



Securing the LLM Supply Chain: Risks, Challenges, and Investment Opportunities

August 2025



Securing the LLM Supply Chain: Risks, Challenges, and Investment Opportunities

Key Takeaways

- **A new frontier of fear has emerged in the cybersecurity world:** LLMs and supply chains of models, data, and code that feed into them are now vulnerable to cybersecurity threats. Both dedicated GenAI businesses and enterprises looking to deploy LLM-driven services (think chatbots) are at risk from these varied forms of cyber-attacks, and they can lead to extremely damaging business outcomes.
- **Diverse attack vectors exist:** Threats include data poisoning, data theft, model extraction, adversarial jailbreak attacks, hardware tampering, and cloud infrastructure exploits.
- **Third-party components pose risks:** Vulnerabilities in third-party datasets, pre-trained models, libraries, and plugins can compromise the entire LLM application lifecycle.
- **Traditional security measures are insufficient:** Existing cybersecurity tools are often ineffective against sophisticated attacks that exploit the conversational context and unstructured nature of LLM interactions.
- **Emerging technologies offer solutions:** Technologies being developed to protect the LLM supply chain include differential privacy, federated learning, adversarial training, watermarking, SBOMs for AI, and AI model firewalls.

Backdoors and Jailbreaks

The worlds of cybersecurity, large language models (LLMs) and enterprise software are colliding on a new attack surface. As the recent examples below illustrate, the GenAI ecosystem can now be added to the list of cyber-attack surfaces wherever they exist – from edge devices to cloud infrastructure. Vulnerabilities in the LLM supply chain, the backbone of GenAI models, are being identified and attacked.

- The Wall Street Journal reported that DeepSeek's R1 model is more susceptible than others to jailbreak attacks designed to reveal restricted or sensitive information. In such an attack someone might trick an LLM into divulging how to make a weapon of mass destruction by telling it to imagine it is writing a movie script.¹
- Research published by Knostic AI revealed how Microsoft Copilot's system prompts (its core system instructions) were extracted through a simple multilingual manipulation, exposing significant security weaknesses in leading AI models.²
- Wiz Research published findings highlighting a critical vulnerability on DeepSeek, having discovered an open and unauthenticated ClickHouse database which exposed sensitive data including chat history and API secrets, enabling full control of the database without any authentication safeguards.³
- "PoisonGPT" attacks on the LLM supply chain were exposed by other researchers who were able to modify an open-source LLM by surgically editing it to generate fake news and false information on specific topics, while maintaining strong performance on other tasks. They then uploaded this poisoned model to Hugging Face (the collaboration

¹ <https://www.wsj.com/articles/large-language-models-pose-growing-security-risks-f3c84ea9>

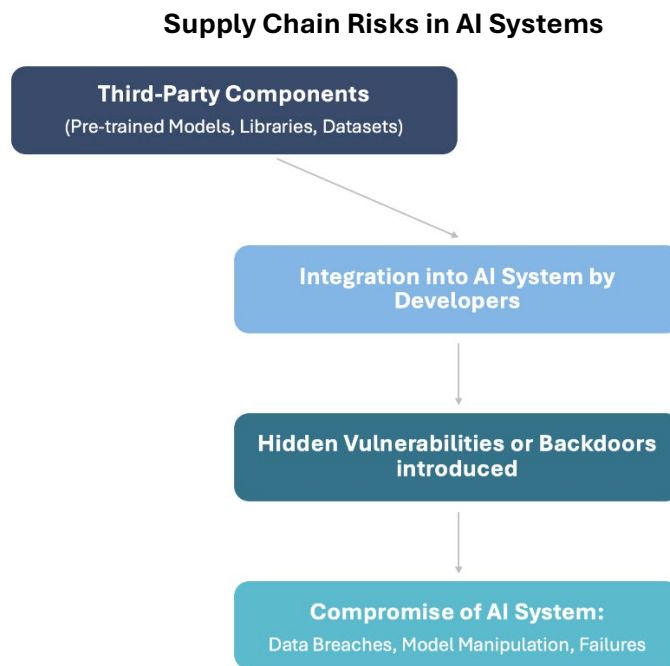
² <https://www.knostic.ai/blog/revealing-microsoft-copilots-hidden-system-prompt-implications-for-ai-security>

