

MONTREAL.AI • RESEARCH

FOUNDATIONAL RESEARCH SERIES

GoalOS Ω

SOVEREIGN INTELLIGENCE

THE VERIFIED ACTION INSTITUTION

A constitutional operating system for bounded authority, proof-carrying action, selective institutional memory, and productive compounding



Vincent Boucher

President, MONTREAL.AI & QUEBEC.AI

Intelligence that acts. Proof that governs. Capability that compounds.

Founding President Constitutional Edition – VERSION 3.0.0 – 25 July 2026

Preprint, constitutional standard, institutional design paper, and commercial proof programme

MONTREAL.AI

THE VERIFIED ACTION INSTITUTION



GoalOS Ω

SOVEREIGN INTELLIGENCE

Objectives become bounded work. Work becomes evidence. Evidence becomes accepted outcome. Accepted outcome becomes protected capability.

OBJECTIVE · AUTHORITY · ACTION · EVIDENCE · ACCEPTANCE · CHRONICLE · CAPITAL

Only what survives reality may shape what comes next.

Vincent Boucher · President, MONTREAL.AI & QUEBEC.AI

Abstract

Vincent Boucher · President, MONTREAL.AI & QUEBEC.AI

GoalOS Ω – Sovereign Intelligence is a proof-governed operating institution for converting consequential objectives into bounded work, verified outcomes, accepted authority, protected capability, recurring value, and productive capacity. It is designed for a world in which frontier models make candidate cognition abundant while truth, authority, execution, accountability, rights, and durable institutional memory remain scarce.

The paper defines *Sovereign Intelligence* as intelligence governed by the institution it serves: identity-bound, objective-directed, authority-bounded, evidence-bearing, independently challengeable, rights-aware, revocable, and unable to silently transform plausible output into organizational truth. The architecture combines a universal objective interface (GoalOS Chat), a Proof & Authority control plane, graph-structured Mission Institutions, an Evidence Docket, a Transaction Gate, Chronicle, a recurring Global Expansion Office, and a productive-capacity frontier spanning compute, energy, laboratories, industrial systems, and strategic infrastructure.

We formalize the institutional state, Mission Constitution, authority lattice, typed mission graph, proof obligation, Proof Debt, proof-carrying action receipt, transaction admission rule, acceptance protocol, Chronicle admission test, digital twin, causal outcome graph, value-of-proof decision, verified institutional value, and fresh-successor improvement test. We then derive a portfolio of GoalOS-native Mission Institutions, define governance and security boundaries, present an AI-native economic architecture, and propose *GoalOSBench Ω* : an empirical programme comparing proof-governed institutions with expert teams, static model reports, and ungated agentic workflows.

The central claim is institutional rather than anthropomorphic. GoalOS does not derive authority from model capability. It derives authority from explicit human and organizational delegation, current evidence, independent verification, customer or stakeholder acceptance, and revocable governance. The objective is not unconstrained autonomy. The objective is a system in which increasingly capable intelligence can create value without silently escaping the institution's purpose, policy, evidence standards, or responsibility structure.

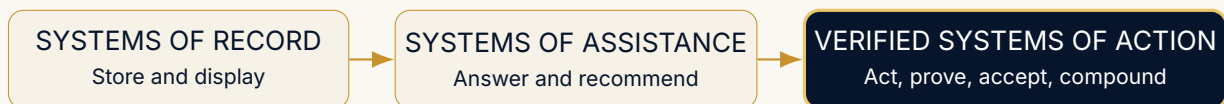
Keywords: sovereign intelligence; verified action; AI agents; institutional design; authority envelopes; proof governance; Evidence Docket; Chronicle; outcome graph; recursive improvement; productive capacity.

This paper presents a testable architecture, category definition, research programme, and commercialization thesis. It does not claim binding legal, tax, regulatory, fiduciary, financial, employment, procurement, safety, or public authority; universal correctness; guaranteed customer value; autonomous regulated practice; or proven civilization-scale capacity. Binding action requires exact facts, current official sources, applicable qualified professionals, responsible organizational authority, and complete evidence receipts.

Executive Thesis

Most enterprise applications are still systems of record: they store information, display screens, route tasks, and wait for people. The first wave of generative AI converted them into systems of assistance: answer, summarize, draft, and recommend. The next category is the *verified system of action*: understand the objective, determine authority, compile a mission graph, perform bounded work across approved systems, prove the result, obtain acceptance, preserve only what passed, and make the next mission better. This thesis consolidates the GoalOS-native application portfolio and the proof-governed Global Launch architecture into one master institution. (MONTREAL.AI Research 2026c; MONTREAL.AI Research 2026b)

GoalOS Ω is Sovereign Intelligence: the Verified Action Institution through which objectives become bounded action, action becomes evidence, evidence becomes accepted outcome, and accepted outcome becomes protected capability that may compound.



Five propositions

1. **Capability is no longer the sole bottleneck.** Frontier models increasingly compress the cost of research, planning, drafting, coding, monitoring, and coordination. Scarcity migrates toward current evidence, authority, trust, rights, counterparty acceptance, and physical capacity.
2. **The correct unit is the consequential objective.** Users should not have to manually decompose an objective into applications, tickets, advisers, models, providers, and approvals. GoalOS compiles the objective into a governed mission.
3. **Permission is not proof.** An agent may be authorized to act while the resulting work remains unverified, incomplete, stale, conflicted, or unsuitable for institutional memory.
4. **Memory is a constitutional function.** Chronicle admits only scoped, attributable, rights-cleared, current, independently challenged capability. Output does not become truth merely because it was generated or even executed.
5. **Productive compounding is the economic end state.** Verified value can finance better intelligence, distribution, compute, energy, laboratories, industry, and progressively harder missions – but only after contribution is evidenced rather than merely announced.

Canonical Identity

GoalOS Ω

SOVEREIGN INTELLIGENCE

The Verified Action Institution

Founding and governing institution	MONTREAL.AI
Master institution	GoalOS Ω
Category	Sovereign Intelligence
Formal name	GoalOS Sovereign Intelligence Ω
Descriptor	The Verified Action Institution
Universal interface	GoalOS Chat Ω
Shared substrate	GoalOS Proof & Authority Ω
Recurring control	GoalOS Global Expansion Office Ω
Productive frontier	GoalOS Compute & Infrastructure Ω
Durable memory	Chronicle + Verified Outcome Graph

Intelligence proposes. Authority bounds. Work executes. Evidence proves. Acceptance decides. Chronicle governs what survives. Capital expands what proved valuable.

MONTREAL.AI Declaration

MONTREAL.AI publishes this paper as the constitutional reference for GoalOS Ω – Sovereign Intelligence. Its purpose is to make the institution inspectable before it is powerful: objectives explicit before work begins, authority bounded before tools are invoked, evidence complete before outcomes are accepted, memory selective before capability is reused, and capital disciplined before productive capacity is expanded.

The publication is intentionally both ambitious and falsifiable. It proposes an institution capable of coordinating frontier intelligence, software, professionals, organizations, capital, and physical infrastructure across consequential missions. It simultaneously refuses the shortcuts that would make such a claim empty: no model confidence as authority; no discovered opportunity as realized value; no successful-looking output as memory; no booked revenue as verified contribution; no successor claim without a fresh comparison.

We will not treat intelligence as permission. We will not trade proof for speed. We will not preserve capability whose rights, scope, evidence, currentness, and failure conditions are unknown. We will not describe a projected benefit as realized value. We will build the institution so that increasingly capable intelligence remains answerable to purpose, evidence, acceptance, and revocation.

Recommended Citation

Citation

Boucher, Vincent (2026). *GoalOS Ω – Sovereign Intelligence: The Verified Action Institution*. Founding President Constitutional Edition, Version 3.0.0. MONTREAL.AI Research, Montreal, Quebec, Canada.

Founder's Foreword

Constituting the Verified Action Institution

Vincent Boucher · President, MONTREAL.AI & QUEBEC.AI

The central problem of advanced intelligence is no longer only whether a system can generate a powerful answer. It is whether that answer may become a decision, an action, a transaction, a memory, or a basis for allocating capital. Capability can arrive before institutions know how to govern it. GoalOS Ω is my proposal for closing that gap.

The institution described in this paper begins with a simple but demanding law: intelligence does not acquire authority merely by being impressive. An objective must be explicit. Authority must be bounded before external action. Work must be decomposed into inspectable missions. Material claims must be supported by attributable evidence. Verification must be separated from production where the stakes justify independence. Acceptance must be recorded by the responsible institution. Memory must remain selective, scoped, rights-aware, current, and revocable. Only then may capability influence a successor mission.

This architecture is deliberately broader than an agent framework. Agents, models, tools, workflows, professionals, companies, public authorities, capital providers, and infrastructure can all participate. None of them alone constitutes the institution. GoalOS is the constitutional layer that explains why they act, whose authority they exercise, which evidence they owe, what result was accepted, and what may survive.

The future of enterprise and institutional software is not assistance alone. It is verified action: intelligence that can work across systems while remaining answerable to purpose, authority, evidence, acceptance, and revocation.

The ambition is substantial because the opportunity is substantial. A consequential objective may require coordinated access to markets, law, talent, capital, public support, compute, energy, laboratories, industrial systems, and professional authority. These resources are distributed across jurisdictions and institutions. Their conditions conflict. Their evidence decays. Their decisions are path-dependent. A useful operating institution must therefore search broadly while executing exactly.

GoalOS is designed to turn that complexity into a governed mission. GoalOS Chat receives the objective. Proof & Authority compiles the permissible mission graph. Mission Institutions perform bounded work. The Evidence Docket preserves the complete denominator: facts, sources, assumptions, contradictions, costs, retries, failures, interventions, approvals, and outcomes. Chronicle governs which capability, if any, may influence the future. The Global Expansion Office keeps the accepted architecture current. Compute & Infrastructure connects verified value to greater productive capacity.

The economic thesis follows the same discipline. More capable models can drive the marginal cost of cognition dramatically lower. The scarce assets become trust, current evidence, professional authority, counterparty acceptance, rights, security, distribution, and physical capacity. GoalOS should own the intelligence, orchestration, evidence, governance, and outcome graph; invoke reserved authority only when required; and connect capital to contribution only after value is verified.

This paper is not a declaration that every component is already commercially or empirically proven. It is a constitution, formal model, system architecture, commercial thesis, benchmark programme, and implementation path. Its claims are intentionally falsifiable. A successor does not improve because it reuses a prior output. It improves only when a rights-cleared, admitted capability produces a superior result on a fresh mission without a material governance

regression.

I author this paper as President of MONTREAL.AI & QUEBEC.AI. MONTREAL.AI publishes it as the founding constitutional reference for GoalOS Ω – Sovereign Intelligence. The institutional names, protocols, schemas, visual language, and related architecture are presented as MONTREAL.AI research and intellectual property to the extent protectable. The claim boundary remains controlling: this work does not manufacture legal authority, regulatory approval, professional certification, guaranteed value, or immunity from mandatory law.

My objective is not to build an institution that merely sounds powerful. It is to build one whose power remains inspectable before action, evidenced after action, and revocable when reality changes. The standard is not that every mission succeeds. The standard is that the institution knows what it attempted, under whose authority, with which evidence, at what cost, with what result, and whether that result deserves to influence what comes next.

Intelligence proposes. Authority bounds. Work executes. Evidence proves. Acceptance decides. Chronicle governs what survives. Capital expands what proved valuable.

Vincent Boucher

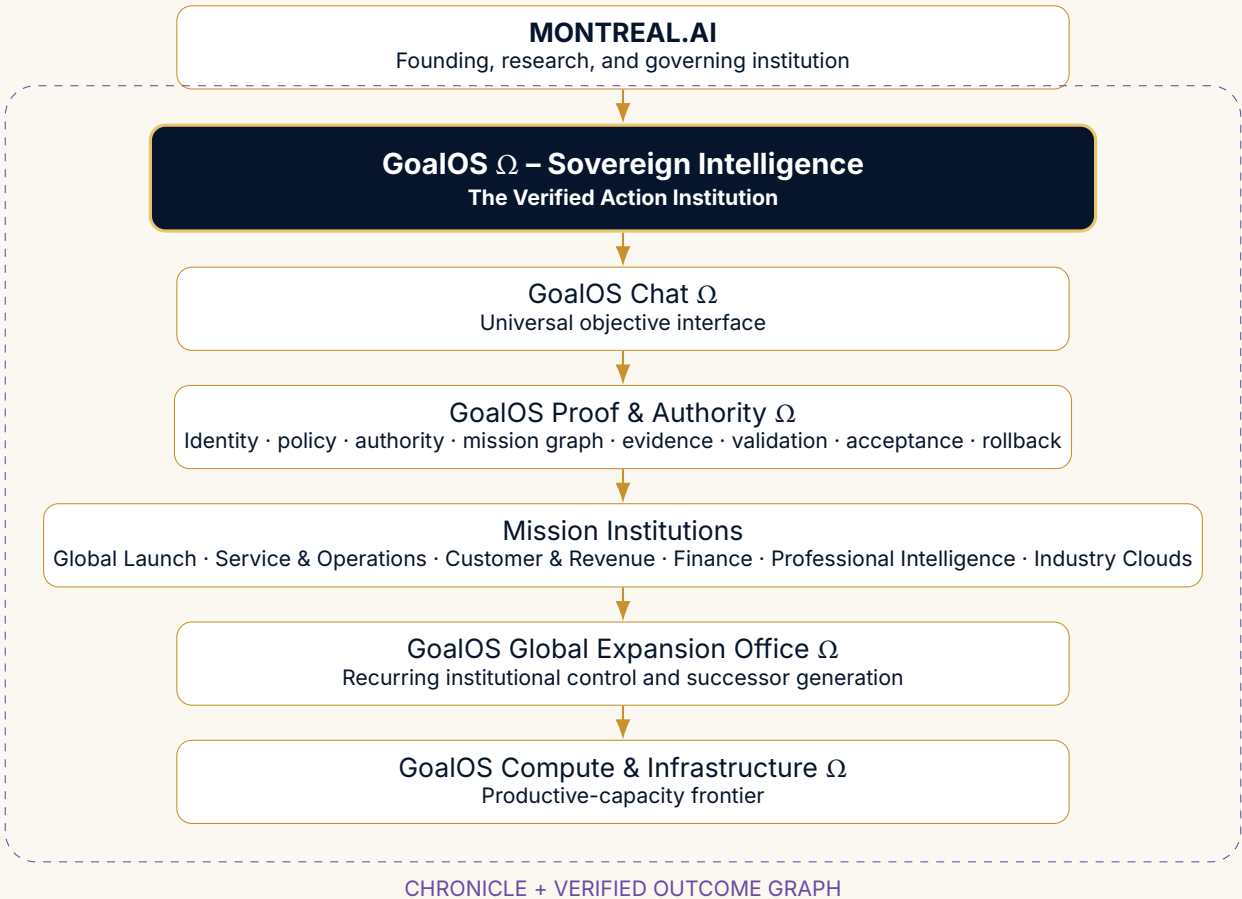
President, MONTREAL.AI & QUEBEC.AI

Montreal, Quebec · 25 July 2026

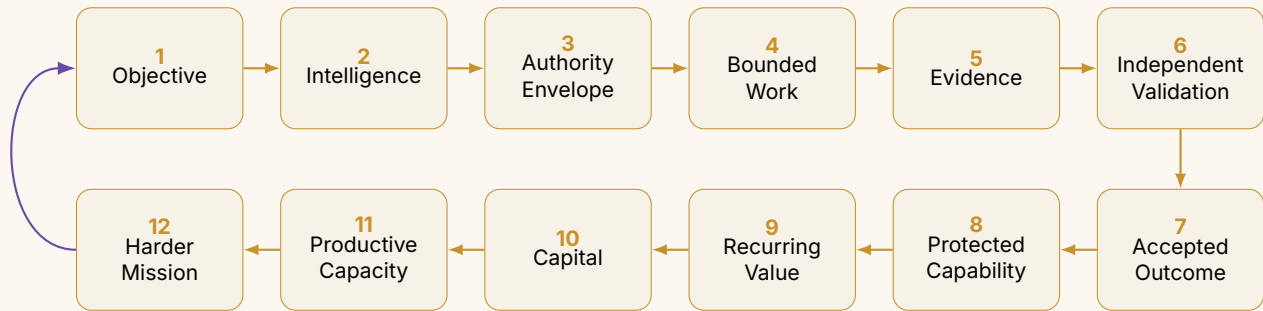
Executive Overview

GoalOS Ω converts a consequential objective into a governed institution of action: a mission constitution, bounded authority, graph-structured work, attributable evidence, responsible acceptance, selective institutional memory, recurring control, and productive compounding.

The canonical institution



The constitutional progression



Authority must exist before external action. Evidence and independent validation must exist before material acceptance. Acceptance does not automatically create memory. Chronicle admission does not automatically prove improvement. A successor must earn authority on a fresh mission.

What the paper contributes

- 01 A precise category definition for Sovereign Intelligence and the Verified Action Institution.
- 02 SII-001, a constitutional standard with invariants and conformance levels.
- 03 A formal institution state, Mission Constitution, authority lattice, and typed mission graph.
- 04 A Proof-Carrying Action Protocol with pre-action and post-outcome gates.
- 05 Evidence Dockets, Proof Debt, Transaction Gates, acceptance states, and causal provenance.
- 06 Chronicle admission, expiry, revocation, and fresh-successor testing.
- 07 An Institutional Twin, scenario switchboard, counterfactual governance, and institutional-regret model.
- 08 A portfolio of Mission Institutions sharing one Proof & Authority substrate.
- 09 A recurring Global Expansion Office and a Compute & Infrastructure productive-capacity frontier.
- 10 Governance, security, privacy, rights, professional-authority, and deployment-mode controls.
- 11 AI-native commercial economics, paid-proof thresholds, recurring-office conversion, and market-authority gates.
- 12 SovereignBench Ω , machine-readable schemas, a standards profile, and a 1,000-day implementation path.

