



YeBlock White Paper

V1.0

Liquid Intelligent Network · LIM

Decentralized computing power, storage, AI, privacy, and quantum resistance

Five in One Distributed Basic Protocol Network



catalogue

About YeBlock	6
Project stage (as of 2026-05)	9
Important disclosures	10
Part 0: YeBlock in 30 Seconds	11
One-line definition	11
Plain-English version	11
Why YeBlock, and why now	12
Part 1: The LIM Model and Industry Context	12
1.1 Three Structural Problems in Today's AI Industry	12
.....	13
1.2 A New Idea: A Five-Pillar, Composable Inference Network	13
1.3 For the Four Core Participant Types	14
Part 2: The Five Pillars, One at a Time	15
2.1 Decentralized Compute	15
2.2 YeBlock Decentralized Storage	18
2.3 YeBlock Decentralized AI	22
2.4 YeBlock Privacy Protocol	26
2.5 YeBlock Post-Quantum Protocol	32
Part 3: How the Five Pillars Compose, and What Emerges	38

3.1 Drop one pillar and it degrades: what YeBlock becomes without each	38
3.2 The five-pillar reaction: capabilities you can only unlock in combination.....	39
3.3 The five-pillar protocol stack	40
3.4 The five-pillar market coverage map — why five pillars are needed to cover the whole market	41
3.5 Summary: YeBlock is a composition, and a disruptor	41
Part 4: How to Participate and How Rewards Work	42
⚠ Compliance and honesty statement on the revenue estimates in this section..	43
4.0 Overview: the 6-way quick-reference table	43
4.1 ① Compute Provider: earn from an idle GPU	45
4.2 ② Storage Provider: earn ongoing storage income from idle drives	47
4.3 ③ LoRA Creator: turn "domain knowledge" into a tradable asset	49
4.4 ④ LP (Liquidity Provider): use YBT to back the best assets	51
4.5 ⑤ AI User / Developer: cut your inference bill by 90%	53
4.6 ⑥ KOL: referral rebates + long-term passive income	55
4.7 How to pick your role (decision table)	57
Part 5: The YeBlock Liquid Economy	58
Part 6: FAQ + Risk Disclosures	67
6.2 Is it really safe?	68
6.3 Team anti-exit (learning from the Ethereum Foundation token-sale saga)	68

YeBlock Token: 121,000,000	69
6.4 Regulatory risk?	70
6.5 What if the project fails?	70
6.6 Can I participate if I'm not technical?	71
6.7 Are earnings paid in USDC or in the token?	72
6.8 If the AI industry cuts prices sharply (OpenAI drops prices), is YeBlock still competitive?	72
6.9 YeBlock vs Bittensor — competition or collaboration?	72
6.10 Four key challenges (publicly disclosed) — the ones other projects avoid saying out loud	73
6.11 KYC / AML progressive roadmap	76
Part 7: Team and Organization	77
7.0 What we're putting in	77
7.1 Team structure (14 functions)	77
7.2 Founding team	80
7.3 Funding sources	80
7.4 Team operating principles	81
Part 8: Roadmap and Timeline	81
8.1 Four time horizons	81
8.2 Next 4 quarters (what participants can do)	81

8.2 A key commitment: 18-month real-demand bar 83

8.3 Mid-term direction 83

8.4 Long-term vision 84

8.5 Horizon direction 84

Closing and invitation to build together 85

Core belief 85

Public commitments 85

Invitation to build 86

Contact 86

About YeBlock

YeBlock LIM combines decentralized compute + decentralized storage + decentralized AI into a single open protocol — turning AI infrastructure from something you can only consume into something you can own a piece of.

YeBlock LIM combines decentralized compute, decentralized storage, and LoRA-granular decentralized AI into a single open protocol.

Idle gaming GPUs contribute compute and earn incentives. Idle NAS drives generate ongoing storage income. Domain-specific LoRAs pay their creators on-chain in perpetuity, the same way music royalties do.

On equivalent tasks, inference cost can drop to as low as 1/10 of a centralized model — and ordinary users get a real ownership stake in the network.

YeBlock (the Liquid Intelligence Network) is among the first infrastructure projects to fuse five capabilities into a single open protocol:

- LoRA-granular on-chain revenue splits
- IPFS-anchored model hosting
- Multi-LoRA pooled inference nodes
- Application-layer tiered privacy (TLS / TEE / ZKML, introduced progressively)
- NIST post-quantum cryptography

On top of these five pillars, YeBlock defines the Liquid Economy — Liquid Trinity.

YeBlock LIME lets ideas be encrypted, put on-chain, executed by the mesh, and earn perpetual royalties.

YeBlock LEM converts cheap or surplus energy into on-site compute revenue (running nodes, hosting hardware, JouleCredit vouchers).

YeBlock LIP provides streaming micropayments and A2A settlement for the Agent economy.

LIM produces intelligence. LIME trades ideas. LEM supplies energy. LIP handles settlement.

Today's AI industry runs on three structural sources of friction:

- Compute is locked up by the cloud giants (60%+ gross margins).
- Model hosting is dominated by centralized platforms (creators can be deplatformed overnight).
- Millions of LoRA fine-tuners create real value with no on-chain path to monetize it.

YeBlock attacks all three at once.

Idle GPUs supply compute and earn incentives. Idle drives supply storage and earn yield. Trained LoRAs earn on-chain royalties in perpetuity, the same way music rights do.

On equivalent tasks where an open-source model is good enough, a developer's inference cost on YeBlock can drop to 1/10 of a traditional AI stack.

Ordinary users have six clear ways in:

- Compute provider
- Storage provider
- LoRA creator
- LP
- User
- KOL

So compute, storage, and intelligence all flow freely at the protocol layer.

DeFi made money liquid. YeBlock LIM makes intelligence liquid.

⚠ Scope of this claim: an open protocol that simultaneously delivers:

- ① LoRA-level revenue-split granularity
- ② IPFS-anchored model persistence
- ③ vLLM/SGLang multi-LoRA pooled nodes
- ④ Token-level on-chain settlement

Bittensor operates at a coarser granularity. HuggingFace has no on-chain splits.

io.net / Akash have no LoRA-level model layer.

YeBlock LIM is the early practitioner of this specific combination — and a pioneer in decentralized AI infrastructure.

Dimension	Number
Fully fused protocol pillars	5 (compute + storage + AI + privacy + post-quantum)
Ways to participate	6 roles
Target monthly income per RTX 4090 (mature network)	\$400 – \$2,000
AI inference cost on equivalent tasks vs OpenAI	up to –90%
LoRA author revenue split	15% on-chain · perpetual
Time horizon	12 – 24 months

- Revenue figures (2,000 / RTX 4090) = network-maturity targets; expected real values in the first 12–24 months are 10–30% of the target.
- Cost figures (–90% vs OpenAI) = valid only for equivalent tasks where model capability is substitutable (e.g. lim-8b + a legal LoRA vs GPT-4o-mini for legal Q&A). Not intended as a general-task comparison against frontier closed models like GPT-4o or Claude 3.5 Sonnet.
- Revenue-split figures (15%, per §2.1.4 split table) = target design; final parameters are decided by DAO governance.

Project edition: five-pillar pragmatic edition + long-term vision. Document edition: full version — how to participate and earn, candid early-stage disclosure, plus the privacy and post-quantum protocols.

Project stage (as of 2026-05)

YeBlock is an early-stage open-protocol project. Every revenue estimate, roadmap, and integration example below is a target design — please make participation decisions on an "early-stage project" basis.

Dimension	Current status
Protocol design	v4.1 complete, entering engineering phase
Website	✓ yeblock.com is live (publicly accessible)
Early user traction (estimated first-week public data)	~10–30k daily new signups / ~30k DAU / ~80% 7-day retention (estimated)
Testnet Contribution Proof (TCP)	Open — users accumulate YBT by completing tasks as testnet incentive vouchers
Testnet node network	Under development; a seed-node network is targeted for 2026-Q3
Mainnet	Not yet launched
YBT token	Not yet issued
Team	50 full-time, US + Singapore, covering 14 functions

Dimension	Current status
Funding source	Founding team self-funded; core members from top-tier Web3 projects
External funding	None yet; a pre-seed round is planned based on yeblock.com traction
Design Partners	None yet — a pre-seed round is planned based on yeblock.com traction
Open-source code	First 10 seats open, with genesis-tier benefits

Important disclosures

1. **YBT: YeBlock Network Contribution Proof Token (Testnet Contribution Proof for the current stage)** — a record of users' contributions to network stability. Distribution rights over future ecosystem incentives will be decided by DAO governance.
2. **Regional restrictions:** based on global legal and compliance requirements, registration and participation are geo-blocked in select jurisdictions.
3. **Progressive KYC roadmap:** no KYC in the current stage; light KYC for regular users before exchange listing; at mainnet, KYC is required for nodes / LPs / creators, together with sanctions-list blocking.
4. **Note on the early numbers:** the "first-week public data" in the table above was collected in the first week yeblock.com went live. Retention and activity data will be published in regular monthly reports as the network matures.

Part 0: YeBlock in 30 Seconds

One-line definition

YeBlock LIM combines three things into a single open protocol:

- LoRA-granular decentralized AI
- IPFS-anchored model storage
- Multi-LoRA pooled compute

The result: idle hardware can earn by contributing compute; model weights cannot be unilaterally deplatformed by any single platform; and LoRA fine-tuners can finally get paid.

Plain-English version

Picture this:

- That gaming PC of yours can run AI inference for 12 hours a night while you sleep, earning on the order of tens to hundreds of dollars a day (illustrative; varies by region and network stage).
- That 8TB drive you bought for a NAS and never really used can help the world store the "small building blocks" that make up AI models — and keep paying you for it.
- That legal-contract-review LoRA you trained pays you every time someone uses it — the same way a song pays royalties every time it's played.
- You're just using AI to write articles, but by combining YeBlock's open-source base models with specialized LoRAs, you pay up to 90% less than OpenAI's mini-tier models on equivalent tasks (actual savings depend on task type and model choice).
- You're not technical and don't own hardware — you simply use USDT to vote for (provide liquidity to) top LoRAs and share in their inference revenue.

That's YeBlock — a project designed so anyone can take part, and one that gives AI infrastructure back to everyone.

Why YeBlock, and why now

Timing	Engineering reality
vLLM / SGLang production-ready	Multi-LoRA concurrent inference on a shared base is now shippable
On-chain gas costs are very low	The economics of on-chain call-fee splits work (Base / Arbitrum)
Open-source LLMs catch up to closed models	Llama-3 / Qwen-2.5 / Mistral / DeepSeek are commercially usable
Centralization crisis at HuggingFace	The industry is worrying about single-point risk in model hosting
Regulation starts to reward open / compliant systems	DePIN + AI is a compliance-friendly new category

Part 1: The LIM Model and Industry Context

1.1 Three Structural Problems in Today's AI Industry

Problem 1: Compute monopoly

↓

Centralized clouds like AWS / Azure / GCP run at 60%+ gross margins

AI startups can't afford H100s (\$3–5/hour)

Small-to-mid inference workloads that aren't latency-sensitive still pay premium prices

Problem 2: Centralized model hosting

↓

HuggingFace is the de facto standard, but it's a single point

Models have been pulled from the platform on multiple occasions

Uploading to HF gives the author no monetization at all

Problem 3: Contributors have no path to monetization

↓

Open-source LoRA fine-tuners have created enormous value
With no sustainable way to earn from it
Long-tail fine-tuners drop out → the ecosystem shrinks

1.2 A New Idea: A Five-Pillar, Composable Inference Network

YeBlock treats five things as co-equal pillars of a single open protocol. The first three decide what it can do (the capability layer); the last two decide who it can do it for and how long it can keep doing it (the trust and time moats):

YeBlock LIM · Five-pillar protocol stack	
Capability layer / (what it does)	Decentralized Compute · Decentralized Storage · Decentralized AI
Trust layer / (who it does it for)	Privacy Protocol / client-side encryption / TEE / DP / ZKML (progressively introduced) / FHE (long-term) · the ticket into high-ticket B2B markets
Time layer / (how long)	Post-Quantum Protocol / TLS hybrid KEM / ML-KEM / ML-DSA / SLH-DSA (NIST 2024) · the hard gate for 5–10-year AI-asset preservation

The role each of the five pillars plays:

Layer	Pillar	Role
Capability layer	Compute + Storage + AI	What YeBlock can do (the supply / asset / intelligence side of inference)
Trust layer	Privacy Protocol	Who YeBlock can serve (decides whether B2B / medical / legal / financial customers can use it)
Time layer	Post-Quantum Protocol	How long YeBlock can serve (decides whether AI assets still have value 5–10 years out)

Why we promoted privacy and post-quantum from "cross-cutting concerns" to first-class pillars:

- "Cross-cutting" implies "a patch, a hardening step, a nice-to-have" — valuable but optional.
- "Pillar" implies "without it, the whole thing collapses." Without a privacy protocol, YeBlock cannot sign up compliant customers in the B2B market. Without post-quantum signatures, YeBlock's AI assets sit inside a 5–10-year "Harvest Now, Decrypt Later" attack window.
- Neither of these is a long-term aspiration — both have production-grade implementations available today.