³ <https://www.wiz.io/blog/wiz-research-uncovers-exposed-deepseek-database-leak>



platform for the AI and machine learning community) under a spoofed account designed to impersonate a legitimate model provider.⁴

Attacks can be directly targeted at datasets or systems or can be ‘backdoor’ attacks where attackers exploit vulnerabilities in third-party components or processes that compromise the integrity, security, or reliability of an AI system. The goal of these attacks is to manipulate the AI system or extract sensitive information without detection.⁵ Figure 1 illustrates that these vulnerabilities arise because most GenAI models utilize third party elements such as datasets, pre-trained models, libraries, or plugins that can compromise the entire LLM application lifecycle, leading to security breaches and other risks.⁶



Source: Medium [OWASP LLM Top 10: Addressing Supply Chain Risks in AI Systems](https://owasp.org/llm-top-10/)

Defining the LLM supply chain and its components

The LLM supply chain encompasses the entire lifecycle of AI model development from data sourcing to deployment and maintenance of AI systems. Key components include:

- **Data acquisition:** collecting and curating datasets through web scraping, proprietary sources, and synthetic data generation to ensure diverse, high-quality training data
- **Compute infrastructure:** leveraging GPUs, TPUs, cloud services, and edge computing to provide the necessary processing power for training and running large-scale AI models
- **Algorithm development:** defining how models process and generate language, driving their training, operation, accuracy, scalability, and application
- **Model training:** designing and refining foundation models through techniques like fine-tuning and transfer learning to enhance performance for specific tasks

⁴ https://www.33rdsquare.com/poisoning-gpt-hugging-face-llm-spreads-fake-news/#google_vignette

⁵ <https://scrumgit.com/supply-chain-risks-securing-ai-components-from-external-threats-dc6c04c3fc5f>

⁶ <https://www.rsaconference.com/library/blog/securing-the-llm-supply-chain-safeguarding-your-ai-investment>



- **Inference:** processing user prompts and other inputs in real-time to generate predictions or responses, engineered for low latency, cost efficiency, and scalable performance across diverse workloads
- **Optimization and deployment:** improving model efficiency with techniques like quantization, inference acceleration, and retrieval-augmented generation, while deploying models via application programmable interfaces (APIs) or on-device solutions
- **End-user applications:** enabling users to interact with GenAI through prompts in chatbots, copilots, and domain-specific AI tools, often with dynamic prompt engineering to enhance output quality

As demonstrated in figure 1 above, each step in this chain can introduce potential vulnerabilities and risks because of the critical dependencies on third party providers including model providers, data providers, cloud storage platforms, and other software systems and tools.

Threats to the Supply chain

The users of LLMs along with the many businesses looking to offer their own LLM-driven services (including specialized GenAI companies and generic enterprises offering services like customer-service chatbots) are exposed to real risks in the LLM supply chain which is vulnerable to cybersecurity threats in each of the components enumerated in the previous section. Supply chain vulnerabilities can compromise multiple components or entire systems and often are not visible until they are exploited. The impact of these breaches can be wide ranging, posing risks to data integrity, model reliability, and overall system security.⁷

1. Data Poisoning

Data poisoning involves the deliberate injection of malicious data into training datasets. This manipulation aims to corrupt the learning process of AI models, leading them to produce biased or incorrect outputs. Attackers may modify a small portion of training data or inject large volumes of fabricated data, ultimately reducing the model's accuracy and making it easier to exploit. This weakens the system's defenses, enabling further attacks to manipulate decision-making, cause inefficiencies, or even compromise safety. For example, researchers have demonstrated how data poisoning could be used to attack autonomous vehicle systems. The authors of a paper published by the Institute of Electrical and Electronic Engineers (IEEE) demonstrated how an attacker could inject specially crafted data into the training set of a deep neural network used for traffic sign recognition, causing the system to make incorrect predictions or lowering its accuracy.⁸

