moving. So, point solutions didn't work. Cholera only stopped when his Lord Bazalgette invented what we now know as

the sewer. Dirty water goes out and clean water comes in. And once you have a system, an infrastructure in place that cleans it at a structural level, systemically, you no longer have cholera. We're in the cholera age of scams and cyber malware. It's everywhere and people are plugging into the internet, and dirty internet is coming in and out. We need to eventually build that infrastructure to make sure that it's cleaned from the pipes. And this goes all the way from the banking infrastructure, the telco infrastructure, the connectivity infrastructure across countries, and the endpoint, the end company infrastructure as well, just needs to be provided in the same way we have clean drinking water.

Taimur Baig

If I'm not mistaken, in the public good article that you wrote, you provided this illustration of this, but the other article that I want to put in our show notes is the Why Singapore's AI Push needs a cyber layer, where you take more of your old job looking at the entire society's cybersecurity perspective and how this defense is as critical as any other layers of defense.

Gaurav Keerthi

Yeah. And like I said earlier, we're going to need this even more, the coming bound, because if the bad guys are using AI to do bad things, and if the good guys are using AI to take greater risks. Who's securing all of this? And somebody needs to step up, and I do think that the government or policymakers in general should think about what levers they have at their disposal to make it less risky. Even the banks may start thinking, and I don't think it's going to happen anytime soon, but the banks may start thinking about this from their own self-preservation perspective. Banks give out a number of loans, car loans, house loans, SME loans. If you default on a car loan, I'll take back your car. If you default on the house loan, I'll take back your house. But if you default on an SME loan, it usually means your SME no longer exists. I take back nothing. So, you've got this huge kind of array of unsecured loans in the SME space sitting on

your risk book, and you're not quite sure what to do about it. You're going to start requiring them to take house insurance, car insurance, and some sort of SME insurance, which will include cyber. And I think it's inevitable that the banks are going to start requiring companies who take loans above a certain amount to show me your cyber first. Show me that you have been a responsible company, protecting yourself against the basic things before I give you this SGD1 million to build your bakery or whatever it is. I think that's also inevitable.

Taimur Baig

Sounds like core infrastructure and public good to me.

Gaurav, what a fascinating discussion and surely, it's not over. There will be a 3rd time, I'm sure about that.

Gaurav Keerthi

I love these conversations because you pick such difficult questions to answer, but they're so much fun.

Taimur Baig

Gaurav, thank you very much for your time and insights.

Gaurav Keerthi

Thank you for having me.

Taimur Baig

Thanks to our listeners as well.

ECONOMICS & STRATEGY

Taimur BAIG, Ph.D.

Chief Economist
Global

taimurbaig@db.com

Mo Ji, Ph.D.

