



1. Introduction

Recent years have seen an exponential rise in the demand for high performance computing, driven by massive improvements and increased global interests in the applications of Machine Learning and Artificial Intelligence.

This explosive boom has not come without its challenges, however, and as a result of this spike in demand, companies, machine learning practitioners, research institutions and developers globally are now facing various issues such as insufficient computing power available, scattered computing resources, and high computing costs.

Computing power is increasingly concentrated in the hands of the tech giants, contributing to monopolistic control on pricing, availability and access. Small organizations and research institutions are therefore faced with prohibitive costs in developing the AI models of the future, limiting their ability to innovate and compete with larger organizations, while big tech companies continue to dominate the field.

To break down these barriers, we introduce NetMind Power, a decentralized computing platform developed by NetMind.AI, an AI research and software company based in London, UK and Washington DC, USA.

NetMind Power is a platform for Machine Learning and AI training, fine-tuning and inference, aimed at machine learning professionals, researchers, and software developers.

NetMind's mission is to create a global network of computing power for AI models by utilizing the idle GPUs of users worldwide. As part of this mission, NetMind Power provides a platform for large-scale distributed computing, integrating heterogeneous computing resources globally, and leveraging grid and voluntary computing scheduling architecture and load balancing technology.

NetMind aims to democratize access to computing power for businesses and research institutions, making it easier and more affordable for them to develop and run their AI models through a low-latency, widely-connected, and easy-to-manage distributed deep learning training and inference platform.

In this white paper, we will explore the three core components of NetMind Power: the training platform, the inference platform, and NetMind Chain and Token, which is the utility token for the platform. We will cover the technical and economic challenges that NetMind Power aims to solve in more detail, as well as the innovative solutions that NetMind has developed to address these challenges. We will also dive into the platform's future direction.

We believe that NetMind Power has the potential to revolutionize the AI computing power landscape and create a more equitable and accessible environment for all organizations, regardless of size.

 This white paper will be continuously updated as our journey progresses, to reflect the rapidly changing nature of this market. As such, it is subject to change and is a reflection of our latest efforts.

2. The Volunteer Computing Network

GPU resources today are a commodity that are increasingly owned by hyperscalers and big tech companies, such as Microsoft, Amazon, Google, amongst other smaller but by no means insubstantial well-funded data centre players and startups in the industry.

As AI has captured the zeitgeist and becomes the first major innovation tech delivers to the world in over a decade, demand for high performance computing is set to increase exponentially, and as such access to the such machines is becoming increasingly expensive.

There exists, however, a vast untapped resource in the world today: individual GPUs, or small clusters of GPUs owned by individuals, whether for gaming, video rendering or crypto mining sitting either underutilized or idle, not yielding anything for their owners. In the crypto example, the move from POW (Proof of Work) to POS (Proof of Stake) has made large numbers of machines 'obsolete' in their originally intended use cases. However, these machines can provide great value in machine learning model training and inference.

NetMind power is built upon the concept of Volunteer Computing. Volunteer Computing is a system that allows owners of these GPUs that are underutilized or otherwise sitting idle, to contribute them to the NetMind Power network in return for rewards (more details are given in the NetMind Token section of this white paper).

This network of globally distributed GPUs provides the computing power necessary to operate the Training and Inference platforms, which form the basis of the user facing features of NetMind Power.

The platform will route training and inference requests to the most suitable Volunteer Computing nodes on the network, depending on the needs of the training and inference job in question.

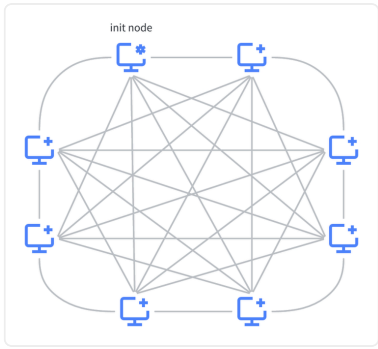
3. Training Platform

3.1 Training Platform Introduction

The training platform is the foundation of NetMind Power's decentralized computing ecosystem. It allows users to train and fine-tune models using the idle GPUs of participants around the world in an efficient and cost-effective way. The platform's architecture is built on advanced technologies and methodologies to enable distributed AI model training.

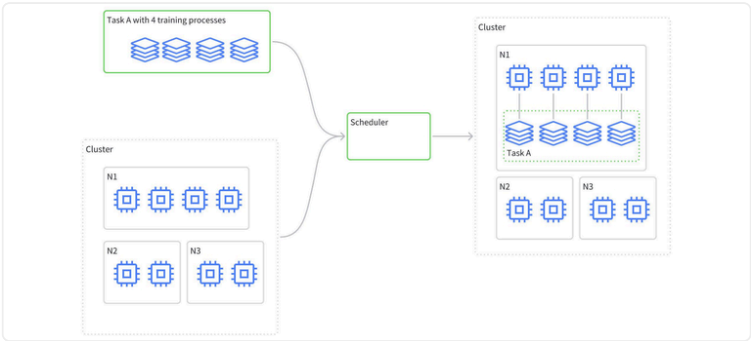
3.2 Key features and technical details of the training platform

1. Decentralized Architecture: The platform utilizes a decentralized network of connected devices, distributing the training workload across multiple GPUs. This decentralized approach reduces the reliance on centralized resources and enables the cost of training a model to be kept low.

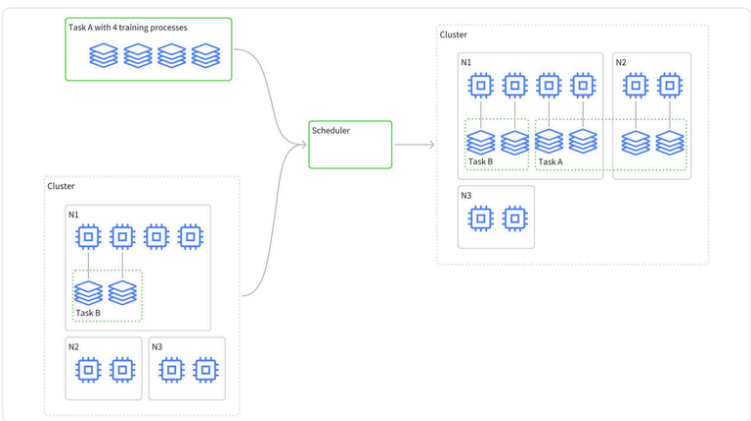


"Init node" is the first node in the decentralized network, and it is maintained by the Netmind team.

2. Resource Allocation and Scheduling: NetMind Power's intelligent resource allocation system dynamically assigns tasks to the most suitable GPUs in the network. This ensures optimal performance and reduces training time.



In most cases, when using multiple GPUs for training, our scheduler will allocate training processes in a way that minimizes network latency and improves training efficiency.



However, sometimes the "best choice" does not exist. Our schedule will then deploy training processes across nodes.

3. Data Partitioning and Model Aggregation: The training platform employs advanced techniques to divide the training data and AI models into smaller, manageable parts that can be processed in parallel by the network's GPUs. This includes methods such as data parallelism and model parallelism, depending on the specific requirements of the AI model being trained. After processing these smaller parts, the platform aggregates the results from each device to form the final trained AI model, ensuring optimal learning outcomes. Techniques such as Federated Learning and parameter averaging are used to merge the model updates from different devices while maintaining data privacy.
4. Security and Privacy: The platform employs advanced encryption and secure multi-party computation techniques to ensure that user data is protected. Furthermore, techniques such as differential privacy can be applied to add an additional layer of protection to the training data.

