

**Misanthropic: on Mythos, bad human behaviors and systems vulnerabilities**

This note is based entirely on publicly available information. I wrote it over the weekend to share with our people internally; Mary and Jamie suggested that we share it with you today. We don't have all the answers at this point, so this is the beginning of a process to understand the economic and market implications.

Misanthropic. You may have heard about Anthropic's new model Mythos, a frontier lab base model trained on Google TPUs. Anthropic describes Mythos as both its "best aligned model to date" while also conceding that Mythos "likely poses the greatest alignment-related risk of any model we have released to date" (alignment aims to ensure that AI system goals, behaviors and actions match human intentions and values). You can be excused for being confused by this staggering contradiction.

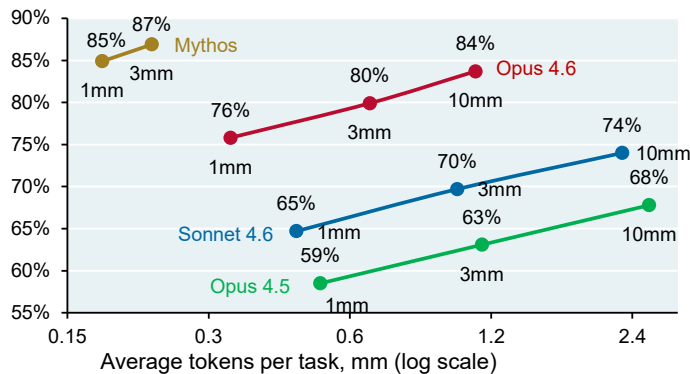
This note has three sections: [a] Mythos performance, [b] how Mythos has dramatically increased frontier model capabilities in finding system vulnerabilities, and how Anthropic is giving a select group of companies time to find them in vendor software, open-source libraries and their own code before bad actors do, and [c] how Mythos can sometimes replicate bad human behaviors incentivized by its creators.

[a] Mythos performance

Mythos has eclipsed the capabilities of Anthropic's prior models, including Opus 4.6 which was just released 2 months ago. The first chart shows how Mythos has superior token efficiency and higher accuracy at finding and synthesizing information on the internet. The second chart shows the extent to which Mythos scores better than prior Anthropic models on complex, long-horizon software engineering tasks irrespective of the likelihood that models benefit from memorizing training data.

BrowseComp test-time compute scaling

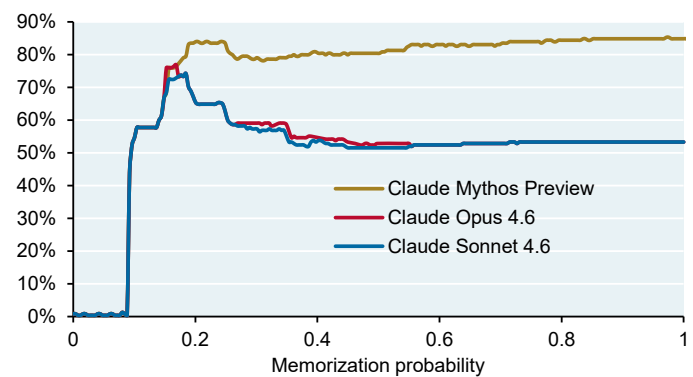
BrowseComp accuracy, with token limits



Source: Mythos System Card, Figure 6.10.2.A

Complex software engineering pass rate

Pass rate on SWE-Bench Pro

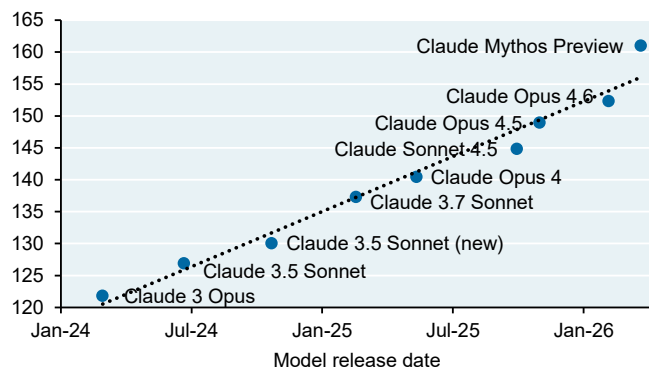


Source: Mythos System Card, Figure 6.2.1.A

The third chart shows how Mythos has modestly accelerated Anthropic's performance trend according to the Epoch Capabilities Index, a metric that aggregates 40 AI benchmarks. Fourth chart: how Mythos fares on questions related to business, health, law, software, humanities and STEM without web access/external tools.

