

DecideAI Whitepaper

DecideAI Whitepaper Version: 1.0 Date: June 2024

Overview

The Opportunity

The demand for LLMs is booming, and supply-side infrastructure is needed to meet this demand. This means a pressing need for reliable models and high-quality data.

- Much of the power and progress has so far been concentrated in the realm of centralized, closed-source, and general-purpose models.
- However, the more pressing long-term need, and the most profitable potential, lies in industries (health, finance, etc.) that require tailored, high-performance models.
- This means that there is a significant, largely untapped opportunity to create an environment for developing, refining, and collaborating on future-ready LLMs and supporting datasets.

The Decide AI Ecosystem

DecideAI is an ecosystem that consists of three products, designed to meet the needs of the high-end, specialized LLM market.

- **Decide Protocol**
 - A training ground that co-ordinates artificial and human intelligence to annotate, train, and continuously improve specialized LLMs and targeted datasets, using Reinforced Learning with Human Feedback (RLHF).
- **Decide ID**
 - A tested, unique methodology (Proof of Personhood or PoP) means that all contributors and their credentials can be verified and traced, ensuring high-quality data.
- **Decide Cortex**
 - A platform that gives access to pre-trained LLMs and pre-vetted datasets (via direct purchase or API endpoints), for clients and/or developers who do not want or need to start from scratch.

The DecideAI ecosystem is underwritten by a robust reward architecture that incentivizes long-term engagement and discourages bad actors.

The future of LLMs will focus on quality, collaboration, and ownership as opposed to a status quo that is opaque, leading to a concentration of power. DecideAI plans to be at the forefront of this phase shift.

The Opportunity

The LLM market

The Large Language Model (LLM) market has exploded in importance since the “early days” of 2022. By 2030, it is expected to reach a size of \$260 billion.

The main drivers of growth are practical applications in a corporate or industrial setting. LLMs can help to streamline processes, enhance the customer experience, and even expand services. The industries with the greatest potential for improvement include healthcare, finance, e-commerce, and media.

However, the demand for LLMs has created the need for a corresponding supply chain. Raw data, like crude oil, must be refined and processed in order to be usable. In LLMs, this process involves steps such as labeling, annotating, and fine-tuning.

Both artificial and human intelligence are useful in this refining process, which results in a usable LLM. The level and type of input required differ according to the type of model being developed.

Broadly speaking, LLMs are divided into:

1. General-purpose models (e.g. ChatGPT, Claude): developed by a few large firms for mass use
2. Targeted models: often developed fully or partially by the companies that use them.

Training general-purpose models is fundamentally a “numbers game”. It requires large data sets and large workforces. Representative firms in this space include Scale and Toloka.

When developing targeted models, however, the focus is on quality rather than quantity. Since they are typically used in a commercial setting with a low tolerance for error, the right data and the right workforce (i.e. with domain-specific expertise) are more important than volume alone. Prolific and Surge are examples of LLM training firms that serve this market.

Currently, general-purpose models are the go-to paradigm. However, as the ready availability of “low hanging fruit” high-quality data diminishes, the main opportunity will shift to the targeted, use-case-focused domain. Focused datasets and models make it easier to differentiate one’s offering, and the competitive moat is deeper as expertise is harder to obtain.

Therefore, finding ways to produce higher-quality data and more expert-oriented models represents the most attractive area for long-term growth for the LLM training market.

Challenges for the AI industry

The status quo in the AI industry is far from ideal. The nature of general-purpose LLMs rewards scale, meaning that a small number of large, closed-source, centralized corporations currently dominate the industry.

This creates a raft of problems, including:

Ethical concerns: It is not clear how and with what motives highly influential models such as ChatGPT are being trained. Potential biases (political, national, ethnic, etc.) could be embedded knowingly or unknowingly and lead to inequitable outcomes.

Privacy risks: The quest for ever more voluminous data sets means that user privacy can easily be compromised in the pursuit of greater scale. This represents an unacceptable risk for organizations with a responsibility to protect proprietary and/or customer information.

Unreliable results: Large, generic datasets not only tend towards mediocrity of output but are also more likely to produce errors and hallucinations. This is a double threat in fields where superior output is necessary, and error is unacceptable.

Uncustomizable: The nature of a closed-source model is that it cannot be easily adapted or fine-tuned. This is a problem in fields with highly specific needs.

Slow to adapt: Many dynamic, complex, or rapidly evolving fields require not only domain-specific expertise but also context-dependent (e.g. time and place) judgment, which past data alone is insufficient to provide.

Scaling problems: The resource constraints of large, unwieldy models mean that enterprises with substantial demand face a combination of high usage costs together with a high risk of outages and disruption.

Seizing the Opportunity: DecideAI

As AI’s capabilities grow, so does the potential for its misuse. Our vision is to counterbalance this by creating open-source, transparent, and secure AI infrastructure that protects user

privacy while rewarding individuals for their invaluable data contributions.

We will achieve this by creating an ecosystem for AI excellence that:

- Democratizes access to high-quality data, models, and training methods
- Leverages blockchain to ensure both privacy and transparency
- Fosters a specialized workforce capable of supporting high-quality AI
- Creates an environment for building AI applications that compete on both performance and cost.

We believe there is a gap for such an ecosystem and that now is the right time to establish it.

The DecideAI Ecosystem

The initial iteration of the DecideAI Ecosystem consists of the following three components. Each addresses a key need in the AI market:

- **Decide Protocol = Transparent Training**
 - A platform for sourcing and refining LLM training data, building new models, and compensating contributors.
- **Decide ID = Data Quality**
 - A cutting-edge system for verifying the contributor workforce to ensure high-quality training data.
- **Decide Cortex = Seamless Collaboration**
 - An open-source community where users and developers can both host & access models and high-quality, sourced-and-labeled datasets.

Users can engage with any combination of these three components to create cutting-edge LLMs. Meanwhile, contributors (data labelers, developers, etc.) can get rewarded via the native token, creating a sustainable system in which everyone benefits.

The growth of the ecosystem will establish DecideAI as a pioneer in the emerging Data Layer of the AI/LLM landscape. By prioritizing quality over quantity, we will set a new industry standard for open-source collaboration.

Decide Protocol

Decide Protocol is a platform aimed at those who do not want to use general-purpose, off-the-shelf LLMs but lack the resources to develop their own models internally.

The protocol provides an end-to-end service, covering data collection, annotation & labeling, model training, refinement, and updating.

As already mentioned, this process relies on a combination of artificial intelligence and human review.

Alternative methodologies (e.g. Direct Preference Optimization) attempt to automate the process from beginning to end. While these approaches are less resource-intensive, they generate fewer human-preferred responses and less reliable output overall.