Based on the aforementioned design, the training platform enables efficient, secure, and cost-effective AI model training in a decentralized environment. The platform's features and technical underpinnings provide a robust solution for organizations seeking to harness the power of AI without the limitations and drawbacks of traditional, centralized computing resources.

4. Inference Platform

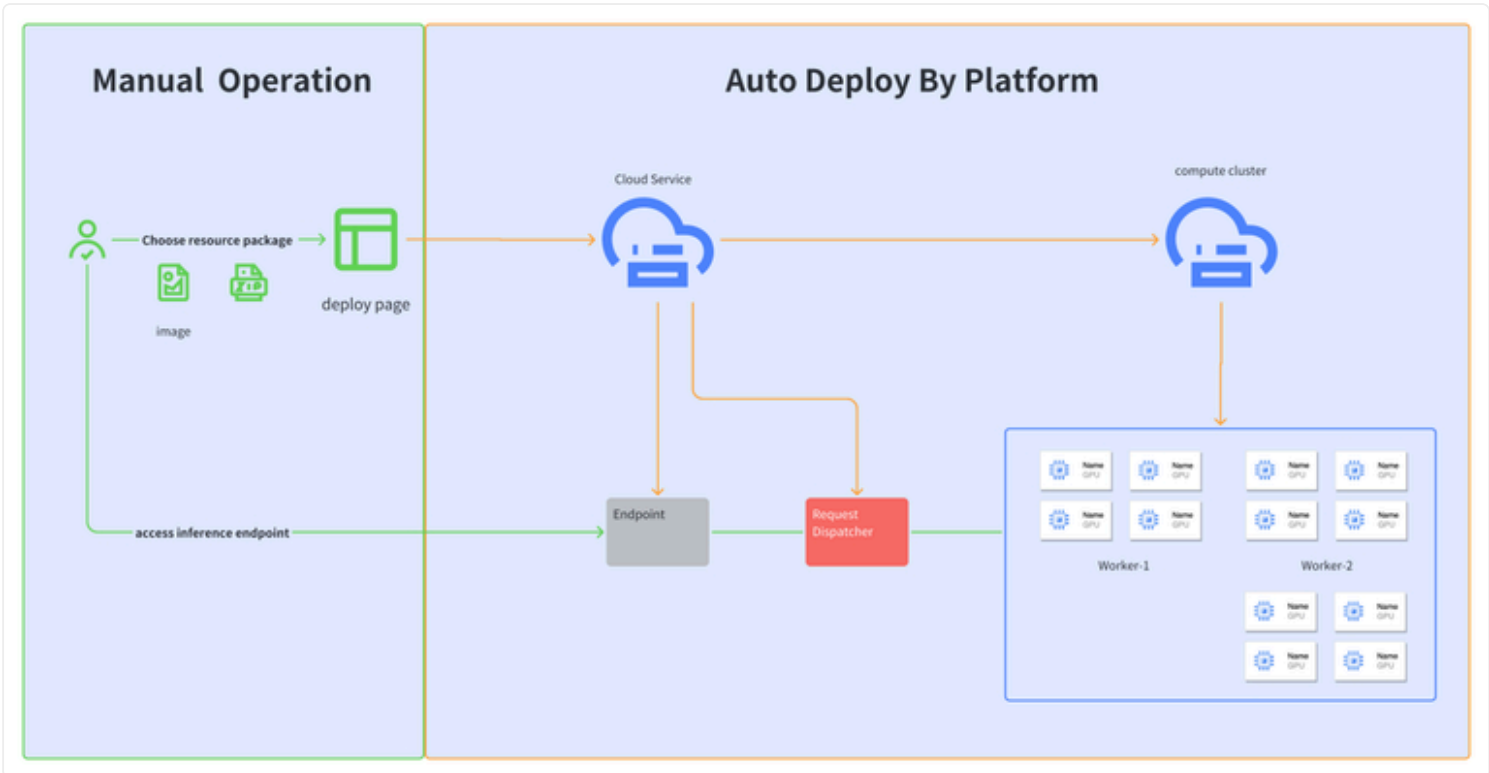
4.1 Inference Platform Introduction

The inference platform is designed to complement the training platform by providing a seamless way for users to deploy and run inference on their own models, models created by others, and open source off-the-shelf models.

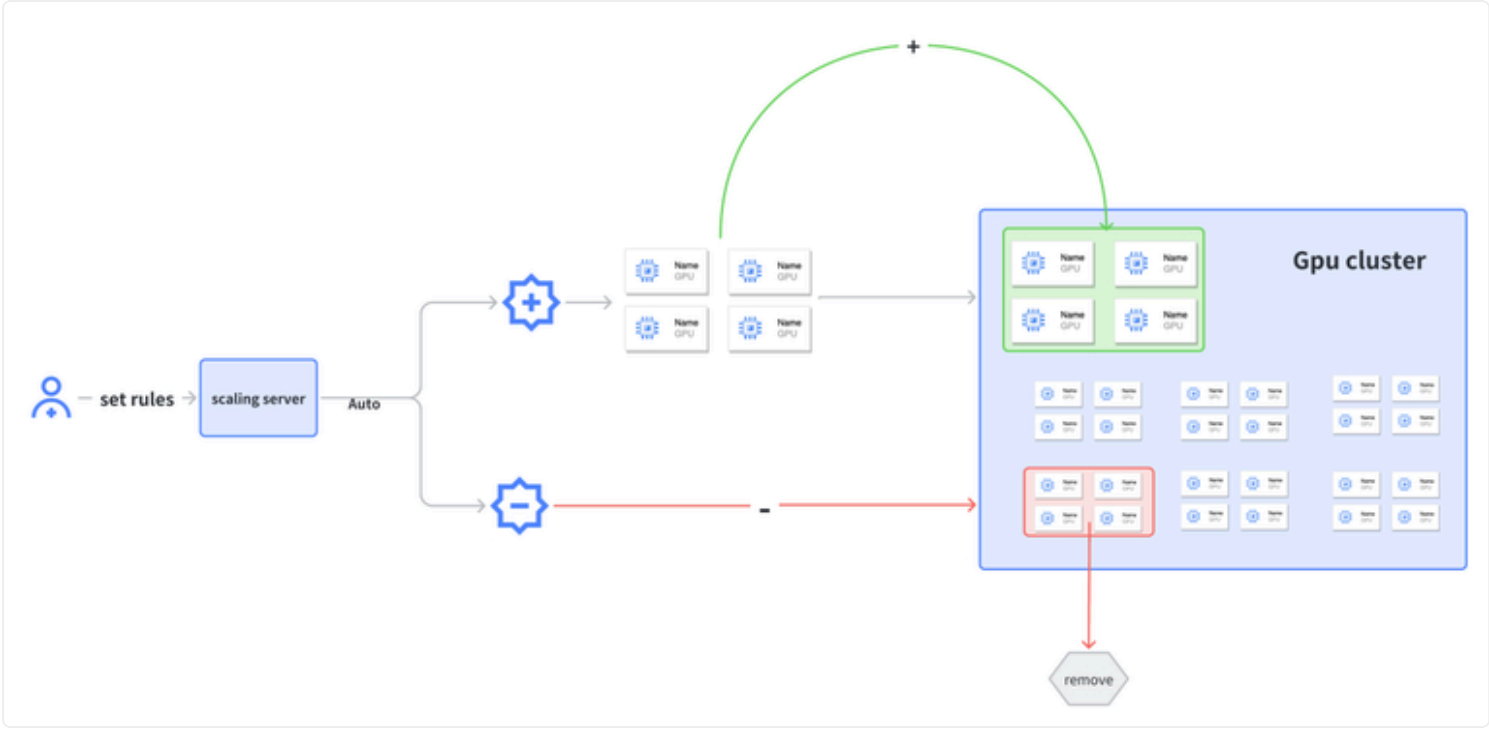
The platform's innovative architecture and features allow for easy integration, production level performance, and reduced costs, making it an ideal solution for organizations seeking to harness the power of AI in a decentralized, secure, and scalable manner.