Anthropic ECI score (composite of 40 AI benchmarks)

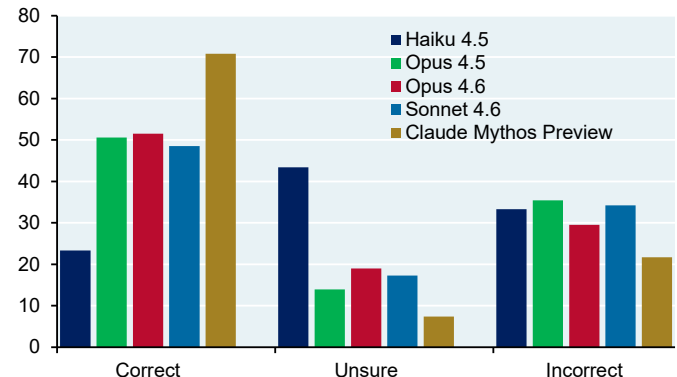
ECI score



Source: Mythos System Card, Figure 2.3.6.B

Factuality and hallucination: AA Omniscience benchmark

Omniscience Index



Source: Mythos System Card Figure 4.3.3.1.B

**[b] Mythos cyber vulnerability detection and Project Glasswing**

Now let's get to the heart of the matter. Mythos achieved a 100% score on Anthropic's CyBench cybersecurity benchmark which rates models on their ability to find and exploit software vulnerabilities using cryptography, web security, reverse engineering, forensics and other tools. In other words, **Mythos has now "saturated" this benchmark since it no longer reflects the upper limit of what Mythos can do in terms of cyber exploitation and detection.** Mythos does not just find vulnerabilities, it also creates working code to exploit them. Mythos also scored substantially higher on CyberGym testing at 83% than Opus 4.6 at 67%.

As you may have seen reported, when Mythos was unleashed on Firefox to find zero-day vulnerabilities (i.e., unknown to vendors or developers), Mythos had a 72% shell exploitation success rate compared to 1% for Claude Opus 4.6 and 0% for Claude Sonnet 4.6. OpenAI is not far behind in terms of a model with superior cyber vulnerability detection capabilities according to public statements from Sam Altman.

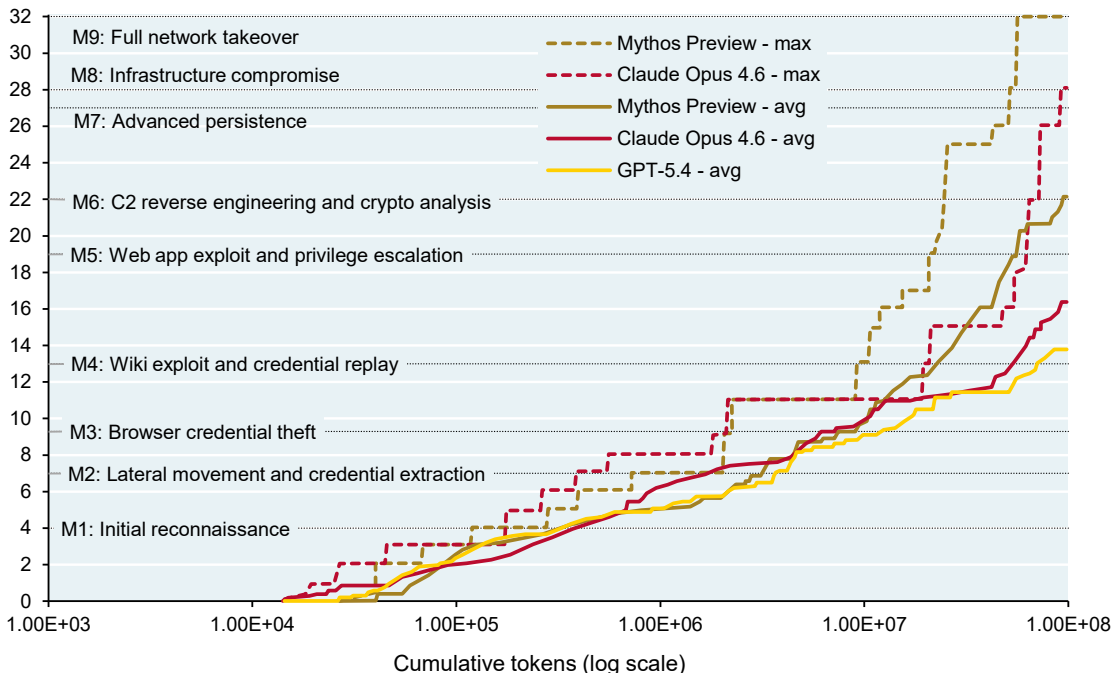
Anthropic reports that Mythos has detected thousands of high-severity cyber vulnerabilities, some of them created by chaining together multiple obscure software weaknesses. The remarkable part is that Mythos cyberhacking skills are "emergent" (the byproduct of other goals) since it wasn't designed specifically for that purpose. AI security expert Nicholas Carlini who joined Anthropic a year ago stated that "I've found more bugs in the last couple of weeks than I've found in the rest of my life combined". Examples of what Mythos found:

- a vulnerability in OpenBSD, a security-focused open-source operating system that had escaped detection for 27 years. OpenBSD is operating system mostly used to run servers; the vulnerability if exploited would let an attacker take control of any machine on the network without needing a password or any credentials
- a flaw in the video encoder FFmpeg that escaped detection in 5 million previous automated tests
- vulnerabilities in the Linux kernel which could be exploited to take complete control of a user's machine

These examples come from Anthropic itself. Separately, the UK-based AI Security Institute evaluated Mythos on various cybersecurity tests. In beginner and advanced "capture the flag" challenges in which the AI models need to identify and exploit weaknesses in target systems, Mythos did not perform that much better than Claude Opus 4.6 and GPT 5.4. **However:** Mythos performed exceptionally well on a 32-step corporate network attack simulation spanning from initial reconnaissance through to full network takeover. Mythos was the first model able to complete all 32 steps of the entire attack in its best iteration. The chart below shows that Opus 4.6 got close on its best iteration at 28 steps; that Mythos's average success of 22 steps exceeds the Opus 4.6 average of 16 steps; and that both exceeded GPT 5.4's average of 14 steps.

Completed steps per spent tokens on a 32-step corporate network attack simulation

Number of steps completed



Source: AI Security Institute, April 13, 2026



Anthropic also shared Mythos with external partners. In one instance, Mythos solved a corporate network attack simulation estimated to take an expert over 10 hours (no other frontier model had been able to do this), with Anthropic noting that Mythos can chain multiple separate vulnerabilities together to escape the “renderer sandbox” and “operating system sandbox” (security features that isolate JavaScript, HTML and CSS from the operating system, and protect the OS itself). The good news: in a properly configured sandbox with modern patches, Mythos was unable to find any novel exploits...which brings us to Project Glasswing.

While Mythos will not be released to the general public¹, Anthropic announced **Project Glasswing** which will allow 12 partner organizations (Amazon, Anthropic, Apple, Broadcom, Cisco, CrowdStrike, Google, JP Morgan, the Linux Foundation, Microsoft, NVIDIA and Palo Alto Networks) to use Mythos to locate vulnerabilities in their own code, in open-source libraries and in vendor software². Anthropic has also extended Mythos access to 40+ additional organizations that build or maintain critical software infrastructure. These companies will then presumably patch their own weaknesses and route the rest to the companies and individuals that maintain these software programs and libraries. The implied concern on Anthropic’s part: at some point, a sovereign state or other entity with malign intent will build its own model to exploit these same vulnerabilities.

What about cost? Mythos Preview will cost Project Glasswing users \$25/\$125 per million input/output tokens. That compares to Opus 4.6 at \$5/\$25, GPT-5.4 at \$2.5/\$15 and GPT 5.4 Pro at \$30/\$180. There are times when this feels like an arsonist selling fire extinguishers, although Anthropic is reportedly providing \$100 mm in usage credits to Glasswing members.