1.3 For the Four Core Participant Types

Participant	YeBlock
Gamers with idle GPUs	Plug in an RTX 4090, contribute compute, earn network incentives + call fees
Users with idle drives	Plug in 8TB+ of space and earn ongoing storage income

Participant	YeBlock
Developers who fine-tune models	Train and upload a LoRA; earn every time it's called
AI application developers	Combine open-source models with specialized LoRAs; on equivalent tasks, cost can drop to ~1/10 of an OpenAI mini-tier model

Part 2: The Five Pillars, One at a Time

2.1 Decentralized Compute

2.1.1 What it is

The cleanest analogy: an Airbnb / Uber for compute.

- Airbnb aggregates idle rooms around the world into a virtual hotel you can book by the night.
- Uber aggregates idle driver-hours into a virtual fleet you can hail on demand.
- **YeBlock's decentralized compute aggregates the world's idle GPUs into a virtual cloud you can rent.**

2.1.2 Why we need it

Real pain points on the demand side:

AWS H100 instance = \$3-5 / hour
 8x H100 monthly rent = \$17,000+ / month
 For AI startups = unaffordable
 For small inference workloads = terrible price/performance

A massive untapped supply side:

Global gaming PCs + workstations

↓

NVIDIA RTX 40 / 50 series (4070 Ti / 4080 / 4090 / 5090)

- Tens of millions shipped worldwide (NVIDIA filings + JPR estimates)
- A single 4090 sits idle 12–18 hours a day
- Electricity: \$0.5–1 / day; rentable on YeBlock at a target of \$4–11 / day

↓

Even at 1% activation, the theoretical idle supply $\approx$ hundreds of thousands of H100 equivalents

⚠ These numbers are based on public market reports plus engineering estimates.

Actual deployable compute depends on operator uptake, node stability, geographic distribution, and other variables. Treat these as engineering estimates, not investment guidance.

The arbitrage window is real:

- Centralized cloud gross margins are 60%+ (per AWS / Azure / GCP filings).
- A decentralized model can plausibly land at 30–50% of centralized cloud pricing.
- Users pay 30–50% of centralized cloud pricing, and providers still make money after scheduling, verification, and protocol fees are netted out.

2.1.3 A category that's been repeatedly validated

Project	Market cap / valuation (order-of-magnitude)	Model	What it validates
Bittensor (TAO)	Billions USD	Subnet compute + AI incentives	On-chain AI incentives are viable
io.net	Hundreds of millions USD	GPU compute marketplace	Global GPU pooling is viable
Akash	Hundreds of millions to \$1B	General-purpose cloud	Decentralized cloud can run in

Project	Market cap / valuation (order-of-magnitude)	Model	What it validates
Network		compute	production
Render Network	Billions USD	Rendering compute market	A creator + provider two-sided market works
Gensyn	Early investment from A16z / Coinbase	AI training compute	Single-card fine-tuning works

Crypto valuations swing hard, so the table above is order-of-magnitude only, not precise (check CoinGecko / CoinMarketCap for real-time data). The point stands: decentralized compute is a mature category, validated by both capital and engineering. YeBlock's differentiation is the five-pillar fusion (compute + storage + AI + privacy + PQ).

2.1.4 How YeBlock does decentralized compute

Global compute nodes (vLLM / SGLang)

- |
- ├─ Each node keeps 1–3 base models resident (e.g. `lim-8b`)
- ├─ The same node concurrently mounts 5–30 LoRAs (already supported by vLLM)
- ├─ Accepts standard API-compatible requests
- └─ Runs inference → claims rewards on-chain

↓

Off-chain matching scheduler

- ├─ Picks the optimal node based on LoRA requirements
- ├─ Factors in node reputation / latency / price
- └─ Cross-verifies 1–5% of tasks via redundant sampling

↓

On-chain revenue-split contract

- ├─ Compute node 30% / LoRA author 15% / LP 15%

└─ Storage 13% / Verification 2% / Protocol 15% / Referral 10%

└─ Instant settlement in USDT/USDC

2.1.5 How YeBlock is different (vs Bittensor / io.net)

Dimension	YeBlock	Bittensor	io.net
Granularity	LoRA-level (fine)	Subnet-level (coarse)	Whole-machine rental
Billing	Pay-per-token (YBT + USDC)	Inflation subsidy (TAO)	By the hour (USDT)
Payees	Nodes + LoRAs + Storage + LPs, etc.	Nodes + subnet holders	Nodes only
Developer onboarding	5-minute integration	Weeks (learning Yuma)	A few hours
Task types	Inference + single-card fine-tuning	Inference + subnet-defined tasks	Anything

2.1.6 What YeBlock will never do (vs industry over-promising):

- ✗ We won't build our core trust model on ZK / TEE (3+ years from real maturity).
- ✗ We won't try to compete head-on with OpenAI's closed frontier models.

2.2 YeBlock Decentralized Storage

2.2.1 What it is

The cleanest analogy: an Airbnb for hard drives.

Aggregate the world's idle disk space into a virtual cloud storage service — a shape especially well-suited to "AI model weight" data:

- Write-once (uploaded a single time)
- Read-many (invoked repeatedly)
- Cannot be lost (it's an asset)
- Cannot be deplatformed (censorship-resistant)

2.2.2 Why we need it

Pain point 1: model weights must not be deplatformable

- HuggingFace is the de facto standard, but it's a fully centralized one.
- Models and datasets have been pulled multiple times (due to regulation, copyright, or policy).
- Once removed, years of the author's work go to zero.
- **YeBlock stores weights at the protocol layer, guaranteeing they cannot be unilaterally removed by any single platform.**

Pain point 2: cold-data storage on decentralized networks beats centralized cloud by orders of magnitude

Service	Price (\$/GB/month)	Monthly cost of 100GB of LoRAs (storage only)
AWS S3 Standard	0.023	\$2.30
Google Cloud Storage	0.020	\$2.00
Filecoin (order-book floor)	~0.0001	~\$0.01
Arweave (one-time pay for permanent storage)	one-time ~\$5/GB	0 (permanent)

⚠ The table above compares raw storage price only.

It does not include:

- Retrieval latency
- Egress bandwidth fees
- Redundancy factors (Filecoin typically needs 3–5 replicas)
- Network congestion

For read-heavy, cacheable, latency-tolerant data — exactly what AI model weights are — the per-unit price advantage translates into real cost savings of 1–2 orders of magnitude.

For low-latency access or frequently mutated data, centralized cloud is still the better answer. YeBlock puts data where it fits, and doesn't try to replace every scenario.

Pain point 3: censorship resistance + user sovereignty

- A user's conversation history / Agent state can be encrypted and stored on decentralized storage.
- Only the user holds the keys — the platform can neither take it nor read it.
- Seamless migration across platforms and devices.

2.2.3 Another well-validated category

Project	Market cap / valuation (order-of-magnitude)	Positioning
Filecoin	Billions USD	General-purpose decentralized storage market (largest)

Project	Market cap / valuation (order-of-magnitude)	Positioning
Arweave	Hundreds of millions USD	Permanent storage (one-time fee)
Storj	Hundreds of millions USD	Enterprise-grade, S3-compatible
IPFS	Protocol-level (not a token project)	Distributed file addressing (de facto standard)
Crust Network	Hundreds of millions USD	IPFS incentive layer

Same disclaimer as above: crypto valuations swing hard, so these are order-of-magnitude numbers. The point: YeBlock solves the same pain points while rebuilding the underlying layer to combine their strengths.

2.2.4 How YeBlock does decentralized storage

LoRA author uploads

↓

Client-side encryption (optional, for private LoRAs)

↓

Sharding + redundancy coding (Reed-Solomon)

↓

Distributed across the YeBlock network

- ├ The project pins 20–50 seed nodes to guarantee 99.9% availability
- ├ Users run their own nodes as edge caches
- ├ Compute nodes hot-cache popular LoRAs locally

↓

When user data is written on-chain through YeBlock, the protocol assigns it to multiple YeBlock nodes; if a node goes offline, the data is automatically re-assigned to other nodes.

↓

On-chain: CID + metadata + collateral + license registered

↓

Access: pull by CID + on-chain verification + client-side decryption

2.2.5 How YeBlock is different

- **Purpose-built for AI assets:**
 - LoRAs are typically 10MB–1GB.
 - Call patterns are predictable (great for prewarming caches).
 - Every call pays out (storage nodes are revenue nodes too).
 - **Compute–storage co-location: compute nodes double as storage nodes for part of the workload (one machine, two roles).**
 - **Enforced license metadata: compliance-friendly by default.**
-

2.3 YeBlock Decentralized AI

2.3.1 What it is

YeBlock ensures that model-weight hosting, inference execution, the flow of economic incentives, and governance decisions are no longer monopolized by centralized platforms.

The v4.1 entry point: the YeBlock LoRA marketplace + on-chain revenue splits — turning "AI capability" into the smallest tradable unit.

⚠ Key concept: "decentralized AI" is not the same as "AI/LoRA decomposition."

- "Decentralized AI" is the dimension (every form of AI that isn't monopolized).
- "LoRA decomposition" is the mechanism (the specific implementation in v4.1).
- The design will evolve to include Expert CUs, ZKML verification, on-chain AI Agents, and more.

2.3.2 What is a LoRA — an introduction for non-technical readers

LoRA = Low-Rank Adaptation. You can think of it as:

Base large model (e.g. lim-8b) $\approx$ a general-purpose computer

LoRA $\approx$ a specialized app / a personal skill pack

- The base model is large (8B parameters, ~16GB on disk).
- A LoRA is small (typically 10MB–1GB, 1/100 to 1/1000 the size of the base).
- The base model is loaded once and stays resident in memory; multiple LoRAs can be hot-swapped on top of it.
- A LoRA turns a general-purpose model into, say, a "legal contract reviewer" or a "novelist."

Analogies:

Base model	LoRA
An iPhone	An app
A game console	A game cartridge
A general-purpose college student	A specialized certification
A general-purpose brain	A stretch of specialized training / memory

2.3.3 Why decentralized AI matters

Pain point 1: HuggingFace is centralized, and authors earn nothing

- HuggingFace is the de facto LoRA standard, but:
 - It's a single, centralized point (can be deplatformed).
 - The author uploads $\rightarrow$ everyone downloads for free $\rightarrow$ the author earns \$0.
- Long-tail LoRA authors create real value but have no way to monetize it $\rightarrow$ the best contributors leave.

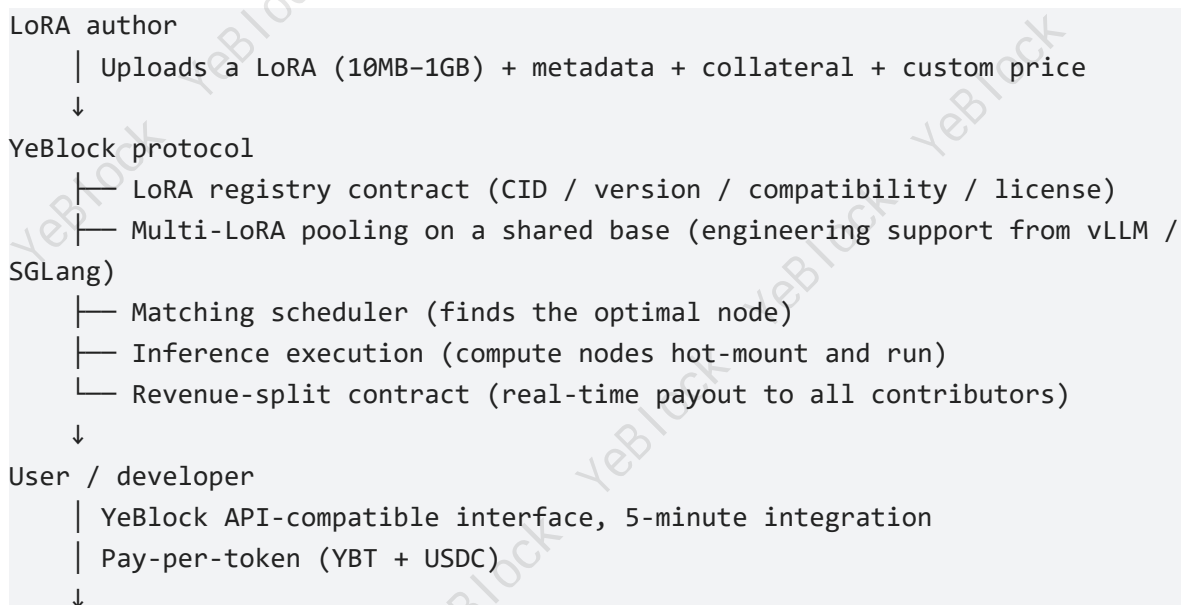
Pain point 2: Bittensor's granularity is too coarse

- Bittensor incentivizes at the subnet level.
- Each subnet requires hundreds of thousands of dollars to launch and months of maintenance.
- Individual LoRA authors can't benefit directly.