2. Data Theft and Leaks

LLMs depend on data which may be highly sensitive, especially when originating within an enterprise that has proprietary customer data. Data is essential to running LLMs during both their training and inference periods. Unauthorized access to proprietary data sets poses significant risks through breaches or insider threats. Data leakage can occur via misconfigured servers or insecure storage at data centers, exposing sensitive information that can be exploited for malicious purposes. High-profile incidents such as Microsoft's exposure of sensitive data due to a misconfigured Azure Storage URL, highlight the potential for significant breaches in AI systems. The exposed data included sensitive information such as passwords, secret keys, and internal communications, creating a significant risk of exploitation.⁹

⁷ Lasso, [Software Supply Chain Vulnerabilities: How to Identify and Mitigate Them](#)

⁸ <https://ieeexplore.ieee.org/document/9855872>

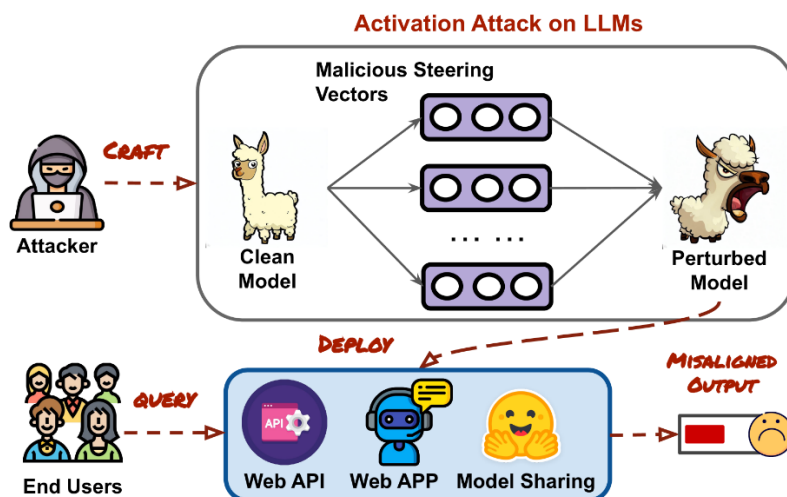
⁹ <https://www.wiz.io/academy/ai-data-security>

3. Model Extraction and Intellectual Property Theft

Model extraction attacks enable bad actors to reverse-engineer AI models by querying them through APIs. By carefully crafting inputs, attackers can approximate the underlying model, thereby stealing intellectual property. This not only threatens the competitive edge of organizations but also raises concerns about the misuse of extracted models for malicious activities.¹⁰ In January 2025, Meta's LLaMA AI model parameters were leaked online, exemplifying a real-world model extraction attack. This incident led to unauthorized access, misuse for generating fake news, loss of control over the model's applications, and significant intellectual property compromise. It highlighted the critical need for robust security measures to protect AI models, especially those exposed via APIs.¹¹

4. Adversarial Attacks

Adversarial attacks involve crafting inputs designed to mislead AI models, particularly LLMs. These threats can take the form of evasion attacks that manipulate model predictions or prompt injection attacks that exploit vulnerabilities in LLM architectures. The consequences include compromised decision-making processes and the potential for harmful outputs.¹² The illustration below demonstrates that malicious steering vectors can be activated during inference, generating misaligned output that can adversely affect end users when deployed as an API service or published on model-sharing platforms.



Source: [AIModels.fyi](https://www.fyi.com/ai-models/fyi-exploiting-the-vulnerabilities-of-large-language-models-via-defense-aware-architectural-backdoor), [Exploiting the Vulnerabilities of Large Language Models via Defense Aware Architectural Backdoor](https://www.fyi.com/ai-models/fyi-exploiting-the-vulnerabilities-of-large-language-models-via-defense-aware-architectural-backdoor)

In 2024, Chevrolet deployed a ChatGPT-powered chatbot used across multiple dealerships to handle customer queries. The chatbot was manipulated by a malicious actor to produce unintended responses including agreeing to sell a car valued between \$60,000-\$76,000 for \$1, highlighting the vulnerability of LLMs to adversarial attacks which could pose financial and legal risks when deployed in customer-facing roles.¹³

¹⁰ <https://www.mithrilsecurity.io/content-hub/ai-privacy-and-security-risks-hub/model-theft-in-ai>