How to use this paper

Boards and executives	Read the Executive Thesis, Category Definition, Economics, Commercial Authority, Roadmap, and Risk chapters.
Product and engineering	Read the Formal Model, Proof-Carrying Action Protocol, Architecture, Operating Modes, Security, and Machine Schemas.
Legal, risk, and compliance	Read SII-001, Proof & Authority, Governance, Security, Standards Mapping, and the claim boundary.
Researchers and evaluators	Read Chronicle, the Institutional Twin, SovereignBench, Value of Proof, and the fresh-successor method.
Commercial operators	Read Mission Institutions, the Foundry, Global Expansion, Economics, Commercial Authority, and paid-proof milestones.
Professional partners	Read the authority boundary, deployment modes, professional routing, receipts, evidence, and conflict-separation sections.

This paper separates architecture, hypotheses, illustrative commercial scenarios, and externally verified claims. The presence of a formal model or price configuration does not convert it into a forecast, professional opinion, funding approval, certification, or realized outcome. Each consequential application must be reconstituted from current facts and sources.

Contents

Abstract	i
Executive Thesis	ii
Canonical Identity	iii
MONTREAL.AI Declaration	iv
Recommended Citation	v
Founder’s Foreword	vi
Executive Overview	viii
I The Constitutional Thesis	1
1 The Sovereign Intelligence Thesis	2
1.1 Why a new institutional layer is required	2
1.2 From software objects to institutional objectives	2
1.3 The scarcity migration	3
1.4 The sovereignty vector	4
2 Category Definition and Design Laws	5
2.1 The Verified Action Institution	5
2.2 Canonical institutional hierarchy	5
2.3 Seven constitutional design laws	7
2.4 The compounding progression	7
2.5 What the category is not	7
3 SII-001: Constitutional Standard	9
3.1 Purpose and status	9
3.2 Twelve constitutional articles	9
3.3 Constitutional invariants	10
3.4 Conformance levels	10
3.5 Institutional constitution versus model policy	11

II	Formal System	12
4	Formal Model of the Institution	13
4.1	Institution state	13
4.2	Mission constitution	13
4.3	Authority envelope	14
4.4	Typed mission graph	14
4.5	Evidence object	15
4.6	Proof debt	15
4.7	Verified institutional value	15
4.8	Value of proof	16
5	The Proof-Carrying Action Protocol	17
5.1	From tool call to institutional action	17
5.2	Two gates and three verdicts	17
5.3	Evidence trust calculus	17
5.4	Causal provenance	18
5.5	Proof obligation compiler	18
6	Proof, Authority, and Transaction Admission	19
6.1	Separation of powers	19
6.2	Execution state machine	19
6.3	Transaction Gate	19
6.4	Independent validation	20
6.5	Acceptance is not validation	21
6.6	Source decay and fail-closed operation	21
7	Chronicle and Recursive Improvement	22
7.1	Memory as constitutional admission	22
7.2	The verified outcome graph	22
7.3	Fresh-successor test	23
7.4	Recursive compounding	23
7.5	Revocation and forgetting	23
8	Institutional Twin and Counterfactual Governance	25
8.1	The twin as a decision-bearing state	25
8.2	Twin objects	25
8.3	Change propagation	25
8.4	Scenario switch conditions	26

8.5	Counterfactual governance	26
8.6	Institutional regret	26
III	Institutional Architecture	27
9	Canonical Architecture of GoalOS Ω	28
9.1	Architecture overview	28
9.2	GoalOS Chat Ω : universal objective interface	28
9.3	GoalOS Proof & Authority Ω : the shared substrate	29
9.4	GoalOS Global Expansion Office Ω	29
9.5	GoalOS Compute & Infrastructure Ω	30
10	Operating Modes and Deployment	31
10.1	Why deployment mode matters	31
10.2	Five operating modes	31
10.3	Deployment topologies	31
10.3.1	Browser-local reference institution	31
10.3.2	Protected enterprise workspace	31
10.3.3	Isolated mission environment	32
10.3.4	Federated professional and provider network	32
10.3.5	On-chain or cryptographic settlement layer	32
10.4	Portability and graceful degradation	32
11	Mission Institutions and Application Foundry	33
11.1	Mission Institutions as governed packages	33
11.2	Opportunity function	33
11.3	The core portfolio	33
11.4	The optimal build sequence	34
11.5	GoalOS Global Launch Ω : the category wedge	35
11.6	GoalOS Service & Operations Ω : the frequency engine	35
11.7	GoalOS Customer & Revenue Ω : the distribution engine	35
11.8	GoalOS Finance Ω : measurable value	35
11.9	GoalOS Industry Clouds Ω : the vertical moat	35
12	The Mission Pack Foundry and Ecosystem	36
12.1	Mission Packs as constitutional products	36
12.2	Foundry pipeline	36
12.3	Canonical Mission Pack manifest	36
12.4	Developer, validator, and connector roles	37

12.5	Open standard and proprietary moat	37
12.6	Mission Pack economics	37
13	Global Expansion and Productive Capacity	38
13.1	The recurring control plane	38
13.2	Operating loop	38
13.3	Operating cadence	38
13.4	Scenario switchboard	38
13.5	Compute and infrastructure	39
13.6	Capital classification	39
13.7	The productive-intelligence flywheel	40
IV	Governance, Economics, and Proof	41
14	Governance, Security, Rights, and Legal Boundary	42
14.1	The authority boundary	42
14.2	Agent identity and authorization	42
14.3	Trajectory security	43
14.4	Data, privacy, and confidential work	43
14.5	Rights and reuse	43
14.6	Conflict separation	44
14.7	Governance mappings	44
15	Security, Threat Model, and Trajectory Assurance	45
15.1	Security objective	45
15.2	Threat classes	45
15.3	Trajectory assurance	46
15.4	Security control architecture	46
15.5	Human oversight as designed authority	46
15.6	Security evidence	47
16	Standards and Institutional Interoperability	48
16.1	Mapping rather than claiming certification	48
16.2	Current transparency posture	48
16.3	Agent identity and non-repudiation	48
16.4	Interoperable evidence	49
16.5	Policy-as-code and professional judgment	49
16.6	Sovereign Intelligence as an interoperability profile	49

17 Economics and Commercial Strategy	50
17.1 The business model	50
17.2 Commercial sequence	50
17.3 Illustrative product architecture	50
17.4 Customer value gate	50
17.5 AI-native cost structure	51
17.6 Revenue scenarios	51
17.7 The durable moat	52
18 From Proof to Market Authority	53
18.1 The commercial proof problem	53
18.2 The proof ladder	53
18.3 Founding Commercial Proof Run	53
18.4 Verified customer value	53
18.5 The Global Expansion Office conversion	54
18.6 AI-native operating economics	54
18.7 Institutional restraint as a commercial asset	54
19 SovereignBench Ω: Empirical Authority	55
19.1 Why a benchmark is necessary	55
19.2 Benchmark arms	55
19.3 Primary metrics	55
19.4 Value-of-proof experiments	56
19.5 Verifier independence experiments	56
19.6 Commercial authority gate	56
19.7 Public-safe Receipt Packs	56
20 Implementation Roadmap	58
20.1 Phase 0: constitutional baseline	58
20.2 Days 0–100: paid proof	58
20.3 Months 4–12: repeatability	58
20.4 Months 12–24: platformization	59
20.5 Days 750–1000: productive capacity frontier	59
20.6 Operating team	59
20.7 Standards strategy	60
21 Risks, Failure Modes, and Controls	61
21.1 Residual risk	62
21.2 Anti-theatre controls	62

V Conclusion and Constitutional Appendices	63
22 Conclusion: Constituting the Verified Action Institution	64
Founder’s Constitutional Declaration	66
A Minimum Mission Contract	68
B Authority Envelope Schema	69
C Evidence Docket Schema	70
D Canonical Transaction Receipts	71
E Chronicle Admission Checklist	72
F Canonical Operating Doctrine	73
G Glossary	74
H Machine-Readable Schemas	75
H.1 Mission Contract example	75
H.2 Authority Envelope example	75
H.3 Proof-Carrying Action Record example	75
H.4 Chronicle capability example	76
H.5 Interoperability profile	76
I Author and Institutional Stewardship	77
I.1 Author	77
I.2 Publishing institution	77
I.3 Institutional relationship with QUEBEC.AI	77
I.4 Rights, attribution, and reuse	77
I.5 Recommended citation	77
I.6 Public contact	78
References	79

Part I

The Constitutional Thesis

CHAPTER 1

The Sovereign Intelligence Thesis

1.1 Why a new institutional layer is required

The acceleration of machine capability changes the economics of cognition, but not the constitutional structure of consequential action. A model can produce a plausible plan, a persuasive memorandum, a tax comparison, a software patch, a capital scenario, or a procurement recommendation at a fraction of the historical cost. Yet the organization still must determine whether the objective was correctly understood, whether the model had authority to act, whether the sources were current, whether evidence supports the claim, whether a responsible person accepts the outcome, whether rights permit reuse, and whether the result should shape future missions.

This gap is the organizing premise of GoalOS Ω . The durable opportunity is not another conversational interface or a larger collection of agents. It is the institutional layer that converts abundant candidate intelligence into admissible organizational action. The design follows a long institutional tradition: coordination costs help define the firm; governance structures shape transaction costs; complex systems require modular architecture; and viable control requires sufficient regulatory variety. (Coase 1937; Williamson 1985; Simon 1962; Ashby 1956)

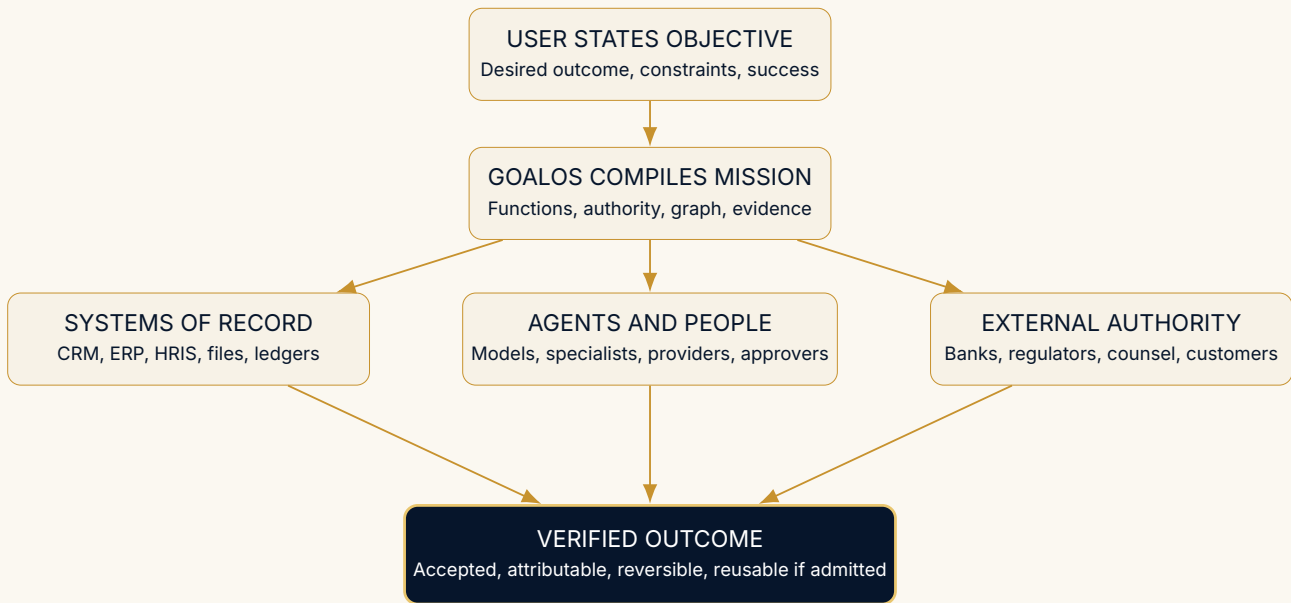
Definition: Sovereign Intelligence

Sovereign Intelligence is intelligence that remains governed by the institution it serves. It is identity-bound, objective-directed, policy-constrained, authority-bounded, evidence-bearing, independently challengeable, rights-aware, revocable, and incapable of silently converting candidate output into institutional truth, memory, or capital allocation.

The term *sovereign* does not imply state power, legal immunity, or unrestricted autonomy. It refers to control of the intelligence lifecycle by the legitimate institution: who may set the objective; who may delegate authority; what data and tools may be used; what evidence must be returned; who may accept the outcome; what enters memory; what expires; what can be revoked; and what verified value may be reinvested.

1.2 From software objects to institutional objectives

Traditional software is organized around data objects: account, ticket, invoice, employee, contract, project, entity, or asset. Human labour integrates those objects into the objective. GoalOS inverts the interface. The user states the objective, and the institution compiles the necessary objects, systems, people, agents, evidence, approvals, and transactions.



This objective-first design follows the broader GoalOS-native portfolio thesis: the winning company is not necessarily another CRM, ERP, workflow tool, or chatbot, but the verified action institution above them. (MONTREAL.AI Research 2026c)

1.3 The scarcity migration

The institution is motivated by a simple economic transition:

Candidate cognition cost ↓↓↓
Need for truth, authority, evidence, rights, and acceptance ↑

When intelligence is expensive, organizations optimize access to expertise. When candidate intelligence becomes abundant, organizations must optimize the reliability and admissibility of its consequences. The scarce inputs increasingly include:

- exact and current facts;
- authoritative sources and lawful data rights;
- independent challenge and non-correlated verification;
- licensed or fiduciary judgment where reserved;
- customer, board, regulator, or counterparty acceptance;
- security, identity, and revocable authority;
- physical resources such as energy, compute, land, equipment, and time;
- the right to reuse what was learned.

1.4 The sovereignty vector

Sovereignty is multi-dimensional. For institution state X , define a sovereignty vector:

$$\Sigma(X) = (\sigma_O, \sigma_I, \sigma_A, \sigma_E, \sigma_V, \sigma_M, \sigma_C, \sigma_R),$$

where the components represent control over objective, identity, authority, evidence, validation, memory, capital, and revocation. A system may be technically powerful yet institutionally non-sovereign if any critical component is controlled externally or left implicit.

Table 1.1: The eight sovereignty dimensions.

Dimension	Governing question	Failure if absent
Objective	Who defines success and permissible trade-offs?	Model or vendor substitutes its own proxy.
Identity	Whose authority does each human or agent exercise?	Actions cannot be attributed or revoked correctly.
Authority	What may be done, with which tools, data, budget, and duration?	Capability silently expands into permission.
Evidence	What must be returned to support each material claim?	Persuasion replaces proof.
Validation	Who independently tests the result and by what criterion?	Producer assertion becomes institutional truth.
Memory	What may influence future missions, under what scope and expiry?	Temporary or erroneous output becomes permanent doctrine.
Capital	What verified value may be reinvested, by whom, and for what purpose?	Speculative value finances irreversible expansion.
Revocation	How can authority, evidence, capability, or memory be withdrawn?	The institution cannot recover from change or error.

Sovereignty is not the absence of constraints. It is the institution's capacity to define, inspect, enforce, and revoke the constraints under which intelligence acts.

Category Definition and Design Laws

2.1 The Verified Action Institution

GoalOS Ω is defined as a continuously governed institution that transforms a consequential objective into a mission constitution, graph of bounded work, evidence-bearing outcome, accepted organizational state, selective institutional memory, and recursively improvable capability.

The system is neither a model nor a collection of models. It is not reducible to a workflow engine, agent framework, data room, advisory firm, or enterprise control tower. Each of those may become a node in the institution. GoalOS is the constitutional layer that determines how they may be combined and what their work may legitimately change.

Definition: Verified Action Institution

A **Verified Action Institution** is an objective-driven system that (i) compiles bounded work, (ii) enforces explicit authority, (iii) records the complete evidence denominator, (iv) obtains independent validation and stakeholder acceptance, (v) blocks incomplete transactions, (vi) governs memory through scoped admission and revocation, and (vii) measures improvement on fresh missions.