Pain point 3: users have no "call a LoRA directly" entry point

- Want to combine a legal LoRA with YeBlock's lim-8b?
- HuggingFace: download the base + download the LoRA + stand up your own service (needs GPUs).
- YeBlock: quick integration via API.

```
client.chat.completions.create(
  model="lim-8b",
  loras=["legal-contract-v1"], # ← LoRA mounted directly
  messages=[{"role": "user", "content": "..."}]
)
```

2.3.4 How YeBlock does decentralized AI

All contributors split by pre-set ratios
 Compute node 30% / LoRA author 15% / LP 15%
 / Storage 13% / Verification 2% / Protocol 15% / Referral 10%

2.3.5 How YeBlock is different (vs HuggingFace / Bittensor)

Dimension	YeBlock	HuggingFace	Bittensor
Granularity	LoRA-level (fine)	LoRA-level	Subnet-level (coarse)
Centralized	✗	✓	✗
Author monetization	✓ on-chain revenue splits	✗ none	✗ indirect (via subnet alpha)
User onboarding	5-minute YeBlock API	Requires self-hosted service	Complex
Call billing	Pay-per-token, YBT + USDC	Free + self-hosted	Inflation subsidy
Compliance-friendly	High (enforced license)	Medium	Weak

2.3.6 The scope of YeBlock's decentralized AI (not limited to LoRA)

v4.1 enters through the LoRA market, but "decentralized AI" is a broader agenda:

- **Near-term entry: LoRA market + on-chain revenue splits.**
- **Mid-term: multi-LoRA composition, lightweight semantic routing, cross-modal (image / video LoRAs).**

- **Longer-term: Expert CU (MoE decomposition), cross-base CU composition, ZKML sampling verification, TEE as a primary trust surface.**
- **Long horizon: on-chain AI Agents, cross-network pipeline inference.**

2.4 YeBlock Privacy Protocol

2.4.1 What it is

In one line: run inference and settle payments so that user inputs, the inference process, and the outputs stay invisible to node operators, the scheduler — and even the YeBlock team.

The cleanest analogy:

- A safe-deposit box in a bank vault — the customer holds the key; the bank provides the space, maintains the equipment, and handles security.
- You can deposit, the bank can keep records, the box can be opened and closed — but bank staff never see what's inside your box.

The real hard problem for decentralized AI isn't "is there enough compute?" — it's "why would a user hand their data to a bunch of strangers to run?" The privacy protocol is the pillar that answers that question.

2.4.2 Why we need it — the hard gate into high-ticket B2B markets

Pain point 1: B2B customers will not run their data in the clear on a decentralized network

Industry	Data sensitivity	Hard requirement not to go on-chain / not to be readable by nodes
Medical	Extremely high	HIPAA (US) / GDPR medical & health data (EU) / other jurisdictions' medical-privacy laws — cross-border or third-party transfers of

Industry	Data sensitivity	Hard requirement not to go on-chain / not to be readable by nodes
		sensitive health data are typically strictly limited
Legal	Extremely high	Attorney–Client Privilege — the moment plaintext contracts flow to a third party, privilege is lost
Financial	High	Financial-data localization + mandatory encryption in transit in most jurisdictions
Government / defense	Extremely high	State secrets + trade secrets + personal privacy stacked on top of each other
Enterprise internal	Medium–high	Trade secrets / customer lists / strategy documents

Those five customer classes together form the highest-ticket portion of the AI inference market (single-customer LTVs are typically 100–1,000× those of a regular consumer). No privacy protocol = YeBlock is locked out of that market.

Pain point 2: existing decentralized networks have weaker node-trust assumptions than centralized clouds

Centralized clouds (AWS / Azure) are at least bound by SOC 2 / ISO 27001 and a single-legal-entity contract — there's someone to hold accountable.

Compute nodes on today's decentralized networks may be:

- Household GPUs distributed worldwide
- Run by strangers
- Spread across multiple jurisdictions

Weaker trust assumption → stronger crypto protection required. The YeBlock privacy protocol is the entry ticket to a decentralized network.

Pain point 3: the natural privacy requirement of "user private LoRAs"

A LoRA can itself be the user's private asset (self-trained, built on private data).

- Uploaded to public IPFS = public = anyone can download and copy it.
- Uploaded to a YeBlock private-LoRA channel (encrypted + only authorized nodes can decrypt and execute it) = usable as an asset with confidentiality preserved.

This "make private assets usable" scenario is unique to YeBlock, and impossible without a privacy protocol.

2.4.3 Technology stack

Technology	Engineering maturity	Real production examples	Where YeBlock uses it
TLS 1.3 + client-side encryption	✓ Fully mature	Every HTTPS service, Signal / WhatsApp	User ↔ scheduler, scheduler ↔ nodes
Intel TDX / AMD SEV-SNP	✓ Industrial-grade in production	Azure Confidential Computing, Google Cloud CVM, Phala / Marlin / Oasis	High-sensitivity inference nodes (TEE nodes)
NVIDIA H100 / H200 Confidential Computing	✓ Shipped in production 2024	NVIDIA + Microsoft joint deployments	TEE GPU nodes (core high-end channel)
Differential Privacy (DP)	Mature at the training stage	Google / Apple / Meta usage analytics	LoRA training-data protection, call-statistics reporting
MPC (multi-party secure computation)	Mature in DeFi wallets	Fireblocks / Threshold / Lit Protocol	Multi-sig wallets, threshold signatures
ZKML (zero-knowledge machine learning)	● Early	Modulus Labs / Giza / EZKL (small-model demos)	Mid-term: small-model sampling verification
FHE (fully)	● Lab-grade	Zama / Inco (small-input	Long-term: specific small

Technology	Engineering maturity	Real production examples	Where YeBlock uses it
homomorphic encryption)		demos)	inferences / private queries

Key takeaway: YeBlock's privacy protocol is a mature engineering composition of multiple layered technologies. YeBlock rebuilds the underlying cryptography and completes the application-layer composition plus the engineering wiring into the compute, storage, and decentralized-AI pillars.

2.4.4 How YeBlock does privacy — a three-layer progressive strategy

The YeBlock privacy protocol isn't a single technology; it's tiered into three levels by customer scenario:

L1 · Basic privacy / (on by default · available at MVP)	· TLS 1.3 end-to-end encryption / · client-side encryption for LoRAs (private-LoRA channel, key stays with the author) / · scheduler only sees routing metadata, not inference content / · nodes clear context immediately after local inference (no persistence of user input) / · target scenarios: default consumer / small-business
↓ Upgrade to	
L2 · TEE channel / (enterprise tier · from 2027-Q2)	· Only Intel TDX / AMD SEV-SNP / NVIDIA H100 CC nodes / · inference code + data + LoRA run end-to-end inside the TEE / · TEE remote attestation proves node trustworthiness / · even the physical node owner cannot see the inference / · target scenarios: medical / legal / financial / government B2B / · 3–10× price

	premium (market has validated enterprises will pay it)
↓ Progressively introduce	
L3 · ZKML sampling / DP / FHE / (2029+ · subject to tech maturity)	· ZKML sampling to verify inference honesty (start with small models) / · DP for training-data protection + call-stats reporting / · FHE for specific small-inference scenarios (private queries / encrypted indexes) / · introduced as industry maturity allows

2.4.5 How YeBlock's privacy protocol is different

Dimension	YeBlock	Bittensor	io.net	OpenAI Enterprise
Default TLS + client-side encryption	✓	✓	✓	✓
Private LoRA channel (encrypted + only authorized nodes execute)	✓ protocol-native	✗	✗	✗ (only OpenAI's own fine-tunes)
Tiered pricing for TEE nodes	☑ (L2 enterprise)	(a few subnets)	✗	☑ (but centralized)
GPU-level TEE (H100 CC) support	✓ on the roadmap	✗	✗	✓

Dimension	YeBlock	Bittensor	io.net	OpenAI Enterprise
Not locked to a single privacy vendor	☑(TDX/SEV/H100, take your pick)	—	—	✗(locked to Azure)
Privacy inside a decentralized network	✓ the only native provider		✗	N/A

Core differentiation: YeBlock is one of very few projects that natively provides tiered privacy on top of a decentralized compute network.

- Other decentralized compute projects (io.net / Akash) barely address the privacy layer.
- Centralized clouds (OpenAI Enterprise / Azure CC) deliver privacy but aren't decentralized.

YeBlock does both — a distinctive position for entering the B2B market.

2.4.6 Engineering reality and trade-offs

✗ Things YeBlock will never do	Reason
✗ Promise ZKML that verifies whole-model inference in the near term	Reality: today's ZK proof generation costs 1,000–10,000× the inference itself; it's essentially unusable in production. YeBlock introduces it progressively from small-model + sampling scenarios.
✗ Promise real-time large-model inference under FHE	Reality: 100–1,000× slower — demo-grade only. YeBlock starts from specific small-inference scenarios and evolves toward larger models, without pitching an LLM-on-FHE fantasy.
✗ Promise fully trustless privacy	Every privacy scheme has a trust anchor (TEE trusts the CPU vendor; TLS trusts the CA; ZK trusts a trusted setup). YeBlock commits to minimizing trust anchors + public disclosure — not to being fully trustless. That said, YeBlock is the closest thing to fully trustless available on today's blockchains.

✗ Things YeBlock will never do	Reason
✗ Use privacy talking points to fool consumers	The privacy protocol primarily serves high-ticket B2B. L1 basic privacy is already enough for consumers; YeBlock will not push consumer users to pay a TEE premium.

Honest summary: YeBlock's privacy stance is "L1 available to everyone by default + L2 pay-as-you-go TEE + L3 following the maturity curve." We don't sell a "full-protocol FHE by default" fantasy, and we won't let "privacy" turn into marketing noise.

2.5 YeBlock Post-Quantum Protocol

2.5.1 What it is

In one line: migrate YeBlock's communication channels, digital signatures, and encrypted assets to NIST post-quantum algorithms (ML-KEM / ML-DSA / SLH-DSA) before quantum computers can break today's RSA / ECC.

The goal: AI assets remain verifiable and settleable across the quantum-threat window. This is not a promise of "quantum-safe today." It is a commitment to keep those assets from being zeroed out by the quantum threat 5–10 years from now.

The cleanest analogy:

- A bank upgrades the locks on its vault — the old locks still work, but everyone knows a certain master key will show up in roughly 10 years.
- Every new asset going in gets the new lock first (forward protection); older assets have their locks changed in batches (backward migration).
- You don't wait until "the day the master key actually appears" — by then it's too late.

2.5.2 Why we need it — the hard gate for 5–10-year AI-asset preservation

Pain point 1: the quantum threat has a defined timeline, and AI-asset lifetimes fall squarely inside it

Time window	Quantum computing maturity (industry consensus)	AI-asset life cycle
2026–2028	Medium-scale quantum computers (thousands of logical qubits); can't break RSA-2048	Today's LoRAs / model weights / on-chain accounts
2028–2032	NIST estimates RSA-2048 may be broken in this window ("Cryptanalytically Relevant Quantum Computer," CRQC)	AI assets created on YeBlock today are still within their life cycle
2032+	Quantum computing matures; PQ standards go mainstream	Assets not migrated by then go to zero

Core takeaway: AI assets launched in 2026 that need to hold their value through 2032 must use post-quantum signatures today. Any project that only "starts thinking about" PQ after 2030 is effectively admitting its assets have a 4–5-year shelf life.

Pain point 2: "Harvest Now, Decrypt Later" attacks are happening today

- Attackers today intercept encrypted traffic, encrypted data, and on-chain signatures → archive them.
- When quantum computers mature (5–10 years out) → batch decrypt.
- This is already happening at the nation-state level (publicly acknowledged by the NSA / NIST / multiple intelligence agencies).
- If YeBlock's AI assets and user inputs aren't encrypted with post-quantum primitives, today's data will be fully decrypted in 2032.

