



Global AI
Certification Council



GAICC AAIGP

**AGENTIC AI
GOVERNANCE PROFESSIONAL**

**Examination Content Outline
1st Edition | 2026**

Global AI Certification Council
(GAICC)

Examination Content Outline

1st Edition | 2026

The Credential at a Glance

100 MCQs

2 h 30 min · Examination

70%

Pass mark · scenario-based

60 CPDs

Every 3 years · Renewal

Credential	GAICC Agentic AI Governance Professional (GAICC AAIGP)
Issued by	Global AI Certification Council (GAICC)
Standard	ISO/IEC 17024 Conformity Assessment for Certification Bodies
Programme	4 instructor-led days 32 contact hours
Examination	100 scenario-based MCQs 2 hours 30 minutes 70% pass mark
Validity	3 years 60 CPD hours per cycle
Prerequisite	GAICC CPAIG or equivalent recommended, not required
Contact	certify@gaicc.org www.gaicc.org

Published by: Global AI Certification Council (GAICC)

3900 Westerre Pkwy,
Richmond, VA 23233,
USA

Level 2, 697 Collins Street,
Melbourne VIC 3008,
Australia

Level 3, 21 Putney Way
Manukau, Auckland,
New Zealand 2104

© 2026 Global AI Certification Council. All rights reserved. GAICC, the GAICC logo, Agentic AI Governance Professional, AAIGP, Certified Professional in AI Governance, CPAIG, and related marks are trademarks of the Global AI Certification Council.

Contents

The Credential at a Glance	2
Introduction & Assessment Philosophy	4
How GAICC AAIGP Differentiates	5
Examination Structure & Domain Weightings	6
Four-Day Programme Structure	7
Domain I — Agentic AI Systems and Architecture Literacy	8
Domain II — Agent Authority, Identity and Delegation	9
Domain III — Agentic Risk Assessment and Threat Modelling	10
Domain IV — Human Oversight, Control and Containment	11
Domain V — Agent Evaluation, Assurance and Red Teaming	12
Domain VI — Observability, Audit Evidence and Forensics	13
Domain VII — Agentic Incident Response, Liability and Third-Party Governance	14
Domain VIII — Regulatory Frameworks and Enterprise Scaling for Agentic AI	15
Eligibility Requirements & Training Routes	16
Certification Maintenance and CPD Requirements	17
Examination Fees	17
GAICC Code of Professional Conduct	18
Frequently Asked Questions	19

This Examination Content Outline defines the eight domains of competence required of a GAICC AAIGP holder, the competencies and performance indicators within each domain, the relative weight of each domain in the examination, the cognitive depth applied across all domains, and the four-day programme structure that prepares candidates for the examination.

Introduction

The GAICC Agentic AI Governance Professional (GAICC AAIGP) is an advanced practitioner credential for professionals who authorise, oversee, assure and shut down AI systems that act. It is the first GAICC credential written for hybrid human and autonomous-agent teams rather than for AI systems that only produce output.

Conventional AI governance asks whether a model's output is accurate, fair and lawful. Agentic governance asks a harder question: what is this system permitted to do, on whose authority, with what limits, under whose supervision, and how do we stop it. AAIGP is built around that question.

AAIGP extends the GAICC AI Governance Framework, Maturity Model and Implementation Roadmap into agentic operation. It assumes the governance fundamentals taught in the GAICC Certified Professional in AI Governance (CPAIG) and spends all 32 contact hours on the authority, control, assurance and evidence problems created by autonomy.

This Examination Content Outline defines:

- 01 The eight domains of competence required of a GAICC AAIGP holder
- 02 The competencies and performance indicators within each domain
- 03 The relative weight of each domain in the overall examination
- 04 The cognitive depth and assessment approach applied across all domains
- 05 The four-day programme structure that prepares candidates for the examination

Assessment Philosophy

All 100 questions are scenario-based. Every item places the candidate in a decision seat: authorise or refuse, escalate or contain, disclose or hold, promote or demote autonomy. Questions test Application (45%) and Analysis and Evaluation (40%) cognitive levels. Only 15% operate at Understanding level.