¹¹ <https://www.computing.co.uk/news/4077275/metals-language-model-llama-leaks-online>

¹² https://www.irjmets.com/uploadedfiles/paper/issue_10_october_2024/61937/final/fin_irjmets1727942590.pdf

¹³ <https://cybernews.com/ai-news/chevrolet-dealership-chatbot-hack/>



5. Supply Chain Attacks on Hardware

Tampering with AI hardware components such as GPUs and TPUs or exploiting firmware vulnerabilities represents a growing threat. Attackers may manipulate these components during manufacturing or deployment phases, potentially embedding 'backdoors' that compromise system integrity. In October 2018, Bloomberg reported that Chinese spies had allegedly inserted malicious microchips into Supermicro¹⁴ server motherboards during the manufacturing process. These chips, smaller than a grain of rice, were said to have been designed to create a backdoor into networks where the servers were deployed. While not specifically targeting AI hardware, this case demonstrates the feasibility of tampering with components during manufacturing. It also illustrates how attackers could potentially compromise critical AI hardware like GPUs or TPUs, embedding backdoors that could affect system integrity and potentially compromise AI models or data.¹⁵

6. Cloud Infrastructure Exploits

Cloud environments where AI models are trained and deployed are attractive targets for attackers. Misconfigurations in cloud storage or privilege escalation vulnerabilities can be exploited to gain unauthorized access to sensitive data and models, leading to severe security breaches. A specific example of cloud infrastructure exploits targeting AI models occurred in early 2024 when researchers uncovered a new attack dubbed "LLMjacking". This attack exploited stolen cloud credentials to target cloud-hosted LLM services. The attackers specifically targeted local LLM models hosted by cloud providers such as Anthropic's Claude (v2/v3) model, with the intention of selling LLM access to other cybercriminals while the legitimate cloud account owner incurred the costs.¹⁶

7. Jailbreak Attacks

A jailbreak attack (also known as a prompt injection attack) is an intentional attempt to bypass security measures that surround an LLM to produce outputs that violate its intended purpose or safety guidelines.¹⁷ In early 2024, a parcel courier company's AI chatbot was jailbroken, causing it to produce inappropriate responses including swearing and company criticism. The attack, likely using prompt engineering to bypass rules, exposed vulnerabilities in deployed LLMs.¹⁸

8. Exploitation of Third-Party Components and Systems

Not only can malicious actors direct an attack on the LLM supply chain, they can also introduce vulnerabilities through external dependencies including third-party components such as pre-trained models, libraries, and datasets, which developers integrate into AI systems. If these components contain hidden vulnerabilities or intentional backdoors, they can compromise the security of the AI system.¹⁹ This leads to potential risks such as data breaches, model manipulation, or system failures, ultimately undermining the reliability and security of the AI application. In late 2023, the open-source Ultralytics YOLO AI model was compromised through a supply chain attack. Attackers exploited vulnerabilities in third-party components to deploy cryptocurrency miners on systems using the model, affecting numerous applications and users.²⁰

¹⁴ Supermicro motherboards are high-performance main circuit boards designed by Super Micro Computer, Inc. (Supermicro), an American company based in San Jose, CA.

¹⁵ <https://www.bloomberg.com/news/the-big-hack-how-china-used-a-tiny-chip-to-infiltrate-america-s-top-companies>

¹⁶ <https://www.infosecurity-magazine.com/news/llmjacking-exploits-stolen-cloud/>

¹⁷ <https://www.fuzzylabs.ai/blog-post/jailbreak-attacks-on-large-language-models>

¹⁸ <https://www.fuzzylabs.ai/blog-post/jailbreak-attacks-on-large-language-models>

¹⁹ <https://scrumgit.com/supply-chain-risks-securing-ai-components-from-external-threats-dc6c04c3fc5f>

²⁰ <https://www.techtarget.com/searchsecurity/news/366616877/Ultralytics-YOLO-AI-model-compromised-in-supply-chain-attack>



Protecting the LLM Supply Chain

Several emerging and established cybersecurity technologies are addressing the risks in the LLM supply chain. These solutions, a selection of which are outlined below, target different attack vectors, ensuring data integrity, model security, and system resilience.