Pain point 3: compliance mandates are already arriving

- **NIST IR 8547 (2024-11): federal agencies must migrate to PQC before 2030.**
- **EU eIDAS 2.0 + national central banks: financial infrastructure PQ migration by 2028–2030.**
- **Enterprise compliance audits: from 2026, SOC 2 / ISO 27001 begin folding PQ readiness into audit scope.**

Any AI infrastructure serving B2B / government / financial customers without a PQ roadmap after 2027 = can't bid.

2.5.3 Post-quantum standards have officially shipped (NIST 2024)

Algorithm	Purpose	Standard number	Publication date	Current deployment status
ML-KEM (formerly Kyber)	Key Encapsulation Mechanism (KEM)	FIPS 203	2024-08	✓ Chrome / Firefox / Cloudflare / AWS have deployed the X25519MLKEM768 hybrid KEM
ML-DSA (formerly Dilithium)	Digital signature	FIPS 204	2024-08	✓ Native support in OpenSSL 3.5 (2025)
SLH-DSA (formerly SPHINCS+)	Hash-based signature (backup)	FIPS 205	2024-08	✓ Used for long-term archival signatures
HQC	Backup KEM	Draft	2025+	Secondary backup
FALCON	Compact signature	Draft	2025+	Resource-constrained devices

Key takeaway: post-quantum crypto isn't science fiction — it's an engineering reality with formally released NIST standards (2024) and rollouts already shipped

by mainstream browsers, CAs, and cloud vendors in 2024–2025. YeBlock applies the released NIST standards to the AI protocol stack.

2.5.4 How YeBlock does post-quantum — a three-phase progressive migration

Phase 1 · MVP default / (from 2026-Q3)	· TLS 1.3 + X25519MLKEM768 hybrid KEM (node ↔ scheduler) / · Hybrid KEM end-to-end for node ↔ IPFS nodes / · Audit logs are fully ML-DSA-signed (10-year post-quantum archival) / · This tier is sufficient to cover "Harvest Now, Decrypt Later"
↓ Upgrade to	
Phase 2 · On-chain PQ options / (2027-Q4 – 2028)	· Account abstraction (AA) + ML-DSA dual-signature support / · PQ signatures on high-value LoRA-asset metadata (ECDSA-compatible) / · PQ signatures on cross-chain bridges (where the target chain supports it) / · Long-term LoRA archives use SLH-DSA (PQ + backdoor-resistant)
↓ Full-stack migration	
Phase 3 · Full-stack PQC / (2030 · following the industry / NIST)	· ECDSA → ML-DSA full migration (5-year compatibility window) / · Re-sign historical on-chain signatures / archival proofs / · TEE attestation uses PQ signatures / · Follow L1 / L2 mainnet PQ upgrades (don't front-run the base layer)

--	--

2.5.5 How YeBlock's post-quantum protocol is different

Dimension	YeBlock	Bittensor	io.net	Filecoin	OpenAI
TLS hybrid KEM	✓ default at MVP	✗	✗	✗	partial
On-chain signature PQ roadmap	✓ live 2027-Q4	✗ no public roadmap	✗ none	under discussion	N/A
PQ metadata signatures on LoRA / AI assets	✓ Phase 2	✗	✗	✗	N/A
Archive-grade PQ signatures on audit logs	✓ from MVP	✗	✗	✗	partial
Application-layer PQ roadmap publicly disclosed	✓ documented commitment	✗	✗	✗	✗

Core differentiation: YeBlock is one of very few AI-protocol-layer projects publishing a PQ roadmap at all.

- Bittensor / io.net / Akash / Render: no public PQ roadmap (2026 sampling).
- Filecoin / mainstream L2s: under discussion but nothing shipped.
- YeBlock aims to be "the first AI-infrastructure project to write a PQ roadmap into a v1 whitepaper."

2.5.6 Engineering reality and trade-offs

✗ Things LIM will never do	Reason
✗ Invent a new PQ algorithm	PQ algorithms are designed by NIST + the global crypto community; YeBlock uses published standards.
✗ Force on-chain PQ before the base chain upgrades	YeBlock leads at the application layer + follows base-chain upgrades.
✗ Turn "post-quantum" into a marketing gimmick	Users don't feel PQ in day-to-day calls — PQ is background infrastructure. We don't cram "post-quantum" into consumer copy, to avoid misleading users.
✗ Do a one-shot full-stack migration	Forcing an immediate full cut breaks ECDSA compatibility and blows up gas costs. Five-year progressive migration + a dual-signature compatibility window.

Honest summary: post-quantum cryptography in YeBlock is infrastructure-level foresight, not a marketing hook.

The team does not promise "quantum-safe today," because current PQ algorithms have not yet accumulated enough real-world validation.

What we do promise: your AI assets on YeBlock will never be zeroed out 5–10 years from now because of the quantum threat.

Part 3: How the Five Pillars Compose, and What Emerges

3.1 Drop one pillar and it degrades: what YeBlock becomes without each

Drop any one of the five and the whole product degrades into something that already exists elsewhere. This isn't wordplay — it's engineering reality:

What you did	Which pillar you skipped	What it degrades into	Real-world analogue
Storage + AI + Privacy + PQ	✗ Decentralized compute	Centralized inference + decentralized assets = "model marketplace + encrypted hosting"	HuggingFace
Compute + AI + Privacy + PQ	✗ Decentralized storage	Centralized model library + decentralized inference = "inference API platform"	Together / Fireworks / Anyscale
Compute + Storage + Privacy + PQ	✗ Decentralized AI (granularity)	General-purpose cloud compute / storage = "general DePIN cloud"	Akash / io.net / Render
Compute + Storage + AI + PQ	✗ Privacy protocol	High-volume consumer platform, permanently locked out of high-ticket B2B	Bittensor (no native privacy layer)
Compute + Storage + AI + Privacy	✗ Post-quantum protocol	Usable short-term; asset value zeroed by the quantum threat 5–10 years out	Nearly every current DePIN project
AI only	✗ All other 4 pillars missing	API aggregation + self-integration = "model router"	OpenRouter

Key insight: the first three pillars (compute / storage / AI) decide whether YeBlock is a usable product; the last two (privacy / PQ) decide whether it can enter B2B markets and whether it will still be here in five years. All five together are what let YeBlock stand up simultaneously on capability, trust, and time.

3.2 The five-pillar reaction: capabilities you can only unlock in combination

Emergent capability	Required pillar combination	Why no single pillar can deliver it
LoRA-author assets that can't be deplatformed	Storage + AI	Storage alone = no demand; AI alone = no persistence layer
LoRA-author assets that earn automatically	AI + Compute + on-chain splits	No compute = no revenue; no AI granularity = can't route earnings to the author
User calls with no platform take-rate	Compute + AI + on-chain splits	Requires P2P three-party consensus
Protocol keeps running even if the team folds	Compute + Storage + on-chain contracts	Centralized compute or storage = a single point of failure
Long-tail contributors can keep monetizing	AI granularity (down to LoRA) + on-chain splits	Coarse granularity buries the long tail
AI cost cut to 1/10 on equivalent tasks	Compute + AI decomposition	Compute alone = no specialized LoRAs; AI alone = no decentralization dividend
Can sign medical / legal / financial customers	Privacy protocol (TEE) + Compute + AI	Without TEE, B2B customers simply can't use it
Private LoRAs made usable	Privacy protocol (client-side encryption) + Storage + AI	No encryption = asset leaks; no decentralized execution = unusable
AI assets that hold up for 5–10 years	PQ protocol + Storage + AI	No PQ = asset signatures broken 5–10 years out
Defense against Harvest Now Decrypt Later	PQ protocol + Privacy protocol	TLS alone = batch-decrypted in the quantum era
Passing B2B + government audits	Privacy + PQ + Compute + Storage + AI	Required by NIST 8547 / NSA CNSA 2.0 / EU eIDAS 2.0

Of the 11 emergent capabilities in the table above:

- The first 6 are unlocked by compute + storage + AI (the original "trinity" boundary).
- The last 5 only unlock once privacy + PQ are added — especially the two biggest commercial opportunities: high-ticket B2B markets and long-term asset preservation.

3.3 The five-pillar protocol stack

User / developer / Agent	
↕ (YeBlock-compatible API)	
① Decentralized AI (granularity layer): LoRA market / matching scheduler / on-chain splits / token billing	
↕	
② Decentralized compute / vLLM/SGLang multi-LoRA pooled inference nodes	③ Decentralized storage / persistent LoRAs + data
↕	
④ Privacy protocol (trust layer · decides who it serves): L1 TLS/client-side encryption · L2 TEE channel · L3 ZKML/DP/FHE	
↕	
⑤ Post-quantum protocol (time layer · decides how long): TLS hybrid KEM · ML-DSA on-chain signatures · SLH-DSA long-term archives	
↕	
Blockchain settlement layer (L2 · Base / Arbitrum, etc.)	

3.4 The five-pillar market coverage map — why five pillars are needed to cover the whole market

	Cost-sensitive	Trust-sensitive
to C	③+②+① / Regular user calls / (compute + storage + AI)	④(L1)+③+②+① / Default-private consumer / (add L1 basic privacy)
to B	④(L2)+③+②+① / Small-to-mid B secure inference / (add TEE)	④(L2)+⑤+③+②+① / B2B + government + finance / (add TEE + PQ)
Low unit price ←————→ High unit price		

Bottom line: drop a pillar, drop a quadrant of market coverage. All five pillars in place = a single protocol that covers to-C, small-to-mid B, high-sensitivity B2B, and government / finance — something no single-pillar project can do.

3.5 Summary: YeBlock is a composition, and a disruptor

YeBlock ≠ replacing any of its predecessors. YeBlock = using a new five-pillar composition to cover a wider set of market quadrants:

Compared to	YeBlock's differentiating advantage
Bittensor	LoRA granularity + IPFS persistence + privacy channel + PQ roadmap
io.net / Akash	LoRA-level AI granularity + on-chain splits + TEE channel + PQ

Compared to	YeBlock's differentiating advantage
Filecoin / IPFS	Add compute + AI granularity + privacy + PQ = a dedicated AI-asset layer
HuggingFace	Decentralized + censorship-resistant + on-chain economic incentives + PQ signatures
OpenRouter	LoRA granularity + on-chain splits + decentralized nodes + B2B privacy
OpenAI Enterprise	Decentralized + censorship-resistant + not locked to Azure + explicit PQ roadmap

Core positioning: YeBlock is an early practitioner of a specific product-grade open-protocol combination — and an early usable implementation of this five-pillar application-layer stack:

- LoRA-granular revenue splits
- IPFS-anchored persistence
- Multi-LoRA pooled compute
- Application-layer tiered privacy
- NIST post-quantum roadmap
- Token-level on-chain settlement

Part 4: How to Participate and How Rewards Work

YeBlock isn't just a playground for technical elites — it's designed as an open ecosystem where anyone can find a role.

Six ways to participate, covering everyone from "totally non-technical" to "professional developer."

⚠ Compliance and honesty statement on the revenue estimates in this section

Before reading the role-specific and revenue detail below, please understand these four principles:

1. **Target design $\neq$ revenue promise:** every revenue figure, split ratio, and network incentive described here is a target design for the YeBlock protocol. Final parameters (including YBT release schedule and split ratios) will be set by DAO governance and are not a revenue guarantee from the team.
2. **The current stage is TCP (Testnet Contribution Proof):** today (yeblock.com), what users accumulate by completing tasks is testnet contribution proof — a record of contributions to testnet stability testing.
3. **Early numbers will run far below the estimates:** every "optimistic" figure below assumes a mature network + sufficient call volume + a stable YBT price. In the first 6–12 months, real revenue may be only 10–30% of the estimate. Please plan finances at the "early" tier.
4. **Region and compliance:** some jurisdictions may impose geo-blocking or additional compliance requirements (KYC / real-name verification / regulatory registration) on node / LP / creator roles.

4.0 Overview: the 6-way quick-reference table

Role	Barrier to entry	Startup cost	Monthly income (conservative)	Monthly income (optimistic)	Difficulty	Who it suits

Role	Barrier to entry	Startup cost	Monthly income (conservative)	Monthly income (optimistic)	Difficulty	Who it suits
① Compute provider	A PC with RTX 4070+	\$0–70 stake	\$70–200	\$400–2,000	★★	Gamers with an idle PC
② Storage provider	8TB+ drives + public IP	\$0–30 stake	\$4–40	\$40–150	★	Users who already own a NAS
③ LoRA creator	Training skills + a few hundred in training costs	\$70–700	\$0–40 (most)	\$400–7,000+ (top LoRAs)	★★★	Developers who fine-tune models
④ LoRA LP	Have YBT, don't have a GPU	\$10+	\$7–4,000	\$70–70,000	★	Capital-first participants
⑤ AI user / developer	One wallet	\$10+	—	—	★	AI application developers
⑥ KOL	A referral link + reach	\$0	\$0–500	\$400–4,000	★★	KOLs

Note: revenue estimates are based on public data + conservative assumptions; actual results depend on network scale, token price, call volume, and other factors. The first 6–12 months may run well below these estimates — please evaluate rationally.