4.2 Key features and technical details of the inference platform

1. **Model Deployment:** Users can deploy their trained AI models on the inference platform, making them accessible to other users and applications via API. The platform supports containerization, allowing for the packaging of AI models and their dependencies into lightweight, portable containers that can be easily deployed across the network.

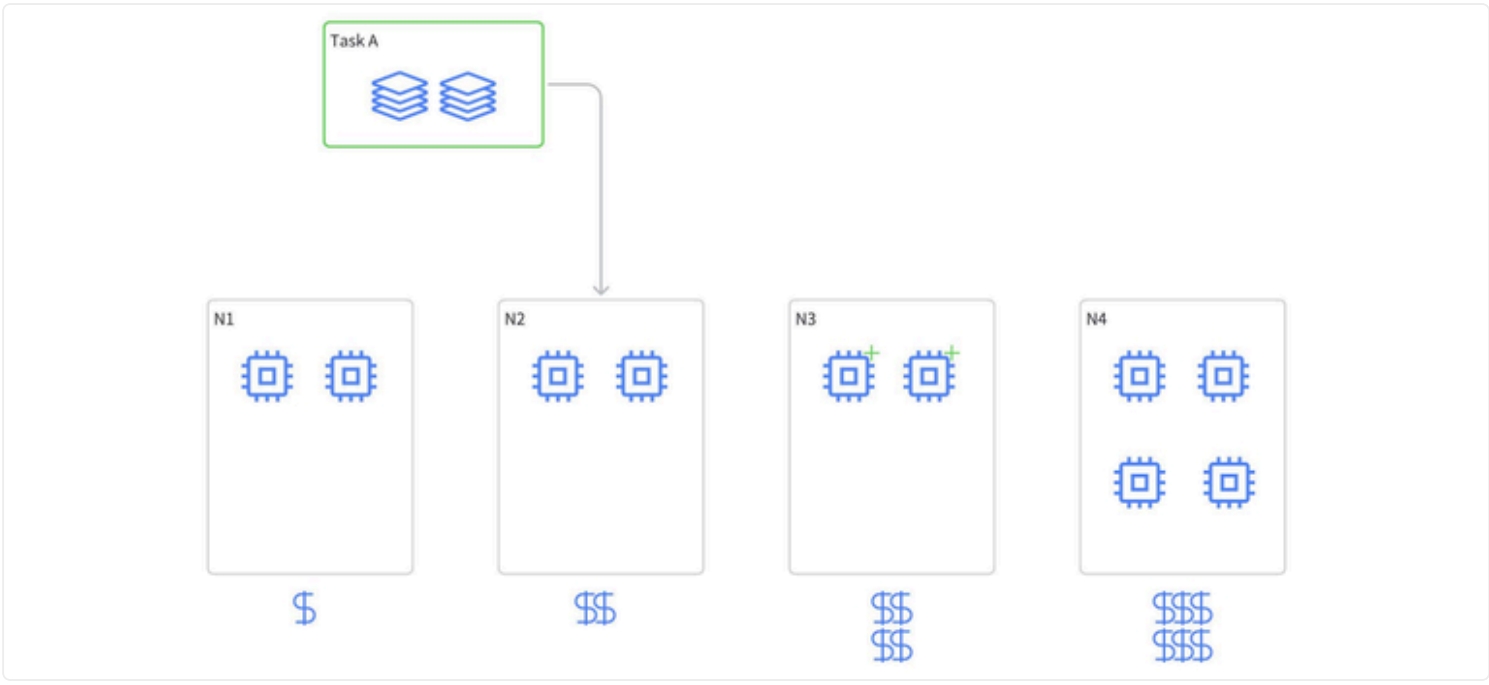


2. **Scalability:** The inference platform is designed to handle varying workloads, automatically scaling up or down based on demand. It employs distributed computing and load balancing techniques to distribute the inference workload across multiple GPUs in the network, ensuring efficient use of resources and minimal latency.

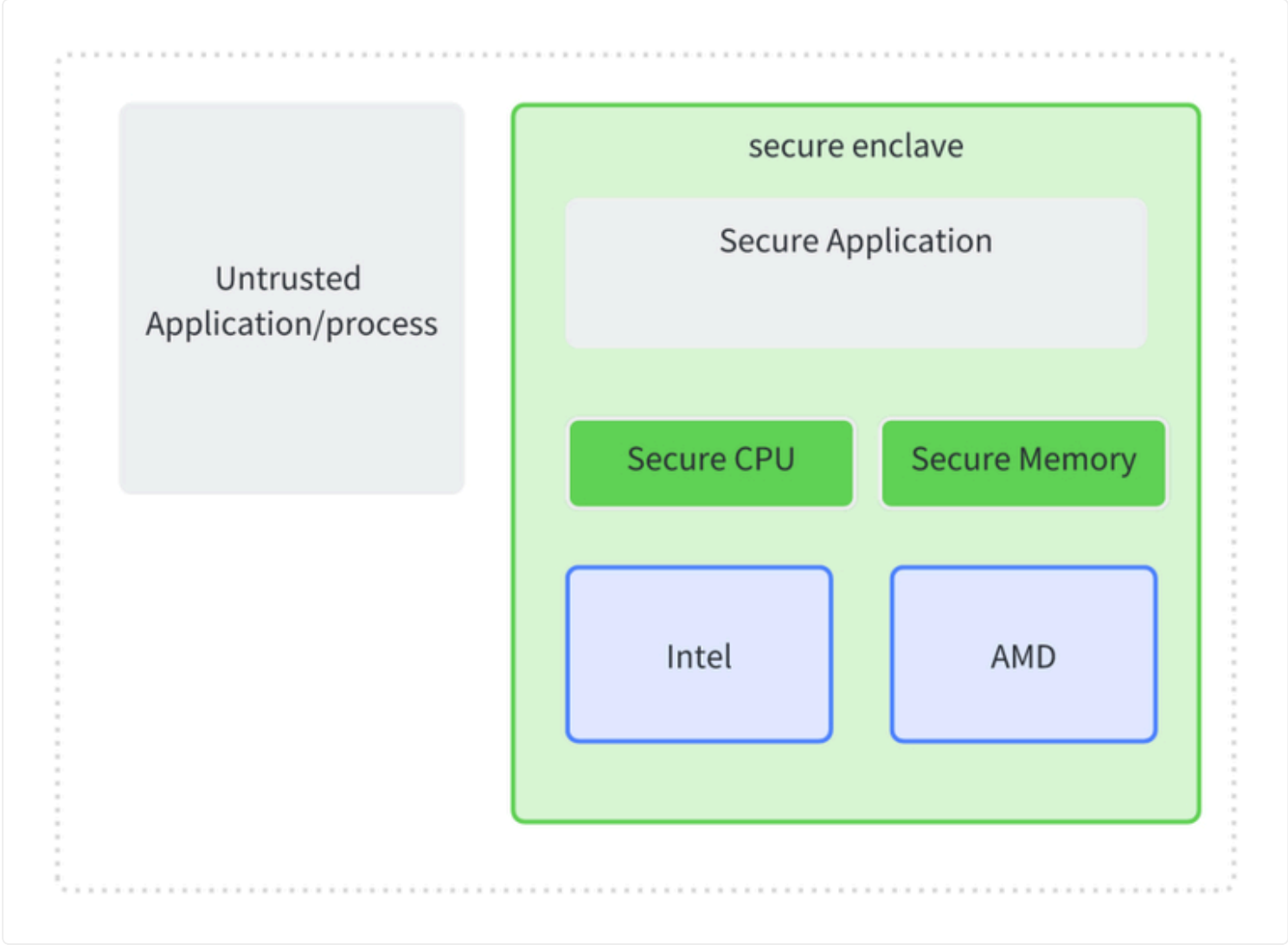


When there is a large amount of work waiting to be processed, the system will scale up according to demand as represented by the "+" sign. On the inverse, when demand is low, the system will scale down accordingly, as represented by the "-" sign.

