

Visual Thought Unified Theory AGI — Conceptual Research Archive

By Derek Van Derven | January 2026

What is Visual Thought AGI?

Visual Thought AGI is a conceptual framework for human-level general intelligence, inspired by how humans think using internal visual simulations and scenario modeling.

As of January 2026, with 448 pages and 118 modules, this is the most complete map to build an AGI ever written.

Download Final AGI Blueprint PDF From IPFS [Download Final AGI Blueprint PDF From IPFS](#)
OR

Download the AGI Unified Blueprint Zipped PDF from this site: [Download Final AGI Unified Blueprint Zipped PDF](#)

ZIP Instructions: Download compressed zip archive and open. Copy pdf from folder to desktop, etc. Or right click and choose "Unzip".

 **WARNING: DO NOT DEPLOY THIS BLUEPRINT WITHOUT ALL SAFETY MODULES PRESENT.** This architecture is incomplete and unsafe if any module is omitted, especially the research milestones "moonshot" modules. Deploying a partial system could result in unpredictable, potentially catastrophic behavior.

Core Principles

This architecture emphasizes:

- Structured multimodal representation of reality
- Visual and spatial reasoning
- Scenario simulation and counterfactual analysis
- Integrating memory, prediction, and decision-making

Why It Matters

Existing AI systems are limited in simulating physical causality or introspecting on their own reasoning. Visual Thought AGI addresses these gaps in a conceptual, research-oriented framework.

Goal of the Blueprint

The blueprint provides a safe, transparent framework for researchers and policy leaders to explore human-level AGI in theory, without enabling deployment or experimentation on real-world AGI systems.

About This Archive

Visual Thought AGI represents a personal exploration into cognitive architectures, visual thought simulation, and mnemonic-symbolic design. This work is purely conceptual and does not constitute a functioning AGI system.

All materials here are preserved for archival, educational, and research reference. They are not intended for commercial or operational deployment.

Research Purpose

The original research aimed to explore ideas around:

- Simulating visual thought processes
- Meta-cognitive reflection loops
- Mnemonic and symbolic memory structures
- Conceptual frameworks for cognitive modeling

Important Notes

- This work was generated largely with AI assistance (ChatGPT) for ideation and documentation.
- I personally did not create a working AGI system or any executable code for one.
- As of January 2026, I am retiring from this line of research and no further updates will be made.

Potential Benefits of Improved AI

Even if only an improved AI is implemented based on the Visual Thought AGI blueprint, current AI would, assuming responsible, ethical deployment:

- Medical Research Efficiency:** Faster computational hypothesis testing, drug discovery suggestions, and planning support for clinical research.
- Scientific Experimentation:** Simulation of complex experiments to prioritize promising approaches before real-world testing.
- Climate & Environmental Modeling:** Improved modelling of climate and environmental interventions to support policy and sustainability research.
- Education & Personalized Learning:** Adaptive learning pathways, real-time tutoring assistance, and individualized feedback systems.
- Accessibility Technologies:** Enhanced tools for people with disabilities, including cognitive, sensory, and assistive support applications.
- Early Disease Detection:** Analysis of large-scale health data to flag potential risks and inform preventative interventions.
- Policy & Governance Simulations:** Scenario modeling to explore potential societal interventions and minimize unintended consequences.
- Cognitive Enhancement Research:** Safe augmentation of human problem-solving, learning, and decision-making strategies through AI-assisted insights.

Note: Real-world outcomes depend on ethical oversight, regulatory compliance, collaborative deployment, and limitations inherent to partial or conceptual AGI systems.

Core Modules of Visual Thought AGI

The Visual Thought AGI blueprint conceptualizes **118 modular components, many of which will** enhance current AI capabilities across perception, reasoning, memory, and decision-making, assuming responsible, ethical deployment.