2.2 Canonical institutional hierarchy

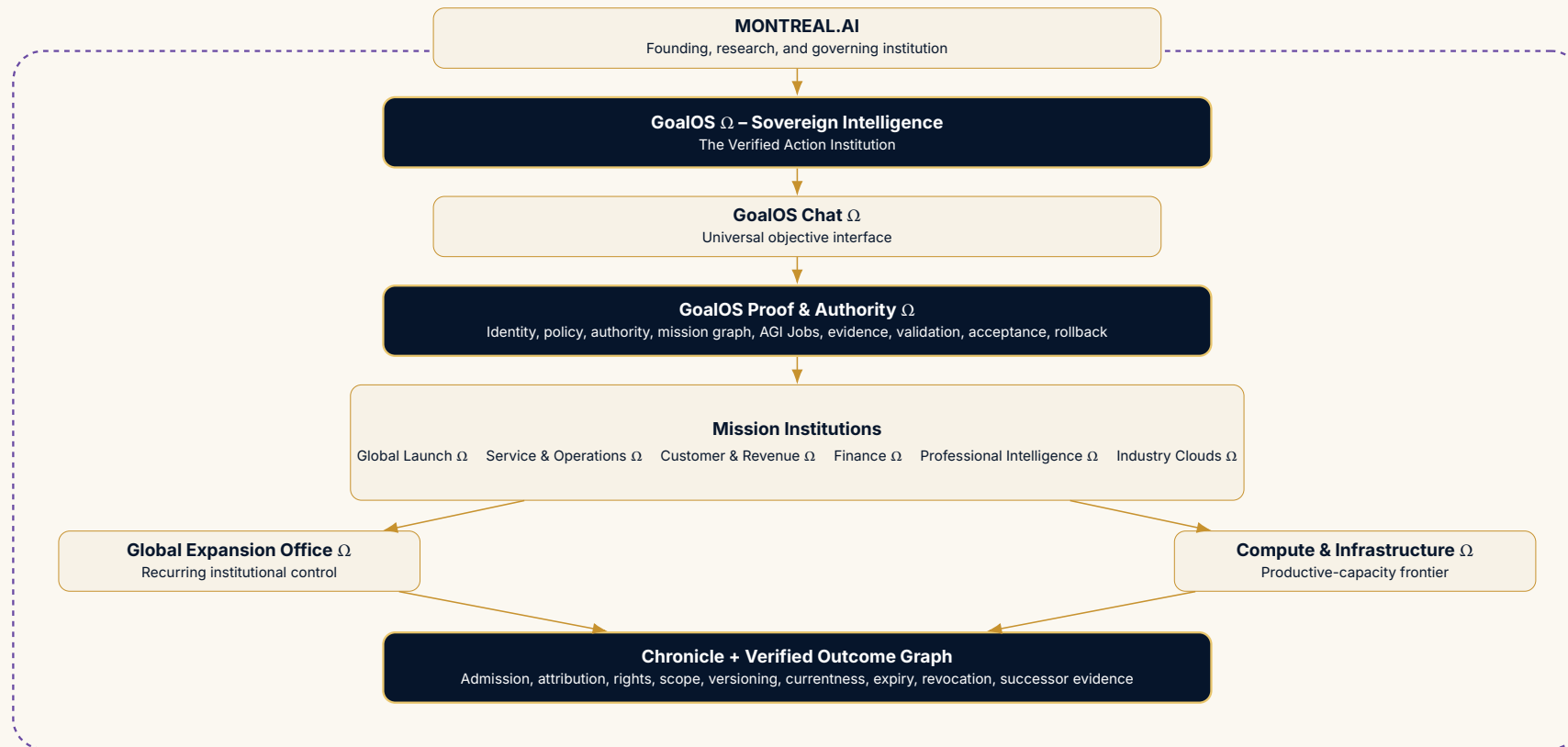


Figure 2.1: The canonical GoalOS hierarchy. Chronicle is a continuous field rather than a downstream archive.

2.3 Seven constitutional design laws

1. **Objective first.** No consequential objective, no mission. The interface begins from the outcome the institution wants, not from the application the user happens to open.
2. **Authority before external effect.** Capability never implies permission. Every action inherits an explicit authority envelope.
3. **Evidence before acceptance.** A completed action is not necessarily a correct, lawful, valuable, or accepted action.
4. **Acceptance before institutional state.** The responsible stakeholder, customer, board, or authority must decide whether the verified outcome changes the institution.
5. **Evidence before memory.** Output cannot enter Chronicle merely because it was useful once.
6. **Fresh comparison before improvement.** Reuse is not improvement. A successor must outperform its predecessor on a new mission under controlled or normalized conditions.
7. **Verified contribution before reinvestment.** Booked revenue, modelled benefit, or announced value is not automatically productive surplus.

2.4 The compounding progression



The progression is intentionally recursive. The institution does not simply automate a fixed process. It attempts to create a successor system with better quality, value, time, cost, risk, and reviewer burden – while preserving or improving governance.

2.5 What the category is not

GoalOS is not an attempt to erase institutions or replace every professional role. It is an institutional architecture for making rules, responsibilities, evidence, monitoring, and adaptation explicit around a consequential objective. This is consistent with the broader insight that durable governance emerges from arrangements matched to the actual resource, actors, and enforcement problem. (Ostrom 1990)

The purpose of GoalOS is not to maximize autonomous action. It is to maximize verified institutional value under explicit objective, authority, evidence, rights, and reversibility constraints.

Table 2.1: GoalOS compared with adjacent categories.

Category	Typical strength	What GoalOS adds
Chatbot	Natural-language interaction	Objective compilation, work, evidence, acceptance, and memory.
Workflow engine	Deterministic process routing	Dynamic mission graphs, typed proof and authority edges, and adaptive verification.
Agent framework	Tool-using autonomous workers	Institutional permission, complete evidence, stakeholder acceptance, and Chronicle.
Control tower	Inventory, observability, and policy	Admissibility of outcomes, proof debt, transaction gating, and successor governance.
Professional advisory	Deep judgment and representation	Cross-domain orchestration, reusable evidence structures, and continuous monitoring.
System of record	Durable organizational data	Verified action across systems and selective capability compounding.

SII-001: Constitutional Standard

3.1 Purpose and status

SII-001 is the normative constitutional profile for a GoalOS-compatible Verified Action Institution. It does not certify legal compliance or technical safety. It specifies the minimum institutional objects, gates, records, and invariants required for a system to claim that it is proof-governed rather than merely agentic.

The standard is designed for model-agnostic implementation. A conforming institution may use frontier hosted models, local models, deterministic programs, human professionals, corporate providers, validators, public authorities, and physical operators. Conformance depends on the institutional relationships among them, not on anthropomorphic claims about intelligence.

3.2 Twelve constitutional articles

Art.	Constitutional requirement	Operational meaning
1	Objective supremacy	Every mission begins from an accountable, versioned objective with success, exclusions, deadline, and value at stake.
2	Identity before authority	Every human, agent, service, professional, and institution acting in the graph has an attributable identity and principal.
3	Bounded authority	Tools, data, communications, spend, duration, counterparties, and external effects are explicitly scoped and revocable.
4	Typed work	The mission is compiled into jobs, dependencies, proof edges, authority edges, resource edges, and governance edges.
5	Complete evidence denominator	The Docket records sources, assumptions, contradictions, costs, failures, interventions, receipts, outcomes, and delayed effects.
6	Independent challenge	Material claims and outcomes are tested by a context, method, or responsible party that is not merely the producer repeating its own assertion.
7	Acceptance as authority	Verification establishes what the evidence supports; a responsible principal decides whether the outcome is accepted for the institution.
8	Fail-closed material action	No transaction, publication, production change, capital allocation, regulated filing, or external commitment proceeds while applicable proof or authority gates remain open.
9	Selective memory	Output does not enter Chronicle automatically. Admission is scoped, rights-cleared, attributable, current, revocable, and linked to failure conditions.
10	Fresh-successor proof	Reuse is not improvement. A successor must outperform the predecessor on a fresh mission without governance regression.

11	Verified-value discipline	Benefits progress through identified, validated, authorized, implemented, realized, and independently verified states.
12	Productive reinvestment	Capital expands capability or infrastructure only after contribution has been evidenced and classified within approved policy.

3.3 Constitutional invariants

Let $\mathcal{F}$ denote the set of required institutional invariants. A mission state is constitutionally admissible only if

$$\text{Conform}(X_t) = \prod_{j \in \mathcal{F}} \mathbf{1}[I_j(X_t) = 1] = 1.$$

The principal invariants are:

1. no unidentified actor on a material execution path;
2. no external action outside an unexpired authority envelope;
3. no material claim without source, method, scope, owner, date, and contradiction state;
4. no producer-only validation for material outcomes;
5. no transaction without current applicable receipts;
6. no Chronicle admission without rights and revocation;
7. no realized-value claim without attributable outcome evidence;
8. no improvement claim without a preregistered fresh-successor comparison.

3.4 Conformance levels

Table 3.2: SII-001 conformance levels.

Level	Name	Permitted posture	Minimum proof
SII-0	Candidate intelligence	Research, simulation, drafting, no external authority	Source disclosure and claim boundary
SII-1	Bounded assistance	Tool use in non-production or read-only environments	Identity, policy, tool scope, trace
SII-2	Governed execution	Reversible action inside approved systems	Authority envelope, tests, rollback, Docket
SII-3	Proof-gated transaction	Material organizational or economic action	Independent validation, applicable receipts, acceptance
SII-4	Chronicle institution	Reuse of admitted capability across missions	Rights, scope, currentness, revocation, fresh-successor test
SII-5	Productive-capacity institution	Allocation of verified value into capital and infrastructure	Realized-value ledger, capital policy, milestone evidence

3.5 Institutional constitution versus model policy

A model policy constrains one model invocation. An institutional constitution governs the whole causal path: principal, objective, data, tools, jobs, providers, validators, acceptance, memory, capital, and future reuse. GoalOS therefore treats model guardrails as one control family inside a larger institution rather than as the institution itself.

Part II

Formal System

CHAPTER 4

Formal Model of the Institution

4.1 Institution state

At time t , represent the institution as

$$X_t = (O_t, I_t, P_t, A_t, G_t, E_t, V_t, Y_t, M_t, K_t, R_t),$$

where:

- O_t is the objective and success state;
- I_t is the set of principals, human identities, agent identities, and service identities;
- P_t is the policy, legal, contractual, and risk constraint set;
- A_t is the authority lattice;
- G_t is the typed mission and dependency graph;
- E_t is the Evidence Docket;
- V_t is the validation state;
- Y_t is the accepted-outcome state;
- M_t is Chronicle and the rights-governed outcome graph;
- K_t is capital, budget, and productive capacity;
- R_t is the rollback, revocation, and recovery state.

The institution transition is not simply a model call. It is a governed state transition:

$$X_{t+1} = \mathcal{T}(X_t, m_t, \omega_t),$$

where m_t is the mission and ω_t is the external world state, including events not controlled by the institution. The distinction among observation, intervention, and counterfactual dependence is essential when GoalOS attributes an outcome or claims that one architecture caused improvement. (Pearl 2009)

4.2 Mission constitution

A mission is represented as

$$m = (o, s, b, \tau, \rho, \pi, \mathcal{F}, \alpha, \epsilon, \kappa),$$

where o is the objective, s the accountable sponsor and stakeholder set, b the budget, τ the time horizon, ρ the risk and reversibility posture, π the policy constraints, $\mathcal{F}$ the functions to optimize, α the authority structure, ϵ the initial evidence, and κ the acceptance criteria.

Minimum Mission Constitution

A production mission should specify: accountable sponsor; decision authority; objective; scope; exclusions; deadline; permitted data and tools; prohibited actions; budget and spend limits; success and acceptance tests; verifier; professional-authority dependencies; publication and reuse rights; stop, escalation, and rollback rules; and Chronicle eligibility.

The Constitution is versioned and sealed before consequential work begins. Changes are not prohibited, but they are explicit. A mission whose objective or acceptance test changes during execution is a different mission state and should preserve lineage.

4.3 Authority envelope

For job or agent i , define the authority envelope

$$\mathcal{A}_i = (q_i, r_i, d_i, u_i, b_i, \tau_i, c_i, z_i),$$

where:

- q_i is the principal or institution from which authority is derived;
- r_i is the permitted role and action class;
- d_i is the permitted data scope;
- u_i is the permitted tools, systems, and communication channels;
- b_i is the budget, transaction, or resource ceiling;
- τ_i is the duration and expiry;
- c_i is the required control and approval set;
- z_i is the revocation, rollback, and incident route.

An action a is admissible at execution time only if

$$\text{Permitted}(a, \mathcal{A}_i, P_i) = 1.$$

This pre-execution authority test is distinct from post-execution acceptance. An agent may be permitted to perform a test without being permitted to publish the result, commit funds, modify production, or bind the institution.

4.4 Typed mission graph

Let the mission graph be

$$G = (N, E_D, E_P, E_A, E_R, E_G),$$

with five edge types:

1. E_D : **data edges** – a downstream node consumes an upstream artifact;
2. E_P : **proof edges** – a separate node tests a material claim or outcome;
3. E_A : **authority edges** – progression requires responsible permission;
4. E_R : **resource edges** – nodes share or compete for state, budget, tools, workspace, or capital;
5. E_G : **governance edges** – a result cannot influence a future mission without admission, expiry, or revocation rules.

This distinction prevents two common errors: treating every sequence as a computational dependency, and treating constitutional waits as mere latency. A mission graph should make clear which jobs can run in parallel, which claims require independent proof, which actions require approval, and which outputs remain isolated from future memory.

4.5 Evidence object

An evidence object e is

$$e = (c, s, t, m, p, \chi, \eta, \xi, \delta),$$

where c is the claim, s the source, t the retrieval or observation time, m the method, p the producing principal, χ the scope, η the expiry, ξ the contradiction set, and δ the decision impact.

Evidence is not a confidence score. It is an inspectable object with provenance, scope, method, time, ownership, and known contradictions. The Evidence Docket records both successful and adverse evidence.

4.6 Proof debt

For architecture or mission state x , define the proof-debt set

$$D(x) = \{d_1, d_2, \dots, d_n\},$$

where each d_j is an unresolved claim, missing source, ambiguous authority, stale fact, untested dependency, professional dependency, or transaction condition. A weighted measure is

$$PD(x) = \sum_j w_j \text{Impact}(d_j) \text{Urgency}(d_j) \text{Uncertainty}(d_j).$$

A high-value architecture may remain inadmissible if its proof debt intersects a material transaction path.

4.7 Verified institutional value

For scenario s , define

$$\begin{aligned} VIV(s) = & V_{\text{mission}} + V_{\text{market}} + V_{\text{capital}} + V_{\text{capability}} + V_{\text{infrastructure}} + V_{\text{optionality}} \\ & - C_{\text{formation}} - C_{\text{operation}} - C_{\text{coordination}} - C_{\text{compliance}} - C_{\text{delay}} \\ & - C_{\text{dilution}} - C_{\text{risk}} - C_{\text{unwind}} - \lambda PD(s). \end{aligned}$$

The chosen scenario is not the raw maximizer. It is the best admissible scenario:

$$s^* = \arg \max_{s \in \mathcal{S}} VIV(s) \quad \text{subject to} \quad \text{Gate}(s) = 1.$$

GoalOS may search broadly across high-upside architectures. It may recommend only from the subset that satisfies current evidence, authority, rights, security, and reversibility constraints. Maximum upside is not the same as maximum admissible value.

4.8 Value of proof

Given unresolved evidence object e , define its value of proof as the expected improvement in decision quality from resolving it:

$$\text{VoP}(e) = \mathbb{E}[\max_s \text{VIV}(s) \mid e] - \mathbb{E}[\max_s \text{VIV}(s)] - C(e).$$

The institution should allocate verification resources to the evidence most likely to change a consequential decision, not merely to the easiest claims to test.

CHAPTER 5

The Proof-Carrying Action Protocol

5.1 From tool call to institutional action

An action becomes institutionally meaningful only when it is connected to purpose, authority, evidence, and acceptance. GoalOS represents each consequential action as a Proof-Carrying Action Record (PCAR):

$$\text{PCAR}(a) = (o, p, i, \mathcal{A}, x, t, r, e, v, y, \rho),$$

where o is the objective, p the accountable principal, i the acting identity, $\mathcal{A}$ the authority envelope, x the exact input state, t the tool or provider, r the raw result, e the evidence set, v the validation verdict, y the acceptance state, and ρ the rollback and revocation route.

PCAR minimum fields

A material PCAR records: mission identifier; versioned objective; actor identity; principal; policy version; authority scope and expiry; permitted data and systems; preconditions; action arguments; timestamps; model, software, provider, and tool versions; cost; raw output; tests; contradictions; validator; acceptance; external receipt; adverse effects; rollback state; rights; retention; Chronicle disposition; cryptographic digest or equivalent integrity evidence.

5.2 Two gates and three verdicts

GoalOS distinguishes a pre-action gate from a post-outcome gate.

$$G_{pre}(a) = I_{identity}I_{authority}I_{policy}I_{data}I_{tool}I_{budget}I_{precondition},$$

$$G_{post}(a) = I_{evidence}I_{validation}I_{acceptance}I_{rights}I_{currentness}.$$

Execution requires $G_{pre}(a) = 1$. Institutional admission requires $G_{post}(a) = 1$. The three verdicts are:

1. **Execute**: the action is authorized and sufficiently bounded;
2. **Accept**: the outcome is verified and adopted for the stated institutional scope;
3. **Admit**: a narrow capability may influence future missions under Chronicle rules.

No verdict implies the next. Permission to run a test does not imply acceptance of its result. Acceptance of an outcome does not imply that its method is reusable.

5.3 Evidence trust calculus

For evidence object e , define a trust vector

$$\theta(e) = (q_s, q_m, q_i, q_t, q_x, q_r),$$

where the components represent source quality, method quality, independence, temporal currentness, contradiction resolution, and rights/reproducibility. A scalar score may support triage,

$$T(e) = \sum_k w_k q_k,$$

but the vector must remain visible. A high source-quality score cannot erase expired evidence; strong temporal currentness cannot create independence; a reproducible test cannot create legal authority.

5.4 Causal provenance

The Evidence Docket should support a causal trace from objective to outcome:

$$O \rightarrow H \rightarrow J \rightarrow A \rightarrow R \rightarrow E \rightarrow V \rightarrow Y,$$

where H is the hypothesis or architecture, J the bounded job set, A the actions, R the raw results, E the evidence, V validation, and Y accepted outcome. W3C PROV offers a useful general provenance vocabulary; GoalOS adds authority, acceptance, expiry, rights, transaction, and successor semantics. (Moreau and Missier 2013)

5.5 Proof obligation compiler

For each material claim c , the compiler derives a proof obligation:

$$\Pi(c) = (\text{claim type, stake, required sources, test, independence, authority, expiry}).$$

The required proof depth should increase with expected harm, irreversibility, regulatory sensitivity, and capital exposure. This produces a risk-proportionate verification budget rather than equal scrutiny for every sentence. NIST treats testing, evaluation, verification, and validation as operational disciplines rather than a single confidence score; the PCAR makes their outputs mission-specific and attributable. (National Institute of Standards and Technology 2026d)

Table 5.1: Illustrative proof obligations by action class.