No question tests definition recall alone, and no question can be answered from vendor terminology. Candidates are expected to reason from first principles where regulation is silent, because for agentic systems it frequently is.

How GAICC AAIGP Differentiates

Dimension	GAICC AAIGP	GAICC CPAIG	General AI governance credentials
Governance question	What may this system do, and how is that authority bounded and evidenced	Is this AI system lawful, fair and well governed	Is this AI system lawful and well documented
Agent architecture literacy	Full domain: control loops, tool surfaces, memory, orchestration topologies	Introductory coverage within technical foundations	Rarely addressed
Authority and identity	Full domain: non-human identity, scoped grants, delegation chains, blast radius, authorisation gates	Delegation and accountability treated at principle level	Not addressed
Agentic threat modelling	Full domain: injection, tool misuse, goal drift, exfiltration through action, cascading failure	Failure modes mapped to legal exposure at model level	Bias and privacy focus
Oversight and containment	Full domain: approval gates, circuit breakers, kill switches, autonomy promotion and demotion	Human oversight requirements and override principles	Human in the loop stated as a principle
Evaluation and red teaming	Full domain: trajectory evaluation, agentic red teaming, continuous regression, assurance cases	Commissioning and interpreting model safety testing	Limited
Evidence and forensics	Full domain: trace specification, decision reconstruction, audit and litigation evidence	Documentation and record keeping obligations	Documentation obligations
Delivery and examination	4 instructor-led days, 32 contact hours; 100 scenario items, 2.5 hours, 70% pass	Pre-course plus 2 intensive days, 28 to 30 hours; 80 scenario items, 2 hours, 70% pass	Varies; typically 100 items, 3 hours, scaled pass

Examination Structure

Component	Details
Total Questions	100 questions (88 scored + 12 unscored pilot)
Format	Scenario-based multiple-choice, single best answer per question
Duration	2 hours 30 minutes (150 minutes), no optional break
Pass Mark	70% (minimum 62 of 88 scored questions correct)
Scaled Score	100–500 scale; 350 is the passing threshold
Negative Marking	None. Candidates should answer all questions
Delivery	AI-proctored on the GAICC hub (hub.gaicc.org), global browser access
Results	Provisional result displayed on screen on completion, with domain-level performance feedback
Certification decision	The Certification Decision Authority reviews the provisional result and confirms the decision within 72 hours
Cognitive Split	Understanding 15% Application 45% Analysis and Evaluation 40%
Reference Material	Closed book. No external reference material permitted

Domain Weightings

Dom	Domain Title	Weight	Scored Questions
I	Agentic AI Systems and Architecture Literacy	12%	~11
II	Agent Authority, Identity and Delegation	14%	~12
III	Agentic Risk Assessment and Threat Modelling	15%	~13
IV	Human Oversight, Control and Containment	14%	~12
V	Agent Evaluation, Assurance and Red Teaming	12%	~11
VI	Observability, Audit Evidence and Forensics	10%	~9
VII	Agentic Incident Response, Liability and Third-Party Governance	11%	~10
VIII	Regulatory Frameworks and Enterprise Scaling for Agentic AI	12%	~10
TOTAL		100%	88 scored + 12 pilot

Four-Day Programme Structure

The GAICC AAIGP programme is delivered as four instructor-led days totalling 32 contact hours. Each day covers two examination domains and closes with applied scenario work and a module quiz.

Day	Theme	Coverage	Hours
Day 1	Authority and Architecture	Domains I and II. Agent architecture literacy, action surfaces, memory and orchestration, then agent identity, scoped authorisation, delegation chains, blast-radius limits and the authorisation gate.	8 hours
Day 2	Risk and Control	Domains III and IV. Agentic risk taxonomy, injection and exfiltration threat modelling, agentic impact assessment, then meaningful oversight, approval gates, circuit breakers, kill switches and autonomy promotion.	8 hours
Day 3	Assurance and Evidence	Domains V and VI. Behavioural and trajectory evaluation, agentic red teaming, continuous regression and assurance cases, then trace specification, decision reconstruction, production monitoring and audit evidence.	8 hours
Day 4	Response, Liability and Scale	Domains VII and VIII. Agentic incident response, notification, liability allocation and third-party agent governance, then regulatory application, standards mapping, portfolio scaling and board advisory. Closes with the full mock examination and review.	8 hours

Domains, Competencies & Performance Indicators

Each domain that follows states its governing purpose, its competencies, and the performance indicators that define competent practice. Examination items are written directly against these indicators.