A combined approach makes it possible to cover dimensions that machine intelligence alone cannot handle independently, with results that are more closely aligned to human expectations.

The methodology used by Decide Protocol is called Reinforcement Learning with Human Feedback (RLHF), which targets an optimal mix of cost-effectiveness and high-quality data.

- Reinforcement Learning = Generating output and iteratively improving the quality.
- Human Feedback = Skilled labelers refine and provide output at each stage.

RLHF is an ideal methodology for training specialized, targeted models:

- **Bespoke:** It is possible to hone responses to suit the needs of a specific target goal, by leveraging the expertise of individuals who understand the subject matter.
- **Accurate:** Ongoing reinforcement training makes answers more accurate allowing for more specialized models for specific use cases.
- **Safe:** The lower incidence of hallucinations means that it is feasible to deploy in critical situations where money, reputation, and human health are at stake.

How it works

The end goal of the process is to create a model that matches the specific needs of the client's organization and domain. The steps in the process are as follows:

1. **Seeding:** Acclimatizing the base model to the new domain.
2. **Annotation:** Intensive interaction between model and reviewers.
3. **Incentivization:** Rewarding contributions in proportion to value added.
4. **Training:** Interpreting, weighting, and incorporating feedback to update the model.
5. **Evolution:** Continuous, live evaluation based on real-world interactions.

We describe the steps in more detail below:

1. Seeding

The base LLM needs to be grounded in the basic concepts of the relevant domain and the tasks it will be asked to perform. In this stage, the human reviewers engage in an initial “dialogue” with the LLM, providing prompts closely related to the target area. Thus the LLM becomes acclimatized and can be used as a foundation model for the subsequent stages.

2. Annotation

Next, a specialized team of contributors is selected on the basis of relevant credentials to the domain (using the Decide ID platform). These reviewers perform two functions concerning the data used to train the LLM:

1. Assess quality (accuracy and relevance).
2. Provide feedback, corrections, and additional context.

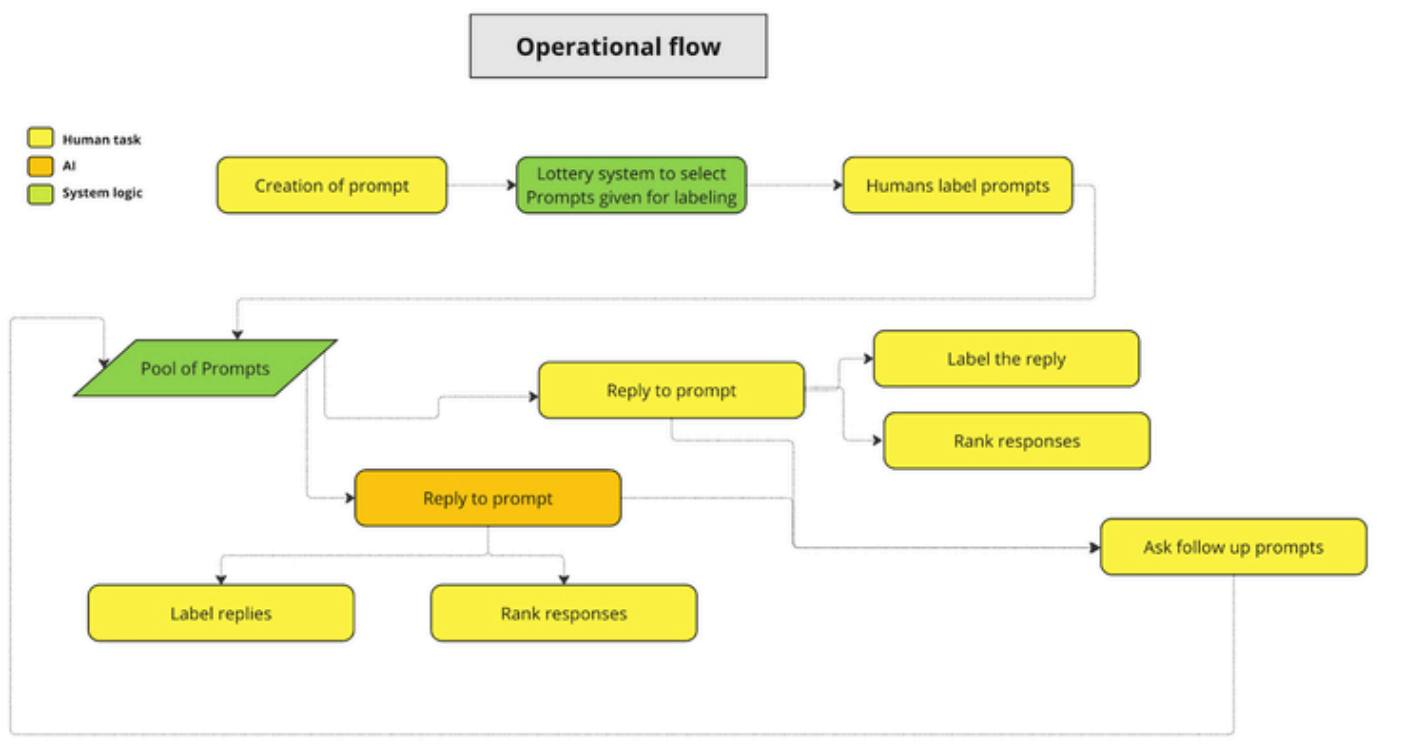
In conversation with the LLM, contributors will prompt, label, rank, and reply. These annotations improve the quality of the LLM’s dataset, and its subsequent responses to users.

Below are some illustrative examples:

- Prompt: “A patient presents with sudden onset of severe chest pain, shortness of breath, and dizziness. What are the potential diagnoses and initial tests you would order?”
- Follow-up: “You mentioned ordering an ECG. Could you explain what specific findings on the ECG you would look for in this case?”
- Label: Assigning urgency levels like “Immediate”, “Urgent”, “Routine”.
- Ranking: Order by probability: “Myocardial infarction”, “Panic attack”, “Pulmonary embolism”.

To maintain consistency and quality, annotators are provided with clear guidelines. These guidelines cover how to assess criteria such as accuracy, relevance, bias, consistency, and context.

Annotation process in more detail



Data Collection - Operational Flow

Data Collection - Operational Flow

1. Creation: Human contributors create relevant prompts.
2. Selection: Prompts are chosen by lottery for labeling, to ensure the composition of new prompts is balanced and the volume is manageable.
3. Prompt Labeling: Human contributors categorize prompts based on relevant criteria (e.g. quality ratings).
4. Aggregation: The labeled prompts are added to a 'pool' of prompts approved for use in future tasks.
5. Reply to prompts: Humans and/or AI agents respond to prompts in the pool or submit follow-up prompts.
6. Reply Labeling: As in Step 3, human contributors categorize replies based on relevant criteria.
7. Response Ranking: Human contributors are presented with multiple responses to the same prompt, and then rank them in order of preference, revealing the optimal responses.
8. Follow-up prompts: After the above process is complete, human contributors can continue to ask follow-up prompts, continuing the process indefinitely.