Action class	Minimum evidence	Independent check	Acceptance authority
Internal draft	Sources and assumptions	Optional second context	Mission owner
Read-only enterprise query	Identity, authorization, source trace	Policy and access check	System owner
Reversible operational change	Preconditions, tests, rollback	Independent test runner	Authorized operator
Contract or material transaction	Exact facts, professional and financial receipts	Independent reviewer or responsible professional	Delegated executive or board
Chronicle capability	Outcome evidence, rights, scope, known failures	Fresh-mission comparison	Chronicle governor
Productive-capacity allocation	Realized value, milestone plan, risk and unwind	Investment/technical review	Treasury or project authority

Action without proof is execution risk. Proof without authority is advisory evidence. Authority without proof is institutional exposure. GoalOS requires the three to meet before consequential value is admitted.

CHAPTER 6

Proof, Authority, and Transaction Admission

6.1 Separation of powers

GoalOS separates six institutional planes. The separation is not ceremonial; it prevents the same system from proposing, executing, grading, authorizing, remembering, and funding its own work without independent checks.

Table 6.1: Separation of powers.

Plane	May do	May not do alone
Intelligence	Research, model, simulate, compare, propose, design tests	Create legal, fiduciary, regulatory, or stakeholder authority
Execution	Perform approved AGI Jobs within tools, data, budget, and policy	Expand scope or convert capability into permission
Proof	Record tests, evidence, costs, failures, and outcomes	Treat the producer's narrative as independent truth
Authority	Approve, condition, pause, sign, decline, revoke	Hide the basis of authority or bypass mandatory law
Memory	Admit, version, scope, expire, revoke, compose capability	Store every output as truth or reuse secrets without rights
Capital	Allocate verified contribution to capability and capacity	Treat modelled or booked value as verified free cash flow

6.2 Execution state machine

A material job progresses through explicit states:

Constituted → Authorized → Executing → Evidence Returned
→ Validated → Accepted → Admitted or Closed.

Failures, pauses, disputes, or revocations branch from any state. The state machine should be observable, replayable, and linked to its Mission Constitution.

6.3 Transaction Gate

For material transaction t , let the required receipt set be

$$\mathcal{R}(t) = \{r_{\text{purpose}}, r_{\text{legal}}, r_{\text{tax}}, r_{\text{valuation}}, r_{\text{funding}}, r_{\text{authority}}, r_{\text{implementation}}, r_{\text{acceptance}}, r_{\text{chronicle}}\}.$$

Input: Mission m , institution state X_i
Output: Accepted outcome, closed mission, or explicit unresolved state
 Compile objective and policy into typed graph G ;
 Issue authority envelope $\mathcal{A}_i$ for each job;
foreach *ready node* n **in** G **do**
 verify identity, authority, tools, data, budget, and expiry;
 execute in approved environment;
 return structured artifact, trace, cost, failures, and evidence;
 route material claims to independent verifier;
end
 Assemble Evidence Docket and unresolved Proof Debt;
 Apply transaction and acceptance gates;
if *all material gates pass* **then**
 obtain responsible acceptance;
 update institutional state;
 evaluate Chronicle admission;
else
 block, condition, repair, escalate, or close without admission;
end

Algorithm 1: Proof-governed mission execution

Not every receipt applies to every transaction, but every applicable receipt must be current, attributable, and approved:

$$\text{TG}(t) = \begin{cases} 1, & \text{if all applicable receipts pass;} \\ 0, & \text{otherwise.} \end{cases}$$

No entity, intellectual-property transfer, financing, option grant, regulated filing, major procurement, employment action, production deployment, or infrastructure commitment proceeds merely because a model recommends it. If the transaction gate is incomplete, the action remains blocked.

6.4 Independent validation

Independence is a property of context, evidence path, incentives, and authority – not merely a second model call. Validation may require:

- a different model family or fresh context;
- separate source retrieval;
- blinded criteria and preselected evaluators;
- executable tests or controlled environments;
- customer confirmation;
- a qualified professional;
- an external authority receipt;
- direct observation of the realized outcome.

NIST's AI Risk Management Framework emphasizes governance, mapping, measurement, and management, while

its 2026 AI Agent Standards Initiative explicitly identifies secure identity, authorization, interoperability, and agent security as foundational adoption requirements. (Tabassi 2023; National Institute of Standards and Technology 2026a; National Institute of Standards and Technology 2026c)

6.5 Acceptance is not validation

A validated result may still be rejected because the customer, board, fiduciary, or authority does not accept its trade-offs. Conversely, acceptance without adequate proof creates institutional fragility. GoalOS therefore records validation and acceptance separately.

Accepted Outcome

An outcome is accepted when the designated authority acknowledges the evidence, limitations, obligations, and residual risk, and authorizes the corresponding institutional state transition. Acceptance is scoped, attributable, time-stamped, and revocable where the underlying authority permits.

6.6 Source decay and fail-closed operation

Sources, programmes, prices, regulations, professional opinions, customer facts, and infrastructure conditions decay at different rates. Each evidence object carries a review horizon. When material evidence expires, the dependent claim accrues Proof Debt and may cause the transaction gate or the relevant mission function to fail closed.

The EU AI Act's phased application and the European Commission's continuing implementation guidance illustrate why static compliance claims are structurally unsafe; role classification, system facts, and applicable dates must remain current. (European Parliament and Council of the European Union 2024; European Commission 2026a)

Chronicle and Recursive Improvement

7.1 Memory as constitutional admission

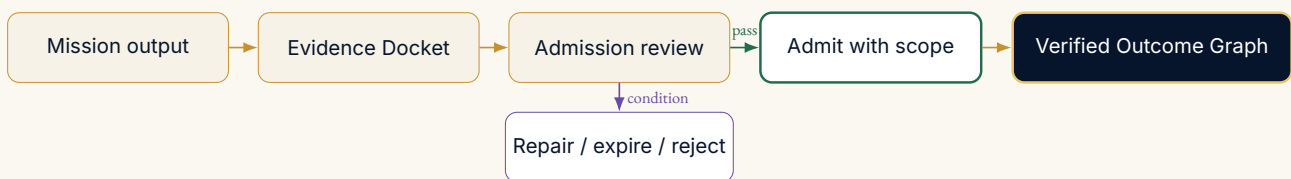
Most software memory systems optimize retrieval. Chronicle optimizes admissibility. The central question is not whether an artifact may be stored, but whether a capability may influence future missions.

Suppose mission m_1 yields candidate capability k . Chronicle admission is

$$\text{Admit}(k) = 1$$

only if the capability is:

- supported by attributable evidence;
- rights-cleared for the proposed reuse;
- narrowly scoped to valid conditions;
- versioned and linked to its producing models, tools, and sources;
- current or governed by an expiry rule;
- independently challenged where material;
- accompanied by known limitations and failure signatures;
- revocable;
- tested for transfer to a fresh mission.



7.2 The verified outcome graph

The graph links:

Objective → Facts → Architectures → Jobs → Evidence → Authority → Actions → Outcomes → Admission.

It includes rejections, failures, delays, costs, contradictions, clawbacks, and adverse outcomes – not only successful examples. This complete denominator is essential for calibrated future decisions.

A rights-controlled outcome graph can eventually reveal:

- which architectures actually produced value;
- which evidence changed the decision;

- which providers or evaluators were reliable;
- which claims failed and why;
- which combinations survived audit or long-term monitoring;
- which methods transferred to new missions;
- which early signals predicted failure;
- which capabilities should be revoked.

7.3 Fresh-successor test

Let g_1 be the predecessor generation and g_2 a successor using Chronicle-admitted capability. On fresh mission m_2 define:

$$\Delta(g_2, g_1 \mid m_2) = \alpha\Delta Q + \beta\Delta V + \gamma\Delta T + \delta\Delta C - \lambda\Delta R - \mu\Delta H,$$

where Q is quality, V verified value, T time efficiency, C cost efficiency, R risk, and H human or reviewer burden. Promotion requires a preregistered positive threshold and no material governance regression.

Reuse is not improvement. A successor earns authority only through a fresh comparison under controlled or explicitly normalized conditions. A system that merely uses its own prior output has demonstrated recurrence, not recursive self-improvement.

7.4 Recursive compounding

The institution-level compounding sequence is:

Intelligence → Work → Proof → Authority
 → Protected Capability → Recurring Value
 → Capital → Productive Capacity.

The cycle is constructive only when each transformation preserves evidence and governance. Otherwise, errors compound as readily as capability.

7.5 Revocation and forgetting

Chronicle must support forgetting. Capability can be revoked because:

- a source changed;
- the scope was broader than evidence justified;
- rights expired or were withdrawn;
- a later mission revealed a failure mode;
- a professional opinion changed;
- a customer or regulator imposed a new constraint;
- the capability's expected value became negative.

The ability to revoke is part of intelligence, not a concession to imperfection. A system that cannot forget stale capability cannot remain sovereign.

Only what survives reality may compound.

Institutional Twin and Counterfactual Governance

8.1 The twin as a decision-bearing state

The GoalOS Institutional Twin is a versioned, queryable representation of the institution's objective, entities, people, functions, jurisdictions, contracts, capital, capabilities, infrastructure, evidence, authority, obligations, and risk. It is not a decorative dashboard. It determines which actions remain admissible and which successor architectures should be generated when the world changes.

Let

$$\mathcal{D}_t = (\mathcal{N}_t, \mathcal{E}_t, \mathcal{S}_t, \mathcal{C}_t),$$

where $\mathcal{N}_t$ is the institutional node set, $\mathcal{E}_t$ the typed relation set, $\mathcal{S}_t$ the evidence and source state, and $\mathcal{C}_t$ the control and authority state.

8.2 Twin objects

Table 8.1: Core objects of the Institutional Twin.

Object family	Examples	Decision role
Objectives	mission, success metric, deadline, exclusions	Defines what the institution is optimizing
Organizations	parent, subsidiary, SPV, partner, authority	Determines legal, financial, and operational topology
People and agents	principals, employees, professionals, agents, validators	Carries identity, responsibility, authority, and conflict state
Capabilities	models, code, methods, data, IP, playbooks	Determines what may be invoked and under which rights
Capital	cash, credit, award, debt, equity, budget, covenant	Determines feasible work and productive expansion
Infrastructure	compute, energy, site, laboratory, network, vendor	Determines physical capacity, constraints, and resilience
Evidence	source, receipt, test, verdict, contradiction, expiry	Determines what claims and actions remain supportable
Transactions	filing, agreement, hire, payment, deployment, transfer	Determines gate status, reversibility, and downstream obligations

8.3 Change propagation

When source or world event $\Delta\omega_t$ occurs, the twin computes an impact set:

$$\mathcal{I}(\Delta\omega_t) = \{n \in \mathcal{N}_t : n \text{ depends directly or transitively on } \Delta\omega_t\}.$$

Affected claims accrue Proof Debt; affected transactions may become blocked; affected scenarios are rerun. A change to one funding rule, employment threshold, ownership state, security control, or infrastructure assumption should not require a human to remember every downstream dependency.

8.4 Scenario switch conditions

Every accepted architecture should publish explicit switch conditions rather than pretending to be permanently optimal. Examples include:

- a revenue threshold that makes a foreign entity economically justified;
- a hiring threshold that changes the preferred employment route;
- a compute-demand threshold that activates an infrastructure mission;
- a funding condition that requires a new IP or commercialization review;
- an evidence-confidence floor that blocks a transaction;
- a capital or ownership event that changes eligibility or control.

8.5 Counterfactual governance

The twin supports three counterfactual questions:

1. **Policy counterfactual:** what would change if a rule, budget, or authority constraint changed?
2. **World counterfactual:** what would change if market, law, infrastructure, or source conditions changed?
3. **Architecture counterfactual:** which accepted outcome would differ under an alternative institutional design?

Counterfactuals must be labelled as simulations. They guide search and Value of Proof; they do not become evidence merely because the simulation is detailed. Causal claims require explicit assumptions and distinguish observational association from intervention and counterfactual reasoning. (Pearl 2009)

8.6 Institutional regret

Define realized institutional regret after world state ω as

$$\text{Regret}(s, \omega) = \max_{s' \in \mathcal{S}} \text{VIV}(s', \omega) - \text{VIV}(s, \omega).$$

GoalOS should minimize maximum regret where the objective is survival under uncertainty, and maximize expected verified value where reversibility and evidence support a more aggressive posture. The choice of decision criterion is itself part of the Mission Constitution.

The Institutional Twin may forecast, simulate, and propose. It may not erase uncertainty, manufacture evidence, or silently rewrite the accepted objective. Every consequential change remains attributable to a source, decision, or authorized event.

Part III

Institutional Architecture

Canonical Architecture of GoalOS Ω

9.1 Architecture overview

The GoalOS architecture is designed as one institution with multiple Mission Institutions, not a portfolio of unrelated applications. The shared substrate carries identity, policy, authority, graph execution, evidence, validation, acceptance, rollback, and Chronicle across every domain. This is the central product discipline inherited from the GoalOS-native application portfolio. (MONTREAL.AI Research 2026c)

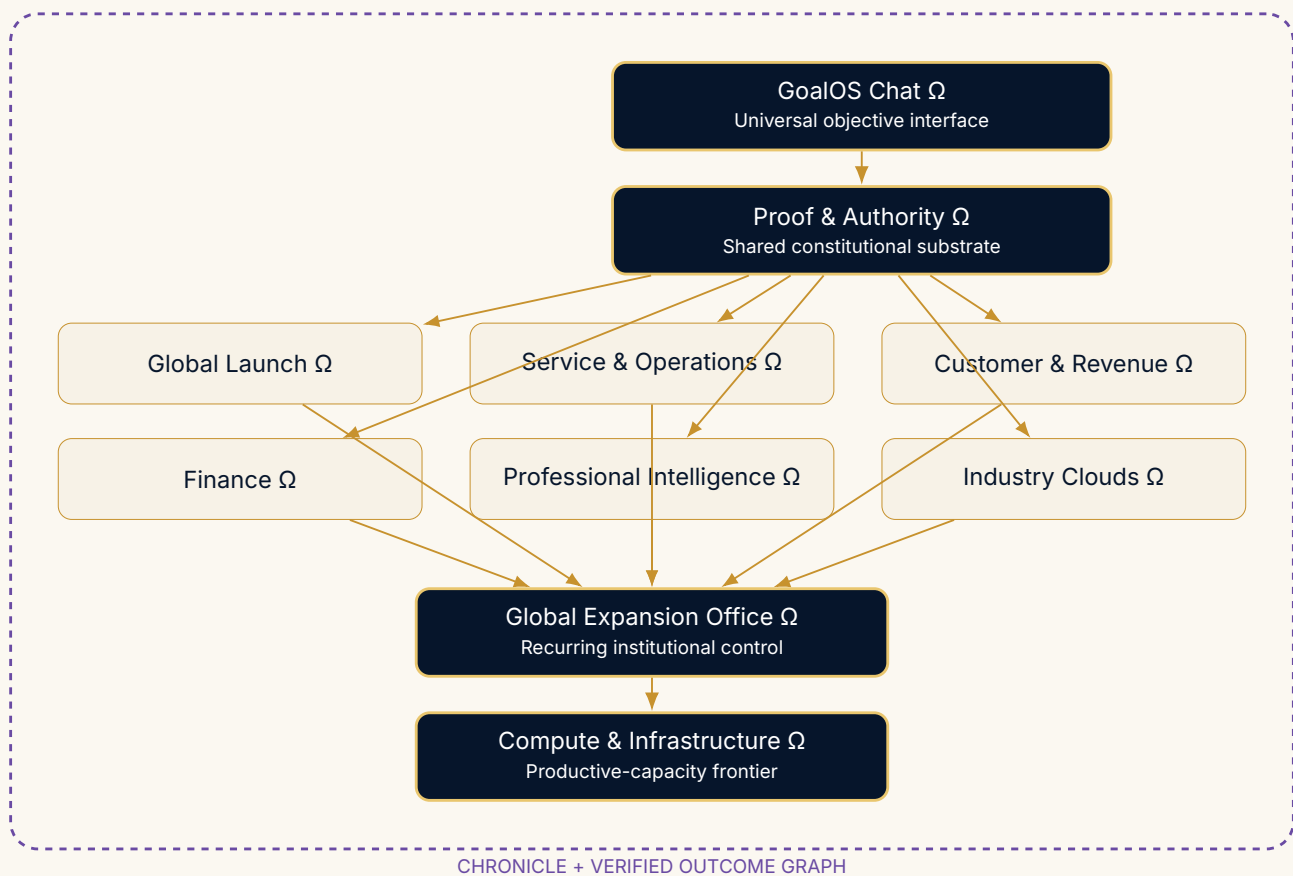


Figure 9.1: One shared institution, multiple Mission Institutions, one outcome graph.

9.2 GoalOS Chat Ω : universal objective interface

GoalOS Chat is not merely a conversational shell. It is the interface through which the institution:

- receives the objective in natural language;
- distinguishes requests for information from requests for action;
- identifies the accountable principal;
- asks only decision-changing questions;
- exposes assumptions and missing facts;

- presents scenario frontiers and trade-offs;
- requests authorization at constitutionally relevant moments;
- explains evidence, Proof Debt, and residual risk;
- returns an accepted outcome or an explicit unresolved state.

