

GAISB

GLOBAL AI STANDARDS BODY

THE PROFESSIONAL STANDARDS LIBRARY — VOLUME III

EDITION 1.0 / 2026

EFFECTIVE Q2 2026

REVIEW CYCLE: TRIENNIAL

JTA-2026-01

THE EMPIRICAL FOUNDATION OF THE CREDENTIAL

The **GAISB** Job Task Analysis

The founding practice analysis that establishes content validity for the Certified in AI Standards credential — task inventory, domain weighting, examination blueprint, capstone rubric, and standard-setting plan, validated by an active practitioner panel.

FOUNDING PRACTITIONER COMMUNITY: 203 MEMBERS · 21 COUNTRIES

Validated by an active-practitioner panel · Aligned to ISO/IEC 17024 & ANSI/NCCA 1100-2024

GOVERNED BY THE GAISB STANDARDS COUNCIL

RATIFIED BY THE GAISB STANDARDS COUNCIL

PUBLIC DOCUMENT · CITE WITH ATTRIBUTION

Pro Bono Publico

FOR THE PUBLIC GOOD

FRONT MATTER

Document Control & Ratification

DOCUMENT TITLE	Job Task Analysis — Founding Practice Analysis for the Certified in AI Standards credential, Edition 1.0
INSTRUMENT CODE	JTA-2026-01
ISSUING AUTHORITY	Global AI Standards Body (GAISB™)
GOVERNING BODY	GAISB Standards Council
EDITION / STATUS	Edition 1.0 — Founding practice analysis, ratified by the GAISB Standards Council
EMPIRICAL BASIS	Prompt Atlas practitioner community (203 members, 21 countries); active-practitioner validation panel
REFERENCE PERIOD	Community fielded through 2025–2026
EFFECTIVE DATE	Effective Q2 2026
NEXT SCHEDULED REVIEW	Triennial — full re-validation Q2 2029; annual practice-review scan in the interim
DISTRIBUTION	Public; reference document for authorized training partners, item writers, and capstone reviewers
RECOMMENDED CITATION	Global AI Standards Body. (2026). Job Task Analysis: Founding Practice Analysis for the Certified in AI Standards credential, Edition 1.0. JTA-2026-01.

RATIFICATION

This founding Job Task Analysis was prepared by the GAISB Standards Division and ratified by the GAISB Standards Council as the controlling practice-analysis instrument for the Certified in AI Standards credential. Full psychometric validation of the item bank — item statistics, reliability, differential-item-functioning, and equating — is conducted on operational data under the plan in Sections 11 and 13.

KIRWIN NARINE — FOUNDER,
GLOBAL AI STANDARDS BODY (GAISB™)

GAISB STANDARDS COUNCIL —
RATIFYING AUTHORITY

DOCUMENT ABSTRACT

This founding Job Task Analysis (JTA) establishes the empirical foundation for the Certified in AI Standards credential issued by the Global AI Standards Body (GAISB™). The task model was developed against the seven-domain Common Body of Knowledge and validated by an active practitioner panel drawn from the Prompt Atlas practitioner community — a founding community of 203 members across 21 countries. The panel reviewed and rated 84 candidate task statements, yielding 67 retained tasks across the seven GAISB domains, each assigned a consensus importance tier. This report documents the methodology, the community and panel, the task inventory, domain weighting, the examination blueprint, the capstone rubric, and the controlling traceability rule linking every assessment item to a validated task.

FRONT MATTER

Founder's Statement

Certified in AI Standards exists because the AI profession needs a standards-grade signal of competence — one that survives scrutiny, equips the holder to do real work, and helps employers cut through the noise of a credential market that has, in the past two years, multiplied without much rigor.

The credential is named deliberately. It is not Certified in AI Strategy. It is not Certified in AI Awareness. It is Certified in AI Standards — because what the modern AI profession lacks is not enthusiasm and not tooling, but a standard: a defensible, measured line between competent practice and improvisation. A holder of this credential has demonstrated, against an explicit task standard, that they design, deploy, and operate AI systems to that standard.

A credential of that caliber cannot be founded on a course outline. It must be founded on a practice analysis: a structured account of what competent AI practitioners actually do, validated by practitioners themselves. That is what this report is — and it is honest about its stage. GAISB is a young body, and this is its founding practice analysis. Its empirical anchor is a real, engaged community of practitioners — the Prompt Atlas community, 203 members across 21 countries, from Trinidad and Tobago to the United Kingdom, the United States, South Africa, Canada, and beyond. The active practitioners in that community form the standing validation panel that reviewed and rated this task model.

We publish the full report rather than an executive summary because transparency is the fastest path to credibility. We are equally transparent about what this edition is and is not: it is a founding, panel-validated practice analysis, not a large-sample probability survey, and the full psychometric validation of the examination — item statistics, reliability, fairness, and equating — is executed on operational data as the item bank goes live. Every assessment item must still trace to a task statement in Appendix B. Items that cannot be traced do not enter the bank.