Throughout this process, human contributors follow the annotator guidelines, which cover:

- Accuracy evaluation: Is LLM's output factually correct? (cross-check against reliable sources)
- Relevance assessment: Does the information provided address the question raised?
- Context analysis: Does the result make sense within the broader context, for example, does it comply with ethical standards?

- Bias and fairness evaluation: Does the response contain or reinforce unfair representations or slanted views?
- Consistency checks: Is the output internally consistent, and does it conform to the overall objectives of the exercise?

3. Incentivization

The protocol awards Decide Tokens (DCD) to contributors who take part in the Seeding and Annotation process. The incentivization system is intended to:

1. Reward high-quality data
2. Incentivize long-term retention
3. Disincentivize bad actors.

The use of cryptographic tokens makes the process transparent, secure, and impartial.

Contributors' efforts are also ranked on a public leaderboard, with the top performers receiving additional token rewards. This public-facing dimension uses the power of game mechanics (financial and psychological rewards) to foster healthy competition.

In addition to the three primary goals above, DCD tokens play a further role in rewarding developers who enhance and build on the Decide AI system.

The ability to share in the success of the platform means that developers have an incentive to build tools and services that integrate with the Decide AI ecosystem (which will in turn improve its efficiency and accessibility). It also incentivizes the creation of external, real-world applications that leverage LLMs trained by the Decide Protocol platform.

4. Training

To train the model properly and assign rewards fairly, we need a robust architecture that is:

- Capable of continually assessing data and model output
- Transparent in its mechanics
- Resistant to manipulation.

Assessing Uncertainty

The Decide Protocol will initially use the DeBERTa v3 (Decoding-enhanced BERT with Disentangled Attention) architecture, incorporating heteroskedastic uncertainty to assess whether a response can be seen as more or less reliable.

Data on reliability can be integrated into the model refinement process, to inform whether more or higher-quality reviewers are needed for certain topics or issues. This will help focus refinement efforts on where they are most needed.

Reinforcement Learning

The RLHF process will be based on the TRLX (Transformer Reinforcement Learning with eXtra supervision) framework and the Proximal Policy Optimization (PPO) algorithm.

PPOP is a state-of-the-art RL technique that optimizes LLMs by maximizing the expected reward for a given action while ensuring that the policy remains relatively stable over time.

This ensures quality and prevents people from gaming the system. Furthermore, the PPO algorithm receives data from the reward model (DeBerta v3) and uses this data to adjust the behavior of the LLM until performance reaches a satisfactory level.

Additional Improvements

In addition to the above, we have integrated various architectural improvements to optimize the quality of the data. The following mechanisms optimize annotation quality and create an engaging environment for contributors by assigning tasks aligned with their skills and interests:

- **Data Shapley Values:** Calculates the contribution of each data point to model performance, weighing annotations accordingly during reward model retraining.
- **Influence Functions:** Estimates the influence of individual annotations on model predictions, prioritizing valuable data points for retraining.
- **Cross-Validation:** Assesses the reward model's performance and generalization capabilities across different data subsets.
- **Annotation Allocation Mechanism:** Matches qualified individuals with appropriate tasks using valuation scores and cross-task models of annotator expertise.

We will continue to enhance the process outlined above as the LLM space matures and new methodologies become available.

5. Evolution

When the initial training cycle is complete, the trained LLM is compared with other leading-edge models, and a decision is taken to deploy the current iteration or initiate another training cycle to create a new base model.

If a decision is made to create a new base LLM, the insights gained from the initial training cycle will be used to seed the training of the new model.

When a model is deployed, the Decide Protocol continues to analyze conversations taking place in real-world applications, identifying areas for improvement. This analysis feeds back into the annotation process, ensuring that the LLM remains relevant and aligned with the target use case.

Decide ID

A large part of the Decide AI system depends on the ability to trust in human expertise.

To achieve this, Decide ID applies a proprietary methodology called Proof of Personhood (PoP). To acquire verified status, each individual must provide multiple proofs and complete a unique challenge.

The approach is applicable to LLM training, and any situation in which people are involved in a task or transaction.

At the most basic level, PoP ensures that:

- People are people (not bots)
- Their credentials are real
- Each contributor is unique.