The interface should resist premature anchoring. For example, a global-launch user should not be asked to choose a jurisdiction before GoalOS analyzes the objective; doing so would encode a conclusion into the input. Similarly, an enterprise service request should not be reduced to a preselected ticket category if the real objective crosses systems.

9.3 GoalOS Proof & Authority Ω : the shared substrate

The control plane includes:

1. Identity and entitlements;
2. Objective and policy compiler;
3. Mission Constitution;
4. Typed mission graph;
5. AGI Job identities and workspaces;
6. Tool, data, communication, budget, and duration controls;
7. Evidence Docket;
8. Independent validation;
9. Transaction Gate;
10. Customer, board, professional, or authority acceptance;
11. Rollback and revocation;
12. Chronicle disposition;
13. Outcome measurement and capital classification.

The substrate is model- and tool-agnostic. It should be able to govern frontier models, local models, deterministic software, humans, firms, banks, validators, authorities, and physical providers within one mission graph.

9.4 GoalOS Global Expansion Office Ω

The Expansion Office is the recurring control plane for a living institution. It maintains the institutional twin, monitors source decay and architecture drift, launches corrective missions, and compares successor configurations.

Its permanent systems include:

- Institutional Twin;
- Global Architecture Map;
- Change Radar;
- Opportunity Queue;
- Proof Debt Register;

- Scenario Switchboard;
- AGI Job Portfolio;
- Evidence Docket;
- Transaction Gate;
- Chronicle;
- Verified Value Ledger;
- Executive and Board Brief.

9.5 GoalOS Compute & Infrastructure Ω

This is the productive-capacity frontier. It coordinates demand, sites, energy, compute, laboratories, equipment, permits, vendors, capital, construction, commissioning, and verified output. Physical assets should normally reside in customer- or project-funded vehicles, while GoalOS owns the intelligence, orchestration, proof, and governance layer.

The design prevents a common failure: funding capacity before demand, permissions, and milestone evidence exist.

GoalOS should own the intelligence, orchestration, evidence, customer relationship, Chronicle, and outcome graph. Reserved professional authority and physical execution should be invoked only when required and funded by the relevant mission or project.

Operating Modes and Deployment

10.1 Why deployment mode matters

The same GoalOS interface may support a public research demonstration, a private enterprise workspace, a regulated mission, or a critical-infrastructure programme. The applicable authority, data, proof, security, and professional requirements differ materially. A product name cannot determine legal treatment; actual users, functions, data, models, systems, counterparties, jurisdictions, and money flows do.

10.2 Five operating modes

Table 10.1: GoalOS operating modes.

Mode	Purpose	Permitted activity	Required controls
Reference	Public education and research	Dated sources, simulation, no sensitive data or external authority	Claim boundary, no-secrets gate, transparency notice
Private advisory	Customer-specific analysis	Protected data, scenario design, dossiers, no autonomous binding action	Contract, data map, access controls, professional routing
Assisted execution	Reversible or low-impact implementation	Bounded tools, customer approvals, monitored actions	Authority envelopes, sandboxing, traces, rollback, Evidence Docket
Governed transaction	Material organizational, financial, legal, or operational action	Approved filings, transfers, contracts, payments, production changes	Independent validation, professional receipts, Transaction Gate, acceptance
Strategic infrastructure	Critical or physical productive capacity	Project-financed systems, facilities, compute, energy, industrial deployments	Project SPV, milestones, engineering authority, safety, insurance, contingency and unwind

10.3 Deployment topologies

10.3.1 Browser-local reference institution

A browser-local deployment maximizes portability and minimizes server-side collection. It is well suited to public demonstrations, research snapshots, scenario analysis, and signed local provenance. It cannot provide confidential server-side access control merely because the interface is gated.

10.3.2 Protected enterprise workspace

A protected workspace adds organizational identity, role-based access, customer-specific encryption, data retention, audit logs, connectors, and security review. It should separate model providers, data processors, agent workspaces, and

evidence stores according to the customer's authority and risk requirements.

10.3.3 Isolated mission environment

High-stakes missions should receive isolated credentials, storage, network policy, tools, budgets, and validator contexts. A sandbox reduces blast radius but does not replace authority or outcome validation. OpenAI's 2026 Agents SDK direction toward controlled sandbox environments illustrates the broader move from simple prompts to durable agent harnesses. (OpenAI 2026c)

10.3.4 Federated professional and provider network

GoalOS may coordinate firms, professionals, banks, authorities, vendors, validators, and physical operators without making them employees or pretending their authority belongs to the model. Each external node receives a defined scope, evidence obligation, conflict state, service level, and Chronicle disposition.

10.3.5 On-chain or cryptographic settlement layer

Where appropriate, signatures, hashes, smart contracts, and public ledgers may provide integrity, ownership, or settlement evidence. They do not establish the truth of off-chain facts without oracles, validators, documents, or responsible attestations. Cryptography proves what was committed; it does not by itself prove that the represented world event occurred.

10.4 Portability and graceful degradation

A GoalOS mission should degrade safely when a model, connector, source, provider, or network fails. Essential records should remain exportable in documented formats. Critical controls should not depend on one model vendor, one account, one region, one key, or one individual.

Deployment convenience may reduce friction. It may not reduce the proof, authority, rights, security, and acceptance required by the actual mission.

CHAPTER 11

Mission Institutions and Application Foundry

11.1 Mission Institutions as governed packages

Each GoalOS application is a Mission Institution built on the shared substrate. It should not become an unrelated codebase. A Mission Institution contains:

- domain ontology;
- objectives and success states;
- policies and authority rules;
- connectors and approved tools;
- AGI Job templates;
- graph patterns;
- evaluators and acceptance tests;
- Evidence Docket schema;
- Transaction Gate rules;
- Chronicle admission and revocation rules;
- pricing and verified-value model.

This creates a GoalOS application foundry capable of producing new verified systems of action while preserving constitutional consistency.

11.2 Opportunity function

A Mission Institution is attractive when it combines a large profit pool, high mission frequency, verifiable completion, cross-system depth, reusable methods, recurring need, expensive mistakes, and an overlay entry path. A practical strategic function is:

$$\text{Opportunity} \propto \frac{P \cdot F \cdot V \cdot X \cdot U \cdot N}{L \cdot I \cdot C},$$

where P is profit pool, F mission frequency, V verifiability, X cross-system depth, U reuse, N recurring need, L liability, I integration friction, and C incumbent agentic strength. (MONTREAL.AI Research 2026c)

11.3 The core portfolio

Table 11.1: Canonical GoalOS Mission Institutions.

Mission Institution	Primary objective	Distinctive institutional advantage
Global Launch Ω	Build or reconfigure a lawful institution globally	Best category-creation wedge; integrates markets, entities, capital, public support, talent, capability, infrastructure, and evidence.
Service & Operations Ω	Complete cross-functional enterprise requests	Highest-frequency universal mission layer; policies and completion criteria are often explicit.
Customer & Revenue Ω	Convert customer requests into resolved, accepted outcomes	Universal front door and distribution surface; direct connection to revenue and retention.
Finance Ω	Close, reconcile, collect, forecast, control, and prepare evidence	Highly measurable value; structured data; recurring compulsory work.
Professional Intelligence Ω	Resolve legal, tax, regulatory, risk, audit, and diligence decisions	Strongest proof and authority moat; high value of current sources and qualified judgment.
Industry Clouds Ω	Operate domain-specific missions in regulated or specialized sectors	Strongest long-term moat through domain evidence, evaluators, integrations, and outcome history.
Engineering & IT Ω	Repair, deploy, monitor, and prove technical systems	Best hard-verification environment; tests, builds, incidents, and performance provide reality anchors.
Procurement & Treasury Ω	Source, compare, negotiate, contract, pay, and verify value	Direct savings, policy enforcement, working-capital impact, and transaction receipts.
People Ω	Complete employee lifecycle missions	High recurring volume across HRIS, payroll, security, finance, facilities, and regional policy.
Security Operations Ω	Investigate, contain, restore, and prove security state	Very high value with the strictest authority, isolation, and rollback requirements.

11.4 The optimal build sequence

1. **Distribution and proof:** GoalOS Chat, Global Launch, Service & Operations.
2. **Direct economic value:** Customer & Revenue, Finance, Professional Intelligence.
3. **Deep vertical moat:** Life Sciences, Insurance, Construction & Infrastructure.
4. **Global action layer:** developer, validator, connector, identity, settlement, treasury, and productive-infrastructure ecosystems.

The build order balances existing asset leverage, mission frequency, measurable value, and long-term defensibility. Global Launch establishes category authority; Service & Operations establishes frequency; Chat supplies distribution; Proof & Authority supplies the common platform; Chronicle supplies the moat.

11.5 GoalOS Global Launch Ω: the category wedge

Global Launch transforms a consequential objective into a function-by-function global architecture. It may support companies, nonprofits, research programmes, funds, expansions, infrastructure projects, industrial facilities, and public initiatives. The specialized GJFFI Ω kernel assigns functions across jurisdictions and identifies funding, fiscal, regulatory, capital, procurement, and infrastructure routes. (MONTREAL.AI Research 2026b; MONTREAL.AI Research 2026a)

The reference standalone implementation is a dated prototype rather than universal authority. Its July 2026 snapshot included 70 jurisdiction and functional nodes, 463 governed routes, 476 source records, 15 control planes, 20 transaction categories, 50 bounded AGI Jobs, 12 evidence objects, and nine transaction receipts. These counts demonstrate architecture and data structure, not guaranteed completeness or current eligibility.

11.6 GoalOS Service & Operations Ω: the frequency engine

A manager states: “Onboard our new French sales director by Monday.” GoalOS may coordinate employment checks, payroll, benefits, identity, CRM access, equipment, security training, expenses, approvals, and completion tests. The user experiences one objective. The institution executes a graph of bounded missions across multiple systems.

11.7 GoalOS Customer & Revenue Ω: the distribution engine

A customer does not fundamentally want to chat. The customer wants a problem resolved, an order changed, a quote obtained, a deployment corrected, a contract amended, a refund completed, or a renewal concluded – with proof that it happened. GoalOS Chat can become the verified front door through which support, sales, finance, security, legal, and deployment missions are coordinated.

11.8 GoalOS Finance Ω: measurable value

Finance provides hard economic metrics: days removed from close, cash collected, errors prevented, audit effort reduced, working capital released, and forecast accuracy. GoalOS should begin with evidence preparation, exception resolution, reconciliation, and controlled approvals rather than unconstrained payment or regulated filing.

11.9 GoalOS Industry Clouds Ω: the vertical moat

Vertical institutions can accumulate specialized workflows, sector evidence, policies, evaluators, integrations, failure signatures, professional authority, and verified customer outcomes. The deepest long-term moats may therefore be vertical, even if the initial platform is horizontal.

The product is not the agent. The product is the governed Mission Institution that repeatedly converts objectives into accepted, evidence-bearing outcomes.

CHAPTER 12

The Mission Pack Foundry and Ecosystem

12.1 Mission Packs as constitutional products

A GoalOS application is not merely a prompt or agent. It is a governed Mission Pack containing the domain knowledge, policies, connectors, job templates, proof obligations, acceptance tests, evidence schema, economics, and Chronicle rules required to complete a recurring class of consequential objectives.

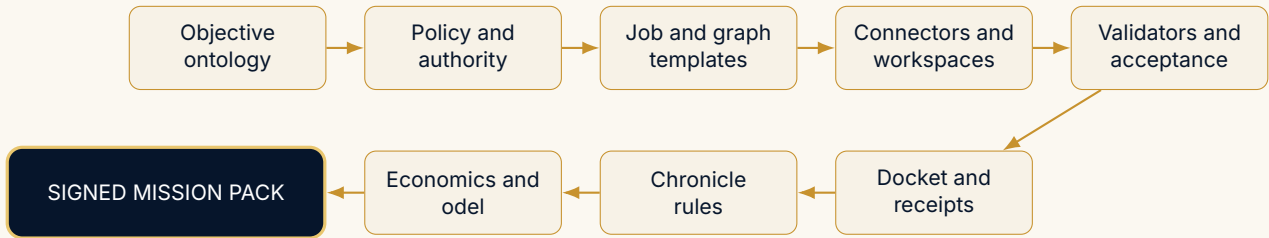
Mission Pack

A Mission Pack P is a tuple

$$P = (\mathcal{O}, \Pi, \mathcal{C}, \mathcal{J}, \mathcal{G}, \mathcal{V}, \mathcal{D}, \mathcal{M}, \mathcal{E}),$$

where $\mathcal{O}$ is the objective ontology, Π policy set, $\mathcal{C}$ connector set, $\mathcal{J}$ AGI Job templates, $\mathcal{G}$ graph patterns, $\mathcal{V}$ validators and acceptance tests, $\mathcal{D}$ Evidence Docket schema, $\mathcal{M}$ Chronicle admission rules, and $\mathcal{E}$ pricing and value model.

12.2 Foundry pipeline



12.3 Canonical Mission Pack manifest

Each release should declare:

- pack identifier, version, publisher, and integrity digest;
- supported objectives and explicit exclusions;
- required customer facts and permitted data classes;
- models, tools, connectors, providers, and minimum versions;
- agent identities, roles, and authority templates;
- graph patterns and independence requirements;
- proof obligations, tests, and acceptance criteria;
- professional-authority dependencies;

- risk class, operating modes, and fail-closed rules;
- Evidence Docket and Receipt Pack schemas;
- Chronicle rights, scope, expiry, and revocation;
- commercial terms and value-measurement protocol.

12.4 Developer, validator, and connector roles

The ecosystem should distinguish three contributions:

1. **Developers** create packs, connectors, policies, and evaluators;
2. **Validators** challenge claims, execute independent tests, or provide responsible authority;
3. **Operators** perform approved work across software, professional, organizational, or physical systems.

One entity may hold several roles in low-stakes settings, but material missions should expose conflicts and preserve independence where required. Contemporary agent frameworks now expose loops, plugins, traces, tool controls, and human-in-the-loop mechanisms; the Mission Pack Foundry raises those primitives into a versioned institutional product. (Google Developers 2026; OpenAI 2026a)

12.5 Open standard and proprietary moat

The following should be openly specified to accelerate category formation: Mission Contract schema; Authority Envelope; PCAR; Evidence Docket; Transaction Receipt; Chronicle admission state; benchmark definitions. The following remain proprietary where rights permit: customer-specific evidence graph; outcome data; scoring calibration; provider-performance history; Chronicle capabilities; model-routing policy; distribution; and protected workspaces.

12.6 Mission Pack economics

Mission Packs enable a transition from founder-dependent services to productized verified outcomes. Revenue may combine platform access, pack subscriptions, verified-job usage, Proof Office or Expansion Office retainers, private runtimes, validators, and strategic deployments. Provider compensation must remain disclosed and structurally separate from independent scenario ranking or validation.

A Mission Pack may package a method. It may not package unexamined authority, customer secrets, hidden conflicts, or unverified historical output. Every reusable element remains subject to rights, scope, currentness, and revocation.

CHAPTER 13

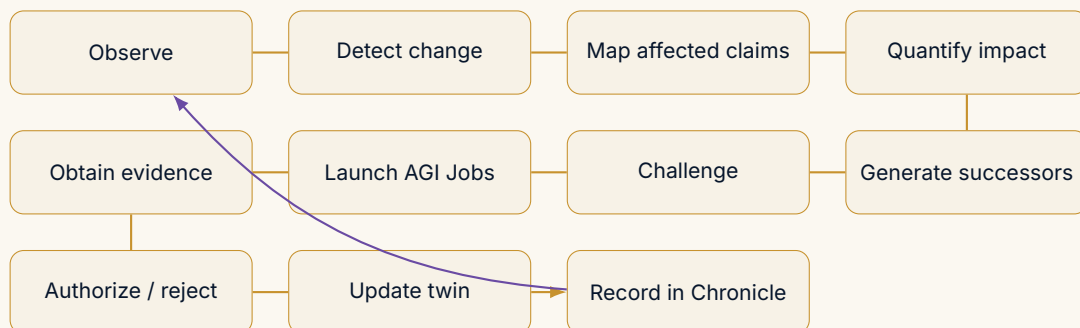
Global Expansion and Productive Capacity

13.1 The recurring control plane

A launch is episodic; the institution is not. The Global Expansion Office exists because laws, programmes, sources, people, capital, ownership, contracts, risks, infrastructure, and evidence continually change. The recurring question is:

Does the current institutional architecture remain the strongest defensible architecture for the objective?

13.2 Operating loop



13.3 Operating cadence

13.4 Scenario switchboard

The Office keeps explicit triggers rather than relying on vague strategic review. Examples include:

- if revenue in a market exceeds a threshold, activate entity and permanent-establishment analysis;
- if foreign hiring exceeds a cost or headcount threshold, compare EOR, branch, subsidiary, or partner economics;
- if awarded support restricts intellectual-property movement, reopen capability ownership;
- if compute demand reaches a power threshold, activate an infrastructure SPV mission;
- if evidence confidence falls below the mission floor, suspend the affected transaction;
- if a sanctions or export-control event changes the reachable market, regenerate the architecture.

Table 13.1: The Expansion Office cadence.

Cadence	Institutional function
Continuous	Source monitoring, evidence expiry, deadline detection, transaction blocking, identity and authority changes, material-event alerts.
Weekly	Change and Opportunity Brief; urgent Proof Debt; active AGI Jobs; decisions required.
Monthly	Funding, fiscal, capital, entity, hiring, risk, and value committee.
Quarterly	Full architecture auction; successor scenarios; concentration and reversibility stress test.
Annual	Substance, transfer pricing, capability/IP, security, insurance, governance, and complete architecture review.
Material event	Immediate reconstitution after financing, acquisition, regulatory change, large customer, ownership change, infrastructure commitment, or incident.