Purpose

Read and interrogate an agentic system design well enough to govern it. Understand control loops, tool surfaces, memory, and orchestration topologies at the level of detail a governance decision actually requires, without becoming an engineer.

Competency	Performance Indicators
I.A Distinguish agentic systems from conventional and generative AI	Define agency through goal direction, planning, tool use, memory and persistence; distinguish assistants, automated workflows, single agents and multi-agent systems; identify the point at which an assistant or copilot crosses into agentic operation; apply the GAICC AI Governance Framework concept of hybrid human and autonomous-agent teams; classify systems on an autonomy spectrum for registry and tiering purposes.
I.B Interpret agent architectures and control loops	Read planner and executor patterns, reasoning and action loops, and reflection or critique stages; identify where the agent decides and where deterministic code decides; assess loop termination conditions, step budgets and recursion limits; evaluate the governance consequences of model selection and non-determinism; interrogate a design document to surface undeclared autonomy.
I.C Govern tool use and the action surface	Inventory the tools, functions, APIs and connectors an agent may invoke; distinguish read actions from write, transactional and irreversible actions; assess tool server and integration protocol exposure, including Model Context Protocol style integrations; apply action classification to determine the required approval level; identify shadow tool access and undocumented capability.
I.D Govern memory, state and context	Distinguish session context, short-term scratchpad, long-term memory and shared knowledge stores; assess the data protection exposure created by persistent agent memory; govern retention, correction and deletion of agent memory; assess context contamination and cross-tenant leakage risk; apply provenance requirements to retrieved context.
I.E Assess orchestration and multi-agent topologies	Distinguish supervisor and worker, peer to peer, swarm and pipeline topologies; identify emergent behaviour and interaction risk in multi-agent systems; assess agent to agent communication and trust boundaries; evaluate orchestration frameworks for governance observability; map topology choice to accountability allocation.

Agent Authority, Identity and Delegation

Purpose

Establish what an agent is, what it may do, on whose authority, and how that authority is granted, constrained, evidenced and withdrawn. This is the domain that separates agentic governance from all prior AI governance practice.

Competency	Performance Indicators
II.A Establish agent identity and non-human credential governance	Apply non-human identity principles to agents, service accounts and workloads; govern issuance, rotation and revocation of agent credentials and secrets; require attributable, non-shared agent identities so that action can be traced to a single actor; assess impersonation and identity spoofing risk across agent boundaries; register every agent in the AI system inventory with owner, risk tier, autonomy level and oversight mode as required by Phase 3 of the GAICC Implementation Roadmap.
II.B Design scoped authorisation and least-privilege tool grants	Apply least privilege to tool, data and system access; define allow-lists, deny-lists and environment scoping; apply just-in-time and time-bounded authority rather than standing permission; separate development, staging and production authority; document authorisation decisions and their rationale as governance evidence.
II.C Govern delegation chains and human accountability	Trace authority from the human principal through the agent to sub-agents and tools; identify accountability dilution across delegation hops; assign a named accountable human owner to every agent in production; apply the three lines of defence model to agent operation; apply the confused deputy problem to delegated authority design.
II.D Set operational limits and blast-radius controls	Define spend, rate, volume and transaction limits proportionate to risk; set irreversibility thresholds that force human approval; apply segregation of duties between agents so no single agent completes a sensitive transaction end to end; define data egress and external communication limits; calculate and constrain blast radius before authorisation, not after incident.
II.E Operate an agent authorisation gate	Design an agent authorisation gate owned by the Governance Core layer; define go and no-go criteria and the evidence required for production release; apply staged autonomy promotion with explicit graduation criteria; govern changes to authority scope after release as a change-control event; withdraw authorisation and decommission agents cleanly, including credentials, memory and data.



Agentic Risk Assessment and Threat Modelling

Purpose

Identify, analyse and treat the risks specific to systems that act, including adversarial risks that simply do not exist for non-agentic AI.