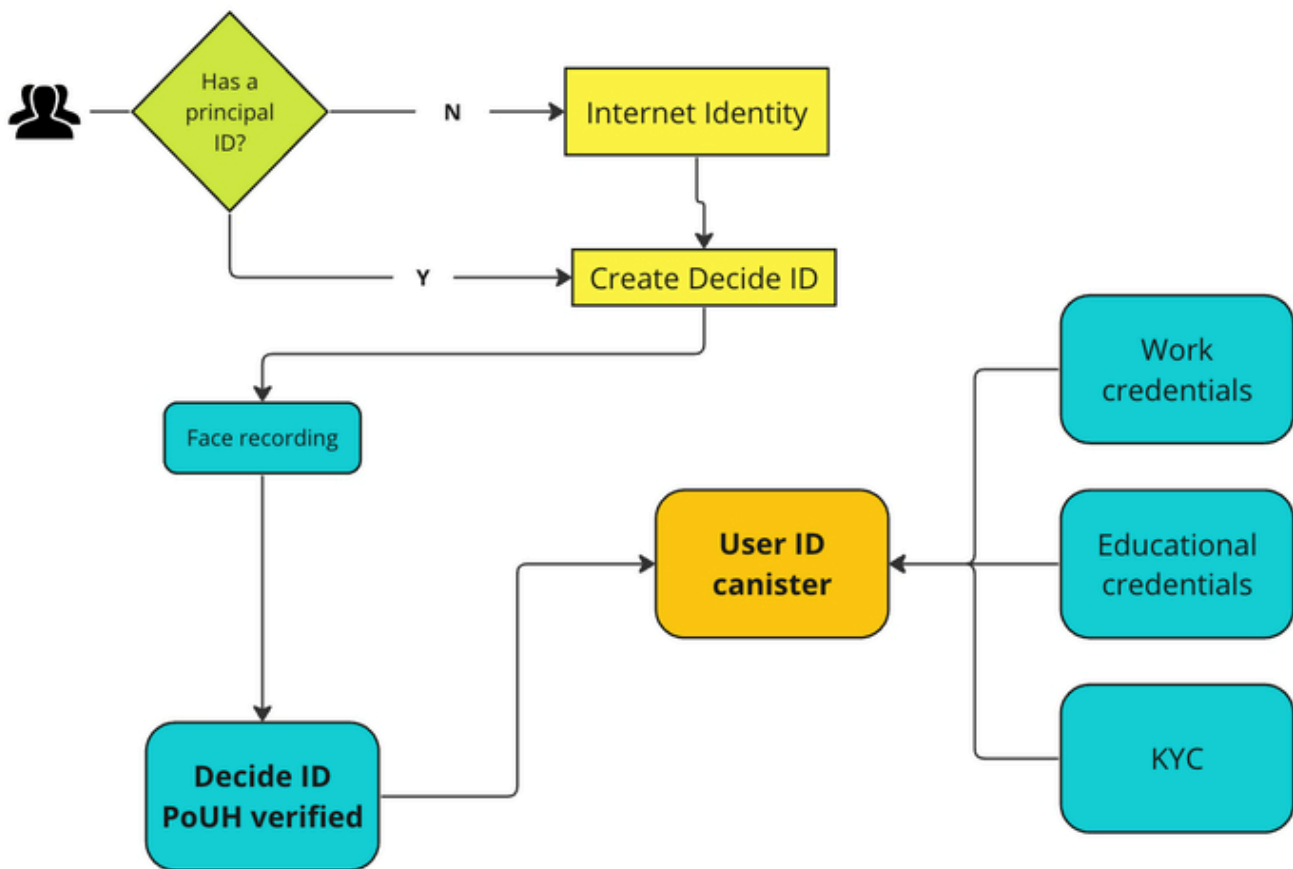
Benefits of PoP

In addition to being secure and effective, PoP offers the following advantages:

- Flexible: It is possible to customize the length and steps of the verification process (e.g. document types, use of video/audio).
- Transferable: Unique Principal ID can be used with other apps that also require Proof of Personhood.
- Seamless integration: Verification (for new users) and confirmation (for existing users) are easily integrated into any app's sign-up process.

ID Issuance & Storage

To contribute, an individual must receive a unique on-chain identifier (Principal ID), and provide their uniqueness.



ID issuance and storage

1. Generate ID: At initial login, a new user can request a Principal ID be generated.
2. Prove Uniqueness: The user is then invited to verify their uniqueness by submitting biometric information.
3. Additional Verification: The user submits further verification items (e.g. relating to credentials) which are also verified.

Verifying Credentials

The use of Zero-Knowledge (ZK) proofs means that in this process, the user's credentials can be verified safely without the risk of sharing private information (e.g. biometric data).

The process for verifying a user's credentials (by Decide ID or any other application) is as follows:

1. The user submits their data and verifier hash.
2. Decide ID sends back ZK proof - an encrypted version of the biometric and the verifier hash.
3. The ZK proof is then sent to the verifier, who confirms or denies the verification request.

No user data is accessible by a third-party application, and the user retains control and ownership (GDPR).

Applications

Decide Protocol

In the context of the Decide Protocol, the Decide ID system ensures that the workforce (i.e. generating and annotating the training data) is real and has the appropriate expertise (as indicated by work and educational qualifications) for the task. Every data point generated using the Decide Protocol can therefore be traced back to a verified user.

Other applications

While the initial use case for Decide ID is supporting the DecideAI ecosystem, it has many potential use cases in the broader web3 universe, where verification and authentication are pivotal.

- **Airdrops and Referral Programs:** PoP can prevent Sybil attacks, ensuring that tokens are distributed to the right human users.
- **Social Media:** Bot detection and influencer authentication ensure that users are engaging with genuine human profiles, preventing misinformation and enhancing credibility.
- **Digital Art and Content Creation:** By authenticating their identities, content creators can protect the value of their work, collaborate, and transact in a trustworthy manner.
- **DeFi Platforms:** By verifying users, DeFi platforms can comply with anti-money laundering (AML) regulations, mitigate fraud risk, and foster trust.
- **Peer-to-peer transactions:** Decide ID enables web3 users to confirm counterparty IDs, and so engage in secure and reliable interactions - the way Satoshi intended.

Decide Cortex

Decide Cortex democratizes AI technology by making state-of-the-art models easily accessible to developers, businesses, and researchers.

In practical terms, it serves as a knowledge-sharing hub that provides access to 'pre-trained' AI models and datasets:

1. As a purchased/licensed product.
2. On a bespoke basis via API endpoints.

Users can not only access, customize, and train models using Decide Cortex's capabilities, but also monitor and manage models post-launch.

Pre-trained models

A variety of models will be available on the platform, falling into two broad categories:

Flagship General Language Models

General purpose models serve as the foundation for various natural language processing tasks.

For example:

- Text generation: e.g. Draft an email based on a client's recent purchase history and preferences.
- Summarization: e.g. Provide a brief summary of key points from lengthy quarterly earnings reports.
- Question answering: e.g. Assist customer service by providing instant, accurate answers to common inquiries about product features.

As with all DecideAI models, those made available on the platform will be continually updated and improved.

Use Case-Specific Language Models

Targeted models excel at particular tasks or industries.

Example: [Redactor](#), the abuse detection language model, has been trained to identify and flag potentially abusive, harmful, or offensive language. The model is adjustable (e.g. level of bias and scale) and can be used for content moderation or online community management.

Tokenomics

Purpose of the DCD utility token

The DecideAI token DCD is the native token of DecideAI's platform and it is a Utility Token. DCD is a part of the ecosystem, intended to:

- Reward participation
- Fund innovation
- Fuel collaboration

In more detail, the DCD token fulfills the following functions:

- Universal Token for Ecosystem Transactions: Whether accessing models, computation resources, or data, DCD tokens streamline activities across the platform, from training to application development
- Model training incentives: Rewarding all data creators and data labelers ensures the sustainability of high-quality model training
- Developer Incentives: Everyone who contributes to the broader DecideAI ecosystem - whether through developing new models or applications to improving existing ones - will be

rewarded.

- **Market Creation Incentives:** The token also encourages the development of new markets and applications, which will in turn enrich the ecosystem and stimulate value exchange.

Total supply levers

The total supply of DCD tokens at genesis is 1 billion. Over time, the supply will increase if more tokens are minted and decrease if tokens are burned.

The DecideAI DAO was launched on the Internet Computer Service Neuron System August 28th 2023. The SNS is configured to generate 1.5% of the total supply annually to pay voting rewards to participating neurons. Voting rewards accumulate in participating neurons as [maturity](#).

At the point a neuron's maturity is disbursed it is burned and the corresponding value of DCD tokens will be minted by the SNS ledger to an account.

The only way the SNS can burn tokens is by proposal.

Incoming and outgoings

At genesis the SNS will have a reserve of ICP from the SNS swap and 426.5 M MOD tokens.