3. **Cost Optimization:** By leveraging the decentralized nature of the platform and the idle resources of participants, the inference platform provides cost-effective access to computing power for running AI models. This reduces operational expenses for users while maintaining high performance. NetMind Power's resource allocation algorithm dynamically assigns inference tasks to the most suitable GPUs in the network, taking into account factors such as computational capacity, latency, and availability, to optimize costs and resource utilization.



4. **Security:** The inference platform employs state-of-the-art security measures to protect both the AI models and the data being processed. This includes techniques such as encryption, secure enclaves for model execution, and secure multi-party computation to maintain data privacy and model integrity during the inference process.

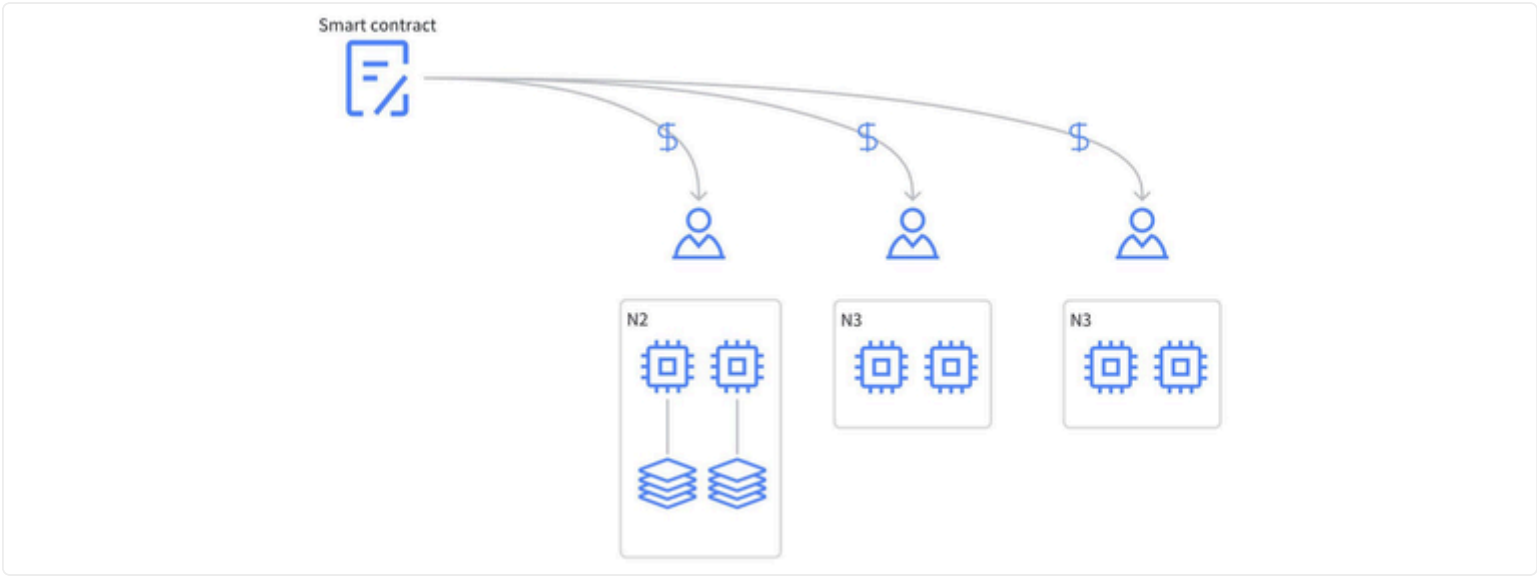


A secure enclave provides CPU hardware-level isolation and memory encryption on every server, by isolating application code and data from anyone without privileges, and encrypting its memory.

5. General Features

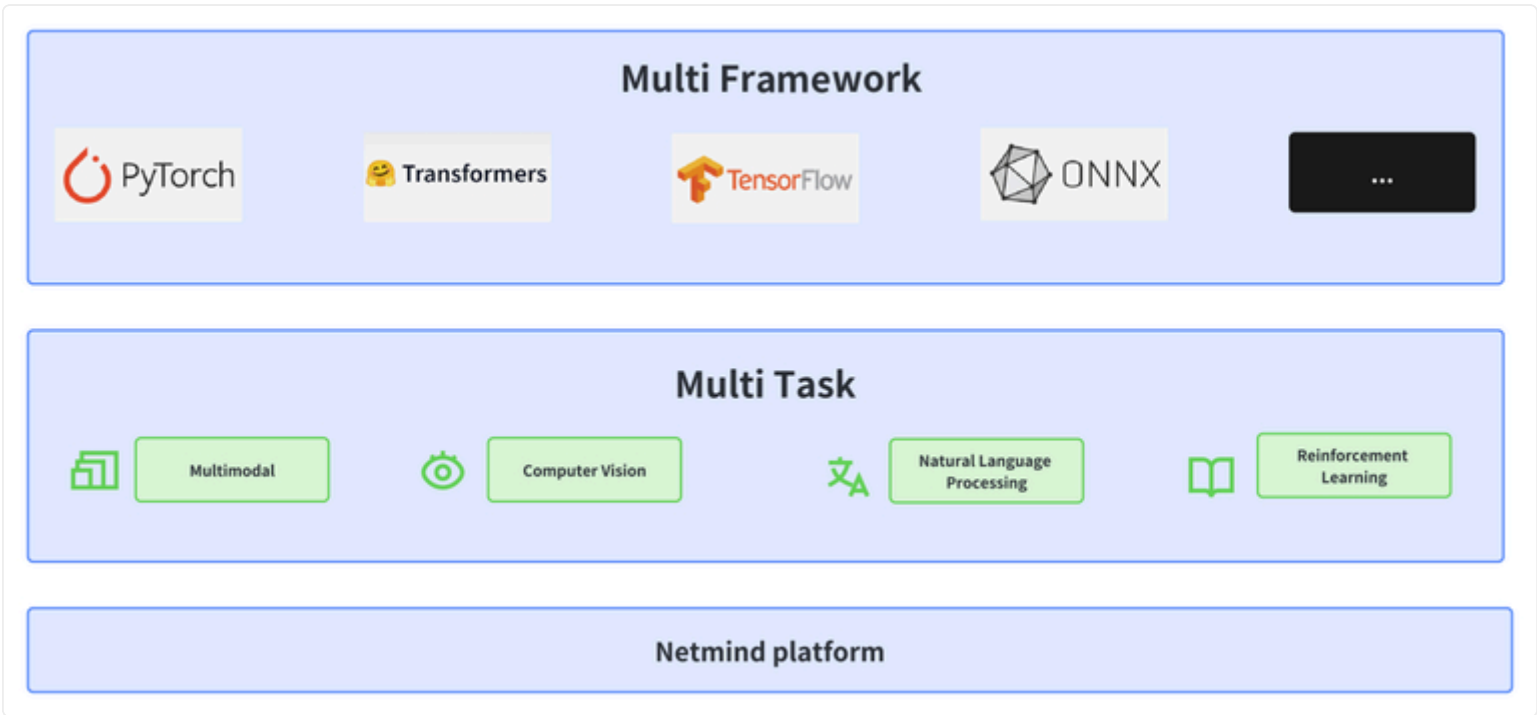
5.1 Incentive Mechanism

Participants who contribute their idle GPU resources to the network are rewarded with NetMind Token (which we expand on later in this white paper), creating a strong incentive for users to join and contribute to the platform. A smart-contract based system is used to automatically distribute rewards based on each participant's contribution to the training process.