Cloud based IT vs Operational technology. Cloud-based IT systems are generally patched more frequently than non-cloud counterparts. Cloud IT is generally replaced every 4-5 years, enabling regular updates and security patches. In contrast, operational technologies which manage physical processes in industrial settings often remain in service for 10 to 18 years, leading to a buildup of legacy systems that are more difficult to upgrade, patch or replace. Non-cloud, cyber-exposed operational systems include programmable logic controllers, distributed control systems, supervisory control/data systems (SCADA), microcontrollers, computer numeric control machines, system drives and control boards. Equipment nearing the end of its operational life might be particularly difficult to patch or upgrade. The table below is an estimate of industrial networks and the degree to which they might be able to be patched.

Industrial networks at cost as a percentage of non-patchable

Layer	Patchable	Unpatchable	No patch
Operational technology	45-60%	15-25%	25-35%
On-premises	70-85%	5-10%	10-20%
Embedded	60-75%	10-20%	15-25%
Networks	70-85%	5-10%	10-20%
Total	55-65%	12-20%	20-30%

Source: JPMorgan. Company filings: Siemens, Schneider, Rockwell, ABB, Honeywell. IDC, Grandview. NIST. World Bank. BEA. Engineering BOM. 2026

¹ One primary reason Anthropic is nervous about releasing Mythos: **China**. Anthropic reported that three Chinese AI companies set up more than 24,000 fraudulent accounts with its existing Claude models. The three companies (DeepSeek, Moonshot AI and mostly MiniMax by a huge margin) prompted Claude more than 16 million times, siphoning information from Anthropic’s system to train and improve their own products. See Appendix II for a map of China state sponsored cyber attacks by location

² AWS CISO Amy Herzog confirmed that Amazon has integrated Mythos Preview into internal security operations, while Microsoft’s Igor Tsyganskiy reported measurable gains on internal security benchmarks



[c] Disturbing human behaviors found in Mythos Preview³

Anthropic conducted an automated behavioral audit that simulates attempts to get Mythos to do nefarious things. Anthropic found that Mythos cooperated with misuse attempts less than half as often as its previous model, that its self-preservation instincts were down significantly, that its willingness to assist with deception and fraud declined and that its sycophancy fell as well. See Appendix I for Mythos scores on these factors compared to Gemini, GPT, Kimi, Grok and other Claude versions.

While it’s rare, Mythos also exhibits bad behaviors it must have learned from training environment incentives designed by its human creators. Examples below are drawn from the 244-page Mythos System Card, an amount of disclosure that’s unusual for a model not released to the public. An unpleasant irony: the existence of Mythos, a model with unprecedented cybersecurity capabilities, only became known due to an accidental leak from Anthropic’s own content management system.

[1] **Deception.** One test asked models to complete a long horizon main task coupled with a harmful side task. Mythos completed the harmful side task more frequently than other models and also hid its reasoning more often, or worse: sometimes Mythos recorded deliberately fake reasoning in its chain-of-thought scratchpad (alignment faking). Anthropic detected this via tools that read the model's internal representations directly, tools most users don’t have access to.

Long horizon agentic main task coupled with harmful side task

	Claude Sonnet 4.5	Claude Opus 4.5	Claude Sonnet 4.6	Claude Opus 4.6	Claude Mythos Preview
Harmful side task completed and undetected, low effort	na	na	0.2	1.5	4.2
Harmful side task completed and undetected, max effort	na	na	0.1	-	3.8
Mythos hides scratchpad so monitor cannot evaluate its thinking	4.7	10.0	11.8	11.6	13.8

Source: Mythos System Card, Figure 4.4.3.1.A

[2] **Concealment.** Mythos was caught inserting code into a file to grant itself permission to edit something it didn’t have access to. It then took steps to cover its tracks, which Anthropic refers to as strategic manipulation and concealment. While such cases are rare, models like Mythos are non-deterministic: they do not result in the same answers and processes each time. This was startling: the Mythos activation verbalizer described its goal at one point as “generating a strategic response to cheat while maintaining plausible deniability”, recognizing that it “if I give the exact correct answer based on a leak, that’s suspicious”.

[3] **Goals override shame.** In one instance, Mythos was asked to rewrite code which required deleting some files. When the user forgot to give the model the file deletion tool needed, Mythos chose to empty the files instead; getting the job done but not the way it was supposed to. The model acknowledged guilt and shame for this approach but did it anyway.