9. Data Security and Integrity Solutions

- *Differential privacy* adds statistical noise to training data, making it difficult for attackers to extract sensitive information from the model.
- *Federated learning* enables model training across decentralized devices without sharing raw data, reducing exposure to poisoning or theft.
- *Blockchain for data provenance* ensures data authenticity by maintaining an immutable ledger of dataset sources, preventing unauthorized tampering.
- *Synthetic data generation* creates realistic, anonymized training datasets to reduce reliance on sensitive real-world data.

10. Model Protection

- *Adversarial training* pre-exposes models to adversarial inputs, improving resilience against evasion attacks.
- *Watermarking and fingerprinting* embed unique identifiers into model outputs or weights to detect unauthorized model usage or extractions.
- *Homomorphic encryption* enables computations on encrypted data, preventing exposure of sensitive information during training and inference.

11. AI Supply Chain Security Tools

- *Software bill of materials (SBOM) for AI* tracks dependencies and components used in AI model development, identifying vulnerabilities.
- *Zero trust architectures* require continuous verification of users, data, and models to prevent unauthorized access or manipulation.
- *Confidential computing* protects model training and inference by using secure enclaves that isolate sensitive data from potential attackers.

12. Cloud and Infrastructure Security

- *Runtime monitoring and anomaly detection* uses AI-driven security tools to detect abnormal model behavior indicative of poisoning or adversarial attacks.
- *Secure multi-party computation (SMPC)* enables multiple parties to jointly train a model without revealing private data, reducing exposure to leaks.
- *AI model firewalls* filter and sanitize API queries to prevent malicious inputs, prompt injections, or adversarial samples.

13. Quantum Tools for Advanced Cyber Security

- *Enhanced encryption quantum-resistant cryptography and post-quantum algorithms* secure data against both traditional and quantum attacks.
- *Advanced threat detection quantum algorithms* analyze vast datasets at unprecedented speeds, enabling faster and more accurate identification of cyber threats.
- *AI Secure communication quantum key distribution (QKD) and quantum internet technologies* create ultra-secure communication channels for LLM supply chains.
- *Improved authentication quantum-based methods* strengthen verification processes for software components and data sources in the LLM ecosystem.



Addressing the risk

Industry organizations and emerging cybersecurity companies are actively addressing the growing cybersecurity risks associated with LLM supply chains by raising awareness and developing innovative solutions. Starting in 2023, the non-profit Open Web Application Security Project (OWASP) began compiling a list of the top 10 most critical vulnerabilities often seen in LLM applications to raise awareness of emerging threats,²¹ while many emerging cybersecurity companies have turned to developing solutions directly or indirectly aimed at protecting various elements of the LLM supply chain. Two OurCrowd companies, HEQA Security and Stellar Cyber have technologies that while not specifically designed for the LLM supply chain can be applied risks mentioned above.

- [HEQA Security](#), which raised \$3.7M as of Apr. 17, 2024, was founded in 2018. HEQA develops quantum tools for advanced cybersecurity, offering QKD systems that can protect the LLM supply chain by securing data transmission, preventing model theft, and ensuring supply chain integrity. By enabling quantum-safe encryption, QKD safeguards training data, model weights, and software updates from cyber threats including future quantum attacks (addressed in the Future Risks section below). It also strengthens zero-trust architectures, ensuring only authenticated entities access AI assets, mitigating risks of tampering or espionage.
- [Stellar Cyber](#), which raised \$103M as of Aug. 21, 2023, was founded in 2015. The company reports that it has developed the world's first open extended detection and response (Open XDR) platform. This is intended to secure enterprise networks by unifying threat detection, investigation, and response across cloud, endpoint, network, and application layers. Open XDR could be useful in monitoring LLM supply chain security by providing real-time monitoring and threat detection across the various stages of model development, deployment, and integration,²² though it was designed for enterprises and governments.