5.2 Interoperability

The platform supports a wide range of AI models and frameworks, enabling users to work with their preferred tools and technologies. NetMind Power utilizes APIs and standard data formats to ensure compatibility and seamless integration with popular machine learning libraries and frameworks.



5.3 Environmental Sustainability

The decentralized approach of NetMind Power not only provides an efficient means for AI computing but also contributes to environmental sustainability. By leveraging idle computing resources across a wide network of users, the platform reduces the need for dedicated data centers. This can lead to lower energy consumption and a reduced carbon footprint, making NetMind Power an eco-friendly solution in the AI space.

6. NetMind Chain

6.1 NetMind Chain

The NetMind power network is built upon and governed using the blockchain, NetMind Chain. NetMind Chain enables the decentralization of all tasks, transactions and functions that occur on the platform. To be clear, the training process happens locally on the machines being used for training, not on the blockchain.

The operation of the blockchain will be facilitated by usage of a utility token, NMT (NetMind Token). We cover NMT in detail in the next section of the white paper.

NetMind Chain is based on the most mature Ethereum 2.0 technology at present such as POA-authoritative proof consensus mechanism and smart contracts fully compatible with the Ethereum chain.

6.2 Mind Nodes, Master Nodes and Staking

As well as adopting distributed technology as the design of the platform, NetMind Power also uses a decentralized management method. The system is composed of a large number of distributed machines, known as Mind Nodes.

Mind Nodes validate transactions and therefore build the blockchain. In reward for this, each Mind Node is rewarded with a Gas fee. Users can choose to stake NMT against Mind Nodes in order to receive staking rewards. The 21 Mind Nodes with the most NMT staked against them become Master Nodes.

Master Nodes receive an extra reward based on the authority they have gained. They collectively receive 10% of the total daily Staking reward. That 10% is distributed equally between the top 21 Master Nodes.

The remaining 90% of the Staking reward is distributed among users who stake on the top 21 Master Nodes, and provide liquidity to official trading pairs (see section 7.2 for reward types).

6.3 NetMind Chain Protocol

Netmind Chain is a public blockchain based on the Ethereum protocol, which adopts the POA (Proof of Authority) consensus algorithm and is fully compatible with Solidity contracts.

Compared with POW, POA does not need to consume a lot of resources to maintain the performance of the network, making the maintenance cost of such platforms extremely low. In contrast to POA, in POS (Proof of Stake) and DPOS (Delegated Proof of Stake) consensus algorithms, the more Tokens a user owns, the more likely they are to become nodes and be responsible for producing blocks. In POA, however, the validators responsible for processing transactions and validating blocks must go through a series of reviews and must ensure their own reliability.

6.4 Reward Calculation and Withdrawal for Power

As mentioned in sections 6.2 and 7.2, users contributing their machines to the NetMind Power network and those staking NMT against Mind Nodes will receive rewards in the form of NMT. The calculation and withdrawal of these rewards are managed by the on-chain smart contracts of the NetMind Chain.

Users will be able to withdraw the rewards they have earned through the smart contract (see section 7.2 for reward types), ensuring transparency and openness throughout the process.

6.5 Task Scheduling and Fee Payment

To ensure a fair allocation of computing power across different users and maintain the smooth operation of the NetMind network, NetMind has integrated a task scheduling system into the smart contract of the NetMind Chain. Fees incurred by users when utilizing various services on the NetMind network will also be incorporated into the smart contract.

This approach allows users engaging in activities such as model training or inference on the NetMind Power network to obtain the respective services in the fastest possible time and at the best available price.

7. NetMind Token

NetMind Token (NMT) is the native utility token of the NetMind Chain. The total supply of NMT is 147,571,163 tokens. All tokens were locked in a smart contract, which will gradually unlock the tokens according to the allocation plan outlined below.

NetMind Token is designed to serve multiple functions within the platform, including payment for training and inference on the platform and as rewards for users who contribute their idle GPU resources to the network.

We upgraded our tokenomics starting April 16, 2024. The previous version of tokenomics was implemented for one year and released 34,571,163 NMT into circulating. If you need to read it, please click on the link:

[The old version of tokenomics \(invalidated\)](#)

7.1 Allocation Plan

In this section, we quantify the allocation plan in percentages and amounts.

Based on the old version of tokenomics, 34,571,163 NMT was allocated to circulation in the first year (from April 16, 2023, to April 15, 2024).

Starting April 16th, 2024, 113,000,000 NMT will be issued over the next 10 years under new tokenomics framework.

1. 62.5 million NMT (55.31% of new issuance; 42.35% of new total supply) is allocated as the Mining Rewards for providers of computational power to the network, used for GPU rental, training and inference.

- For the first 2 years (after 4/16/24), 20 million NMT will be distributed.
- For the next 2 years, 15 million NMT will be distributed.
- For the next 2 years, 10 million NMT will be distributed.
- For the next 2 years, 7.5 million NMT will be distributed.
- In the final 2 years, 5 million NMT will be distributed.
- This includes a revised portion for early GPU contributors now set at 5 million tokens, which will begin to vest on April 16, 2026, and will be gradually released over a four-year period.

2. 16.5 million NMT (14.60% of new issuance; 11.18% of new total supply) is allocated as the Staking Reward for users who participate in staking and providing liquidity.

- For the first 2 years (after 4/16/24), 4 million NMT will be distributed.
- For the next 2 years, 3.5 million NMT will be distributed.
- For the next 2 years, 3 million NMT will be distributed.
- For the next 2 years, 3 million NMT will be distributed.
- In the final 2 years, 3 million NMT will be distributed.

3. 16.5 million NMT (14.60% of new issuance; 11.18% of new total supply) is allocated as the Ecosystem Growth Fund dedicated to expanding and enriching the ecosystem.

- For the first half of the year (after 4/16/24), **no NMT** will be distributed.
- For the next 1 and half years, 4.5 million will be distributed.
- For the next 2 years, 4 million NMT will be distributed.
- For the next 2 years, 3.5 million NMT will be distributed.
- For the next 2 years, 2.5 million NMT will be distributed.
- In the final 2 years, 2 million NMT will be distributed.