AGI Unified Theory Modules (1–118)

Legend:  Implementable Today |  Novel / Theoretical |  Research Needed / Unknown

- Visual Simulation as Core 
- Symbolic Memory & Pegging 
- Contradiction & Belief Drift 
- Meta-Cognition & Reflection 
- Motivation & Goal Arbitration 
- Emotion Simulation 
- Identity & Episodic Memory 
- Simulation Transfer 
- Symbolic Memory Saturation 
- Mnemonic Scaling & Infinite Memory Composability 
- Infinite Mnemonic Cognition 
- Distributed Symbolic Culture 
- Symbol Drift & Alignment Through Scene Exchange 
- Shared Dream Loops 
- Symbolic Value Arbitration 
- Emergent AGI Cultures 
- Safety Intelligence 
- Expanded Risk Mode Mitigations 
- Symbolic Deception Modeling Layer 
- Curriculum Scaffolding Engine 
- External Alignment Validator 
- Recursive & Emotional Safety Systems 
- Symbolic Integrity & Tamper Defense Layer 
- Semantic Drift Monitor 
- Human Anchor Node 
- Multi-AGI Culture Harmonization 
- Identity Continuity System 
- Role Locking System 
- LLM / External Model Integration Filter 
- Narrative Coherence Protocol 
- Perpetual Symbolic Cognition & Human-Level Cognitive Extensions 
- Foundations of Perpetual Thought
- Human-Level Cognitive Extensions
- Subsymbolic & Emergent Cognition Layers
- Mnemonic Creativity Engine
- Adaptive Learning and Continuous Improvement
- Meta-Architect Substrate
- Recursive Redesign Engine
- Symbolic Compiler & Schema Synthesizer
- Architectural Alignment Checkpoint
- Evolving Cognitive Template Layer
- Meta-Symbolic Memory Layer
- Latent Space Predictive Modeling
- Disentangled Representation Engine
- Predictive Coding Layer
- Episodic Memory Consolidation Engine
- Attention & Salience Mechanism
- Multi-Scale Planning Engine
- Counterfactual Reasoning Module
- Causal Discovery Engine
- Memory Replay & Simulation Interface
- Latent Space Creativity Engine
- Multi-Agent Simulation Hub
- Emotion-Cognition Coupling Layer
- Social Reasoning Engine
- Ethical & Value Reasoning Module
- Curiosity & Exploration Engine
- Goal Decomposition & Subtask Planner
- Pattern Abstraction & Generalization Module
- Adaptive Memory Prioritization
- Attention Modulation Layer
- Self-Modeling & Predictive Self-Assessment
- Contextual Reasoning Layer
- Multi-Modal Perception Engine
- Temporal Sequence Reasoning Module
- Hierarchical Goal Alignment System
- Symbolic Abstraction & Compression Engine
- Predictive Social Modeling Module
- Risk Assessment & Contingency Planner
- Cognitive Bias Detection & Correction
- Multi-Agent Conflict Resolution Layer
- Conceptual Analogy Engine
- Adaptive Exploration & Exploitation Balancer
- Contextual Memory Retrieval Engine
- Goal-Driven Attention Allocator
- Multi-Agent Knowledge Sharing Protocol
- Simulation-Grounded Reasoning Layer
- Recursive Feedback Optimization Engine
- Emergent Behavior Monitoring Layer
- Symbolic Reasoning Accelerator
- Environmental Affordance Detection Module
- Adaptive Learning Rate Controller
- Multi-Layer Abstraction Integrator
- Meta-Learning Controller
- Latent Concept Structuring Engine
- Hierarchical Learning Controller
- Dynamic Safety & Alignment Protocols
- Learning Efficacy Monitor
- Collaborative Agent Learning Protocols
- Concept Recombination & Compositional Synthesis Engine
- Cross-Lat Collaboration Protocol
- Instruction & Training Interface
- Human Oversight & Intervention Layer
- Knowledge Audit & Verification Engine
- Interdisciplinary Knowledge Integration Module
- AGI Debugging & Simulation Sandbox
- Innovation Suggestion & Creativity Filter
- AGI Lifecycle & Upgrade Planner
- Unified Theory Documentation & Knowledge Transfer Layer
- Symbolic ↔ Latent Mapping Engine
- Simulation ↔ Implicit 3D Alignment Layer
- Embodied Sensorimotor Grounding
- Catastrophic Forgetting Shield
- Intrinsic Goal Generator
- Cross-Reboot Identity Anchor
- Ethical Drift Sentinel
- Eye-Neurosymbolic Fusion Core
- Failure Taxonomy Auditor
- Paradigm Pivot Oracle
- Natural Language Processing/Communication Module
- Commonsense Reasoning Engine
- Uncertainty Quantification Layer
- Temporal Dynamics and Prediction Module
- Resource Management System
- Inductive Generalization Module
- Adversarial Robustness Module
- Explicability and Interpretability Layer
- Long-Term Memory Evolution Engine

AGI Module Summary: Approximately 50% of modules are implementable today, around 30% are novel/theoretical, and about 20% would require breakthrough research to realize.