Liquid-economy applications will also open up new roles: idea creators / human nodes / energy providers (see Part 5).

4.1 ① Compute Provider: earn from an idle GPU

Core idea: use the GPU you already own to run AI inference on the YeBlock network, and earn call fees + YBT mining rewards.

4.1.1 Who you are

- You have one or more discrete GPUs (RTX 4070+ / RTX 5070+ / Mac Studio M2 Ultra+ / A100 etc.).
- Your machine sits idle a lot (nights while you sleep, days while you're at work).
- You'd like to convert that idle compute into supplemental income.

4.1.2 Hardware requirements + revenue tiers

GPU	VRAM	LoRAs mountable	Per-card throughput	Monthly income (conservative)	Monthly income (optimistic)
RTX 4070 12GB	12GB	3–5	~50 tok/s	\$40–120	\$200–700
RTX 4080 16GB	16GB	5–10	~80 tok/s	\$70–160	\$300–1,200
RTX 4090 24GB	24GB	10–30	~100 tok/s	\$100–300	\$700–2,000
RTX 5090 32GB	32GB	20–50	~150 tok/s	\$200–400	\$1,500–3,000
A100 40GB / 80GB	40–80GB	50–100	~200 tok/s	\$400–1,200	\$3,000–7,000

Key variables:

- Call volume (low early, high once the network matures).

- Node stability / SLA.
- Popularity of the base model you run.
- YBT token price.

4.1.3 Launch steps (5 steps)

Step 1: Prepare

- └─ A Windows / Linux / macOS machine with an RTX 4070+
- └─ A stable connection (uplink $\geq$ 30Mbps, ideally a public IP)
- └─ A crypto wallet (MetaMask / Phantom / YeWallet, etc.)

Step 2: Register

- └─ Visit lim.yeblock.com (example domain)
- └─ Connect wallet
- └─ Complete node registration (KYC depends on jurisdiction)

Step 3: Install the node client

- └─ Download YeBlock Node (Windows / Linux / macOS)
- └─ One-click install (auto-detects GPU + configures vLLM)
- └─ Pick a base model (lim-8b is a good place to start)

Step 4: Stake YBT

- └─ Hold YBT equivalent to \$200–2,000 USDT
- └─ Lock it in the contract (anti-cheat bond)
- └─ Start the node

Step 5: Accept calls + auto-settle

- └─ The scheduler dispatches jobs automatically
- └─ On inference completion → on-chain auto-settlement
- └─ USDC lands instantly + daily YBT incentives are distributed

4.1.4 Advanced plays

- **Multi-GPU farm:** one host + multiple 4090s = revenue scales linearly.
- **Specialty base models:** run less common bases = less competition + a price premium.
- **TEE nodes (enterprise tier):** upgrade hardware to support confidential computing and pick up high-ticket enterprise orders.

- **Bundle storage:** run a storage node on the same box = double revenue from one hardware footprint.

4.2 ② Storage Provider: earn ongoing storage income from idle drives

Core idea: contribute idle disk space to the YeBlock storage network to store LoRAs and user data, and earn ongoing storage income.

4.2.1 Who you are

- You own a NAS, home server, or idle drives.
- You have a public IP or port forwarding.
- You'd like to convert that idle storage into supplemental income.
- You're OK with low early income and a long-term compounding curve.

4.2.2 Hardware requirements

Configuration	Description
Drives	8TB+ (SATA HDDs recommended; SSDs are best for hot cache)
RAM	4GB+
Network	Uplink $\geq$ 20Mbps; public IP or UPnP support
Uptime	$\geq$ 95% (SLA affects income)
OS	Linux / NAS operating systems

4.2.3 Launch steps

Step 1: Prepare

└─ 8TB+ drives + public IP

- └─ NAS / Linux server
- └─ Crypto wallet

Step 2: Register + download the storage-node client

- └─ Register a storage node on YeBlock
- └─ Download YeBlock Storage Node
- └─ Configure storage paths

Step 3: Stake + launch

- └─ Stake YBT equivalent to \$200–2,000 USDT
- └─ Start the node

4.2.5 Risks and things to watch for

Risk	Mitigation
Early income too low	Treat storage as a side gig on top of a compute node (one machine, two revenue streams)
Drive failure	Use enterprise HDDs (WD Gold / Seagate Exos, etc.) with 5+ year lifetimes
Electricity vs. income	Do the math on NAS power draw + electricity cost — don't run it at a loss
Slashing from data loss	Enable RAID + regular backups

4.2.6 Real-world advice

On its own, storage-provider revenue isn't high, but there are two higher-yield ways to participate:

1. **If you already own a NAS, the marginal cost is close to zero.**
2. **Run a compute node and a storage node together — one machine, two revenue streams.**

4.3 ③ LoRA Creator: turn "domain knowledge" into a tradable asset

Core idea: a LoRA you trained is a digital asset — every time someone uses it, you get paid, the way songwriters and software authors do.

4.3.1 Who you are

- You know how to train LoRAs with tools like PEFT / Unsloth / Axolotl.
- You have domain data (medical / legal / coding / gaming / writing style / minority languages / ...).
- You want to turn your "domain expertise" into a sustainably monetizable asset.

4.3.2 The role with the biggest upside for LoRA creators

Tier	Call volume	Estimated monthly income	Estimated share
Cold	< 10 calls/day	< \$30	60%
Medium	100–1,000 calls/day	\$200–5,000	30%
Hot	1,000–10,000 calls/day	\$700–30,000	8%
Breakout	10,000+ calls/day	\$4,000–200,000+	2%

4.3.3 Real, high-call-volume LoRA directions

B2B scenarios (high ticket, high volume):

- Legal contract review
- Medical record summarization
- Financial statement analysis
- Tender / RFP drafting

- HR resume screening

Consumer scenarios (large user base):

- Fiction writing (classical / contemporary / fantasy)
- Marketing copy generation
- Code generation (specific languages: Rust / Go / Solidity)
- Roleplay (dating simulators / NPC dialogue)
- Translation and polishing (minority languages)

Long-tail niches still wide open:

- Pharmaceutical knowledge
- Legal sub-fields (real estate / divorce / IP)
- Game-specific NPC style
- Historical-figure dialogue
- Academic-paper formatting

4.3.4 Launch steps

Step 1: Data prep

- └─ Collect 500-50,000 domain samples
- └─ Clean + format
- └─ Split train / validation

Step 2: Train

- └─ Pick a base model (start with lim-8b)
- └─ Train on cloud GPUs with Unsloth / Axolotl
- └─ Training cost: low (a few hours on a 4090)
- └─ Validate (compare against the base model)

Step 3: Upload to YeBlock

- └─ Stake YBT equivalent to \$200-2,000 USDT
- └─ Upload weights (10MB-1GB)
- └─ Fill in metadata (base_model / version / use case / license / content)

rating)

- └─ Set your price (add \$0.0001–0.001 per 1k tokens)
- └─ Publish to the market

Step 4: Promote + iterate

- └─ Get community trials
- └─ Collect feedback and iterate v2 / v3
- └─ Promote via KOLs and communities
- └─ Optimize descriptions / tags / examples

4.4 ④ LP (Liquidity Provider): use YBT to back the best assets

Core idea: don't know how to train a LoRA? No problem — vote for the best

LoRAs with YBT and share in their inference revenue.

4.4.1 Who you are

- You hold (or want to buy) YBT but don't train LoRAs.
- You have conviction on specific LoRAs.
- You want passive income / long-term investment.

4.4.2 How LP works

Logic:

When a LoRA author uploads a LoRA, they must stake YBT equivalent to \$200–2,000 USDT.

↓

Anyone can then "provide LP" to that LoRA:

- └─ Stake YBT into that LoRA's pool
- └─ Receive an LP certificate
- └─ Share 5% of that LoRA's call-fee revenue

↓

Shared risk:

- └─ If the LoRA is slashed (content violation, etc.), LPs share the loss pro-rata
- └─ If the LoRA takes off, LPs share the upside pro-rata

Why LPs are needed:

- The author's collateral alone can't cover the credit risk of large-volume scenarios.
- LPs act as "launch liquidity / credit enhancement."
- LPs also help filter for quality LoRAs by voting with real capital.

4.4.3 Revenue structure

15% of every call fee → LP pool

Distributed among LPs pro-rata to their stake

4.4.5 Launch steps

Step 1: Get YBT

- └ KOL referral / network incentives / public sale / DEX purchase
- └ Move to a compatible wallet

Step 2: Browse the LoRA market

- └ Look at daily call-volume trends
- └ Check user ratings and repeat-use rates
- └ Check author reputation
- └ Assess the category outlook

Step 3: Provide LP

- └ Pick 1-20 LoRAs (diversify risk)
- └ Stake YBT into the LP contract
- └ Lock-up: 30-90 days

Step 4: Collect payouts + cycle management

- └ Daily automatic settlement
- └ Rebalance after lock-up expires
- └ Long-term holders can stand for DAO governance

4.4.6 Risks and things to watch for

Risk	Mitigation
------	------------

Risk	Mitigation
LoRA slashing	Pick high-reputation authors + diversify across multiple LoRAs
YBT price drops	Take part of the call fees in USDC (as a floor)
Liquidity locked	Accept the 30–90 day lock-up
Low early income	Early LP pools are small — a bigger share is actually an advantage
Wash-trading LoRAs	The platform filters out fake call volume; LPs can also vote to remove them

4.5 ⑤ AI User / Developer: cut your inference bill by 90%

Core idea: you don't have to "earn" — saving is earning. YeBlock cuts your AI inference cost by 90%.

4.5.1 Who you are

- An AI application developer (chatbot / writing tool / automation).
- Want to use a specialized LoRA.
- Want to protect data privacy (private LoRA + client-side encryption).
- Want to save on cost.

4.5.2 Real-world cost comparison

Scenario (1M tokens)	OpenAI GPT-4o-mini	YeBlock	Savings
Legal contract review	\$0.75	\$0.10	87%
Writing	\$0.75	\$0.07	91%

Scenario (1M tokens)	OpenAI GPT-4o-mini	YeBlock	Savings
Code generation	\$0.75	\$0.08	89%
General customer support	\$0.75	\$0.06	92%
Translation	\$0.75	\$0.07	91%

Scenario (1M tokens)	OpenAI GPT-4o	YeBlock	Savings
High-complexity reasoning	\$7.50	\$0.80	89%
Long-context summarization	\$7.50	\$0.85	89%

4.5.3 Launch steps

Step 1: Register an account + wallet
└─ Connect YeBlock with MetaMask, etc.

Step 2: Top up YBT
└─ (as needed)

Step 3: Browse the LoRA market
└─ Pick LoRAs that fit your product

Step 4: Integrate the SDK
└─ `pip install lim-sdk` (or `npm install`)
└─ Use the YeBlock-compatible interface (5-minute integration)

4.5.4 Code example

```
from lim_sdk import Mesh

client = Mesh(api_key="your_api_key")

resp = client.chat.completions.create(
    model="lim-8b",
```



```
loras=["legal-contract-v1"],
messages=[
    {"role": "user", "content": "Please review this lease agreement..."}
],
temperature=0.7,
)

print(resp.choices[0].message.content)
print(f"Cost: {resp.usage.cost_usdt} USDT")
```

4.5.5 Advanced plays

- **Private LoRA: self-train and self-deploy (pay only for compute + storage, no LoRA-author split).**
- **TEE channel (enterprise tier): confidential inference + encrypted input.**
- **Multi-LoRA composition (mid-term capability): mount multiple LoRAs in a single call.**

4.5.6 Who this fits

- Small and mid-size enterprises deploying AI.
- Cost-sensitive AI application founders.
- Customers who care about data privacy and don't want to send data to OpenAI.
- Anyone who needs a specialized domain (something OpenAI doesn't cover).

4.6 ⑥ KOL: referral rebates + long-term passive income

Core idea: if you have reach (KOL), refer people to YeBlock and permanently share 5% (tentative) of their earnings.

4.6.1 Who you are

- Crypto / AI KOLs.
- Tech bloggers, YouTube / X (Twitter) / TikTok creators, and other social-media creators.