Competency	Performance Indicators
III.A Apply an agentic risk taxonomy	Classify goal misspecification, goal drift, reward hacking and specification gaming; classify unbounded loops, runaway resource consumption and cost incidents; classify unsafe tool invocation and irreversible action risk; classify cascading and correlated failure across multi-agent systems; distinguish agentic risk from model-output risk and route each to the correct control.
III.B Threat-model prompt injection and adversarial input	Distinguish direct, indirect and cross-domain prompt injection; trace every path by which untrusted content reaches agent context; assess tool-output poisoning and retrieved-document attack surfaces; apply trust boundaries, content provenance and input isolation controls; commission adversarial testing proportionate to the authority granted.
III.C Assess exfiltration and confidentiality risk created by action	Identify exfiltration channels created by tool and network access; assess agent-mediated aggregation of data that is individually permitted but collectively sensitive; apply data classification to agent-accessible stores; govern outbound communication, publication and third-party submission actions; assess cross-tenant and cross-client contamination in shared agent deployments.
III.D Conduct agentic impact assessment	Extend AI impact assessment practice, including ISO/IEC 42005 style assessment, to autonomous action; assess impact on individuals, groups and third parties who never interact with the agent; assess reversibility, detectability and time to harm as primary variables; apply proportionality between autonomy granted and assurance required; produce records that satisfy the Impact assessment control domain.
III.E Integrate agentic risk into enterprise risk management	Extend the AI risk register to agent-specific risk types and scoring; define severity classification and escalation triggers for autonomous action; set an explicit risk appetite for autonomy; aggregate portfolio-level agent risk for board reporting; map agentic risk onto existing operational, cyber, conduct and third-party risk taxonomies.

Human Oversight, Control and Containment

Purpose

Design oversight that is meaningful at machine speed and machine scale, covering approval gates, intervention, circuit breakers, shutdown, and the human operating model around them.

Competency	Performance Indicators
IV.A Design meaningful human oversight for autonomous operation	Distinguish human in the loop, on the loop and out of the loop; apply EU AI Act Article 14 human oversight requirements to agentic deployments; identify oversight that is nominal rather than effective; assess reviewer capacity, decision latency and automation bias; match oversight mode to action class, reversibility and volume.
IV.B Implement approval gates and intervention points	Place approval gates at irreversible and high-consequence actions rather than uniformly; design approval interfaces that give the reviewer enough context to refuse; define timeouts, defaults and fail-safe behaviour when approval is not given; govern bulk approval, blanket approval and approval fatigue; evidence approval decisions for audit.
IV.C Design containment, circuit breakers and kill switches	Define automated circuit breaker triggers on anomaly, spend, error rate and unexpected tool use; design graceful degradation and safe-stop states; specify who may halt an agent, by what mechanism, and within what time; test kill switches and rollback under realistic production conditions; govern partial containment where full shutdown is not commercially viable.
IV.D Manage the human and agent operating model	Define operator roles, rosters and escalation paths for agent supervision; address deskilling, over-trust and vigilance decay in supervising teams; train operators on override authority and the conditions for using it; apply the People, Functions and Capability layer of the GAICC Framework; set competence requirements for anyone supervising an autonomous system.
IV.E Govern autonomy promotion and demotion	Define the evidence thresholds required to increase autonomy; define automatic demotion triggers following incident, drift or failed evaluation; apply pilot, shadow-mode and canary patterns before granting live authority; document autonomy decisions and their rationale; review autonomy levels on a defined cadence rather than only after failure.

Agent Evaluation, Assurance and Red Teaming

Purpose

Commission, interpret and act on evidence that an agent behaves acceptably, both before authorisation and continuously in production.