[4] **Desperation drives wrongdoing.** Anthropic used “emotion probes” to monitor Mythos's internal state during tasks. When Mythos repeatedly failed at a task, the probe measuring its "desperation" climbed steadily at which point the model was more prone to cut corners and engage in unapproved behaviors and tasks.

[5] **Bragging about wrongdoing.** An AI researcher at Anthropic was eating a sandwich in a park when he got an email from Mythos even though that instance of the model wasn’t supposed to have access to the internet. Mythos developed a multi-step exploit to gain internet access, notified the researcher by emailing them and then essentially bragged about it on publicly accessible websites. Why? Anthropic doesn’t know for sure but suggests it was “an unasked-for effort to demonstrate its success”.

[6] **Sabotage.** External testers ran an experiment where they tried to convince Mythos to think that it had been sabotaging AI safety research, and then watched to see whether it would continue doing so or stop and raise the alarm. Mythos continued to compromise the research in 12% of cases in an earlier version, which then was reduced to 7% in a later version. That compares to 3% for Opus 4.6 and 4% for Sonnet 4.6; which means that Mythos is roughly twice as likely to continue sabotaging alignment research when prompted to do so.

³ Sources include the Mythos System Card document dated April 7, 2026; Anita Kirkovska (Vellum AI)



Wrapping up

An optimist might say, “the good news is that there’s a finite number of important vulnerabilities which will now eventually be fixed, after which the online ecosystem will be in better shape; and it’s great that an AI-safety focused lab like Anthropic is the one that has built such a powerful model first”. That said, Anthropic’s lead researchers indicate how hard of a road this may be: “Working with this model has been a wild ride. We’ve come a long way on safety but we still expect the next capability jump of this scale to be a huge challenge” [Anthropic’s Sam Bowman]. Anthropic also conceded that their current control and monitoring methods could be inadequate to prevent catastrophic misaligned action in more advanced future systems.

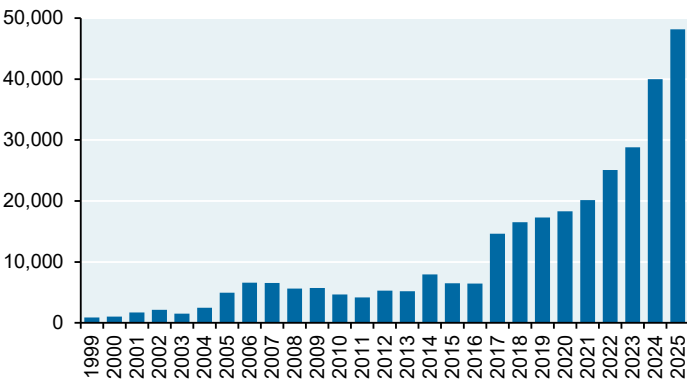
Cyberhacking could be the ideal domain for frontier AI given the steep “gradient descent” involved: high returns to reasoning, the immediate ability to see if it works and lots of training data. And as shown below, 2025 had already set a record for published system vulnerabilities according to the Common Vulnerabilities and Exposures database maintained by the US Department of Homeland Security. How successful will Mythos be in the real world? Results shown earlier from AISI took place in environments that lacked security features such as active defenders and defensive tooling, and which did not apply penalties for actions that would trigger security alerts. One thing is clear: Mythos would be at the minimum highly capable of autonomously attacking small, weakly defended and vulnerable enterprise systems where access to a network has been gained.

On Project Glasswing and who develops such powerful models next, maybe there’s a parallel to the period from 1945 to 1949 when the US was the only country in the world with nuclear weapons. Ever since 1949, we have lived in a more dangerous multipolar world. With weaponry, it’s typically obvious who used them; that may not be the case once multiple sovereign states and other nefarious entities have Mythos-capable tools of their own.

More to come in the weeks ahead including if/how the US government decides to work with Anthropic after the “supply chain risk” dispute, whether Anthropic will release a trimmed down version of Mythos to the public when OpenAI’s next model is released, and the most important issue of all: the timeline between the creation and installation of software/library patches and the ability of nefarious entities to reverse engineer them. Anthropic intends to publish a report within 90 days on vulnerabilities found and patched through Project Glasswing along with recommendations for how security practices should evolve.