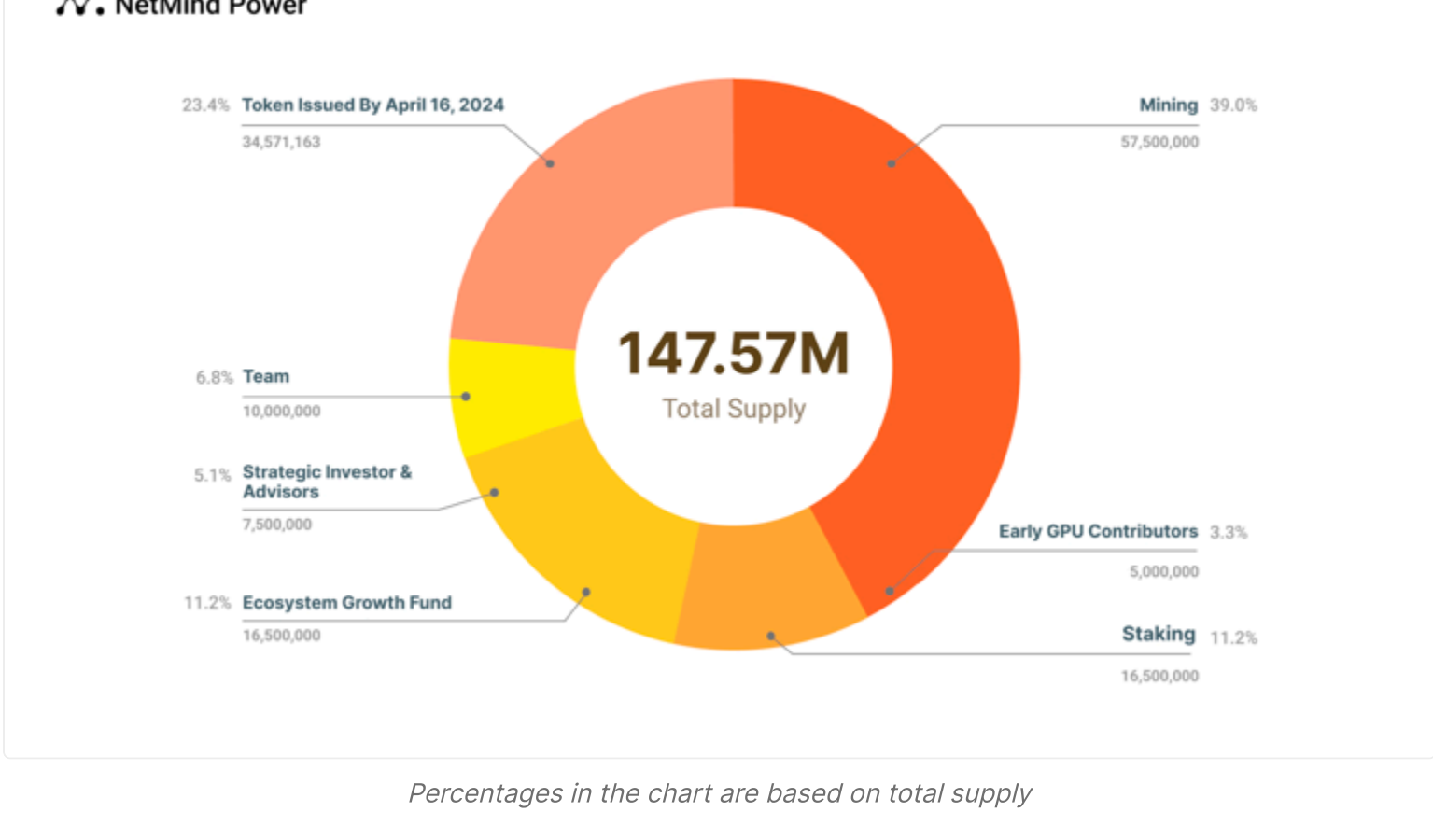
4. 7.5 million NMT (6.64% of new issuance; 5.08% of new total supply) is reserved for strategic investors & advisors, which is a new category with 7,500,000 tokens aimed at reinforcing relationships with key strategic partners.

- For the first half of the year (after 4/16/24), 3 million NMT will be distributed.
- For the next 4 and half years, 4.5 million will be distributed.

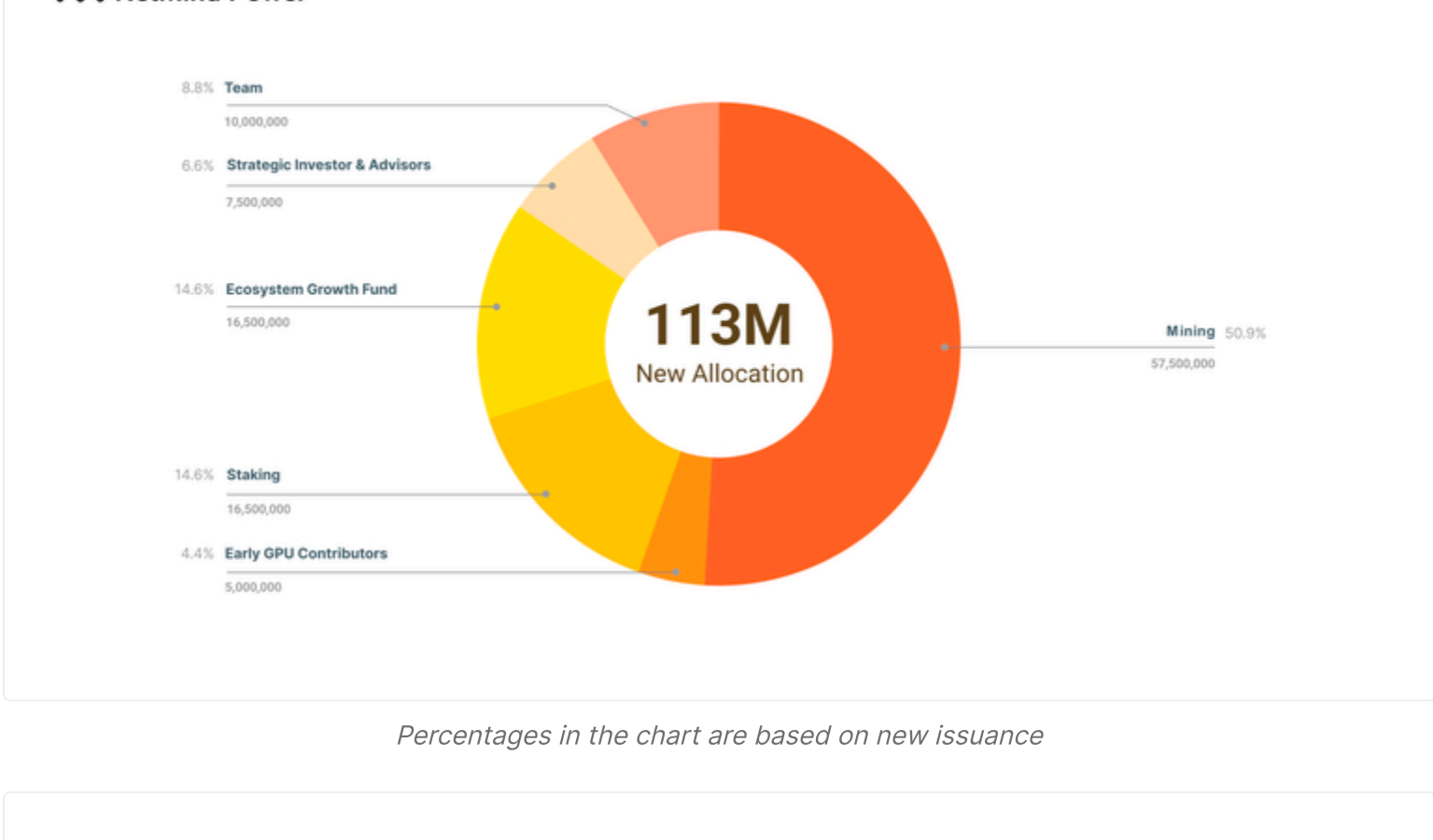
5. 10 million NMT (8.85% of new issuance; 6.78% of new total supply) is allocated as the Team Fund reserved for the NetMind technical team.

- For the first half of the year (after 4/16/24), **no NMT** will be distributed.
- For the next 4 and half years, 6 million will be distributed.
- In the final 5 years, 4 million NMT will be distributed.

The duration unit "year" in the reward rules is calculated as 365 days, regardless of leap years.