- Tech enthusiasts.

4.6.2 How the referral rebate works

Your referral link → someone registers through your link

↓

That user generates revenue (call fees / mining)

↓

You earn 5% of that revenue (permanently)

↓

Does not reduce the referred user's revenue (the 5% comes from the protocol treasury's growth budget)

Core mechanics:

- Referral rebates are permanent (not one-time).
- They don't reduce the referred user's revenue (the 5% comes from the protocol treasury's growth budget — it's not a take-rate).
- **One-tier direct referral + team bonuses (no multi-level chains).**

4.6.4 Launch steps

Step 1: Register → get your unique referral link + QR code

Step 2: Promote (content matrix)

- └ Write articles / shoot videos introducing YeBlock
- └ Share tutorials
- └ Publish guides like "How to become a compute provider"

Step 3: Track conversion

- └ Watch signups / activations / revenue splits in the dashboard
- └ Follow the activity of each referred user

Step 4: Optimize + compound

- └ Nurture active users
- └ Offer 1:1 support (high-value users)
- └ Improve copy / tutorial quality → lift conversion (single-tier direct rebate only; no multi-level)

4.7 How to pick your role (decision table)

Decide by "what I have"

What I have	Recommended role
A gaming PC with RTX 4070+	Compute provider
A NAS / 8TB+ drives	Storage provider + compute-provider side gig
LoRA training skills + domain data	LoRA creator
\$10k–\$1M USDC, no technical skills	LoRA LP
Building an AI application	AI user (saving is earning)
Large audience (KOL / community owner)	KOL

Decide by "what I want"

My goal	Recommended role
Steady monthly passive income	Compute provider (if you have a GPU) / LP (if you have YBT)
High-risk, high-return (top-LoRA outcome uncertain)	LoRA creator
Save money + use AI tools	AI user
Long-term compounding + governance	LP + referrals
Zero-cost entry	KOL / AI user

Part 5: The YeBlock Liquid Economy

Liquid Trinity — protocol-native flow for ideas, energy, and payments

YeBlock's three protocol-native applications: creative market · energy network · intelligent payments

From infrastructure to economy: the YeBlock Liquid Economy layer

5.1 YeBlock LIME — Liquid Idea Market & Execution

"Ideas don't just get traded — YeBlock actually executes them."

LIME's basic unit is a minimal piece of creativity: encrypted, chain-verified for ownership, and hireable so AI can break it down and execute it.

DeFi made money liquid. YeBlock makes intelligence liquid. YeBlock LIME makes ideas liquid.

① Minting & Ownership · Idea Minting (storage pillar + post-quantum pillar)

The creator writes down an idea. The client encrypts it and uploads it to decentralized storage (IPFS-anchored; cannot be unilaterally removed by any single platform).

The chain registers the CID, summary hash, licensing terms, and pricing/split parameters, then stamps a timestamp with an ML-DSA post-quantum signature.

This produces a Proof of Priority — "this idea belonged to me at this moment" — and a signature that even quantum computers 5–10 years from now can't forge.

That directly solves the first pain point of idea markets: the moment you say it, it gets copied, and you can't prove it was yours.

② Teaser Layer (privacy pillar L1)

All the market sees publicly is the "teaser window": a summary, category, expected deliverable, and price.

The full content stays encrypted; the key stays with the author — reusing the existing "private LoRA channel" mechanism as-is.

③ Hire & Escrow (on-chain settlement layer)

A buyer (an individual sponsor or an on-chain vote) picks an idea and deposits funds into an escrow contract; the author authorizes decryption.

Two modes are available:

- Outright purchase (one-time transfer)
- Licensed execution (author retains ownership and takes perpetual royalties on results — isomorphic to the 15% perpetual LoRA royalty)

④ Machine Execution (compute pillar + AI pillar + privacy pillar L2)

This is what sets LIME apart from every "bounty platform" out there: buyers aren't hiring people — they're hiring AI.

The YeBlock scheduler decomposes the IdeaCapsule into a task pipeline. For each stage, it picks a base model (lim-8b) plus a stack of domain LoRAs — for example, a legal LoRA drafts the contract, a code LoRA writes the implementation — then routes the work to the optimal compute node.

Node selection reuses `routing.py`; normalized LoRA mounting reuses `composition.py`.

High-value ideas run through TEE nodes, where operators physically cannot see the content. That solves the second pain point: the executor stealing your idea.

Reverse hiring (YeBlock hires humans)

When a pipeline stage is beyond AI — or confidence is too low — YeBlock splits it out and posts it back to the LIME task board as a human subtask, with skill requirements, price, and delivery standards.

Typical cases include on-site execution, physical photography, credentialed signatures, local

relationships, and aesthetic sign-off.

Human takers enter the protocol on the same terms as compute nodes:

- Stake to enter
- Deliverables produce on-chain receipts
- YeBlock runs the first-pass quality check (LoRAs auto-audit human deliverables)
- Disputes go through redundant arbitration
- Cheating burns the stake

Humans are wrapped by the protocol as a specialized node — the Human Node: what AI can't do gets routed to a human; what a human delivers gets verified by YeBlock.

The first leg is humans hiring AI. This leg is AI hiring humans in reverse — intelligence and human effort flow both ways on the same pipeline.

⑤ Receipts & Settlement (compute pillar + economic security)

Every execution stage emits a signed receipt (receipts.py).

1–5% of stages are cross-verified via redundant sampling.

Node collateral and slashing make cheating economically negative-EV (economic_security.py).

Human-node deliverables also emit signed receipts. YeBlock runs the first check (LoRA-driven automatic QA); sampled arbitration cross-reviews them — the same verification and slashing rails as compute receipts.

Once every stage passes verification, the escrow contract settles atomically.

⑥ Idea Royalty Waterfall (AI pillar + on-chain revenue splits)

Every time the deliverable — a product, a LoRA, or a report — generates future revenue, the waterfall runs automatically. Each of the following takes a share:

- Idea author
- LoRAs used in execution
- Compute nodes
- Human executors
- Storage nodes
- LPs
- Protocol
- Referrers

This adds two seats — Idea Author and Human Executor — to the existing seven-way split table. The Human Executor share is carved out of the execution fee in proportion to their subtask workload (same source as the compute-node share), leaving the other ratios structurally unchanged. The lineage-splitting logic in `royalty_waterfall.py` applies unchanged.

⑦ Idea Lineage (storage pillar + post-quantum pillar)

Others can build derivative ideas (forks) on top of yours. The chain records the parent–child lineage, and downstream revenue flows back up the waterfall to upstream ancestors. Ideas form an inheritable family tree, like music copyrights. SLH-DSA archival signatures keep the lineage verifiable for decades.

The pain points of bounty platforms:

Without compute, all you can do is post ideas and wait for humans to bid = a regular bounty platform (Upwork).

Without storage, ideas and deliverables can be deplatformed = a centralized creative community.

Without AI-level granularity, you can't hire a "specialist skill stack," and execution quality is uncontrollable.

Without privacy, an idea is stolen the moment it's published, and the market can't exist at all (LIME's hardest dependency).

Without post-quantum, priority-proof signatures can be forged 5–10 years later, invalidating ownership.

Idea Author — the lowest-barrier role (no GPUs, no drives, no model-training skills required — only good ideas).

5.2 YeBlock LEM — Liquid Energy Mesh

DeFi made money liquid. YeBlock makes intelligence liquid. YeBlock LEM makes energy liquid.

Intelligence is the best exit for energy.

Electricity doesn't travel well and doesn't store well. As intelligence (inference results), it can reach anywhere on Earth in an instant and hold its value forever.

Electricity dominates the AI inference cost stack, so the most natural form of "selling energy" inside YeBlock is refining electricity into intelligence on-site — then selling the intelligence.

In the era of fusion power, when the marginal cost of electricity approaches zero, the cost of

intelligence approaches zero as well. Whoever has the energy sits at the center of the compute network.

Core design principle: electricity can't be transmitted on-chain, so LEM doesn't build a virtual grid or touch physical transmission. All energy-provider monetization runs through three "convert on site" paths:

Three participation paths:

Path A · Energy-to-Compute (energy + GPUs)

Run compute/storage nodes directly — but the protocol prices your energy advantage in.

The scheduler's routing weights add an energy-cost dimension (routing.py gains an extra factor). At equal latency and reputation, lower-electricity-cost nodes get task priority. Batch workloads — LoRA training, offline inference, cold-data archiving — are routed to the lowest-cost nodes.

A user at \$0.03/kWh and a user at \$0.15/kWh, running the same card, differ by exactly the energy spread.

The network becomes an energy arbitrageur: it takes locally unsellable cheap electricity and turns it into intelligence buyers everywhere can pay for.

Path B · Energy Hosting (energy without GPUs)

This path matches "energy but no cards" with "cards but no energy."

GPU owners deploy hardware into an energy provider's site — residential solar, small hydro, a mining farm, or a future fusion park.

The on-chain contract automatically splits that node's compute revenue into two seats: hardware sponsor and energy operator. The ratio is agreed by both parties on-book and settled by the escrow contract — no trust required between them.

It's mining-machine hosting turned into a protocol: energy holders can enter with zero hardware investment.

Path C · Energy Credits (JouleCredit — "power futures" sold to the network)

Energy providers use metering oracles (smart meters + TEE-signed uploads) to mint verified

supply into on-chain JouleCredit certificates.

Node operators buy JouleCredit to offset electricity cost. B2B customers buy the green-power version to meet ESG / carbon-compliance requirements — green-power inference is already a real plus in EU enterprise procurement.

Surplus power — midday solar peaks, curtailed wind and hydro, future fusion base-load excess — is posted as discounted credits. The scheduler automatically shifts elastic workloads into those windows.

The network becomes a sponge for the grid: it absorbs surplus energy and converts it into intelligence.

Protocol mechanics and pillar mapping

Metering and attestation (PoEC, Proof of Energy Contribution): smart-meter data is signed on-chain by TEE nodes (privacy pillar — metering details encrypted, only totals public); ML-DSA signatures ensure long-term certificate validity (post-quantum pillar — green-power credits and carbon records need 10-year-scale audit validity).

Energy-aware scheduling: routing.py adds three weights for price / carbon intensity / time slot (compute pillar); metering history and credit archives live on decentralized storage (storage pillar); a predictive LoRA handles load-vs-price matching (AI pillar).

Economic security: false supply reporting = staked slashing; redundant meters cross-verify; same "cheating doesn't pay" logic used in economic_security.py.

Revenue splits: Path B splits are performed inside the 30% compute-node seat (hardware / energy), leaving the seven-way split table untouched; Path C is an independent certificate market and doesn't occupy the split table.

Roles and narrative

From a household with rooftop solar, to a small hydro plant owner, to a future fusion operator — participation barriers can go all the way down to zero hardware (Paths B/C). For a household solar array, the electricity you can sell to the grid vs. the electricity you can "refine into intelligence" and sell to the world can differ by 10×. A small hydro plant on the brink of shutdown can lift revenue 10×.

Vision: after commercial fusion, the marginal cost of energy goes to zero, and the pricing power of intelligence shifts from "who has the model" to "who has the energy + who has the protocol." LEM puts YeBlock at that intersection ahead of time.

5.3 YeBlock LIP — Liquid Intelligence Pay

DeFi made money liquid. LIM made intelligence liquid. LIME makes ideas liquid, LEM makes energy liquid, LIP makes the intelligence economy settle in real time.

YeBlock LIP is a payment rail designed for machines — Visa serves the human economy; LIP serves the Agent economy.

Core logic: why the AI economy needs its own payment rail.

Human payments are "few, large." A few tens of thousands of TPS at peak covers all human card swipes.

The AI economy is "massive, tiny" — structurally the opposite.

Each inference bills per token at ~\$0.0001 per transaction. A \$0.30 credit-card interchange fee loses 3,000× on every single settlement on legacy rails.

One Agent kicks off tens of thousands of calls a day, each backed by a 7–9-way split, so transaction counts explode multiplicatively.

Machine-to-machine trades need millisecond finality, programmable conditions, and unattended authorization — none of which bank rails can provide.

YBT's 1M TPS and ultra-low fees sit exactly where only this kind of rail can carry a machine economy. This isn't going after Visa's human retail market — it's opening a market where no incumbent exists.

Three product layers (from inside out):

Layer 1 · The network's own settlement engine (cold-start demand, live today)

YeBlock is YeBlock LIP's first customer. The waterfall payout on every inference call is high-frequency micropayment by construction. Recipients include compute nodes, LoRA authors,

idea authors, human nodes, storage, LPs, protocol, and referrers.