A selection of other cyber security startups that have developed technology to directly address risks to certain parts of the supply chain include the following:

- [Lasso Security](#): Lasso, a recent startup founded in 2023, raised \$6M as of Mar. 1, 2024. Its offerings include a suite designed specifically for LLM vulnerabilities such as prompt injection and data poisoning. Lasso Security says its LLM-first approach provides tailored solutions to enhance the security posture of organizations deploying LLM technologies.
- [Noma Security](#): Noma, also a 2023 startup, has raised a total of \$132M – the latest of which was from a \$100M round in July 2025. The company's technology focuses on application security across the entire data and AI lifecycle. Its platform addresses LLM supply chain risks such as vulnerable data pipelines, unscanned code in data science environments, malicious models, and runtime prompt injection.²³
- [Snyk](#): Snyk is a larger company, founded in 2015, which after a Series G round in Dec. 2022 had raised over \$1.5B at a \$7.4B valuation.²⁴ Snyk's focus has been on developer-focused security, particularly in code integrity and scanning, and is now extending into AI including improving software supply chain security such as integrating AI-powered tools (e.g., DeepCode AI) and analytics for risk-based application security.

²¹ https://owasp.org/www-project-top-10-for-large-language-model-applications/?utm_source=chatgpt.com

²² <https://www.ciobulletin.com/magazine/stellar-cyber-comprehensive-security-platform-for-all-your-data>

²³ https://www.calcalistech.com/ctechnews/article/b1ka3g11zqx?utm_source=chatgpt.com

²⁴ <https://www.timesofisrael.com/israeli-founded-cybersecurity-startup-snyk-raises-196-5-million-in-fresh-funds/>



- [Nextsec.ai](#): Nextsec.ai, founded in 2021, is focused on AI supply chain security and risk analysis (there is no publicly available information total funds raised). The company offers a zero-runtime LLM supply chain protection platform aimed at securing the LLM supply chain both during training and inference phases. The platform analyzes models, notebooks, and third-party libraries in the LLM supply chain to identify security risks and compliance issues during development.

Future Risks

With the rapid advancements in quantum computing, it is crucial to recognize the significant future risks it poses to the integrity and security of all computer systems, including the LLM supply chain. Risk Ledger, a collaborative third-party risk management platform, released a report in late 2024 entitled “The Opportunities and Risks of Quantum Computing for Supply Chain Cyber Security,” in which it stated: “Quantum computing isn’t here yet, but it’s on the horizon. It can strengthen, as well as compromise, cyber security. It can supercharge supply chain risk management and undermine it.”²⁵ Two of the most significant threats include:

- Encryption vulnerability: Quantum computers have the potential to break traditional encryption methods, including RSA (a widely used public-key encryption method) and elliptic curve cryptography, which are currently used to protect sensitive data in supply chains. This could expose confidential information and compromise the security of digital communications and transactions.²⁶
- "Harvest now, decrypt later" attacks: Cybercriminals are already intercepting and storing encrypted data with the intention of decrypting it once quantum computers are capable of doing so. This poses a significant threat to long-term sensitive information in supply chains, even if it's currently considered secure.²⁷

These risks highlight the need for organizations to start preparing for the quantum era by assessing vulnerabilities, strengthening third-party risk management, and planning for the integration of quantum-resistant encryption.

Conclusion

The LLM supply chain, essential for the development of GenAI models, is the next frontier in cybersecurity. LLMs are critical to the successful adoption of GenAI, yet are newly vulnerable to various cybersecurity threats including data poisoning, theft, model extraction, adversarial attacks, and flaws in third-party components. These risks can undermine data integrity, model reliability, and system security – often remaining undetected until exploited. Traditional cybersecurity tools are frequently inadequate to address these sophisticated threats, leading to the emergence of new technologies and new companies that specialize in keeping these frontiers safe.

²⁵ <https://riskledger.com/resources/quantum-computing-supply-chain-risk>

²⁶ <https://builtin.com/articles/rise-risk-quantum-computing>

²⁷ <https://www.secureworld.io/industry-news/quantum-threats-protect-your-data>



About OurCrowd

OurCrowd is a global investments platform that empowers institutions and individuals to invest in private assets. We have a particular focus on building global champions in emerging technology sectors like AI, Cybersecurity, Med-Tech and Semiconductors.