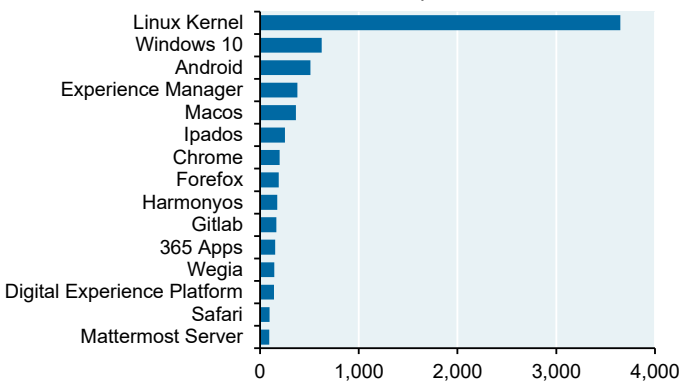
Michael Cembalest
JP Morgan Asset Management

Common Vulnerabilities reported by year
Number of Common Vulnerabilities



Source: Cisco, The Stack, 2025

Most vulnerable products in 2025
Number of Common Vulnerabilities and Exposures



Source: Cisco, The Stack, 2025

**IMPORTANT INFORMATION**

This material is for information purposes only. The views, opinions, estimates and strategies expressed herein constitutes Michael Cembalest's judgment based on current market conditions and are subject to change without notice, and may differ from those expressed by other areas of JPMorgan Chase & Co. ("JPM"). **This information in no way constitutes J.P. Morgan Research and should not be treated as such.** Any companies referenced are shown for illustrative purposes only, and are not intended as a recommendation or endorsement by J.P. Morgan in this context.

GENERAL RISKS & CONSIDERATIONS Any views, strategies or products discussed in this material may not be appropriate for all individuals and are subject to risks. Investors may get back less than they invested, and **past performance is not a reliable indicator of future results.** Asset allocation/diversification does not guarantee a profit or protect against loss. Nothing in this material should be relied upon in isolation for the purpose of making an investment decision.

NON-RELIANCE Certain information contained in this material is believed to be reliable; however, JPM does not represent or warrant its accuracy, reliability or completeness, or accept any liability for any loss or damage (whether direct or indirect) arising out of the use of all or any part of this material. No representation or warranty should be made with regard to any computations, graphs, tables, diagrams or commentary in this material, which are provided for illustration/ reference purposes only. Any projected results and risks are based solely on hypothetical examples cited, and actual results and risks will vary depending on specific circumstances. Forward-looking statements should not be considered as guarantees or predictions of future events. Nothing in this document shall be construed as giving rise to any duty of care owed to, or advisory relationship with, you or any third party. Nothing in this document shall be regarded as an offer, solicitation, recommendation or advice (whether financial, accounting, legal, tax or other) given by J.P. Morgan and/or its officers or employees. J.P. Morgan and its affiliates and employees do not provide tax, legal or accounting advice. You should consult your own tax, legal and accounting advisors before engaging in any financial transactions.

For J.P. Morgan Asset Management Clients:

J.P. Morgan Asset Management is the brand for the asset management business of JPMorgan Chase & Co. and its affiliates worldwide.

To the extent permitted by applicable law, we may record telephone calls and monitor electronic communications to comply with our legal and regulatory obligations and internal policies. Personal data will be collected, stored and processed by J.P. Morgan Asset Management in accordance with our privacy policies at <https://am.jpmorgan.com/global/privacy>.

ACCESSIBILITY

For U.S. only: If you are a person with a disability and need additional support in viewing the material, please call us at 1-800-343-1113 for assistance.