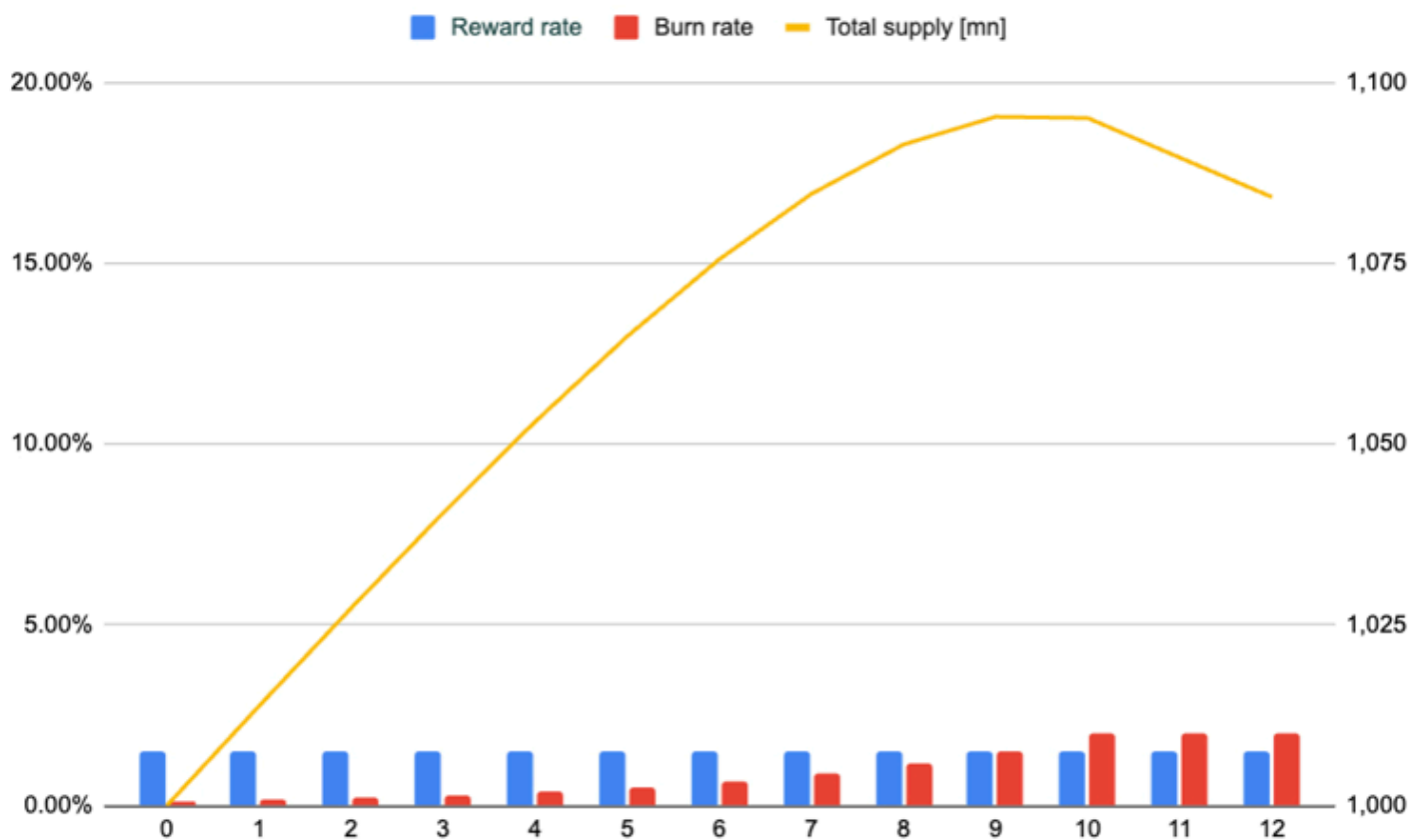
The SNS will receive DCD tokens from partner applications rewarding users for labeling services.

The SNS will have various outgoings. It will use ICP to pay the DecideAI dApp hosting costs (cycles), 3rd parties for services, and the DecideAI team. It will use DCD to reward users and for community bounties.

In the early years, outgoings will outstrip incomings, DCD reserve will largely be used to provide user rewards and community bounties. The expectation is that as incoming grows over time it will eventually balance outgoings. A higher rate of incoming would also allow a higher rate of user rewards and bounties to encourage higher growth in users and usage.

As the DecideAI DAO sees fit it can choose to burn DCD tokens to reduce the total supply. The expectation is that over several years the SNS will be able to afford to burn DCD at an increasing rate until the burn rate exceeds the minting rate from voting rewards and the total supply starts decreasing.

The following diagram depicts a projection of the total supply of DCD over time. For the purposes of this projection it is assumed that the reward rate will remain at a constant 1.5% and that the burn rate will start at 0.125%, increasing by a factor of 1.32 each year, until it overtakes the reward rate.



Total DCD supply over time projection

Liquidity Management

After the decentralization swap, participants will receive a basket of neurons of varying dissolve delays with only 1/5 being immediately liquid. The voting reward rate, initialized to 1.5%, is expected to encourage token holders to lock up a certain proportion of tokens thus, at least temporarily, removing them from the liquid supply. In the case of the seed funders, their neurons have vesting periods from 6-36 months before they can even start dissolving. In the case of the founding team, the vesting periods are from 1 year cliff with 36 month linear vesting.

There are various tokenomics parameters which can affect the proportion of DCD that is locked up. These include the max dissolve delay, the dissolve delay bonus, min dissolve delay to vote, max age, max age bonus, and voting reward rate. We have carefully chosen [initial values](#) for these parameters which we believe provide a good balance of incentives but these are all levers available to the DAO to allow it to influence the total and liquid supply and therefore the price if so desired.

Consider the SNS reserve of DCD tokens. These tokens are liquid but are only being trickled out (in percentage terms) to the community as user rewards and bounties, and then only some proportion will find their way onto the market (DEXes). It is a similar story for the portion of DCD held in the NNS reserve - it is liquid but will not enter the market unless the DAO decides to conduct a future swap.