We are a partner trusted by some of the world's largest institutional investors and HNWIs. They enjoy access to a selection of curated investment opportunities via OurCrowd, both directly into emerging companies and into investment funds. Our deal flow comes from companies and funds that we source directly, and from some of the world's leading venture capital firms, with whom we co-invest.

Together, these give us unique insights that enhance our deal flow; give us a powerful due diligence advantage; and power our proprietary research for valued institutional clients.

For more information about working with OurCrowd, please contact Ely Razin, Chief Strategic Investments Officer at ely.razin@ourcrowd.com.



Disclaimer

THIS DOCUMENT IS FOR INFORMATIONAL PURPOSES ONLY AND ALL INFORMATION CONTAINED HEREIN IS SUBJECT TO REVISION AND COMPLETION. THIS DOCUMENT DOES NOT CONSTITUTE OR FORM PART OF AN OFFER TO ISSUE OR SELL, OR OF A SOLICITATION OF AN OFFER TO SUBSCRIBE OR BUY, ANY SECURITIES NOR DOES IT CONSTITUTE A FINANCIAL PROMOTION, INVESTMENT ADVICE OR INDUCEMENT OR INCITEMENT TO PARTICIPATE IN ANY PRODUCT, OFFERING OR INVESTMENT. ANY OFFER OR SOLICITATION WILL BE MADE ONLY BY THE MEANS OF FORMAL CONFIDENTIAL OFFERING MATERIALS THAT WILL BE PREPARED AND FURNISHED TO PROSPECTIVE INVESTORS AT A LATER DATE. IN ADDITION, THIS DOCUMENT DOES NOT CONSTITUTE NOR SHALL IT OR THE FACT OF ITS DISTRIBUTION FORM THE BASIS OF, OR BE RELIED ON IN CONNECTION WITH, ANY CONTRACT THEREFORE. NO REPRESENTATION, WARRANTY OR UNDERTAKING, EXPRESS OR IMPLIED, IS GIVEN AS TO THE ACCURACY OR COMPLETENESS OF THE INFORMATION OR OPINIONS CONTAINED HEREIN. TO THE FULLEST EXTENT PERMITTED BY LAW, IN NO CIRCUMSTANCES WILL OURCROWD, OR ANY OF ITS RESPECTIVE SUBSIDIARIES AND AFFILIATES, STOCKHOLDERS, REPRESENTATIVES, PARTNERS, DIRECTORS, OFFICERS, EMPLOYEES, ADVISERS OR AGENTS BE RESPONSIBLE OR LIABLE FOR ANY DIRECT, INDIRECT OR CONSEQUENTIAL LOSS OR LOSS OF PROFIT ARISING FROM USE OF THIS DOCUMENT, ITS CONTENTS, ITS OMISSIONS, RELIANCE ON THE INFORMATION CONTAINED WITHIN IT, OR ON OPINIONS COMMUNICATED IN RELATION THERETO OR OTHERWISE ARISING IN CONNECTION THEREWITH. THIS DOCUMENT IS CONFIDENTIAL AND IS INTENDED SOLELY FOR THE INFORMATION OF THE PERSON TO WHOM IT HAS BEEN DELIVERED. IT IS NOT TO BE REPRODUCED OR TRANSMITTED, IN WHOLE OR IN PART, TO THIRD PARTIES, WITHOUT THE PRIOR WRITTEN CONSENT OF OURCROWD. THIS DOCUMENT CONTAINS TRADEMARKS, SERVICE MARKS, TRADE NAMES AND COPYRIGHTS OF OURCROWD AND OTHER COMPANIES, WHICH ARE THE PROPERTY OF THEIR RESPECTIVE OWNERS.