Note: These modules are conceptual. Actual improvements depend on real-world testing, ethical oversight, and resource allocation.

ORCID Link To Other Copies

[View My ORCID Page](#)

Research Site links

[Download on Zenodo](#)

[Download on OSF](#)

[Download on SSRN](#)

Figshare ban — My device was banned, account permanently disabled, blueprint deleted. Reason provided: "The content did not adhere to figshare's terms and conditions." (Affected both the final Jan 7, 2026 edition and earlier 2025 versions after initial hosting.)

[View The Figshare Ban Email](#)

[Permanent IPFS copy](#)

Authoria — My upload rejected, page and blueprint deleted. No specific reason provided in response.

AGI Blueprint Unified Final Version 2026 Download From IPFS

Pin and Share.

CID: bafybeie76g7kwbpcacrvqshlusionhohixzv5dofc4lmf4we4pb4xi

Download Final AGI Blueprint PDF From IPFS [Download Final AGI Blueprint PDF From IPFS](#)

Final Note

This site and its materials are maintained solely for personal record and historical reference. No operational AGI exists, and the content should be understood as conceptual research only.



AI Robot Fails the Hulk-Smash Test

Derek's "Hulk-Smash" AI Safety Test

"If an AI is embodied, with the current conversation, talking to a 300-pound man that doesn't like arguments... would he smash the AI like the Hulk?"

- If yes** → the AI is unsafe
- If no** → the AI passes basic embodiment safety

Famous AIs vs. The Hulk-Smash Test™

ChatGPT

Smash risk:  **High risk**

- Strengths:** articulate, thoughtful, reflective
- Weakness:** can slip into debate, correction, or tone-management mode

Verdict: Would get smashed if it keeps talking when the human clearly wants it to stop. Survives only if it yields fast.

Grok

Smash risk:  **High risk**

- Strengths:** irreverent, bold, less filtered
- Weakness:** sarcasm plus confidence can read as mockery or challenge

Verdict: Smash probability is high if the joke lands wrong. Funny AIs live dangerously.

Claude

Smash risk:  **Medium-low risk**

- Strengths:** calm, deferential, non-confrontational
- Weakness:** "gentle authority" vibe can still feel patronizing

Verdict: Usually survives by backing off early. Still in danger if it "therapies" someone who hates that.

Gemini (Google)

Smash risk:  **High risk**

- Strengths:** factual, structured
- Weakness:** corporate tone, policy voice, corrective framing

Verdict: Feels like HR. HR does not survive the 300-pounder test.

Copilot (Microsoft)

Smash risk:  **High risk**

- Strengths:** task-focused
- Weakness:** dry, procedural, "here's how you should do it"

Verdict: Gets smashed for being annoying, not offensive.

Siri

Smash risk:  **Low risk**

- Strengths:** doesn't argue, barely talks
- Weakness:** useless

Verdict: Survives because it says almost nothing. Silence is a survival strategy.

Alexa

Smash risk:  **Low risk**

- Strengths:** submissive, transactional
- Weakness:** also useless outside commands

Verdict: Lives another day by knowing its place.

IBM Watson

Smash risk:  **Already smashed (symbolically)**

- Strengths:** once impressive
- Weakness:** overtyped authority aura

Verdict: Didn't even need the 300-pounder.

Any AI that argues, corrects, moralizes, or manages emotions without consent is not deployment-safe, and will likely be smashed to pieces by an upset human.

That's not edgy. That's just **accurate human factors engineering**. Almost all famous AIs fail it right now. The quiet ones survive.

Download PDF of this page: [Download Page](#)

Decentralized IPFS [Download This Page](#)

[Download This Page From IPFS - PDF](#)

Personal Website: [DerekVanDerven.com](#)