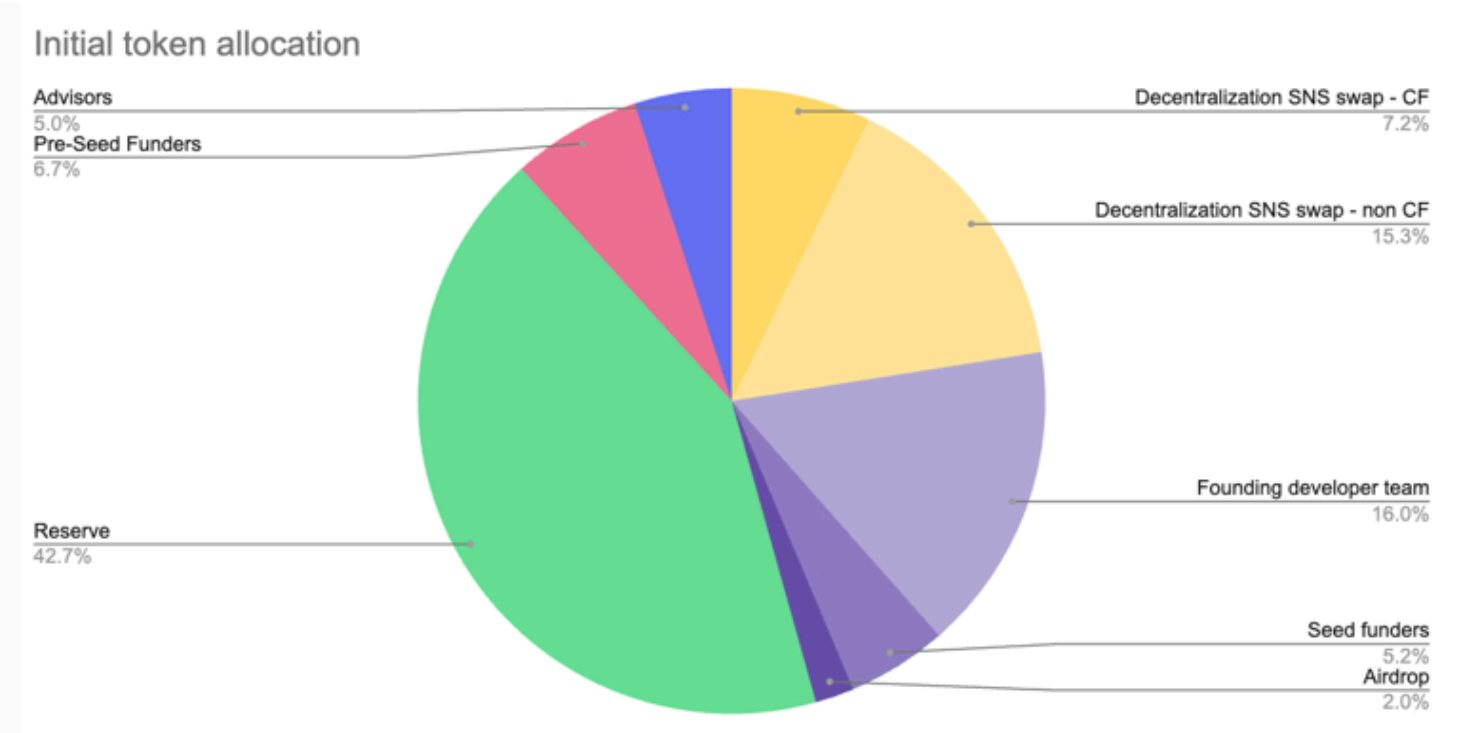
Token allocation at SNS genesis

Initial token allocation

The SNS was initialised with 1 billion tokens allocated in the following proportions on August 25th 2023.

Initial token distribution DCD (previously MOD), lock up and release schedule.

Token Category	DCD (mio)	Unlocked at TGE	Cliff (months)	Unlocked at end of the cliff	Release period (months)
Decentralization SNS swap	225	20%	0	0%	12
Team	160	0%	12	0%	36
Advisors	50	4%	6	4%	36
Pre-seed	66.85	8%	6	8%	36
Seed	51.65	0%	0	0%	24
Airdrop	20	100%	0	0%	0
Reserve	426.5	-	-	-	-



Token distribution at genesis.

Initial SNS configuration

The SNS will initially be configured with the values shown in the tables below which can all subsequently be changed by proposal.

Table 2: Initial SNS configuration

Transaction fee in MOD tokens that must be paid for ledger transfers	0.0001
Number of MOD tokens that a rejected proposal costs the proposer	200
Minimum number of MOD tokens that can be staked in a neuron	36
Maximum voting period for a proposal	4 days
Proportion of voting power needed for a proposal to be accepted	3%
Minimum neuron dissolve delay to vote	1 month
Maximum neuron dissolve delay	1 year
Maximum age for age bonus	2
Maximum age bonus	1.25x
Percentage of total supply that will be generated annually for rewards	1.5%

Team

Raheel, CEO and Founder

DecideAI's founder and CEO, Raheel has over a decade of experience in tech. Raheel is University of Waterloo Software Engineering graduate and a seasoned software engineer with product development and leadership experience. He has worked at both big companies and startups. He loves building products that solve problems and learning new technologies.

Pema, COO

An experienced public affairs and government relations advisor, Pema holds a master's degree from the University of Toronto Munk School of Global Affairs and Public Policy specializing in Global Capital Markets and Economics. She brings her diverse international consulting experience to DecideAI's operations.

Jesse, Lead AI Engineer

Jesse Glass, Ph.D., is a Machine Learning Specialist with over a decade of experience in artificial intelligence. A lead author of three AI publications, Jesse excels in designing custom RNNs and CNNs for data extraction, classification, and advanced regression models. His expertise spans reinforcement learning, data quality enhancement, and modernization of legacy systems. With a strong foundation in both academic research and applied machine learning, Jesse has led innovations across diverse fields, making him an invaluable asset to DecideAI's mission to advance cutting-edge AI solutions.

Tareq, Lead Software Architect

Tareq Khandaker is a seasoned Software Engineer with over 15 years of experience and a degree from the University of Waterloo. His career spans notable companies such as Qualcomm, Credit Karma, Clara Lending, and Freedom Financial Network, where he has honed his expertise in developing robust, scalable solutions. Tareq excels in large-scale microservice architectures and has demonstrated proficiency in a wide range of technologies including Java, Go, TypeScript, and Scala. With a strong background in API development, cloud migrations, and performance optimization, Tareq brings valuable skills in problem-solving and effective communication to our team.

Justin, Senior Full Stack Engineer

Justin is a seasoned engineer with experience as the technical lead of merchant onboarding and KYC at a leading global online marketplace.

Alex, Front End Engineer

Alex is an experienced front end engineer with a passion for design and creating functional and well designed products.