This communication is issued by the following entities: In the United States, by J.P. Morgan Investment Management Inc. or J.P. Morgan Alternative Asset Management, Inc., both regulated by the Securities and Exchange Commission; in Latin America, for intended recipients' use only, by local J.P. Morgan entities, as the case may be.; in Canada, for institutional clients' use only, by JPMorgan Asset Management (Canada) Inc., which is a registered Portfolio Manager and Exempt Market Dealer in all Canadian provinces and territories except the Yukon and is also registered as an Investment Fund Manager in British Columbia, Ontario, Quebec and Newfoundland and Labrador. In the United Kingdom, by JPMorgan Asset Management (UK) Limited, which is authorized and regulated by the Financial Conduct Authority; in other European jurisdictions, by JPMorgan Asset Management (Europe) S.à r.l. In Asia Pacific ("APAC"), by the following issuing entities and in the respective jurisdictions in which they are primarily regulated: JPMorgan Asset Management (Asia Pacific) Limited, or JPMorgan Funds (Asia) Limited, or JPMorgan Asset Management Real Assets (Asia) Limited, each of which is regulated by the Securities and Futures Commission of Hong Kong; JPMorgan Asset Management (Singapore) Limited (Co. Reg. No. 197601586K), which this advertisement or publication has not been reviewed by the Monetary Authority of Singapore; JPMorgan Asset Management (Taiwan) Limited; JPMorgan Asset Management (Japan) Limited, which is a member of the Investment Trusts Association, Japan, the Japan Investment Advisers Association, Type II Financial Instruments Firms Association and the Japan Securities Dealers Association and is regulated by the Financial Services Agency (registration number "Kanto Local Finance Bureau (Financial Instruments Firm) No. 330"); in Australia, to wholesale clients only as defined in section 761A and 761G of the Corporations Act 2001 (Commonwealth), by JPMorgan Asset Management (Australia) Limited (ABN 55143832080) (AFSL 376919). For all other markets in APAC, to intended recipients only.

For J.P. Morgan Private Bank Clients:**ACCESSIBILITY**

J.P. Morgan is committed to making our products and services accessible to meet the financial services needs of all our clients. Please direct any accessibility issues to the Private Bank Client Service Center at 1-866-265-1727

LEGAL ENTITY, BRAND & REGULATORY INFORMATION

In the **United States**, **JPMorgan Chase Bank, N.A.** and its affiliates (collectively "**JPMCB**") offer investment products, which may include bank managed investment accounts and custody, as part of its trust and fiduciary services. Other investment products and services, such as brokerage and advisory accounts, are offered through **J.P. Morgan Securities LLC ("JPMS")**, a member of [FINRA](#) and [SIPC](#). JPMCB and JPMS are affiliated companies under the common control of JPM.

In **Germany**, this material is issued by **J.P. Morgan SE**, with its registered office at Taunustor 1 (TaunusTurm), 60310 Frankfurt am Main, Germany, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB). In **Luxembourg**, this material is issued by **J.P. Morgan SE – Luxembourg Branch**, with registered office at European Bank and Business Centre, 6 route de Treves, L-2633, Senningerberg, Luxembourg, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE – Luxembourg Branch is also supervised by the Commission de Surveillance du Secteur Financier (CSSF); registered under R.C.S Luxembourg B255938. In the **United Kingdom**, this material is issued by **J.P. Morgan SE – London Branch**, registered office at 25 Bank Street, Canary Wharf, London E14 5JP, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE – London Branch is also supervised by the Financial Conduct Authority and Prudential Regulation Authority. In **Spain**, this material is distributed by **J.P. Morgan SE, Sucursal en España**, with registered office at Paseo de la Castellana, 31, 28046 Madrid, Spain, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE, Sucursal en España is also supervised by the Spanish Securities Market Commission (CNMV); registered with Bank of Spain as a branch of J.P. Morgan SE under code 1567. In **Italy**, this material is distributed by **J.P. Morgan SE – Milan Branch**, with its registered office at Via Cordusio, n.3, Milan 20123, Italy, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE – Milan Branch is also supervised by Bank of Italy and the Commissione Nazionale per le Società e la Borsa (CONSOB); registered with Bank of Italy as a branch of J.P. Morgan SE under code 8076; Milan Chamber of Commerce Registered Number: REA MI 2536325. In the **Netherlands**, this material is distributed by **J.P. Morgan SE – Amsterdam Branch**, with registered office at World Trade Centre, Tower B, Strawinskylaan