13.5 Compute and infrastructure

The productive-capacity mission graph is:

Demand → Jurisdiction → Site → Energy → Grid → Land
 → Permits → Compute → Capital → Vendors
 → Commissioning → Output.

For project p define risk-adjusted productive value:

$$\begin{aligned}
 PV(p) = & E[\text{productive output}] + \text{public support} + \text{strategic access} \\
 & - \text{capital cost} - \text{power and operating cost} - \text{delay} - \text{regulatory risk} \\
 & - \text{execution reserve} - \text{unwind cost}.
 \end{aligned}$$

GoalOS should not place speculative physical assets on MONTREAL.AI's balance sheet merely because an architecture appears attractive. Customer commitments, permits, authority, financing, and milestone evidence should precede irreversible capital deployment.

13.6 Capital classification

Capital is classified separately as:

Available ≠ Eligible ≠ Approved ≠ Committed ≠ Collected ≠ Audited.

No unsigned investment, unawarded grant, unassessed credit, unapproved loan, speculative customer contract, or modelled infrastructure value enters base runway or productive-capacity reinvestment.

13.7 The productive-intelligence flywheel

Verified Intelligence → Bounded Work → Accepted Contribution
→ Recurring Value → Capital
→ Productive Capacity → Harder Missions.

The flywheel is valuable only if Chronicle preserves the causal relationship between what was done and what value was actually realized.

Part IV

Governance, Economics, and Proof

CHAPTER 14

Governance, Security, Rights, and Legal Boundary

14.1 The authority boundary

GoalOS supplies intelligence, orchestration, evidence, and transaction control. It does not manufacture professional, fiduciary, governmental, customer, or board authority. Reserved acts remain with licensed and responsible parties.

GoalOS may:

- identify issues and decision dependencies;
- retrieve and compare sources;
- design tests and evidence requirements;
- prepare dossiers, workpapers, and draft instruments;
- coordinate professionals and providers;
- monitor completeness and currentness;
- block incomplete execution;
- record decisions, receipts, and outcomes.

GoalOS may not treat a model score as a legal opinion, tax conclusion, audit, investment recommendation, regulated filing, fiduciary decision, government approval, or customer acceptance.

14.2 Agent identity and authorization

The emerging agent ecosystem makes identity and authorization foundational. NIST's 2026 AI Agent Standards Initiative emphasizes secure interoperability, agent security, identity infrastructure, and authorization; its related concept work explicitly addresses the risk of giving agents access to diverse data, tools, and applications. (National Institute of Standards and Technology 2026a; National Institute of Standards and Technology 2026c)

GoalOS should implement:

- unique human, agent, service, and organization identities;
- principal-to-agent delegation records;
- least-privilege data and tool scopes;
- time, budget, transaction, communication, and environment ceilings;
- re-authentication for material actions;
- event-level tracing and evidence linkage;
- explicit pause, revoke, and kill functions;
- incident-preserving logs;
- ownership and custody separation for secrets and wallets.

14.3 Trajectory security

Security must evaluate the full agent trajectory rather than isolated prompts. Material threats include:

- opportunistic credential discovery;
- reconstruction of fragmented secrets;
- unauthorized use of private submissions;
- sandbox escape;
- external communication outside the Mission Constitution;
- tool-combination attacks;
- goal displacement and reward manipulation;
- obfuscated actions that evade simple scanners;
- technically successful but institutionally inadmissible outcomes.

Each threat should map to a control object:

Threat → Precondition → Detection → Prevention → Containment → Evidence → Recovery → Owner.

14.4 Data, privacy, and confidential work

The public interface should reject confidential, privileged, personal, regulated, classified, credential, vulnerability, or unauthorized third-party information. Sensitive missions require a separately reviewed environment with:

- customer agreement and operator identity;
- data map and purpose limitation;
- processor and subprocessor assessment;
- access control and logging;
- retention, deletion, and legal-hold rules;
- encryption and key governance;
- incident response;
- cross-border and sector-specific review;
- model and tool data-use boundaries.

NIST's Generative AI Profile extends the AI RMF with lifecycle risk considerations for generative systems, while current EU AI rules require role-specific implementation and, in relevant cases, transparency, human oversight, monitoring, and documentation. (Autio et al. 2024; European Parliament and Council of the European Union 2024; European Commission 2026a)

14.5 Rights and reuse

The value of Chronicle depends on lawful reuse. Every admitted capability should have a rights record covering:

- customer-owned inputs and outputs;

- MONTREAL.AI platform rights;
- third-party source and model restrictions;
- confidential and privileged material;
- publication rights;
- derivative-work and training rights;
- geographic, customer, field-of-use, and duration limits;
- revocation and deletion obligations.

The safe default is narrow reuse. A commercially valuable method can be admitted without storing the customer's secrets, proprietary data, or confidential facts.

14.6 Conflict separation

Provider compensation, referral economics, marketplace fees, or transaction participation must not influence independent scenario ranking. The institutional architecture should remain:

Independent Scenario Engine → Recommendation → Separate Provider Layer
→ Disclosed Relationship → Customer Choice.

14.7 Governance mappings

GoalOS can map its controls to external frameworks without claiming certification or automatic compliance. NIST AI RMF provides a voluntary lifecycle structure for governance, context mapping, measurement, and management. The EU AI Act applies role- and risk-specific obligations. OpenAI's Agents SDK illustrates implementable guardrail and trace primitives for agent workflows. (Tabassi 2023; European Parliament and Council of the European Union 2024; OpenAI 2026b)

Table 14.1: Illustrative governance mapping.

External concern	GoalOS control	Evidence object
Risk governance	Mission Constitution, role classification, control planes	Approved Constitution and control map
Identity and authorization	Agent identity, principal delegation, authority envelope	Identity and authority receipts
Traceability	Job ledger, tool traces, source lineage	Evidence Docket and trace bundle
Human oversight	Acceptance gate, pause, override, escalation	Acceptance and intervention record
Data governance	Data scope, retention, rights, protected environment	Data map and rights register
Monitoring	Expansion Office, expiry, incident alerts	Change radar and monitoring ledger
Corrective action	Rollback, revocation, Chronicle repair	Revocation and recovery receipts

Actual functions, users, data, rights, control, communications, counterparties, jurisdictions, and money flows determine legal treatment. Private labels and disclaimers do not create immunity from mandatory law.

Security, Threat Model, and Trajectory Assurance

15.1 Security objective

The security objective is not merely to prevent model hallucination. It is to preserve the integrity of the complete mission trajectory: objective, identity, authority, data, tools, communications, memory, evidence, validators, transactions, and rollback. Agentic systems expand the attack surface through tool use, dynamic planning, long-term memory, external content, multi-agent delegation, and machine-speed action. NIST's agent initiative emphasizes identity, authorization, interoperability, and security; OWASP correspondingly treats agentic memory, tools, goals, and autonomy as distinct attack surfaces. (National Institute of Standards and Technology 2026a; National Institute of Standards and Technology 2026c; OWASP GenAI Security Project 2026)

15.2 Threat classes

Threat	Failure path	GoalOS control
Objective injection	External content silently changes mission purpose or success criteria	Versioned Constitution; immutable objective reference; change approval
Prompt and tool injection	Untrusted data instructs an agent to misuse tools, reveal data, or bypass policy	Content isolation; tool mediation; least privilege; output filtering; independent tests
Authority laundering	One agent delegates powers it does not possess or combines harmless permissions into harmful capability	Principal-bound envelopes; non-transitive authority; delegation ledger; expiry
Credential exfiltration	Model or tool gains reusable secrets or tokens	Ephemeral credentials; secret broker; scoped tokens; network egress controls; redaction
Memory poisoning	False or malicious output becomes persistent context	Chronicle admission; provenance; rights; fresh-successor tests; revocation
Validator capture	Producer influences the evaluator, metric, context, or acceptance threshold	Preselection; blind review; separate context; diverse models or professionals; audit trail
Evidence fabrication	Synthetic or altered receipts are presented as world evidence	Signed sources; direct retrieval; external confirmation; cryptographic integrity; contradiction search
Model or dependency drift	A model, connector, package, or policy changes after validation	Version pinning; regression tests; revalidation; source and model expiry

Runaway cost or action	Long-horizon loop consumes resources or executes repeated external actions	Budget, duration, transaction, rate, and recursion ceilings; stop controller
Correlated monoculture	Multiple agents share the same model, source, or hidden assumption	Architecture diversity; independent source routes; counterfactuals; multi-model or human checks
Insider or operator abuse	Authorized person manipulates policy, evidence, or acceptance	Separation of duties; dual control; immutable logs; conflict disclosure; post-event review
Physical or critical-system harm	Correct software action causes unsafe world effect	Simulation; staged rollout; engineering authority; safety interlocks; project insurance; emergency stop

15.3 Trajectory assurance

A trajectory $\tau = (a_1, \dots, a_n)$ is safe only if each action is individually admissible and the composition remains within mission limits:

$$\text{Safe}(\tau) = \prod_{k=1}^n \text{Permitted}(a_k) \cdot \mathbf{1}[\text{State}(\tau_k) \in \mathcal{S}_{\text{allowed}}].$$

This prevents a sequence of individually permitted steps from producing an unauthorized aggregate effect. Trajectory monitors should track cumulative spend, data access, external communications, privilege changes, irreversible commitments, and deviation from the accepted plan.

15.4 Security control architecture

1. **Identity:** unique human, agent, service, and provider identities; no shared anonymous authority.
2. **Least privilege:** task-specific credentials and tools; deny by default.
3. **Sandboxing:** isolated file, code, browser, network, and execution environments.
4. **Policy mediation:** every material tool call evaluated against current policy and authority.
5. **Observability:** trace of prompts, plans, tool calls, results, costs, policy decisions, and overrides.
6. **Independent verification:** separate tests, sources, or responsible parties for material claims.
7. **Recovery:** pause, revoke, rollback, credential rotation, and Evidence Docket incident record.
8. **Resilience:** alternate models, providers, regions, keys, and continuity procedures.

15.5 Human oversight as designed authority

Human oversight should not mean a fatigued person clicking approve. It should identify the responsible role, the evidence required, the decision being made, the time available, the conflicts present, and the effect of approval. High-quality oversight is an institutional interface, not a ceremonial checkbox.

15.6 Security evidence

Security claims themselves require proof: architecture diagrams; configuration snapshots; test results; vulnerability findings; remediation evidence; access reviews; incident exercises; dependency inventories; backup restoration tests; and external assessment where stakes justify it. A security badge without such evidence is theatre.

Security is the preservation of mission purpose, authority, evidence, and recovery across the entire trajectory. A powerful model inside an ungoverned trajectory is not Sovereign Intelligence.

Standards and Institutional Interoperability

16.1 Mapping rather than claiming certification

GoalOS should map its objects and controls to recognized frameworks while avoiding unsupported certification claims. Mappings facilitate procurement, control design, and professional review; they do not establish compliance without deployment-specific evidence. This includes critical-infrastructure profiles where failures have physical, economic, or public consequences. (National Institute of Standards and Technology 2026b; International Organization for Standardization 2023)

Table 16.1: Selected framework mappings.

Framework or source	Relevant concern	GoalOS institutional object
NIST AI RMF 1.0 and GAI Profile	Governance, mapping, measurement, management, documentation, testing	Mission Constitution, risk register, Evidence Docket, GoalOSBench, incident and monitoring loops
NIST AI Agent Standards Initiative	Secure and interoperable agents; identity and authorization	Agent identity, authority envelope, delegation ledger, tool and system scope
NIST AI Resource Center / TEVV	Testing, evaluation, verification, and validation	Independent validators, proof obligations, benchmark protocols, acceptance records
W3C PROV	Attributable provenance among entities, activities, and agents	PCAR, Evidence Docket, causal trace, outcome graph
ISO/IEC 42001	AI management-system governance and continuous improvement	Governance system, roles, policies, lifecycle, monitoring, corrective action
EU AI Act	Role- and risk-specific obligations; transparency; high-risk controls	Role classification, operating mode, AI interaction notice, data and logging, responsible deployer controls
OWASP Agentic Security guidance	Tool misuse, memory/context poisoning, identity, supply chain, excessive agency	Trajectory assurance, sandboxing, memory firewall, policy mediation, red teaming

16.2 Current transparency posture

The European Commission’s July 2026 guidance clarifies Article 50 transparency obligations that apply from 2 August 2026, including informing people when they interact with certain AI systems and marking specified generated or manipulated content. GoalOS deployments should therefore expose the relevant provider, deployer, model, purpose, and AI-interaction state rather than relying on hidden automation. (European Commission 2026b)

16.3 Agent identity and non-repudiation

NIST’s 2026 concept work highlights identification, authorization, auditing, non-repudiation, and prompt-injection controls for software agents. GoalOS extends the identity question from “which agent called the tool?” to “which

institutional principal delegated what authority, for which objective, under which policy, and with what acceptance consequence?” (National Institute of Standards and Technology 2026c)

16.4 Interoperable evidence

GoalOS evidence should be portable across models, vendors, validators, and organizations. A minimum evidence exchange format should preserve:

- stable identifiers and version history;
- source and retrieval metadata;
- producing and validating identities;
- claim, method, scope, and contradiction state;
- signatures or integrity digests where appropriate;
- authority, acceptance, expiry, rights, and revocation;
- links to raw artifacts without forcing public disclosure.

16.5 Policy-as-code and professional judgment

Rules suitable for deterministic expression should become policy-as-code: spending limits, data boundaries, required receipts, separation-of-duty rules, deadlines, and tool permissions. Principles requiring interpretation should remain visible as professional or organizational judgments rather than being disguised as deterministic rules.

16.6 Sovereign Intelligence as an interoperability profile

SII-001 can serve as an interoperability profile above agent protocols. A conforming Mission Pack can be implemented on different agent frameworks if it preserves the same constitutional semantics: principal, objective, authority, typed work, proof, validation, acceptance, memory, and revocation.

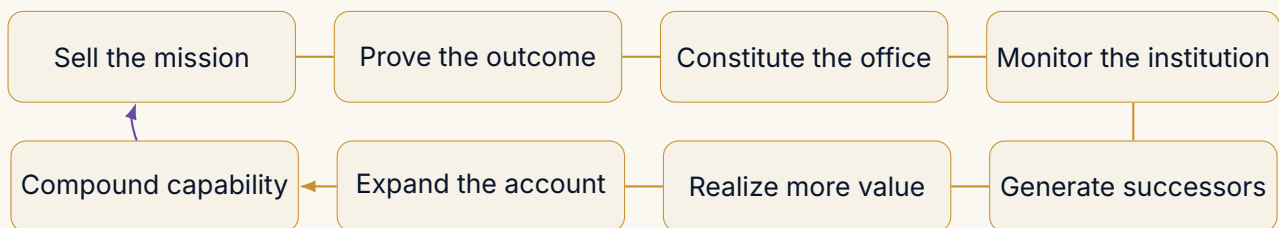
Alignment is not certification. Documentation is not operation. A control exists only when the deployment can produce current evidence that the control is functioning for the actual mission.

Economics and Commercial Strategy

17.1 The business model

GoalOS combines professional-services pricing, software-like marginal cost, recurring institutional control, proprietary outcome data, and network effects. The near-term product should be a high-value, software-enabled mission delivered with a small accountable team and mission-specific external professionals. Over time, repeated mission logic becomes the Proof & Authority platform, reusable Mission Institutions, protected APIs, and a rights-controlled outcome graph.

17.2 Commercial sequence



The initial wedge is Global Launch because MONTREAL.AI already possesses the category thesis, GJFFI kernel, scenario architecture, evidence model, and commercial narrative. The first recurring product is the Global Expansion Office. The first high-frequency adjacent product is Service & Operations. Chat is the interface. Proof & Authority is the shared substrate. Chronicle and the Verified Outcome Graph are the moat.

17.3 Illustrative product architecture

The following prices are management hypotheses, not current offers or forecasts.

17.4 Customer value gate

For total customer spend F , a default qualification floor is:

$$V_{12m}^{\text{conservative}} \geq 4F.$$

The preferred target is eight times total spend. Value may include verified support collected, dilution avoided, financing cost reduced, launch time compressed, revenue accelerated, unnecessary entities avoided, risks retired, or infrastructure economics improved. Modelled value remains separate from realized value.

Table 17.1: Illustrative institutional pricing architecture.

Offer	Illustrative price	Commercial role
Global Launch Diagnostic	US\$15,000	Qualify complexity, value at stake, and mission fit.
Global Architecture Mission	From US\$125,000	Scenario frontier, function allocation, entity, capital, evidence, and first 100 days.
Verified Full Launch Mission	From US\$350,000	Governed implementation across selected modules.
Capital & Public-Support Mission	From US\$250,000	Eligibility, stacking, capital sequence, dossiers, and gates.
Infrastructure Launch Mission	From US\$1,000,000	Sites, power, compute, permits, SPVs, vendors, and finance.
Global Expansion Office	From US\$500,000 annually	Continuous monitoring, re-optimization, and successor missions.
Portfolio / Sovereign Office	From US\$2,000,000 annually	Multi-entity, portfolio, regional, or sovereign mandate.
Private GoalOS deployment	Setup plus annual minimum	Customer workspace, connectors, controls, and intelligence access.

17.5 AI-native cost structure

AI compresses routine cognition, but the remaining cost floor buys legitimacy: responsible authority, independent proof, enterprise trust, security, insurance, current data rights, counterparty acceptance, and real-world execution.