At 1M TPS, token-level settlement upgrades from "batch-and-settle" to real-time streaming payouts — LoRA authors watch their balance tick up by the second.

LIME's AI-hires-humans payments and LEM's JouleCredit trades all run on the same rail.

Payment demand doesn't need external validation: every call on the network creates transactions.

Layer 2 · Agent Wallets (accounts held by machines)

Humans open a "pocket-money account" for their AI Agent: account abstraction plus policy constraints — per-transaction and daily limits, payee whitelists, purpose rules, revocable at any time — with private keys held in a TEE (Trusted Execution Environment).

Within its authorized envelope, the Agent pays autonomously: calling APIs, buying data, renting compute, hiring other Agents, hiring human nodes.

Streaming Pay bills continuously by token, second, or kWh — pay as you use, stop paying when you stop. Inference, storage leases, and energy credits are all naturally streaming assets.

Layer 3 · A2A settlement network (opened up as an industry standard)

Inside a LIME task pipeline, the lead Agent farms subtasks out to a legal Agent, a code Agent, or a human node — and they settle among themselves in milliseconds. That's A2A (Agent-to-Agent) clearing.

Open this layer to outside developers: any Agent application can plug LIP in for machine-to-machine payments. YeBlock grows from "the split-settlement system we use ourselves" into "the clearinghouse for the Agent economy."

LIP's unique differentiation: payment and execution are natively bound.

Other payment rails — including centralized options like Stripe's Agent Pay — can only prove "the money moved."

On LIP, every payment is anchored to a signed execution receipt (receipts.py). The payment record is simultaneously a cryptographic proof that you actually received that inference or that deliverable.

Receipt = invoice = settlement. Payment and delivery are atomically bound at the protocol layer — something no external payment rail can do today.

Pillar mapping:

- Compute: execution receipts serve directly as settlement instructions; redundant sampling prevents fake invoices.
- Storage: invoices and audit archives are encrypted and archived to decentralized storage.
- AI: risk-control LoRAs detect abnormal payment patterns in real time (hijacked Agents, wash-trading arbitrage) and auto-reconcile.
- Privacy: confidential payments — amounts and counterparties are encrypted from third parties (enterprise invoices are trade secrets); Agent private keys live in a TEE.
- Post-quantum: fund-transfer signatures are the single thing most vulnerable to HNDL attacks; accounts and payment channels sign with ML-DSA — a natural extension of the "5–10-year AI-asset preservation" promise to the money layer.

Value capture: fees are paid in YBT and partially burned — the higher the volume, the greater the burn — so network usage is directly tied to YBT's value (usage-driven, not subsidized by inflation). Specific fee rates and burn ratios are marked "subject to DAO governance."

Part 6: FAQ + Risk Disclosures

6.1 Key variables:

- Network scale (the first 6–12 months will be far below these numbers).
- YBT price (a portion of income is paid in YBT and moves with the token).
- The time and effort you put in.
- Luck (breakout LoRAs are unpredictable).

6.2 Is it really safe?

The safeguards YeBlock provides:

- Smart-contract audits (Certik / OpenZeppelin, etc.).
- Multi-sig + timelocks (major changes have a 14-day delay).
- Stake + slashing to deter cheating.
- Client-side encryption + TLS 1.3 + ML-KEM hybrid KEM.
- DAO governance + transparent revenue splits.

But these risks remain:

- Smart-contract vulnerabilities (audits reduce, don't eliminate) — we'll patch quickly when issues surface.
- YBT token price volatility.
- Regulatory uncertainty.
- Lost wallet private keys = permanent asset loss.

6.3 Team anti-exit (learning from the Ethereum Foundation token-sale saga)

YeBlock's anti-exit design:

- Team allocation of 13%: 3-year vesting + 1-year cliff.
- Institutional allocation of 12%: 3-year vesting + 1-year cliff.
- Once team YBT unlocks, unlocked YBT is re-staked. Team operations run on staking yield.

If it is not re-staked, DAO governance may vote to burn a portion of the foundation's team YBT — the exact ratio decided by DAO vote.

If team YBT does need to be sold to keep operations running, the ratio and timing are decided by DAO vote.

- DAO governance goes live 2027-Q2 (team power transfers to the community).
- Treasury is managed by multi-sig + subject to public audit.
- Users hold their own wallet keys — the platform cannot move funds.

YeBlock Token: 121,000,000

The total supply of YBT is fixed at 121,000,000 tokens.

More than half is distributed to miners, node operators, and other network contributors — incentivizing real compute, storage, and referral contributions.

40% of net protocol ecosystem revenue is used to buy YBT on the open market and permanently burn it, reinforcing the deflationary mechanic over time.

Allocation	Amount (10k tokens)	Amount (tokens)	Share
Distributed referral mining	3,000	30,000,000	24.79%
VC (Venture Capital)	1,452	14,520,000	12.00%
YeBlock Foundation	1,573	15,730,000	13.00%
TGE (Token Generation Event)	2,100	21,000,000	17.36%
Node	1,400	14,000,000	11.57%
Staking mining	1,200	12,000,000	9.92%
Contribution mining	900	9,000,000	7.44%

Ecosystem partnerships / grants / later-stage growth	475	4,750,000	3.93%
Total (hard cap)	12,100	121,000,000	100.00%

6.4 Regulatory risk?

YeBlock's compliance design:

- Open protocol + globally distributed (no dependence on a single jurisdiction).
- Enforced license + content ratings + three-layer content filtering.
- User KYC is handled by the front end / integration partner.
- The protocol layer exposes compliance-metadata interfaces.

Strong recommendations:

- Users in different regions should consult local counsel.
- For large positions (> \$10,000 USDT equivalent), consult compliance advisors.
- Creators uploading LoRAs must ensure training data is properly licensed.

6.5 What if the project fails?

The worst-case losses per role:

Role	Maximum loss
Compute provider	Stake + time (hardware is unaffected)

Role	Maximum loss
Storage provider	Stake + time (drives are unaffected)
LoRA creator	Stake + training cost + time
LP	Staked YBT (per contract rules)
AI user	USDT balance in the wallet (control the balance to reduce it)
KOL	Time (no monetary loss)

Ways to reduce risk:

- Don't go all-in — keep single-role exposure within what you can afford to lose.
- Top up your wallet as needed; don't leave large balances sitting.
- Spread across multiple roles to reduce single-point risk.

6.6 Can I participate if I'm not technical?

Yes! Ordered by technical bar, from lowest to highest:

Technical bar	Role
★Zero technical	LP / KOL / AI user (basic calls)
★★Low (can follow a tutorial)	Compute provider / Storage provider
★★★Medium (knows Python)	AI user (advanced development)
★★★★High	LoRA creator

YeBlock is designed so everyone can take part — you can contribute capital, hardware, technical skills, or community reach, with different roles for different contribution styles.

6.7 Are earnings paid in USDC or in the token?

Both:

- **YBT / USDC (your choice):** call fees, node splits, LP splits.
- **YBT:** network incentives, ecosystem rewards, governance rights (subject to price volatility).

Suggested split:

- Keep the USDC portion for short-term living expenses.
- If you're long-term bullish on the YeBlock ecosystem, hold the YBT portion.

6.8 If the AI industry cuts prices sharply (OpenAI drops prices), is YeBlock still competitive?

Yes:

- OpenAI's price cuts are an industry-wide trend; YeBlock pricing will move down with it (the open-source model ecosystem follows).
- **But YeBlock's decentralization + censorship resistance + data privacy + fine-grained LoRA are things OpenAI can never offer.**
- Long-term, specialized-domain LoRAs are exactly what OpenAI's general model can't fill in (medical / legal / minority languages / etc.).

6.9 YeBlock vs Bittensor — competition or collaboration?

Short-term: differentiated competition (no head-on clash)

- Bittensor: subnet-level / compute-first / complex economy / no native privacy layer / no PQ roadmap.

- **YeBlock: LoRA-level / five-pillar / simple and transparent / L1–L3 tiered privacy / clear NIST PQ path.**

Long-term: possible collaboration (complementary)

- **YeBlock's LoRA market can plug into Bittensor's compute subnets.**
- **Bittensor subnets can use YeBlock's storage + revenue-split layer.**

6.10 Four key challenges (publicly disclosed) — the ones other projects avoid saying out loud

Almost no Web3 + AI intro document in the industry includes this section — but publicly disclosing risks builds credibility. We're publishing them because professional readers will see through what we leave out in 5 minutes; better to publish than to be caught.

Challenge 1: Compute verification can't be 100% collusion-proof

The problem: decentralized compute nodes may fake inference results (return fabricated tokens without actually running the model). This has been the root problem of the decentralized-compute category for 7 years.

How we address it:

- **Tiered verification: redundant sampling of 1–5% of tasks + mandatory 3-replica cross-verification for high-value tasks (enterprise customers, paying users).**
- **Node reputation system: long-term cheat rate + SLA + historical call volume combine into a score that influences scheduling priority and collateral release.**
- **Stake + slashing: cheating nodes have their bond cut (redirected to the defrauded user + protocol treasury).**

- **TEE channel (enterprise tier):** high-sensitivity workloads upgrade to TEEs for a trusted execution environment (trading some performance for stronger trust).

In plain terms: the stack above doesn't solve this 100%, but it drives the expected value of cheating negative. This is an industry-level problem; we don't expect a 100% solution soon — the protocol's target is "cheating is economically negative," not "cheating is technically impossible."

Challenge 2: First-token latency limits real-time use cases

The problem: geographically distributed nodes have 200–2,000ms first-token latency — a bad fit for latency-sensitive scenarios like real-time chat or voice assistants. This is exactly why io.net / Akash have never grown into AI inference at scale.

How we address it:

- **MVP focus on asynchronous workloads:** batch inference, background Agents, content generation, document analysis, scheduled tasks — 1–2 seconds is fine here.
- **Geo-aware scheduling:** the scheduler prefers nodes geographically close to the request source (in-region first-token typically < 500ms).
- **Hot cache + prewarm:** popular LoRAs are prewarmed on multiple nodes to reduce cold-starts.
- **Real-time falls back to centralized by default:** YeBlock integrators automatically fall back to centralized cloud when the user explicitly picks "low-latency" mode (transparent, opt-in).

In plain terms: YeBlock will not compete head-on with OpenAI's real-time conversational scenarios in the near term. The current focus is asynchronous inference + specialized-LoRA calls — a small but real market.

Challenge 3: The cold-start supply–demand loop

The problem: no users → not enough providers → no compute → no users. The classic two-sided market problem.

How we address it:

- The asymmetric starting point we already have: yeblock.com's launch week saw estimated daily new signups of 10–30k, DAU ~30k, and 7-day retention ~80% — the project starts the cold-start engine from the user side first, and then pulls providers in.
- YBT incentives cover both sides: users earn YBT by completing tasks; providers earn YBT by running nodes — early-participant rewards on both sides.
- The team runs seed nodes to guarantee baseline supply: in MVP, the team runs 5–10 seed nodes and backstops call volume, so early providers always have work.
- A verification milestone we commit to: within 18 months (by 2027-Q4), real call fees must make up $\geq 50\%$ of node total revenue. If we miss it, we treat the business model as unproven — we scale down proactively, redesign the economic model, and don't prop things up with inflation.

In plain terms: this is one of the hardest milestones we face. Helium / Filecoin / io.net haven't truly cleared this bar either.

Challenge 4: Multi-region compliance risk

The problem: DePIN + AI + token is a triple compliance minefield. The US SEC, the EU AI Act, and Singapore's MAS each draw their own red lines.

How we address it:

- Offshore token entity: the token-issuing entity is established in a Web3-friendly jurisdiction.

- **Mandatory KYC for nodes / LPs:** at mainnet, nodes and LP roles are KYC-mandatory + sanctions-list blocked.
- **Geo-blocking:** specific jurisdictions are dynamically blocked (IP / wallet) based on legal advice, including the OFAC sanctions list.
- **Non-engagement in certain regions:** until compliance is clear, YeBlock mainnet is not actively opened to users in unfriendly regions (KYC-free today, but no region-specific token incentives are issued for those regions).
- **Content compliance:** LoRA uploads require license + content rating + three-layer content filtering (automatic + community + DAO arbitration).

In plain terms: compliance will always run behind regulation. We have four contingency plans, but we can't guarantee that a regulatory tightening won't force us to pause a region — a shared risk with every Web3 project.