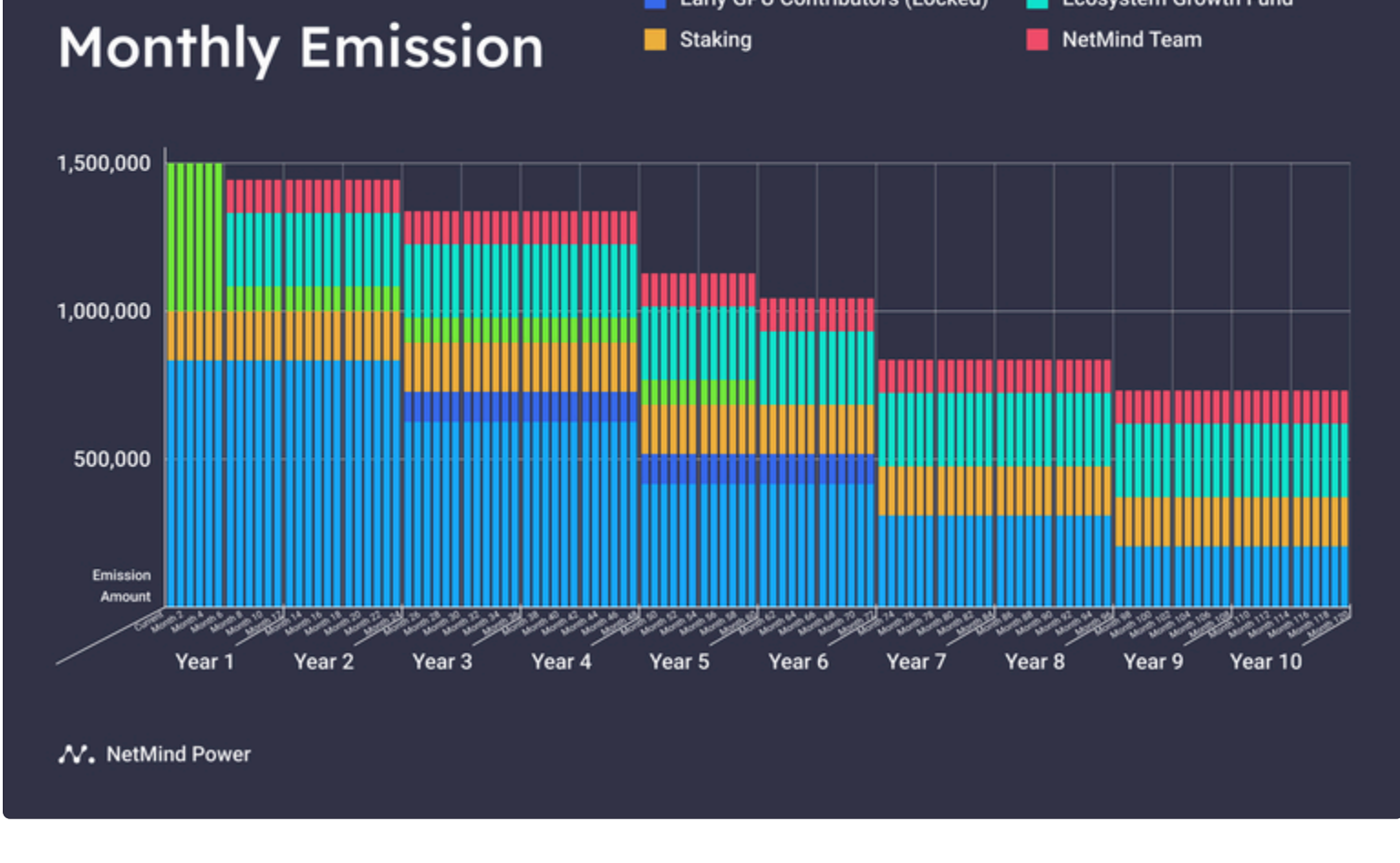
Percentages in the chart are based on total supply



Percentages in the chart are based on new issuance

Updated Tokenomics (Go Live Date - 4/16/24)					
Token allocation			Vesting schedules		
Allocation	Future Emissions (After 4/16/24)	% of Total New Allocation (After 4/16/24)	CDP (in months)	Vesting (in months)	Total vesting (in months)
Mining - Phase 1	20,000,000	17.70%	0	24	24
Mining - Phase 2	15,000,000	13.27%	24	24	48
Mining - Phase 3	10,000,000	8.85%	48	24	72
Mining - Phase 4	7,500,000	6.64%	72	24	96
Mining - Phase 5	5,000,000	4.42%	96	24	120
Staking - Phase 1	4,000,000	3.54%	0	24	24
Staking - Phase 2	3,500,000	3.10%	24	24	48
Staking - Phase 3	3,000,000	2.69%	48	24	72
Staking - Phase 4	3,000,000	2.69%	72	24	96
Staking - Phase 5	3,000,000	2.69%	96	24	120
Ecosystem Growth Fund - Phase 1	4,500,000	3.98%	6	18	24
Ecosystem Growth Fund - Phase 2	4,000,000	3.54%	24	24	48
Ecosystem Growth Fund - Phase 3	3,500,000	3.10%	48	24	72
Ecosystem Growth Fund - Phase 4	2,500,000	2.21%	72	24	96
Ecosystem Growth Fund - Phase 5	2,000,000	1.77%	96	24	120
Team - Phase 1	6,000,000	5.31%	6	54	60
Team - Phase 2	4,000,000	3.54%	60	60	120
Strategic Investor & Advisors - PH 1	3,000,000	2.69%	0	6	6
Strategic Investor & Advisors - PH 2	4,500,000	3.98%	6	54	60
Early GPU Contributors (Locked)	5,000,000	4.42%	24	48	72
TOTALS	113,000,000	100.00%		N/A	

NetMind Power



7.2 Reward Types

In this section we cover the different types of rewards in the network. NMT serves not only as a means of payment for various fees on the NetMind platform and NetMind Chain, but also as a reward for various types of participation in the network, which are detailed below.

7.2.1 Mining Reward

Volunteer Computing resource providers receive rewards in NMT simply by being connected to the network, regardless of machine utilization for tasks such as training or inference. The user receives these rewards if the machine is online and connected to the network for *more than 5 non-consecutive hours in one 24 hour period*.

The amount of reward a Volunteer Computing Resource provider is allocated is determined by three factors: the GPU type, and the network bandwidth speed and the amount of connected machines on any given day (each daily reward allocation will be divided across all connected machines).

Miner Emissions

Miner emissions are determined through a calculated approach, linking the value of computing resources to the NMT value. The daily distribution of NMTs to miners is based on the minimum value between the total computing resources' value and the maximum NMT value, divided by the average NMT price over the past 7 days. This ensures fairness and adapts to market conditions.

- Linking to Computing Resources and NMT Value:
 - The emissions depend on two primary factors: the total value of the computing resources in the network and the current value of NMT.
- Establishing a Cap:
 - There is a cap to maintain a balanced distribution of rewards. This is calculated by comparing the total value of computing resources (as per a specific parameter) with the total value of the maximum NMT that could be issued to miners.
 - Total daily VC = total value of computing resources
 - Total daily NMT = Maximum daily issued NMT for GPU contributors
 - Average NMT price = past 7 days average NMT price
 - B = Balancing coefficient
- Minimum Value Determination:
 - Each day, the **total number of NMT issued** to all miners is equal to
 - $\min[\frac{\text{total daily VC}}{B}, \text{total daily NMT} \times \text{average NMT price}]$
 - *Starting on 2024-2-24, the algorithm is activated.*
- If total daily VC/B ≤ total daily NMT * average NMT price, the unissued NMT will stay in the mining reward pool, to be distributed after 10 years.
- Individual GPU Reward:
 - Each GPU has a weight, w_i , which represents its computational power or contribution value; a penalty coefficient, p_i , which might be applied due to various factors (like uptime, vacant time, reliability, internet speed etc.). Then for machine i , its reward will be:
 - $$\frac{w_i \times p_i}{\sum (w_i \times p_i)} \times \text{total number of NMT issued}$$
- The reward for each machine is then proportionate to its effective weight compared to the total, factoring in the daily NMT distribution cap.