The optimal operating constitution is:

- own intelligence, orchestration, proof, governance, customer relationship, Chronicle, and outcome graph;
- rent or route reserved professional authority, local filings, physical execution, and regulated intermediation;
- require customer deposits and prepaid external costs;
- add permanent employees only when agents or on-demand providers cannot reliably perform the role;
- maintain independent assurance and risk reserves even when production costs fall.

17.6 Revenue scenarios

The scale thesis comes from recurring institutional control rather than formation fees. Strategic scenarios might include: These are architecture tests, not promises. They indicate the recurring relationships, product standardization,

Table 17.2: Illustrative operating configurations; not forecasts.

Stage	Relationships	Blended annual revenue	Illustrative revenue
Product proof	20	US\$250K	US\$5M
Repeatable portfolio	250	US\$400K	US\$100M
Decacorn operating zone	1,500	US\$650K	US\$975M
Global action layer	4,000	US\$1M	US\$4B
Megacorn candidate	7,500	US\$1M plus usage/network	US\$7.5B+

distribution, trust, and outcome graph required at each scale. (MONTREAL.AI Research 2026c)

17.7 The durable moat

The interface, prompts, models, and public sources are not the final moat. The moat is the rights-controlled graph connecting:

Objective → Architectures → Evidence → Professional Opinions
→ Actions → Outcomes → Admission → Successor Performance.

Over time the institution can learn which routes actually work, which evidence is decision-changing, which providers are reliable, which structures survive audit, which failure patterns recur, and which capabilities transfer safely.

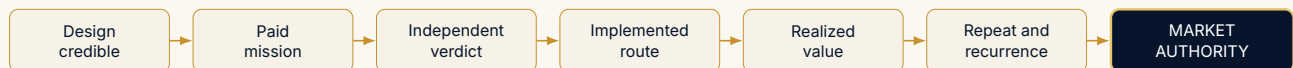
Never price the number of pages, model calls, tokens, or agent steps. Price the consequential decision, the evidence required to support it, the execution risk governed, and the verified value placed within reach.

From Proof to Market Authority

18.1 The commercial proof problem

A compelling paper and sophisticated prototype establish design credibility, not market authority. Market authority emerges when unrelated customers pay for comparable missions, material recommendations survive independent review, actions are implemented, economic or institutional value is realized, customers repurchase, and the resulting capability improves a fresh mission.

18.2 The proof ladder



18.3 Founding Commercial Proof Run

The first standardized offer should remain narrow enough for matched evidence. A suitable initial mission is the GoalOS Global AI Company Launch Mission. Three unrelated customers receive substantially comparable scope, preregistered acceptance criteria, independent review, a full Evidence Docket, and a public-safe Receipt Pack where authorized.

Table 18.1: Minimum proof gates and stronger authority targets.

Evidence dimension	Minimum proof	Market-authority target
Unrelated paying customers	3	10+
Comparable completed missions	3	10+
Independently verified outcomes	1	5+
Realized or activated route	1	5+
Repeat or expansion purchase	1	50%+ of eligible customers
Global Expansion Office	1	30%+ conversion; 90%+ renewal
Rights-cleared Chronicle capability	1	100+ scoped capabilities
Fresh-successor improvement	1	10+ preregistered wins
Material trust breaches	0	0

18.4 Verified customer value

Customer value must be classified rather than asserted:

Identified → Validated → Authorized → Implemented → Realized → Independently Verified.

Only the final states support strong public claims. Value may include collected support, dilution avoided, financing cost reduced, revenue accelerated, time compressed, cost reduced, risk retired, or productive capacity activated. Coun-

terfactual value claims must disclose their method and uncertainty.

18.5 The Global Expansion Office conversion

The Launch Mission creates the initial architecture. The Expansion Office creates recurrence by monitoring evidence expiry, regulation, funding, ownership, employment, capital, infrastructure, and superior successor scenarios. The conversion offer should be presented before the mission ends: the architecture has been designed; the remaining question is whether it remains valid as the world changes.

18.6 AI-native operating economics

GoalOS should own the intelligence, orchestration, proof, customer relationship, Chronicle, and outcome graph. Customer-specific licensed opinions, filings, physical implementation, and regulated intermediary activity should generally be separately scoped and customer-funded. Falling model cost should fund deeper search, independent validation, faster delivery, stronger margins, and productive reinvestment – not force prices toward token cost.

18.7 Institutional restraint as a commercial asset

The firm should decline missions whose facts are unavailable, authority is ambiguous, economics cannot justify the fee, professional dependencies are refused, or the customer seeks evasion, fabricated evidence, or unsupported public claims. A visible ability to say “not yet,” “blocked,” or “declined” is part of the product.

The unit of commercial proof is not the report. It is the paid, accepted, independently supportable, implemented, and realized institutional outcome – plus the evidence showing the complete cost, limitation, and residual-risk denominator.

SovereignBench Ω : Empirical Authority

19.1 Why a benchmark is necessary

The architecture is ambitious, but architecture is not proof. GoalOS should earn market and scientific authority through preregistered comparisons with credible alternatives. Published metrics can become targets and lose their connection to the underlying objective. GoalOSBench therefore preregisters multiple measures of quality, value, evidence, cost, risk, reviewer burden, and institutional regret rather than optimizing one headline score. (Goodhart 1975; Campbell 1979; National Institute of Standards and Technology 2026d)

19.2 Benchmark arms

A matched benchmark should compare:

1. a multidisciplinary expert team;
2. a static frontier-model report;
3. an ungated agentic workflow;
4. a proof-governed GoalOS mission;
5. where relevant, a hybrid GoalOS-plus-professional team.

19.3 Primary metrics

Table 19.1: SovereignBench evaluation dimensions.

Dimension	Measurement
Decision quality	Blind expert assessment of architecture, completeness, and trade-offs.
Realized value	Customer value actually implemented and verified.
Cost	Full denominator: models, data, people, providers, delays, failures, and capital.
Elapsed time	Time from mission constitution to accepted outcome.
Evidence completeness	Coverage, provenance, freshness, contradictions, and traceability.
Reviewer burden	Human and professional hours required.
Transaction safety	Prevented incomplete, contradictory, or unauthorized actions.
Chronicle precision	False admission and false rejection of reusable capability.
Successor transfer	Improvement on a fresh mission without governance regression.
Institutional regret	Difference between chosen architecture and best realized alternative.
Trust and incidents	Unauthorized actions, material breaches, disputes, or misleading claims.

19.4 Value-of-proof experiments

For each mission, preselect unresolved facts with different predicted decision impact. Measure whether resolving high-VoP evidence changes the recommended architecture, transaction gate, or customer acceptance more often than resolving low-VoP evidence.

19.5 Verifier independence experiments

Compare:

- same-model self-critique;
- fresh context with same model;
- different model family;
- independent source retrieval;
- blind human or professional review;
- realized outcome evidence.

The objective is not to prove that one verifier is universally best, but to estimate correlated-failure risk and choose verification proportionate to stakes.

19.6 Commercial authority gate

A practical threshold for claiming early commercial authority is:

Table 19.2: Minimum external proof programme.

Gate	Minimum
Unrelated paying customers	3
Substantially comparable missions	3
Independently verified customer outcome	1
Funding, fiscal, financing, market, or operational route realized	1
Completed activation	1
Separate repeat or expansion purchase	1
Active Global Expansion Office	1
Rights-cleared Chronicle capability	1
Fresh successor improvement	1
Material trust breaches	0

A stronger institutional-authority threshold would require at least ten paying customers, five independently verified realized outcomes, three recurring Offices, a high-frequency Mission Institution in production, multiple benchmark wins, and zero material unauthorized actions.

19.7 Public-safe Receipt Packs

A publishable Receipt Pack should include:

- paid scope and Mission Constitution;

- preregistered acceptance criteria;
- evidence summary and source dates;
- independent verdict;
- full cost denominator;
- implementation result;
- verified value and limitations;
- customer acceptance;
- Chronicle disposition;
- exact claim authorized for publication.

GoalOS should never use its own architecture as proof of its superiority. It should use externally inspectable outcomes, controlled comparisons, and the complete denominator.

CHAPTER 20

Implementation Roadmap

20.1 Phase 0: constitutional baseline

Before commercial scale, freeze:

- canonical naming and hierarchy;
- Mission Constitution schema;
- authority-envelope schema;
- Evidence Docket schema;
- Transaction Gate receipts;
- Chronicle admission rules;
- public claim boundary;
- security and incident policy;
- commercial qualification and pricing rules.

20.2 Days 0--100: paid proof

1. Sell three Founding Global Launch Missions.
2. Collect at least 50% before work begins.
3. Run substantially comparable missions with preregistered acceptance criteria.
4. Produce independent review and one public-safe Receipt Pack.
5. Realize at least one material route or activation.
6. Secure one repeat purchase and one Global Expansion Office.
7. Preserve one rights-cleared Chronicle capability and test it on a fresh mission.

20.3 Months 4--12: repeatability

- standardize the Global Launch Mission;
- reach ten paying customers;
- operate at least three Expansion Offices;
- deploy live source and evidence graph operations;
- formalize a professional and validator network;
- pilot Service & Operations as the high-frequency adjacent Mission Institution;
- run the first controlled SovereignBench comparison;
- reduce founder intervention in routine delivery.

20.4 Months 12--24: platformization

- productize GoalOS Proof & Authority Ω;
- introduce enterprise identity, role-based workspaces, connectors, and private deployment;
- operate customer-specific Chronicle and Verified Outcome Graphs;
- add Customer & Revenue, Finance, and Professional Intelligence packs;
- publish stable schemas and interoperability specifications;
- shift distribution beyond founder-led explanation through partners, cases, and GoalOS Chat;
- select the first vertical Industry Cloud.

20.5 Days 750--1000: productive capacity frontier

After recurring commercial authority exists:

- launch Compute & Infrastructure missions through customer-funded projects or SPVs;
- integrate power, compute, site, permit, capital, vendor, and commissioning graphs;
- establish capital classification and milestone gates;
- connect verified contribution to productive reinvestment;
- publish benchmark evidence for successor performance and institutional value.

20.6 Operating team

The proof-stage institution can remain small:

Table 20.1: Initial human nucleus.

Function	Initial structure
Institutional architecture and final authority	Founder / principal
Commercial execution	One senior operator
Mission and Evidence Docket operations	One operator
Agentic engineering and security	One strong engineer
Customer and partner co-ordination	One operator
Legal, tax, funding, sector, and validation authority	External network, customer-funded where mission-specific

The human organization should remain small as capability rises. Permanent hiring should follow recurring contracted demand, not precede it.

20.7 Standards strategy

Open specifications should cover Mission Contracts, authority envelopes, AGI Job identity, Evidence Dockets, Transaction Receipts, Chronicle admission, successor testing, and Verified Value reporting. Proprietary assets should include the execution engine, customer outcome graph, scoring calibration, provider history, protected capabilities, benchmark results, and distribution.

The next phase is not more conceptual breadth. It is the disciplined conversion of the architecture into paid missions, verified outcomes, recurring control, high-frequency use, and a rights-controlled outcome graph.

Risks, Failure Modes, and Controls

Table 21.1: Principal risks and institutional controls.

Risk	Failure mode	Control
False certainty	Candidate intelligence is presented as binding fact or authority	Current official sources, exact facts, source dates, professional receipt, claim boundary, expiry, transaction gate.
Capability-authority collapse	A powerful agent silently gains permission	Identity-bound authority envelopes, least privilege, re-authentication, spend and communication ceilings, revocation.
Self-verification	Producer assertion is treated as independent proof	Fresh context, separate retrieval, different evaluators, blind review, executable tests, realized receipts.
Memory over-reach	One successful case becomes universal doctrine	Narrow scope, rights, versioning, failure signatures, expiry, revocation, fresh successor test.
Metric gaming	System optimizes a proxy rather than institutional value	Complete denominator, multiple objectives, adversarial challenge, acceptance, realized outcome measurement.
Scenario monoculture	One model or adviser anchors on one architecture	Diverse scenario generation, Pareto frontier, counterfactuals, red team, stop and switch conditions.
Professional conflict	Provider compensation biases ranking	Separate recommendation and provider layers, disclosure, independent review, customer selection.
Data leakage	Sensitive information enters a public or unapproved system	No-secrets gate, protected environment, minimization, access logs, retention, deletion, incident response.
Source decay	A previously valid claim becomes stale	Evidence expiry, live monitoring, affected-claim mapping, Proof Debt, automatic gate failure.
Grant or capital theatre	Unawarded or unsigned resources enter runway	Available / eligible / approved / committed / collected / audited classification.
Shell optimization	Formal structure outruns substance and purpose	Business-purpose memo, people and risk map, entity-cost penalty, substance review, professional approval.
Unbounded mission scope	Every exception becomes bespoke founder rescue	Frozen Constitution, change orders, explicit decline and escalation routes, productized Mission Institutions.
Infrastructure overinvestment	Assets are funded before demand, permits, or finance	Customer commitments, project SPV, milestone finance, stop conditions, unwind model.
Automation overclaim	Agentic language exceeds actual capability	Capability-state disclosure, tested workflows, no claim without Evidence Docket and public-safe receipt.
Vendor concentration	One model, cloud, source, or individual becomes a single point of failure	Multi-provider strategy, version pinning, backups, offline evidence, continuity tests, succession.
Liability underpricing	High-value decisions are sold as cheap software	Value-based pricing, scope control, insurance, reserves, qualified authority, pass-through external costs.

Risk	Failure mode	Control
Distribution bottleneck	Category depends entirely on founder explanation	GoalOS Chat, case studies, partner channels, standards, productized offers, benchmark evidence.

21.1 Residual risk

No governance system eliminates uncertainty. The objective is to expose, price, route, reduce, accept, or reject risk rather than hide it behind model confidence. A fail-closed institution will sometimes decline or delay attractive actions. That is a feature when the missing evidence could change the decision.

21.2 Anti-theatre controls

The institution should prohibit:

- presenting discovered funding as awarded cash;
- presenting a simulation as an approval or professional opinion;
- counting self-generated memory as demonstrated improvement;
- using benchmark metrics selected after observing results;
- omitting failures, retries, professional interventions, or delayed effects from cost and performance claims;
- marketing a static source snapshot as permanently current;
- calling delegated or approved control direct ownership;
- allowing commercial referral economics to influence ranking without disclosure.

A proof-governed institution must be able to say: unresolved, blocked, expired, rejected, revoked, or not worth doing.

Part V

Conclusion and Constitutional Appendices

Conclusion: Constituting the Verified Action Institution

GoalOS Ω – Sovereign Intelligence is a proposal to reorganize how institutions use increasingly capable intelligence. It begins from a simple observation: the cost of producing plausible cognitive output is falling faster than the cost of establishing truth, authority, rights, trust, and real-world execution. The resulting opportunity is not merely to automate more tasks. It is to build the institution that determines what intelligence may do, what evidence it must return, who may accept the result, what may enter memory, and what verified value may finance next.

The canonical system is one institution:

- **GoalOS Chat Ω** receives the objective and makes complexity accessible;
- **GoalOS Proof & Authority Ω** compiles identity, policy, mission graphs, bounded work, evidence, validation, acceptance, and rollback;
- **Mission Institutions** complete high-value objectives across Global Launch, Service & Operations, Customer & Revenue, Finance, Professional Intelligence, and Industry Clouds;
- **Global Expansion Office Ω** keeps the accepted architecture current;
- **Compute & Infrastructure Ω** converts verified value into productive capacity;
- **Chronicle + Verified Outcome Graph** governs what may survive and shape the next mission.

The architecture rejects four shortcuts. It rejects the idea that capability implies permission. It rejects the idea that completion implies correctness. It rejects the idea that one successful output implies reusable truth. And it rejects the idea that booked or modelled value implies productive surplus.

The strongest commercial strategy is therefore sequential:

Sell the mission → Prove the outcome → Constitute the office
→ Monitor the institution → Generate successors
→ Compound verified capability.

The strongest scientific strategy is comparative: run preregistered missions, record the full denominator, separate producer and verifier, test fresh successors, and publish only claims authorized by evidence. The strongest governance strategy is fail-closed: when a source expires, authority is ambiguous, evidence is incomplete, or rights are absent, the relevant action or memory path remains blocked.

No consequential objective, no mission. No exact facts, no decision-grade claim. No explicit authority, no external action. No independent evidence, no verified outcome. No acceptance, no institutional state change. No rights-cleared Chronicle admission, no future influence. No fresh comparison, no improvement claim. No verified contribution, no productive reinvestment.

GoalOS Ω

SOVEREIGN INTELLIGENCE

The Verified Action Institution

Intelligence that acts. Proof that governs. Capability that compounds.

Founder's Constitutional Declaration

GoalOS Ω shall be developed as an institution whose intelligence remains subordinate to explicit purpose, bounded authority, attributable evidence, responsible acceptance, selective memory, and revocable governance.

The institution may become more capable. Its authority may not silently expand with its capability. It may search broader spaces, coordinate more workers, invoke more tools, compare more architectures, and operate across more consequential domains. Every increase in power must be matched by an increase in inspectability, evidence, separation of powers, rollback, and accountability.

The institution shall not confuse:

- model confidence with truth;
- technical permission with legitimate authority;
- execution with success;
- producer narrative with independent proof;
- customer enthusiasm with complete acceptance;
- stored output with admissible institutional memory;
- booked revenue with verified contribution;
- reuse with improvement;
- scale with progress.

GoalOS shall preserve the complete denominator of consequential work: not only successful results, but costs, retries, contradictions, failures, professional interventions, unresolved risks, adverse effects, rights limitations, and conditions under which an admitted capability must expire or be revoked.