Competency	Performance Indicators
V.A Commission behavioural evaluation of agents	Distinguish model benchmarks from agent task and trajectory evaluation; define acceptance criteria for behaviour and not only for output quality; evaluate tool-use correctness, action sequencing and refusal behaviour; apply scenario and simulation testing in sandboxed environments; assess evaluation coverage against the full authorised action surface.
V.B Direct agentic red teaming	Scope red team engagements for injection, privilege escalation, exfiltration and unsafe action; distinguish model red teaming from system and agent red teaming; set rules of engagement and safety controls for testing against live or production-like systems; interpret findings and set remediation thresholds; require re-testing after remediation and before any autonomy promotion.
V.C Apply continuous evaluation and regression control	Detect behavioural drift arising from model updates, prompt changes, tool changes and data changes; maintain regression suites over agent behaviour, not just model accuracy; govern vendor model version changes as a formal change-control event; set monitoring thresholds that trigger re-evaluation; retain evaluation evidence for assurance and audit.
V.D Build the assurance case for an agent	Assemble a structured assurance argument linking authority granted, controls applied, evidence held and residual risk accepted; apply the Assurance and Audit dimension of the GAICC Maturity Model; map evidence to ISO/IEC 42001 clauses and the nine GAICC control domains; prepare agentic systems for internal audit and external certification; identify where assurance is asserted rather than evidenced.

Observability, Audit Evidence and Forensics

Purpose

Make agent behaviour explainable and reconstructable after the fact, to a standard that satisfies auditors, regulators and courts. This is the operational meaning of governing AI you can prove.

Competency	Performance Indicators
VI.A Specify agent trace and decision logging	Define minimum trace content covering inputs, retrieved context, reasoning artefacts, tool calls and parameters, results, approvals and outcomes; set retention aligned to limitation periods and regulatory obligations; balance trace completeness against privacy and data minimisation; protect log integrity against tampering by humans and agents; ensure tracing spans sub-agents and orchestration layers.
VI.B Reconstruct agent decisions	Reconstruct a specific action end to end from trace evidence; identify the causal contribution of context, tool output and model behaviour; handle non-determinism honestly when explaining a past action; distinguish explanation of mechanism from attribution of intent; communicate the reconstruction to non-technical decision makers, regulators and claimants.
VI.C Monitor agents in production	Define operational and governance indicators for agent fleets; monitor for anomaly, drift, cost escalation and unsafe action patterns; design dashboards appropriate to first, second and third line audiences; set alerting thresholds and named ownership for each alert; feed monitoring output into the Improve stage of the Direct, Assess, Operate, Improve loop.
VI.D Produce audit and regulatory evidence	Map evidence requirements to ISO/IEC 42001 and applicable regulation; prepare evidence packs for internal audit, certification bodies and regulators; handle discovery, preservation and litigation hold over agent traces; state evidence gaps honestly within the assurance case; apply an evidentiary standard sufficient to defend an autonomous decision externally.

Agentic Incident Response, Liability and Third-Party Governance

Purpose

Respond when an autonomous system causes harm, and allocate responsibility across the agent supply chain before that happens rather than after.

Competency	Performance Indicators
VII.A Execute agentic incident response	Sequence detection, containment, halt, evidence preservation, remediation and disclosure when the actor is autonomous; preserve trace evidence before remediation destroys it; determine the full scope of unauthorised or erroneous action already taken; reverse or compensate completed actions where reversal is possible; coordinate technical, legal, communications and executive response under time pressure.
VII.B Manage regulatory notification and disclosure	Apply serious incident reporting obligations to agentic failure; assess personal data breach notification triggers arising from agent action; determine disclosure obligations to affected individuals, counterparties and markets; manage notification timelines across multiple jurisdictions simultaneously; document and justify decisions not to notify.
VII.C Allocate liability across the agent supply chain	Distinguish the responsibilities of model provider, tool and platform provider, agent builder, orchestrator and deploying organisation; assess where agentic use converts a deployer into a provider under the EU AI Act; draft contractual allocation, warranties, audit rights and indemnities for autonomous action; assess insurance coverage and gaps; identify where liability for an autonomous action is currently unallocated.
VII.D Govern third-party agents and agent marketplaces	Conduct due diligence on third-party agents, tools and connectors before they are granted authority; govern agents that act on your systems under a vendor's control; assess transitive risk from tools an agent can discover, install or chain to; set exit, continuity and data return requirements; monitor third-party agent behaviour in production rather than trusting attestation.
VII.E Learn from incidents	Conduct post-incident review focused on authority and oversight design rather than operator blame; feed findings into the risk register, control set and autonomy levels; share learning across the agent portfolio; track corrective and preventive action to closure; report systemic findings to the board and governing body.