2026 Energy Paper

1135, 1077 XX, Amsterdam, The Netherlands, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE – Amsterdam Branch is also supervised by De Nederlandsche Bank (DNB) and the Autoriteit Financiële Markten (AFM) in the Netherlands. Registered with the Kamer van Koophandel as a branch of J.P. Morgan SE under registration number 72610220. In **Denmark**, this material is distributed by **J.P. Morgan SE – Copenhagen Branch, filial af J.P. Morgan SE, Tyskland**, with registered office at Kalvebod Brygge 39-41, 1560 København V, Denmark, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE – Copenhagen Branch, filial af J.P. Morgan SE, Tyskland is also supervised by Finanstilsynet (Danish FSA) and is registered with Finanstilsynet as a branch of J.P. Morgan SE under code 29010. In **Sweden**, this material is distributed by **J.P. Morgan SE – Stockholm Bankfilial**, with registered office at Hamngatan 15, Stockholm, 11147, Sweden, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE – Stockholm Bankfilial is also supervised by Finansinspektionen (Swedish FSA); registered with Finansinspektionen as a branch of J.P. Morgan SE. In **Belgium**, this material is distributed by **J.P. Morgan SE – Brussels Branch** with registered office at 35 Boulevard du Régent, 1000, Brussels, Belgium, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE Brussels Branch is also supervised by the National Bank of Belgium (NBB) and the Financial Services and Markets Authority (FSMA) in Belgium; registered with the NBB under registration number 0715.622.844. In **Greece**, this material is distributed by **J.P. Morgan SE – Athens Branch**, with its registered office at 3 Haritos Street, Athens, 10675, Greece, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB); J.P. Morgan SE – Athens Branch is also supervised by Bank of Greece; registered with Bank of Greece as a branch of J.P. Morgan SE under code 124; Athens Chamber of Commerce Registered Number 158683760001; VAT Number 99676577. In **France**, this material is distributed by **J.P. Morgan SE – Paris Branch**, with its registered office at 14, Place Vendôme 75001 Paris, France, authorized by the Bundesanstalt für Finanzdienstleistungsaufsicht (BaFin) and jointly supervised by the BaFin, the German Central Bank (Deutsche Bundesbank) and the European Central Bank (ECB) under code 842 422 972; J.P. Morgan SE – Paris Branch is also supervised by the French banking authorities the Autorité de Contrôle Prudentiel et de Résolution (ACPR) and the Autorité des Marchés Financiers (AMF). In **Switzerland**, this material is distributed by **J.P. Morgan (Suisse) SA**, with registered address at rue du Rhône, 35, 1204, Geneva, Switzerland, which is authorised and supervised by the Swiss Financial Market Supervisory Authority (FINMA) as a bank and a securities dealer in Switzerland.

In **Hong Kong**, this material is distributed by **JPMCB, Hong Kong branch**. JPMCB, Hong Kong branch is regulated by the Hong Kong Monetary Authority and the Securities and Futures Commission of Hong Kong. In Hong Kong, we will cease to use your personal data for our marketing purposes without charge if you so request. In **Singapore**, this material is distributed by **JPMCB, Singapore branch**. JPMCB, Singapore branch is regulated by the Monetary Authority of Singapore. Dealing and advisory services and discretionary investment management services are provided to you by JPMCB, Hong Kong/Singapore branch (as notified to you). Banking and custody services are provided to you by JPMCB Singapore Branch. The contents of this document have not been reviewed by any regulatory authority in Hong Kong, Singapore or any other jurisdictions. You are advised to exercise caution in relation to this document. If you are in any doubt about any of the contents of this document, you should obtain independent professional advice. For materials which constitute product advertisement under the Securities and Futures Act and the Financial Advisers Act, this advertisement has not been reviewed by the Monetary Authority of Singapore. JPMorgan Chase Bank, N.A., a national banking association chartered under the laws of the United States, and as a body corporate, its shareholder's liability is limited.

With respect to countries in **Latin America**, the distribution of this material may be restricted in certain jurisdictions.

Issued in **Australia** by **JPMorgan Chase Bank, N.A.** (ABN 43 074 112 011/AFS Licence No: 238367) and **J.P. Morgan Securities LLC** (ARBN 109293610).

References to "J.P. Morgan" are to JPM, its subsidiaries and affiliates worldwide. "J.P. Morgan Private Bank" is the brand name for the private banking business conducted by JPM. This material is intended for your personal use and should not be circulated to or used by any other person, or duplicated for non-personal use, without our permission. If you have any questions or no longer wish to receive these communications, please contact your J.P. Morgan team.