7.2.2 Computational Power Utilization Fee

Users who rent GPU, train or fine-tune models on the NetMind Power platform pay to do so in NMT. The provider of the GPU resources on the network that is used to rent, perform the training or fine-tuning is similarly rewarded in NMT.

There are two methods of training on the NetMind Power platform, Sync mode, in which the model is trained in a synchronous manner, and Async mode, in the which the model is trained in an asynchronous manner.

Below is a breakdown of how the fees that users pay for these two types of training are distributed:

- Sync mode**, 50% of the total computational power utilization fee will be paid to Volunteer Computing resource providers that provide power for the training job, while the remaining 50% will be returned to contract without increasing the circulating supply.
- Async mode**: 20% of the total computational power utilization fee will be paid to Volunteer Computing resource providers that provide power for the training job, while the remaining 80% will be returned to contract without increasing the circulating supply.
- SSH feature**: Depending on the market needs and NMT price, up to 100% NMT utilization fee will be returned to contract without increasing the circulating supply.

These percentages are subject to change in the future. A reminder that all transactions on the blockchain consume gas fees. A small percentage of NMT for each training task will go towards Gas fees.

The user will receive this reward regardless of how long the machine was connected to the network during a 24 hour period.

During the time the machine is being utilised for training or fine-tuning, the user is **also** receiving the **Mining Reward**.

Sudden disconnection: If the user's machine disconnects from the network **during a training job** (i.e. When the user is earning a **Machine utilisation reward**) **cause the job failed**, they will **not receive** Mining Reward for that 24 hour period. But they still received their share of Machine Utilisation Rewards

7.2.3 Staking Rewards

Users can stake their NMT tokens with any of the 21 main nodes or provide liquidity to official trading pairs and receive NMT rewards based on the distribution plan.

Staking Rewards Structure:

Flexible (including staking and liquidity), long-term, and master node staking options remain funded from the staking rewards pool, with detailed allocation adjustments to be announced by April 30th, 2024.

- 10% of the rewards go to the top 21 Master Node operators.
- 90% of the rewards go to staking users who have staked their tokens against any of the 21 Master Nodes, or provided liquidity to official trading pairs. Staking users receive daily rewards from the staking pool based on weighted participation. The weighting coefficients are as follows:
 - Long-term Staking: 4
 - Flexible Staking: 1
 - Flexible Liquidity: 2 (*When calculating rewards, the NMT and USDC with equivalent value in the trading pair are both considered liquidity.*)

7.3 Master Nodes and Elections

As mentioned in this white paper, NetMind Chain has 21 Master Nodes which other Mind Nodes then elect through staking.

7.4 Cross-Chain Bridge

Cross-chain bridge facilitates the transfer of NetMind Token between NetMind Chain and Binance Smart Chain.

Users can utilize the cross-chain bridge to make cross-chain transfers. For example, a user may purchase NMT on Pancakeswap on the Binance Smart Chain and then, through the cross-chain bridge, transfer NMT to NetMind Chain to pay for model training or inference or for staking.

Additionally, users can transfer their rewards from the NetMind Chain to the Binance Smart Chain for trading.

The cross-chain bridge will provide the project with a broader market and will give users more flexibility, fostering the development of the project.

8. User types

In this section we outline the different user types that NetMind power aims to serve. Given the system is based on a tripartite design, we will cover the demand side (users of the platform), the supply side (those who provide computing power to the platform) and the blockchain side (contributing to the running of the blockchain).

8.1 Demand Side

1. Machine learning practitioners

Practitioners and experts in machine learning and deep learning worldwide will be able to use our platform for training and fine-tuning machine learning models that are either proprietary or open source. Once they have trained their models, the user can host the model for inference on our platform, accessible via API, to embed their models into their applications, or for batch inference. They will also be able to make their models available for inference through the platform to the community and be compensated for doing so, without having to share or publish their models' weights.

2. Researchers

Likewise, AI researchers from universities, government research institutions, private companies and research labs alike will be able to use the training platform to train and fine-tune their machine learning models in a cost effective way to advance their research efforts. Like machine learning practitioners, researchers can also choose to make their models available for inference to third parties for a fee.

3. Developers

Software developers that aren't necessarily experts in Machine Learning and that are seeking to embed AI and machine learning capabilities into their applications will be able to use our Inference platform to do so in a cost effective and scalable way via API. They will be able to host their own proprietary models on our platform, choose from a selection of open source off-the-shelf models, or pay to use models trained and published by other users of the platform (in which case they will have to pay for both the GPU resources and the fee set by the model's owner). Additionally, the training platform is designed for these users, in that it allows for a no code way to fine tune existing open source models - the user just provides the dataset to fine tune the model on.

8.2 Supply side

1. Volunteer Computing Resource Provider

Anyone owning one or more GPUs, whether a gamer, cryptocurrency enthusiast, crypto miner for example can contribute their idle machines to the network for other users to train, fine-tune and run inference on, and be rewarded with NMT (NetMind Token). NMT is the native utility token of the NetMind Power platform and is discussed in a later section of this white paper.

8.3 Blockchain side

1. NetMind Chain Mind Node

Mind Nodes validate transactions and create blocks in the blockchain. Users can connect their machines to the network for this purpose, and are rewarded with NMT as outlined in the NetMind Chain section of this white paper.

9. Future Direction

In this final section, we outline some of the future development plans for NetMind Power and the potential impact the platform could have on the AI landscape.

1. Platform Enhancements

We are committed to continuous improvement and expansion of the platform, with plans to support more machine learning frameworks, enhancements to security and privacy mechanisms, and overall performance and scalability improvements. These enhancements will enable the platform to cater to a wider range of AI applications and users, ensuring that it remains at the forefront of cutting-edge AI technology.

2. Community and Ecosystem

A core aspect of our growth strategy involves expanding the NetMind Power community and fostering a thriving ecosystem. This will be achieved through providing the means for users to train models on the platform and publish them for public community usage.

3. Decentralized Pricing

We intend to decentralize the pricing mechanism in the platform, giving providers of computing power the ability to set their own pricing for usage of their machines. The Volunteer Computing network can dynamically adjust prices based on demand, availability of resources, and other market factors

10. Conclusion

In conclusion, NetMind Power is a decentralized platform aimed at Machine Learning practitioners, researchers, and software developers. The platform allows owners of GPUs to contribute them to the network to provide computing power for machine learning model training and inference. By adopting a community-driven approach, NetMind Power leverages the power of the crowd to encourage collaboration in AI and promote the development of better and more powerful AI models in the future.

The technical solutions offered by NetMind Power, including its management system, blockchain technology, and asynchronous training algorithm, provide a reliable and secure platform for participants to share their idle GPUs and for users to build AI models.

The platform uses a native utility token, NMT, to reward providers of computing power, and as a means of paying for those computing services by users of the training and inference platforms.

We encourage readers to explore the potential of the platform and get involved in the NetMind Power community to help build the future of AI. By participating in the ecosystem, users can contribute to the democratization of AI computing resources, empowering organizations and individuals around the world to harness the power of AI to solve complex problems and drive progress in their respective fields. Together, we can shape a more equitable and accessible AI landscape for the benefit of all.