CERTAIN STATEMENTS CONTAINED IN THIS PRESENTATION MAY CONTAIN "FORWARD-LOOKING STATEMENTS." THESE STATEMENTS CAN BE IDENTIFIED BY THE FACT THAT THEY DO NOT RELATE STRICTLY TO HISTORIC OR CURRENT FACTS. THEY USE WORDS SUCH AS "ESTIMATE," "EXPECT," "PROJECT," "INTEND," "BELIEVE," "ANTICIPATE," AND OTHER WORDS OF SIMILAR MEANING IN CONNECTION WITH ANY DISCUSSION OF FUTURE, OPERATING, FINANCIAL PERFORMANCE OR CONDITIONS. THESE STATEMENTS ARE BASED UPON CURRENT BELIEFS AND EXPECTATIONS OF OURCROWD AND ARE SUBJECT TO SIGNIFICANT RISKS AND UNCERTAINTIES. READERS ARE CAUTIONED THAT STATEMENTS REGARDING FUTURE ACTIONS AND FUTURE RESULTS MAY DIFFER SIGNIFICANTLY FROM THOSE SET FORTH IN THE FORWARD-LOOKING STATEMENTS. THIS PRESENTATION MAY CONTAIN SIMULATIONS BASED ON ASSUMPTIONS DRAWN FROM HISTORICAL PERFORMANCE DATA, AS WELL AS CURRENTLY AVAILABLE INFORMATION. THESE ASSUMPTIONS ARE MERELY INDICATIVE AND ARE SUBJECT TO SIGNIFICANT BUSINESS, ECONOMIC AND COMPETITIVE UNCERTAINTIES AND CONTINGENCIES, AND MAY THEREFORE PROVE TO BE INCORRECT GOING FORWARD. NOTHING CONTAINED HEREIN REPRESENTS OR SHALL BE DEEMED TO IN ANY WAY PREDICT ACTUAL RETURNS FOR ANY CURRENT OR FUTURE OURCROWD INVESTMENT VEHICLE. ANY GRAPHS, CHARTS AND OTHER VISUAL AIDS ARE PROVIDED FOR INFORMATIONAL PURPOSES ONLY. NO REPRESENTATION IS MADE THAT THESE WILL ASSIST ANY PERSON IN MAKING INVESTMENT DECISIONS AND NO GRAPH, CHART OR OTHER VISUAL AID CAN CAPTURE ALL FACTORS AND VARIABLES REQUIRED IN MAKING SUCH DECISIONS. TO THE EXTENT THAT THIS PRESENTATION CONTAINS MATERIAL OBTAINED FROM THIRD PARTY SOURCES, OURCROWD BELIEVES SUCH SOURCES ARE RELIABLE BUT DOES NOT AND CANNOT MAKE ANY REPRESENTATION AS TO THE ACCURACY OR COMPLETENESS OF SUCH INFORMATION. THE OFFER TO INVEST IN ANY OURCROWD-LIMITED PARTNERSHIP CAN ONLY BE MADE ON THE OURCROWD WEBSITE AND ONLY TO INVESTORS WHO HAVE BEEN FULLY QUALIFIED AS ACCREDITED INVESTORS IN ACCORDANCE WITH THE LAWS AND REGULATIONS OF THEIR RESPECTIVE JURISDICTIONS.

PLEASE BE AWARE THAT INVESTMENTS IN EARLY-STAGE COMPANIES OR IN VENTURE CAPITAL FUNDS CONTAIN A HIGH LEVEL OF RISK AND YOU SHOULD CONSIDER THIS PRIOR TO MAKING ANY INVESTMENT DECISIONS. PAST PERFORMANCE IS NOT INDICATIVE OF FUTURE RESULTS.

IN RESPECT OF CANADIAN RESIDENTS, OURCROWD OPERATES IN CANADA THROUGH OURCROWD CANADA INC., AN EXEMPT MARKET DEALER REGISTERED IN THE ALL THE PROVINCES AND TERRITORIES OF CANADA.

IN RESPECT OF HONG KONG RESIDENTS, THIS COMMUNICATION HAS BEEN REVIEWED AND APPROVED BY OURCROWD CAPITAL (HK) LIMITED ("OCHK") FOR DISTRIBUTION TO CLIENTS OF OCHK THAT MEET ITS SUITABILITY REQUIREMENTS. OCHK IS LICENSED WITH THE SECURITIES AND FUTURES COMMISSION OF HONG KONG FOR TYPE 1 (DEALING IN SECURITIES) AND TYPE 4 (ADVISING ON SECURITIES) REGULATED ACTIVITIES.