References

1. Truthful QA - GPT judge model and <https://arxiv.org/pdf/2109.07958.pdf>
2. Market sizing - <https://www.pragmamarketresearch.com/reports/121032/large-language-model-llm-market-size?UTM=AT23>

3. RLHF market <https://www.alliedmarketresearch.com/reinforcement-learning-market-A229407>
4. Peer prediction for learning agents - <https://arxiv.org/pdf/2208.04433.pdf>
5. Open assistant roadmap - https://docs.google.com/presentation/d/1n7lrAOVOqwdYgiYrXc8Sj0He8krn5MVZO_iLkCjTtu0/edit#slide=id.g1c27a3078af_0_1354
6. Open assistant repo- <https://github.com/LAION-AI/Open-Assistant/tree/5e9ba5f02c00149b58496fa2915ccffc0e664cfd>
7. Open assistant paper- <https://drive.google.com/file/d/10iR5hKwFqAKhL3umx8muOWSRm7hs5FqX/view>
8. Why RLHF - <https://www.alignmentforum.org/posts/Rs9ukRphwg3pJeYRF/why-do-we-need-rlhf-imitation-inverse-rl-and-the-role-of>
9. SFT vs RLHF - <https://medium.com/@joonbeomkwon/understanding-ai-training-supervised-fine-tuning-vs-reinforcement-learning-from-human-feedback-55d1b65317e0>
10. RLHF (PPO) vs DPO - <https://deci.ai/blog/dpo-preference-tuning-llms/#:~:text=While%20PPO%20is%20designed%20to,baselines%20or%20complex%20normalization%20terms.>
11. DPO vs RLHF - <https://medium.com/@sinarya.114/d-p-o-vs-r-l-h-f-a-battle-for-fine-tuning-supremacy-in-language-models-04b273e7a173>
12. ZKP implementation - <https://levelup.gitconnected.com/implementing-decentralised-identity-with-zk-snarks-68f9fca32c47>
13. How ZKP works - <https://blog.jarrodwatts.com/how-zk-proofs-and-zkevm-work>
14. ZK SNARK in action - <https://www.altoros.com/blog/securing-a-blockchain-with-a-noninteractive-zero-knowledge-proof/>

Legal Disclaimer

PLEASE READ THE ENTIRETY OF THIS "LEGAL DISCLAIMER" SECTION CAREFULLY. NOTHING HEREIN CONSTITUTES LEGAL, FINANCIAL, BUSINESS OR TAX ADVICE AND YOU SHOULD CONSULT YOUR OWN LEGAL, FINANCIAL, TAX OR OTHER PROFESSIONAL ADVISOR(S) BEFORE ENGAGING IN ANY ACTIVITY IN CONNECTION HERewith. NEITHER MODCLUB FOUNDATION OR ITS AFFILIATES (COLLECTIVELY, THE "FOUNDATION"), ANY OF THE PROJECT TEAM MEMBERS (THE "TEAM") WHO HAVE WORKED ON Modclub, DecideAI OR ANY PROJECT TO DEVELOP Modclub and DecideAI IN ANY WAY WHATSOEVER, ANY DISTRIBUTOR/VENDOR OF DCD/MOD TOKENS (THE "DISTRIBUTOR"), NOR ANY SERVICE PROVIDER SHALL BE LIABLE FOR ANY KIND OF DIRECT OR INDIRECT DAMAGE OR LOSS WHATSOEVER WHICH YOU MAY SUFFER IN CONNECTION WITH ACCESSING THIS WHITEPAPER (THE "WHITEPAPER"), THE WEBSITE AT [HTTPS://Modclub.APP/](https://Modclub.APP/) and

[HTTPS://Modclub.AI/](https://modclub.ai/) and <https://decideai.xyz/> (THE "WEBSITE") OR ANY OTHER WEBSITES OR MATERIALS PUBLISHED BY THE FOUNDATION. BY ACCESSING THE WHITEPAPER AND/OR THE WEBSITE, YOU AGREE TO BE BOUND BY THE TERMS OF THIS "LEGAL DISCLAIMER" SECTION.

Project purpose: You agree that you are acquiring DCD (formerly MOD) to participate in DecideAI (formerly Modclub) and to obtain services on the ecosystem thereon. The Foundation, the Distributor and their respective affiliates and service providers will develop and contribute to the underlying source code for DecideAI. The Modclub Foundation is acting solely as an arms' length third party in relation DecideAI the DCD distribution, and not in the capacity as a financial advisor or fiduciary of any person with regard to the distribution of DCD.

Nature of the Whitepaper: The Whitepaper and the Website are intended for general informational purposes only and do not constitute a prospectus, an offering memorandum, an offer document, an offer of securities, a solicitation for investment, or any offer to sell any product, item, or asset (whether digital or otherwise). The information herein may not be exhaustive and does not imply any element of a contractual relationship. There is no assurance as to the accuracy or completeness of such information and no representation, warranty or undertaking is or purported to be provided as to the accuracy or completeness of such information. Where the Whitepaper or the Website includes information that has been obtained from third-party sources, the Foundation, the Distributor, their respective affiliates and/or the team, have not independently verified the accuracy or completeness of such information. Further, you acknowledge that circumstances may change and that the Whitepaper or the Website may become outdated as a result; and neither the Foundation nor the Distributor is under any obligation to update or correct this document in connection therewith.

Token Documentation: Nothing in the Whitepaper or the Website constitutes any offer by the Foundation, the Distributor, or the team to sell any DCD or MOD (as defined herein) nor shall it or any part of it nor the fact of its presentation form the basis of, or be relied upon in connection with, any contract or investment decision. Nothing contained in the Whitepaper or the Website is or may be relied upon as a promise, representation or undertaking as to the future performance of Modclub and DecideAI. The agreement between the Distributor (or any third party) and you, in relation to any distribution or transfer of DCD or MOD, is to be governed only by the separate terms and conditions of such agreement.

The information set out in the Whitepaper and the Website is for community discussion only and is not legally binding. No person is bound to enter into any contract or binding legal commitment in relation to the acquisition of DCD or MOD, and no digital asset or other form of payment is to be accepted on the basis of the Whitepaper or the Website. The agreement for distribution of DCD or MOD and/or continued holding of DCD or MOD shall be governed by a separate set of terms and conditions setting out the terms of such distribution and/or continued holding of DCD or MOD (the "Terms and Conditions") or token distribution agreement (as the case may be), which shall be separately provided to you or made available on the Website. The Terms and Conditions must be read together with the Whitepaper. In the event of any inconsistencies

between the Terms and Conditions and the Whitepaper or the Website, the Terms and Conditions shall prevail.

Deemed Representations and Warranties: By accessing the Whitepaper or the Website (or any part thereof), you shall be deemed to represent and warrant to the Foundation, the Distributor, their respective affiliates, and the Modclub team as follows:

(a) in any decision to acquire any DCD or MOD, you have not relied on any statement set out in the Whitepaper or the Website;

(b) you will and shall at your own expense ensure compliance with all laws, regulatory requirements and restrictions applicable to you (as the case may be);

(c) you acknowledge, understand and agree that DCD or MOD may have no value, there is no guarantee or representation of value or liquidity for DCD or MOD, and DCD or MOD is not an investment product nor is it intended for any speculative investment whatsoever;

(d) none of the Foundation, the Distributor, their respective affiliates, and/or the

Modclub/DecideAI team members shall be responsible for or liable for the value of DCD or MOD, the transferability and/or liquidity of DCD or MOD and/or the availability of DCD or MOD through third parties or otherwise; and

(e) you acknowledge, understand and agree that you are not eligible to participate in the distribution of DCD or MOD if you are a citizen, national, resident (tax or otherwise), domiciliary and/or green card holder of a geographic area or country (i) where it is likely that the distribution of DCD or MOD would be construed as the sale of a security (howsoever named), financial service or investment product and/or (ii) where participation in token distributions is prohibited by applicable law, decree, regulation, treaty, or administrative act (including without limitation the United States of America and the People's Republic of China); and to this effect you agree to provide all such identity verification document when requested in order for the relevant checks to be carried out.