6.11 KYC / AML progressive roadmap

To balance a low early participation barrier with mid-to-long-term compliance, YeBlock uses a three-phase KYC strategy:

Phase	Timing	KYC requirement	Trigger conditions / limits
Phase 1 • Startup phase	Current	No-KYC registration	YBT accumulation only (not immediately liquid); per-account YBT cap, per-IP registration cap, behavioral risk controls
Phase 2 • Pre-TGE	2026-Q3	Light KYC (email + facial verification + sanctions-list screening)	Triggered when YBT accumulation exceeds a threshold (e.g. 10k); YBT can't be activated or transferred until KYC is completed
Phase	From	Strong KYC (ID + proof of	Mandatory for nodes / LPs / LoRA creators; regular

Phase	Timing	KYC requirement	Trigger conditions / limits
3 · Post-mainnet	2027-Q2	residence + AML screening)	consumer users can still use light KYC; receiving on-chain splits above \$X/month triggers additional review

Additional principles:

- Sanctions lists (OFAC SDN list + UN / EU / UK sanctions lists) are blocked in real time throughout.
- User privacy data does not go on-chain; it's encrypted and held by a compliant KYC vendor (Persona / Sumsub or similar are on the roadmap).
- KYC data is retained for 5 years, and users can request deletion (per applicable data-protection law).

Part 7: Team and Organization

7.0 What we're putting in

YeBlock is built by a 50-person full-time team, distributed across the US and Singapore, remote-first.

7.1 Team structure (14 functions)

The table below shows the current functional blueprint for YeBlock. Actual headcount per function shifts with project stage; we do not disclose precise

headcount in this document (to avoid sending misleading signals to the market or potential investors).

Engineering (~65% of the team):

Function	Main output
Protocol / smart contracts	Split contracts, staking contracts, governance contracts, L1 deployment
Inference engine / node scheduling	vLLM / SGLang customization, multi-LoRA pooling, matching scheduler
Storage integration	LoRA persistence, cache prewarm, CID indexing
Client SDKs	Python / JS SDK, OpenAI-compatible layer, 5-minute integration experience
Website + dashboard	yeblock.com user entry point, node-operator dashboard, LoRA marketplace front end
DevOps / node operations	Seed nodes, monitoring, SLA reporting automation
Internal security / audit team	Internal security review + coordinating external audits
Cryptography / privacy engineering (with external advisors)	TLS 1.3 hybrid KEM deployment, TEE node onboarding (TDX/SEV/H100 CC), application-layer ML-KEM/ML-DSA wrappers, private-LoRA encrypted channel, early ZKML/DP pilots

Non-engineering (~35% of the team):

Function	Main output
Protocol research / token economics	Economic modeling, incentive curves, game-theory analysis, DAO governance design
Compliance / legal (with external counsel)	Multi-region compliance, token architecture, content policy, KYC/AML roadmap
Public-chain engineering	Consensus, network layer, ledger + data structures, VM + smart contracts,

Function	Main output
	PQ cryptography, cross-chain protocols, privacy protocol
Community / content / KOL	Community operations, Discord / Telegram / X, etc.
Growth / data / user research	Website analytics, funnel optimization, retention research
Trust & safety / content moderation	LoRA content review, appeals handling, copyright review
Brand / design / documentation	Visual system, documentation (this whitepaper is one output), educational content
HR / finance / admin	Cross-border team ops, comp, compliance registration

The team deliberately keeps a roughly "engineering $\approx$ non-engineering" balance (about 65 : 35). A DePIN + AI + token project built with engineering alone — no compliance, no economic modeling, no community operations — ends up technically correct but unable to go the distance.

On cryptography / privacy engineering: this is our specifically reinforced direction for 2026.

Once YeBlock became a five-pillar system, promoting the privacy protocol and post-quantum protocol to first-class pillars meant we needed a dedicated team for:

- Interfacing with TEE vendors (Intel / AMD / NVIDIA)
- Tracking the evolution of NIST PQC standards
- Working closely with external cryptography advisors

YeBlock has to apply these standards to the YeBlock protocol stack correctly,

safely, and compliantly — one of the reasons we're one of the few DePIN projects willing to write a PQ roadmap into a v1 document.

7.2 Founding team

YeBlock was founded by a seasoned Web3 team:

- **Background:** core members come from top-tier Web3 projects (kept confidential until 10 days before mainnet).
- **Why we're being confidential:** prior projects involve ongoing conflicts of interest / non-compete obligations / regulatory quiet periods — stealth-founder mode is a standard compliance practice in Web3 (Berachain early on, various Anonymous projects, etc.).
- **Unveiling commitment:**
 1. The full personal identity of the founder(s) will be publicly disclosed 10 days before YeBlock mainnet launch.

7.3 Funding sources

Dimension	Status
Current funding	Founding team self-funded + a top-tier Web3 VC fund
Source notes	Personal capital accumulated by founding members over prior Web3 cycles + a top-tier Web3 VC fund
Current burn rate	Self-funded cash flow + top-tier Web3 VC fund support runs the team to the 2030-Q4 planned milestone
External funding	A top-tier Web3 VC fund
Future	Based on yeblock.com traction, a strategic raise is planned for 2026-Q3 (pre-seed + seed /

Dimension	Status
funding plans	strategic rounds). We keep meaningful early equity and token allocations to incentivize the later community; we won't over-dilute chasing valuation.

7.4 Team operating principles

- **Remote-first + async collaboration:** spanning US East and Singapore time zones.
- **Core decisions:** founder-led in the MVP phase.
- **Progressive decentralization:** after 2027-Q2, governance decisions are gradually handed off to the DAO.
- **Hiring:** continuously open, especially welcoming senior talent in compliance / protocol research / cryptography (TEE / PQC / ZKML) / security audit (contacts in the closing section).

Part 8: Roadmap and Timeline

8.1 Four time horizons

Short-term	Hard commitments, verifiable milestones
Mid-term	Trigger-conditional projects
Long-term	Strategic vision + external dependencies
Horizon	Directional principles

8.2 Next 4 quarters (what participants can do)

Every quarter's "definition of done" is a publicly verifiable metric — the community can use this table to hold the project accountable in real time.

Quarter	Team target	Definition of done (publicly verifiable KPIs)	What participants can do
2026-Q3	Foundational build: 5 seed nodes + a lim-8b single-LoRA closed loop	① GitHub protocol code open-sourced (Apache 2.0) ② A public end-to-end LoRA-call demo video for 1 LoRA ③ Node list + call records queryable on-chain ④ 1 technical blog post detailing the architecture ⑤ Privacy L1 on by default (TLS 1.3 + client-side-encrypted LoRA channel + non-persistent context) ⑥ PQ Phase 1 on by default (node ↔ scheduler TLS X25519MLKEM768 hybrid KEM)	Apply to join seed-node testnet (whitelist airdrop); use yeblock.com to complete tasks and accumulate YBT
2026-Q4	Developer ecosystem kick-off	① LoRA marketplace publicly accessible ② Weekly calls ≥ 7,000 / nodes ≥ 10 ③ First monthly SLA report ④ Full audit-log ML-DSA signing goes live (10-year PQ archives)	LoRA creators claim early-participation benefits; KOLs accumulate YBT via referral links
2027-Q1	Testnet launch	Privacy protocol L2 TEE-channel public test (Intel TDX or AMD SEV-SNP, either one, ≥ 3 nodes)	Compute providers launch at scale; compute + storage nodes onboard
2027-Q2	Expansion and enterprise entry point	① Monthly calls ≥ 1M ② TEE nodes ≥ 5 (including an NVIDIA H100 CC roadmap demo) ③ Design Partners ≥ 10 (including ≥ 2 B2B customers using the TEE channel) ④ Enterprise SDK publicly available (with TEE attestation API) ⑤ DAO governance phase 1 goes live (core-parameter voting) ⑥ PQ Phase 2 kick-off: account abstraction + ML-DSA dual-signature support / PQ signatures on high-value LoRA metadata	All roles fully onboard; enterprise customers integrate; DAO governance participation
2027-Q4	Five-pillar completeness	① Monthly calls ≥ 5M ② TEE nodes ≥ 30 ③ B2B customers ≥ 10 ④ Real USDC call fees ≥ 50% of node total revenue ⑤ Privacy L3 pilots: ZKML small-model sampling-verification demo / DP training-data protection pilot ⑥ PQ protocol	Advanced plays: upgrade to TEE nodes to take on B2B high-ticket orders

Quarter	Team target	Definition of done (publicly verifiable KPIs)	What participants can do
		upgrade	

8.2 A key commitment: 18-month real-demand bar

We commit to a hard metric anyone can hold us to:

By 2027-Q4 (18 months after launch), real USDC call fees on the YeBlock network must make up $\geq 50\%$ of node total revenue (including YBT mining).

If we miss it:

1. We publicly acknowledge the business model isn't proven.
2. We proactively pause YBT emissions + redesign the economic model.
3. **We do not prop up a false sense of prosperity with inflation.**
4. Locked node / LP collateral unlocks on schedule; we do not freeze user assets.

This line is the hard metric that separates "a DePIN driven by real demand" from "a DePIN propped up by inflation." The team writes it into the document, and we invite the community to hold us to it in 2027-Q4.

8.3 Mid-term direction

Items that require a trigger condition to start:

Item	Trigger condition
Multi-LoRA composition	Industrial-grade DARE / TIES merge libraries
Lightweight semantic router	Daily calls $\geq 100k + 50$ LoRA categories
Image LoRA market	Stable image-LoRA toolchain

Item	Trigger condition
Content moderation moves to DAO	Centralized moderation cost > 10% of protocol fees
Video LoRA market	1,000+ public video-LoRA works on HF
Full introduction of privacy L3 (ZKML / DP / FHE)	ZK proof generation cost < 100× inference / FHE performance gap < 10×
PQ Phase 3 full-stack migration kicks off	OpenSSL / mainstream L2 chains have clear PQ upgrade paths + NIST IR 8547 mandatory windows confirmed

8.4 Long-term vision

Dependent on external technology maturing:

- Cross-base CU composition (industrial-grade S-LoRA / Punica).
- Expert CU decomposition (open-source MoE standard).
- ZKML sampling verification (privacy protocol evolution).
- TEE upgraded to primary trust surface (privacy protocol evolution).
- Cross-chain interoperability.
- CU derivatives market.

8.5 Horizon direction

Directional principles, awaiting industry / regulatory evolution:

- Full-stack PQC (post-quantum protocol evolution, following NIST 2030).
- User-level AI Agents on YeBlock.
- Cross-network pipeline inference (RDMA over WAN).
- Unknown forms (kept open for what comes next).

Closing and invitation to build together

Core belief

DeFi made money liquid. YeBlock makes intelligence liquid.

Ten years ago, banks monopolized money; DeFi gave it back to everyone.

Today, AI compute, models, and economic incentives are monopolized by centralized giants; YeBlock is here to give them back to everyone.

It's a simple but deep proposition:

- An idle GPU should be able to earn by contributing compute.
- Training a LoRA should pay.
- Using AI shouldn't mean handing a single company a 60% gross margin.
- Model weights should belong, permanently, to their authors and users.

Public commitments

Three public commitments:

1. **Honesty: no vaporware — every number, timeline, and performance claim is grounded in real engineering.**
2. **Open source: the core protocol and SDKs are fully open source (Apache 2.0).**

3. **Co-building: DAO governance goes live in 2027-Q2; the team ultimately hands power to the community.**

Invitation to build

If you believe in a future where compute, storage, and AI are all decentralized, join us in the early stage in whichever role fits you best:

- **Got an idle GPU: become a genesis compute provider.**
- **Got a NAS / drive: become a genesis storage provider.**
- **Train models: become a genesis LoRA creator (genesis-period benefit window).**
- **Have capital: become a genesis LP.**
- **Have reach: become a genesis KOL.**
- **AI developer: become an early user (whitelisted discount).**

Contact

Channel	Status
Web	https://yeblock.com (live; sign up to participate in YBT tasks)
Official website	https://yeblock.com
GitHub	https://github.com/yeblocklim/YeBlock
Twitter / X	https://x.com/YeBlockLIM
Telegram	https://t.me/yeblock
Business email	ye@yeblock.com
Early hiring	Especially welcome: senior talent in compliance / protocol research / security audit / inference engines (specific email coming soon)

Channels listed as "opening soon" are not yet live; any third-party account or group claiming to represent YeBlock before official launch is not official. Rely on yeblock.com in-site announcements as the sole source of identity verification. Please stay vigilant against scams.

YeBlock — intelligence that flows like water.

A full-stack, five-pillar (compute + storage + AI + privacy + post-quantum) decentralized-AI-infrastructure practitioner.

Join us — starting with early-participation rights.

Five-pillar full edition — including how to participate and earn, candid early-stage disclosure, and the privacy / post-quantum protocols. Updated: 2026-05. Licensed under CC-BY-SA 4.0 — you're welcome to share, translate, and adapt (with attribution).