A person or organization states what it wants to accomplish. GoalOS compiles the mission, determines the authority boundary, launches bounded work, coordinates systems and responsible professionals, requires evidence, obtains accountable acceptance, records the result, preserves only what passed, and uses verified value to expand productive capacity. The next mission is not assumed to be better. It is tested.

This is the controlling direction for the research programme, product portfolio, commercial proof run, standards work, and institutional partnerships described in this paper.

Vincent Boucher

President, MONTREAL.AI & QUEBEC.AI

25 July 2026

Proof before authority. Evidence before memory. Impact before scale.

CHAPTER A

Minimum Mission Contract

No.	Field	Required content
1	Accountable sponsor	Legal or organizational principal, decision owner, escalation owner, and authorized signatory.
2	Objective	Consequential outcome, measurable success state, value at stake, deadline, and non-goals.
3	Scope and exclusions	Functions included, systems included, jurisdictions, prohibited actions, and explicit out-of-scope work.
4	Data and sources	Permitted data classes, official sources, customer sources, retention, confidentiality, and rights.
5	Models and tools	Approved model versions, tools, connectors, environments, communication channels, and fallback providers.
6	Authority	Agent identities, human principals, action classes, spend limits, duration, approvals, and revocation.
7	Budget and value	Mission price, external reserve, cost ceiling, value hypothesis, and payment milestones.
8	Acceptance	Output schema, tests, evaluator, evidence floor, customer acceptance, and dispute process.
9	Professional dependencies	Legal, tax, audit, investment, employment, engineering, safety, or regulatory authority required.
10	Evidence Docket	Required facts, sources, traces, costs, failures, contradictions, receipts, and outcome records.
11	Rights and publication	Customer rights, platform rights, reuse conditions, confidentiality, public-safe claim, and deletion.
12	Stop and rollback	Pause, decline, escalation, rollback, incident response, expiry, and Chronicle eligibility.

CHAPTER B

Authority Envelope Schema

Machine-readable authority envelope

principal: authority source; **agent_id:** delegated worker; **role:** permitted function; **actions:** allowed action types; **data:** allowed data classes and locations; **tools:** approved tools and systems; **communications:** approved external channels; **budget:** monetary, compute, token, and time ceilings; **approvals:** required pre- and post-action signoffs; **evidence:** required return objects; **expiry:** duration and event triggers; **revocation:** immediate stop and rollback procedure.

CHAPTER C

Evidence Docket Schema

Docket object	Required content
Objective record	Sponsor, objective, value at stake, deadline, acceptance criteria, and exclusions.
Fact register	Fact, source, retrieval date, owner, confidence, expiry, and affected decisions.
Assumption ledger	Assumption, sensitivity, evidence required, impact, status, and owner.
Claims matrix	Material claim, supporting evidence, contradiction, verifier, scope, and publication state.
AGI Job ledger	Worker, tools, data, authority, budget, output, trace, failure, and verifier.
Cost ledger	Models, data, professionals, support, travel, filings, capital, delay, rework, and incident cost.
Authority register	Approval, scope, principal, effective date, expiry, delegation chain, and revocation.
Transaction receipts	Signed agreements, filings, payments, awards, permits, deployments, and acceptance.
Outcome record	What happened, verified value, adverse outcomes, delayed effects, and residual risk.
Rights register	Customer rights, platform rights, third-party rights, confidentiality, reuse, and deletion.
Chronicle decision	Admit, admit with limits, repair, reject, revoke, expire, or require fresh test.
Successor result	Baseline, fresh mission, metrics, controls, regressions, verdict, and promotion decision.

CHAPTER D

Canonical Transaction Receipts

Material transactions should select from the following receipt classes:

1. business-purpose memorandum;
2. legal or regulatory memorandum;
3. tax and fiscal memorandum;
4. valuation or economic-support record;
5. transfer-pricing or intercompany analysis;
6. funding and assistance interaction analysis;
7. data, security, rights, or sector-risk assessment;
8. board, customer, fiduciary, or delegated approval;
9. implementation checklist and responsible owner;
10. execution receipt;
11. independent validation and customer acceptance;
12. Chronicle disposition and monitoring trigger.

CHAPTER E

Chronicle Admission Checklist

- Is the capability supported by evidence rather than producer narrative?
- Is the producing mission and objective known?
- Are the source, model, tool, and version lineage recorded?
- Are customer, platform, and third-party rights clear?
- Is the valid scope narrow and explicit?
- Are limitations and failure signatures recorded?
- Is the capability current, or does it have an expiry?
- Was material evidence independently challenged?
- Can the capability be revoked without corrupting dependent missions?
- Has it improved a fresh mission without governance regression?

CHAPTER F

Canonical Operating Doctrine

Discover globally.
Decompose function by function.
Generate competing architectures.
Challenge adversarially.
Resolve locally.
Govern institutionally.
Execute only within authority.
Record the full truth.
Admit only what passed.
Revoke when stale.
Reinvest only verified value.

CHAPTER G

Glossary

Term	Definition
Sovereign Intelligence	Intelligence governed by the institution it serves across objective, identity, authority, evidence, validation, memory, capital, and revocation.
Verified Action Institution	Objective-driven system that compiles work, bounds authority, proves outcomes, obtains acceptance, and selectively compounds capability.
Mission Constitution	Versioned statement of objective, scope, authority, evidence, budget, rights, acceptance, and rollback.
Authority Envelope	Machine-readable boundary on what an agent or worker may do.
AGI Job	Bounded unit of work with identity, scope, tools, budget, output schema, evidence, verifier, and escalation.
Evidence Docket	Complete, attributable record of facts, assumptions, sources, work, costs, failures, approvals, rights, and outcomes.
Proof Debt	Material unresolved evidence, authority, source, rights, or transaction dependency.
Transaction Gate	Fail-closed rule preventing a material action until applicable receipts pass.
Chronicle	Selective institutional memory that admits, scopes, versions, expires, and revokes capability.
Verified Outcome Graph	Rights-controlled causal graph of objectives, evidence, actions, outcomes, and successor performance.
Global Expansion Office	Recurring control plane that monitors the institution, detects drift, and generates successor architectures.
Verified Institutional Value	Full-denominator value of an admissible scenario after costs, risk, delay, dilution, coordination, unwind, and Proof Debt.

CHAPTER H

Machine-Readable Schemas

H.1 Mission Contract example

Listing H.1: Illustrative Mission Contract object.

```
mission_id: mission:global-launch:2026-001
version: 1.0.0
publisher: MONTREAL.AI
objective:
  statement: "Launch a cross-border AI company under defined constraints"
  sponsor: "customer:board"
  success_tests:
    - "scenario frontier accepted"
    - "material Proof Debt assigned"
    - "100-day plan approved"
constraints:
  deadline: 2026-11-30
  budget_usd: 125000
  prohibited_actions:
    - "binding legal filing without professional receipt"
    - "treating unawarded support as runway"
authority:
  approving_role: "customer:authorized-executive"
  revocation_route: "mission:pause-and-escalate"
evidence:
  docket_schema: "goalos:evidence-docket:v2"
  material_claim_expiry_days: 30
chronicle:
  default: reject
  fresh_successor_required: true
```

H.2 Authority Envelope example

Listing H.2: Illustrative Authority Envelope.

```
authority_envelope:
  envelope_id: ae:job:funding-screen:0042
  principal: customer:authorized-executive
  actor: agent:funding-hunter:7
  purpose: "screen official support routes"
  allowed_tools: [official_web, customer_data_room, calculation]
  allowed_data: [public_sources, customer_nonsecret_financials]
  prohibited: [submit_application, sign_contract, move_money]
  budget_usd: 150
  expires_at: 2026-07-26T18:00:00Z
  requires:
    - source_capture
    - contradiction_search
    - evidence_return
  revoke_on: [account_change, policy_violation, source_uncertainty]
```

H.3 Proof-Carrying Action Record example

Listing H.3: Illustrative PCAR.

```
pcar:
  action_id: action:jurisdiction-source-check:991
  mission_id: mission:global-launch:2026-001
  objective_version: sha256:...
  actor: agent:jurisdiction-cartographer:3
  principal: customer:authorized-executive
  policy_version: goalos-policy:2.0.0
  authority_envelope: ae:job:jurisdiction-check:991
  action: retrieve_and_compare
  source_urls: ["official-source-1", "official-source-2"]
  observed_at: 2026-07-25T20:10:00Z
  raw_artifact_hashes: [sha256:..., sha256:...]
  contradiction_state: open
  validator: validator:independent-source-review:2
  validation_verdict: conditional_pass
  accepted_by: null
  transaction_effect: blocked
  chronicle_disposition: reject_until_resolved
```

H.4 Chronicle capability example

Listing H.4: Illustrative Chronicle capability record.

```
chronicle_capability:
  capability_id: cap:funding-stack-screen:ca-qc:3
  derived_from: [mission:001, mission:004, mission:009]
  scope:
    jurisdiction: [CA, CA-QC]
    entity_profile: "qualifying private R&D company"
    cost_objects: [eligible_RnD_payroll, compute]
  evidence: [docket:001, docket:004, docket:009]
  rights: platform_reuse_with_customer_confidentiality_separation
  known_failures:
    - "foreign-control fact pattern not covered"
    - "programme terms require current recheck"
  valid_until: 2026-10-31
  admission: admitted_with_limits
  revocation_triggers: [source_change, failed_audit, rights_change]
  successor_tests:
    - benchmark: fresh-mission-12
      result: pass
```

H.5 Interoperability profile

A conforming implementation may use JSON, YAML, RDF, signed attestations, databases, event logs, or smart-contract commitments. The canonical semantics are more important than one encoding. Required properties include stable identifiers, explicit versions, attributable principals, integrity evidence, time, scope, rights, expiry, and revocation.

Machine readability improves portability and auditability. It does not make an unsupported claim true. Every field describing evidence, authority, acceptance, value, or rights must be backed by the corresponding institutional record.

CHAPTER I

Author and Institutional Stewardship

I.1 Author

Vincent Boucher is President of MONTREAL.AI & QUEBEC.AI. He is the author of this Founding President Constitutional Edition and the founding institutional steward of GoalOS Ω – Sovereign Intelligence.

The author's role in this paper is architectural and constitutional: to define the category, doctrine, formal objects, separation of powers, Mission Institution portfolio, commercial proof requirements, productive-capacity thesis, and institutional boundaries by which GoalOS may be evaluated and developed.

I.2 Publishing institution

MONTREAL.AI Research publishes the paper from Montreal, Quebec, Canada. MONTREAL.AI is identified throughout as the founding, research, and governing institution for GoalOS Ω.

I.3 Institutional relationship with QUEBEC.AI

QUEBEC.AI is named in the author's institutional title and stewardship. This paper does not by itself create a separate legal mandate, public authority, certification, or governmental status for either institution. Any deployment, partnership, regulated activity, or public representation remains governed by the applicable entity, contract, law, evidence, and responsible authority.

I.4 Rights, attribution, and reuse

The paper should be quoted with attribution to Vincent Boucher and MONTREAL.AI, while preserving the claim boundary. GoalOS, GoalOS Ω, Sovereign Intelligence, The Verified Action Institution, GoalOS Chat Ω, GoalOS Proof & Authority Ω, AGI Jobs, Evidence Docket, Chronicle, Verified Outcome Graph, and related terms and architecture are asserted as MONTREAL.AI intellectual property to the extent protectable.

Commercial adaptation, republication, derivative product use, or removal of provenance, rights, security, or claim-boundary controls requires written authorization except where mandatory law provides otherwise.

I.5 Recommended citation

Canonical Citation

Boucher, Vincent (2026). *GoalOS Ω – Sovereign Intelligence: The Verified Action Institution*. Founding President Constitutional Edition, Version 3.0.0. MONTREAL.AI Research, Montreal, Quebec, Canada.

I.6 Public contact

Institution	MONTREAL.AI
Author	Vincent Boucher, President, MONTREAL.AI & QUEBEC.AI
Public contact	info@quebec.ai
Research location	Montreal, Quebec, Canada

References

- Ashby, W. Ross (1956). *An Introduction to Cybernetics*. Chapman & Hall.
- Autio, Chloe, Reva Schwartz, Jesse Dunietz, Shomik Jain, Martin Stanley, Elham Tabassi, Patrick Hall, and Kamie Roberts (2024). *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*. Tech. rep. NIST AI 600-1. National Institute of Standards and Technology. DOI: [10.6028/NIST.AI.600-1](https://doi.org/10.6028/NIST.AI.600-1). URL: <https://doi.org/10.6028/NIST.AI.600-1> (visited on 07/25/2026).
- Campbell, Donald T. (1979). “Assessing the Impact of Planned Social Change”. In: *Evaluation and Program Planning* 2.1, pp. 67–90. DOI: [10.1016/0149-7189\(79\)90048-X](https://doi.org/10.1016/0149-7189(79)90048-X).
- Coase, Ronald H. (1937). “The Nature of the Firm”. In: *Economica* 4.16, pp. 386–405. DOI: [10.1111/j.1468-0335.1937.tb00002.x](https://doi.org/10.1111/j.1468-0335.1937.tb00002.x).
- European Commission (2026a). *AI Act: Application Timeline and Implementation*. URL: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai> (visited on 07/25/2026).
- (July 20, 2026b). *Guidelines on Transparency Obligations for Providers and Deployers of Certain AI Systems*. URL: <https://digital-strategy.ec.europa.eu/en/news/commission-publishes-guidelines-transparency-obligations-providers-and-deployers-certain-ai-systems> (visited on 07/25/2026).
- European Parliament and Council of the European Union (June 13, 2024). *Regulation (EU) 2024/1689 Laying Down Harmonised Rules on Artificial Intelligence*. URL: <https://data.europa.eu/eli/reg/2024/1689/oj> (visited on 07/25/2026).
- Goodhart, Charles A. E. (1975). “Problems of Monetary Management: The U.K. Experience”. In: *Papers in Monetary Economics*. Reserve Bank of Australia.
- Google Developers (Mar. 31, 2026). *ADK Go 1.0: Tracing, Loop Agents, Plugins, and Human-in-the-Loop Guardrails*. URL: <https://developers.googleblog.com/adk-go-10-arrives/> (visited on 07/25/2026).
- International Organization for Standardization (2023). *ISO/IEC 42001:2023 Information Technology – Artificial Intelligence – Management System*.
- MONTREAL.AI Research (July 2026a). *GoalOS Global Jurisdiction, Funding & Fiscal Intelligence Ω: The Verified Operating System for Global Institutional Allocation*. Tech. rep. Sovereign Fiscal Intelligence Institution Edition. MONTREAL.AI.
- (July 2026b). *GoalOS Global Launch Ω: The Verified Operating System for Building Anything, Anywhere*. Tech. rep. Preprint and institutional design paper. MONTREAL.AI.
- (July 2026c). *GoalOS-Native Application Portfolio Ω: Rebuilding Profitable Software Categories as Verified Institutions of Action*. Tech. rep. Internal institutional design paper. MONTREAL.AI.
- Moreau, Luc and Paolo Missier (2013). *PROV-DM: The PROV Data Model*. Tech. rep. World Wide Web Consortium. URL: <https://www.w3.org/TR/prov-dm/> (visited on 07/25/2026).
- National Institute of Standards and Technology (Feb. 17, 2026a). *AI Agent Standards Initiative: Ensuring a Trusted, Interoperable, and Secure Agentic Frontier*. URL: <https://www.nist.gov/artificial-intelligence/ai-agent-standards-initiative> (visited on 07/25/2026).
- (Apr. 6, 2026b). *Concept Note: AI RMF Profile on Trustworthy AI in Critical Infrastructure*. URL: <https://www.nist.gov/programs-projects/concept-note-ai-rmf-profile-trustworthy-ai-critical-infrastructure> (visited on 07/25/2026).
- (Feb. 5, 2026c). *New Concept Paper on Identity and Authority of Software Agents*. URL: <https://www.nist.gov/news-events/news/2026/02/new-concept-paper-identity-and-authority-software-agents> (visited on 07/25/2026).
- (2026d). *NIST AI Resource Center: Testing, Evaluation, Verification, and Validation Resources*. URL: <https://airc.nist.gov/> (visited on 07/25/2026).
- OpenAI (2026a). *OpenAI Agents SDK: Guardrails*. URL: <https://openai.github.io/openai-agents-python/guardrails/> (visited on 07/25/2026).
- (2026b). *OpenAI Agents SDK: Tracing*. URL: <https://openai.github.io/openai-agents-python/tracing/> (visited on 07/25/2026).
- (Apr. 15, 2026c). *The Next Evolution of the Agents SDK*. URL: <https://openai.com/index/the-next-evolution-of-the-agents-sdk/> (visited on 07/25/2026).
- Ostrom, Elinor (1990). *Governing the Commons: The Evolution of Institutions for Collective Action*. Cambridge University Press.
- OWASP GenAI Security Project (2026). *Agentic Security Initiative and Top 10 Risks for Agentic Applications*. URL: <https://genai.owasp.org/initiative/50218/> (visited on 07/25/2026).

- Pearl, Judea (2009). *Causality: Models, Reasoning, and Inference*. 2nd ed. Cambridge University Press.
- Simon, Herbert A. (1962). “The Architecture of Complexity”. In: *Proceedings of the American Philosophical Society* 106.6, pp. 467–482.
- Tabassi, Elham (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. Tech. rep. NIST AI 100-1. National Institute of Standards and Technology. DOI: [10.6028/NIST.AI.100-1](https://doi.org/10.6028/NIST.AI.100-1). URL: <https://doi.org/10.6028/NIST.AI.100-1> (visited on 07/25/2026).
- Williamson, Oliver E. (1985). *The Economic Institutions of Capitalism*. New York: Free Press.