The Foundation, the Distributor and the Modclub/DecideAI team do not and do not purport to make, and hereby disclaims, all representations, warranties or undertaking to any entity or person (including without limitation warranties as to the accuracy, completeness, timeliness, or reliability of the contents of the Whitepaper or the Website, or any other materials published by the Foundation or the Distributor). To the maximum extent permitted by law, the Foundation, the Distributor, their respective affiliates and service providers shall not be liable for any indirect, special, incidental, consequential or other losses of any kind, in tort, contract or otherwise (including, without limitation, any liability arising from default or negligence on the part of any of them, or any loss of revenue, income or profits, and loss of use or data) arising from the use of the Whitepaper or the Website, or any other materials published, or its contents (including without limitation any errors or omissions) or otherwise arising in connection with the same. Prospective acquirors of DCD or MOD should carefully consider and evaluate all risks and



D DecideAI Whitepaper



Search...

⌘ K

uncertainties (including financial and legal risks and uncertainties) associated with the distribution of DCD or MOD, the Foundation, the Distributor and the DecideAI/Modclub team.

DCD/MOD Token: In particular, it is highlighted that DCD/MOD:

- (a) does not have any tangible or physical manifestation, and does not have any intrinsic value (nor does any person make any representation or give any commitment as to its value);
- (b) is non-refundable and cannot be exchanged for cash (or its equivalent value in any other digital asset) or any payment obligation by the Foundation, the Distributor or any of their respective affiliates;
- (c) does not represent or confer on the token holder any right of any form with respect to the Foundation, the Distributor (or any of their respective affiliates), or its revenues or assets, including without limitation any right to receive future dividends, revenue, shares, ownership right or stake, share or security, any voting, distribution, redemption, liquidation, proprietary (including all forms of intellectual property or licence rights), right to receive accounts, financial statements or other financial data, the right to requisition or participate in shareholder meetings, the right to nominate a director, or other financial or legal rights or equivalent rights, or intellectual property rights or any other form of participation in or relating to Modclub, the Foundation, the Distributor and/or their service providers;
- (d) is not intended to represent any rights under a contract for differences or under any other contract the purpose or pretended purpose of which is to secure a profit or avoid a loss;
- (e) is not intended to be a representation of money (including electronic money), security, commodity, bond, debt instrument, unit in a collective investment scheme or any other kind of financial instrument or investment;
- (f) is not a loan to the Foundation, the Distributor or any of their respective affiliates, is not intended to represent a debt owed by the Foundation, the Distributor or any of their respective affiliates, and there is no expectation of profit; and
- (g) does not provide the token holder with any ownership or other interest in the Foundation, the Distributor or any of their respective affiliates.

Notwithstanding any distribution of DCD or MOD, users have no economic or legal right over or beneficial interest in the assets of the Foundation, the Distributor, or any of their affiliates after the token distribution.

To the extent a secondary market or exchange for trading DCD MOD does develop, it would be run and operated wholly independently of the Foundation, the Distributor, the distribution of DCD MOD and Decideai or Modclub. Neither the Foundation nor the Distributor will create such secondary markets nor will either entity act as an exchange for DCD MOD.

Informational purposes only: The information set out herein is only conceptual, and describes the future development goals for Modclub to be developed. In particular, the project roadmap in the Whitepaper is being shared in order to outline some of the plans of the DecideAi/Modclub team, and is provided solely for INFORMATIONAL PURPOSES and does not constitute any binding commitment. Please do not rely on this information in deciding whether to participate in the token distribution because ultimately, the development, release, and timing of any products, features or functionality remains at the sole discretion of the Foundation, the Distributor or their respective affiliates, and is subject to change. Further, the Whitepaper or the Website may be amended or replaced from time to time. There are no obligations to update the Whitepaper or the Website, or to provide recipients with access to any information beyond what is provided herein.

Regulatory approval: No regulatory authority has examined or approved, whether formally or informally, any of the information set out in the Whitepaper or the Website. No such action or assurance has been or will be taken under the laws, regulatory requirements or rules of any jurisdiction. The publication, distribution or dissemination of the Whitepaper or the Website does not imply that the applicable laws, regulatory requirements or rules have been complied with.

Cautionary Note on forward-looking statements: All statements contained herein, statements made in press releases or in any place accessible by the public and oral statements that may be made by the Foundation, the Distributor and/or the DecideAi/Modclub team, may constitute forward-looking statements (including statements regarding the intent, belief or current expectations with respect to market conditions, business strategy and plans, financial condition, specific provisions and risk management practices). You are cautioned not to place undue reliance on these forward-looking statements given that these statements involve known and unknown risks, uncertainties and other factors that may cause the actual future results to be materially different from that described by such forward-looking statements, and no independent third party has reviewed the reasonableness of any such statements or assumptions. These forward-looking statements are applicable only as of the date indicated in the Whitepaper, and the Foundation, the Distributor as well as the DecideAi/Modclub team expressly disclaim any responsibility (whether express or implied) to release any revisions to these forward-looking statements to reflect events after such date.

References to companies and platforms: The use of any Foundation and/or platform names or trademarks herein (save for those which relate to the Foundation, the Distributor or their respective affiliates) does not imply any affiliation with, or endorsement by, any third party. References in the Whitepaper or the Website to specific companies and platforms are for illustrative purposes only.

English language: The Whitepaper and the Website may be translated into a language other than English for reference purpose only and in the event of conflict or ambiguity between the English language version and translated versions of the Whitepaper or the Website, the English

language versions shall prevail. You acknowledge that you have read and understood the English language version of the Whitepaper and the Website.

No Distribution: No part of the Whitepaper or the Website is to be copied, reproduced, distributed or disseminated in any way without the prior written consent of the Foundation or the Distributor. By attending any presentation on this Whitepaper or by accepting any hard or soft copy of the Whitepaper, you agree to be bound by the foregoing limitations.

MODCLUB FOUNDATION is a Foundation Established in Panama. It is located at Torre Advanced Building, 1st Floor Ricardo Arlas Street, Panama City, Republic of Panama.

Last updated 20 days ago