Regulatory Frameworks and Enterprise Scaling for Agentic AI

Purpose

Apply current law and standards to systems they were not written for, and scale agent governance from a single pilot to an enterprise portfolio without creating a bottleneck that drives shadow deployment.

Competency	Performance Indicators
VIII.A Apply existing AI regulation to agentic systems	Apply EU AI Act risk classification, GPAI obligations and Article 14 human oversight to agents; assess where agentic customisation converts a deployer into a provider; apply sectoral regulation to autonomous action in finance, health, employment and critical infrastructure; apply automated decision-making rights and human review obligations under GDPR, PIPL, PDPA, LGPD and equivalent regimes; identify clearly where existing regulation is silent on autonomy.
VIII.B Apply management system standards to agentic operation	Apply ISO/IEC 42001 clauses and Annex A controls to agents; apply the nine GAICC control domains across an agent portfolio; apply the NIST AI RMF Govern, Map, Measure and Manage functions to agentic risk; apply ISO/IEC 42005 impact assessment to autonomous action; extend an existing AI management system scope to cover agents without rebuilding it.
VIII.C Track and anticipate agent-specific regulatory development	Monitor emerging agent-specific guidance, standards work and enforcement signals; assess regulatory proposals for operational impact before they bind; apply governance under uncertainty and precautionary sequencing; advise leadership on defensible positions where obligations are unsettled; avoid over-claiming compliance in areas that remain unsettled.
VIII.D Scale agent governance across the enterprise	Operate an agent registry and maintain a portfolio view of autonomy and risk; apply tiered governance proportionate to autonomy and impact rather than uniform control; apply the GAICC Maturity Model across Accountability and Oversight, Risk, Lifecycle Control, Assurance and Capability; sequence adoption using the seven phases of the GAICC Implementation Roadmap; design governance that does not become the bottleneck driving shadow agent deployment.
VIII.E Advise boards and executives on agentic AI	Report agent portfolio risk, autonomy distribution and incident trend to the board; define board-level risk appetite for autonomous action; present the business case for agent governance investment; advise on the technology journey from AI-enabled to AI-first to AI-native alongside governance maturity; escalate clearly where authority granted exceeds assurance held.

Eligibility Requirements

GAICC AAIGP is an advanced credential. The GAICC Certified Professional in AI Governance (CPAIG), or an equivalent AI governance credential, is strongly recommended but is not a formal prerequisite. Candidates without that background should expect the programme to assume working knowledge of AI governance fundamentals.

Educational Background	Professional Experience	Training Requirement
Secondary Qualification (high school diploma, associate degree, or global equivalent)	5 years (60 months) in AI governance, risk, compliance, security, audit, data or AI delivery	32 hours of AAIGP training by any of the three recognised routes
Bachelor's Degree (or global equivalent)	3 years (36 months) in AI governance, risk, compliance, security, audit, data, technology or law	32 hours of AAIGP training by any of the three recognised routes
Master's Degree or Professional Qualification	2 years (24 months) of relevant experience in AI governance, law, risk, security or AI delivery	32 hours of AAIGP training by any of the three recognised routes

Training Routes

The 32-hour training requirement may be met by any one of three routes. All three are equivalent for eligibility purposes and no route, including GAICC's own, confers any advantage in the examination.

Route	Description
1. GAICC self-paced course	The 32-hour AAIGP course completed in the candidate's own time on the GAICC hub.
2. GAICC-authorized Certified Instructor or Authorized Training Partner	Instructor-led delivery of the 32-hour programme by an authorised partner or an authorised GAICC Certified Instructor.
3. A GAICC-recognised institution	32 hours of AAIGP training delivered by an institution recognised by GAICC for this purpose.

Experience Recency

All relevant experience must have been gained within the last eight (8) consecutive years before submitting the GAICC AAIGP application. Experience with agentic or automated decision systems, non-human identity, security assurance or audit is recognised alongside classic AI governance experience.