Chief Economist
China/HK SAR
mojim@db.com

Nathan CHOW

Senior Economist
China/HK SAR
nathanchow@db.com

Radhika RAO

Senior Economist
Asean/EU/South Asia
radhikarao@db.com

Tieying MA, CFA

Senior Economist
North Asia
matieying@db.com

Han Teng Chua, CFA

Senior Economist
Asean
hantengchua@db.com

Byron LAM

Economist
China/HK SAR
byronlamfc@db.com

Eugene LEOW

Senior Rates Strategist
G3 & Asia
eugeneleow@db.com

Philip WEE

Senior FX Strategist
Global
philipwee@db.com

Wei Liang CHANG

FX & Credit Strategist
Global
weiliangchang@db.com

Samuel TSE

Rates Strategist
Asia
samueltse@db.com

Sherilyn Chew

Multi-asset Strategist
Global
sherilync Chew@db.com

Mervyn Teo

Senior Credit Analyst
USD, SGD, AUD
mervynteo@db.com

Dexter CHUN

Credit Analyst
USD
dexterchun@db.com

Tracy Li Jun LIM

Credit Analyst
USD, SGD
tracylimt@db.com

Ian Haan Chui

Credit Analyst
USD
ianchui@db.com

Amanda SEAH

Credit Analyst
USD, SGD, AUD
amandaseah@db.com

Teng Chong LIM

Credit Analyst
USD, SGD, AUD
tengchonglim@db.com

Joel SIEW, CFA

Credit Analyst
USD, SGD, AUD
joelsiew@db.com

Iris GAO

Credit Analyst
USD
irisgao@db.com

Lilian LV

Credit Analyst
USD
lilianlv@db.com

Daisy SHARMA

Analyst
Data Analytics
daisy@db.com

Violet LEE

Associate
Publications
violetleeyh@db.com

GENERAL DISCLOSURE/ DISCLAIMER (For Macroeconomics, Currencies, Interest Rates, Digital Assets or Commodities)¹

The information herein is published by DBS Bank Ltd and/or DBS Bank (Hong Kong) Limited (each and/or collectively, the "Company"). It is based on information obtained from sources believed to be reliable, but the Company does not make any representation or warranty, express or implied, as to its accuracy, completeness, timeliness or correctness for any particular purpose. Opinions expressed are subject to change without notice. This research is prepared for general circulation. Any recommendation contained herein does not have regard to the specific investment objectives, financial situation and the particular needs of any specific addressee. The information herein is published for the information of addressees only and is not to be taken in substitution for the exercise of judgement by addressees, who should obtain separate legal or financial advice. The Company, or any of its related companies or any individuals connected with the group accepts no liability for any direct, special, indirect, consequential, incidental damages or any other loss or damages of any kind arising from any use of the information herein (including any error, omission or misstatement herein, negligent or otherwise) or further communication thereof, even if the Company or any other person has been advised of the possibility thereof. The information herein is not to be construed as an offer or a solicitation of an offer to buy or sell any securities, futures, options or other financial instruments or to provide any investment advice or services. The Company and its associates, their directors, officers and/or employees may have positions or other interests in, and may effect transactions in securities mentioned herein and may also perform or seek to perform broking, investment banking and other banking or financial services for these companies. The information herein is not directed to, or intended for distribution to or use by, any person or entity that is a citizen or resident of or located in any locality, state, country, or other jurisdiction (including but not limited to citizens or residents of the United States of America) where such distribution, publication, availability or use would be contrary to law or regulation. The information is not an offer to sell or the solicitation of an offer to buy any security in any jurisdiction (including but not limited to the United States of America) where such an offer or solicitation would be contrary to law or regulation.

[#for Distribution in Singapore] This report is distributed in Singapore by DBS Bank Ltd (Company Regn. No. 196800306E) which is Exempt Financial Advisers as defined in the Financial Advisers Act and regulated by the Monetary Authority of Singapore. DBS Bank Ltd may distribute reports produced by its respective foreign entities, affiliates or other foreign research houses pursuant to an arrangement under Regulation 32C of the Financial Advisers Regulations. Where the report is distributed in Singapore to a person who is not an Accredited Investor, Expert Investor or an Institutional Investor, DBS Bank Ltd accepts legal responsibility for the contents of the report to such persons only to the extent required by law. Singapore recipients should contact DBS Bank Ltd at 65-6878-8888 for matters arising from, or in connection with the report.

DBS Bank Ltd., 12 Marina Boulevard, Marina Bay Financial Centre Tower 3, Singapore 018982. Tel: 65-6878-8888. Company Registration No. 196800306E.

DBS Bank Ltd., Hong Kong Branch, a company incorporated in Singapore with limited liability. 18th Floor, The Center, 99 Queen's Road Central, Central, Hong Kong SAR.

DBS Bank (Hong Kong) Limited, a company incorporated in Hong Kong with limited liability. 11th Floor, The Center, 99 Queen's Road Central, Central, Hong Kong SAR.

¹ This disclaimer may not apply if the applicable assets fall within the definition of 'financial instruments' that are set out in Article 2(1) EU MAR (e.g. financial instruments that are traded on a regulated market, MTF or OTF, etc.). Section C of Annex I of MiFID2 specifies these 'financial instruments'.