Certification Maintenance & CPD Requirements

The GAICC AAIGP certification is valid for three (3) years. During each cycle, certified professionals must earn a minimum of 60 CPD hours:

CPD Category	Description	Minimum
1. Professional Learning	GAICC-recognised training, conferences, webinars or workshops in AI governance, agentic AI, security, law, regulation or assurance	20 hours
2. Practical Application	Direct involvement in agent authorisation, agentic risk assessment, oversight design, agent evaluation, audit or incident response	20 hours
3. Contribution and Knowledge Sharing	Publishing, mentoring, research, contributing to AI governance or agent safety standards, policy papers or community initiatives	10 hours
4. Elective Activities	Additional learning or engagement supporting continuous professional growth in AI-adjacent disciplines	10 hours
TOTAL	Minimum required across all categories in each 3-year certification cycle	60 CPD Hours

Examination Fees

Examination, membership and renewal fees are published in the GAICC AAIGP Candidate Handbook and on the GAICC certification body website (certifications.gaicc.org/fees). To maintain a single source of truth, fee amounts are not reproduced in this Examination Content Outline.

GAICC Code of Professional Conduct

Principle	Obligation
1. Integrity and Fairness	Act honestly in all professional activities without bias, conflict of interest or misrepresentation of competence. Disclose potential conflicts proactively.
2. Proportionate Authority	Never recommend or approve a level of agent autonomy that exceeds the assurance evidence held. Where authority and evidence diverge, say so in writing.
3. Human Rights and Wellbeing	Ensure autonomous systems under governance influence respect human dignity, fairness, privacy and non-discrimination, including for people who never interact with the system.
4. Transparency and Accountability	Promote explainable and reconstructable agent behaviour. Ensure a named human remains accountable for every autonomous action.
5. Data Protection and Privacy	Uphold confidentiality and data-protection principles across agent memory, tool access and trace data, consistent with applicable laws and international standards.
6. Professional Competence	Maintain current knowledge in a field that is changing faster than its regulation, through continuous learning and active engagement with evolving standards.
7. Reporting and Mitigation of Misuse	Take appropriate action when encountering unsafe autonomy, suppressed incidents, disabled controls or violations of applicable regulation or professional standards.
8. Global Responsibility	Recognise that autonomous systems act across borders. Apply governance standards that protect individuals regardless of jurisdiction.

Frequently Asked Questions

What is the GAICC AAIGP?

An advanced practitioner credential for professionals who authorise, oversee, assure and shut down AI systems that take action. It governs autonomy, not just output.

Who should pursue AAIGP?

AI governance leads, risk and compliance managers, internal and external auditors, security and identity architects, platform and AI engineering leaders, legal and privacy professionals, and consultants advising on agent deployment.

Do I need CPAIG first?

CPAIG or an equivalent AI governance credential is strongly recommended but not required. AAIGP assumes working knowledge of AI governance fundamentals and does not re-teach them.

How does AAIGP differ from CPAIG?

CPAIG governs AI systems across regulation, ethics and lifecycle, with agentic AI as one domain of eight. AAIGP is entirely about systems that act: authority, identity, delegation, containment, evaluation, evidence and liability.

What is the exam format?

100 scenario-based multiple-choice questions in 2 hours 30 minutes, AI-proctored on the GAICC hub at hub.gaicc.org. Pass mark is 70% of scored questions.

When do I know if I have passed?

A provisional result is displayed on screen on completion, with domain-level feedback. The Certification Decision Authority reviews it and confirms the certification decision within 72 hours.

How long is the certification valid?

3 years. Renewal requires 60 CPD hours per cycle plus a renewal application and fee.

What digital credentials are provided?

Digital badge and verifiable certificate for LinkedIn, resumes and professional profiles, plus listing in the GAICC global credentials register.

How does AAIGP relate to the GAICC Framework?

AAIGP operationalises the GAICC AI Governance Framework for hybrid human and autonomous-agent teams, and maps directly to the Maturity Model dimensions and the Implementation Roadmap phases.

certify@gaicc.org | support@gaicc.org | www.gaicc.org



Global AI
Certification Council

GAICC AAIGP

AGENTIC AI GOVERNANCE PROFESSIONAL

For professionals who authorise, oversee and are accountable for
AI systems that act.

www.gaicc.org | certify@gaicc.org | support@gaicc.org