GAISB is built to a single standard: the holder consistently designs, deploys, and operates AI systems with sound judgment under realistic constraints, and holds the line on responsible practice when it would be easier not to. Everything in this document is in service of that standard.

KIRWIN NARINE
FOUNDER, GLOBAL AI STANDARDS BODY (GAISB™)

TABLE OF CONTENTS

Contents

FRONT MATTER

—	Document Control & Ratification
—	Founder's Statement
—	Executive Summary

REPORT

1	Introduction & Purpose
2	Standards Alignment
3	Methodology
4	The Practitioner Community & Validation Panel
5	Common Body of Knowledge — Domains & Sub-Domains
6	Task Inventory & Panel Consensus
7	Domain Weighting & Panel Concordance
8	CBK Validation Outcomes
9	GAISB Assessment Architecture
10	Capstone Rubric
11	Pass-Score & Standard-Setting Plan
12	Examination Security & Administration
13	Reliability, Fairness & Limitations
14	Recommendations & Maintenance Plan

APPENDICES

A	Panel Rating Instrument
B	Full Task Inventory (67 tasks)
C	

Sample Examination Items with Task Traceback

D Capstone Rubric — Full Anchors

E Comparative Positioning

F Community Data

G Panel Method Notes

H Glossary

I Sources & References

FRONT MATTER

Executive Summary

This founding Job Task Analysis (JTA) documents the empirical foundation for the Certified in AI Standards credential issued by the Global AI Standards Body (GAISB™). The JTA establishes content validity for the GAISB examination and capstone by tying every assessment item to tasks validated by working AI practitioners. The task model was developed against the seven-domain Common Body of Knowledge (CBK 1.0) and validated by an active practitioner panel drawn from the Prompt Atlas community — a founding community of 203 members across 21 countries, anchored in Trinidad and Tobago with substantial representation across the United Kingdom, the United States, South Africa, Canada, Europe, and the Asia-Pacific. The panel reviewed 84 candidate task statements and retained 67 across the seven domains, assigning each a consensus importance tier.

203	21	67	7
COMMUNITY MEMBERS	COUNTRIES	VALIDATED TASKS	CBK DOMAINS

KEY FINDINGS

- All seven CBK 1.0 domains were validated by the panel; none was eliminated and none added. Two sub-domains were rescoped following panel review.
- The two highest-weighted domains — Prompt Engineering & LLM System Design (18%) and AI Agents & Agentic Workflows (18%) — together account for 36% of the examination and hold the largest share of Essential-tier tasks, matching the panel's consensus that production prompt and agent design define current practice.
- The panel's task-importance consensus was concordant with the published CBK 1.0 domain weights; the Standards Council retained the published weights as final.
- The GAISB knowledge examination is a 75-item blueprint administered over a three-hour session; the AIS-400 Capstone is evaluated against a calibrated six-dimension rubric. Both must be passed for the credential to be conferred.
- Pass scores are established by a modified-Angoff procedure for the knowledge examination and a calibrated-rubric procedure for the capstone, conducted before the first operational administration.

SCOPE OF THIS EDITION

This is a founding, panel-validated practice analysis. Task importance is reported as practitioner-panel consensus, not as a large-sample probability survey; population inferential statistics are therefore not claimed. Full psychometric validation of the item bank (item statistics, reliability, DIF, equating) is executed on operational data per Sections 11 and 13.

SECTION 1 · INTRODUCTION & PURPOSE

Introduction & Purpose

1.1 Why a Job Task Analysis?

A Job Task Analysis is the cornerstone of any defensible professional certification. It links the content of an examination to the actual work performed by competent practitioners. Without a current practice analysis, an examination has no demonstrated content validity and cannot withstand legal, accreditation, or stakeholder scrutiny.

The AI-practitioner role is consolidating in real time. Practitioners cross over from product management, consulting, transformation, engineering, and operations. A practice analysis grounded in real practitioners is the only credible basis for a credential that aspires to set a professional standard for that role.

1.2 Scope

This JTA covers tasks performed by AI practitioners with operational responsibility (or substantive contribution) for one or more of: designing AI systems, including LLM systems and agentic workflows; selecting and deploying generative-AI models; engineering production prompts and retrieval patterns; designing and operating AI agents; embedding ethics, data, and responsible-AI controls into system design; leading AI transformation through deployed systems; and managing applied-AI innovation pipelines.

Out of scope: pure ML research roles with no deployment responsibility; data engineering without applied-AI duties; and sales, procurement, or pure marketing roles without applied-AI responsibilities.

1.3 Intended Audience

- Examination item writers and capstone rubric reviewers authorized by GAISB.
 - Authorized training partners building GAISB preparation curricula.
 - Employers using GAISB for hiring, role definition, and AI-workforce planning.
 - Accreditation reviewers and standards bodies evaluating the credential.
 - Candidates seeking transparency into the empirical basis of the assessment.
-

1.4 Use & Maintenance

The JTA is the controlling content reference for examination construction and capstone rubric calibration. As the community and item bank grow, the panel base broadens and the psychometric evidence accumulates. An annual practice-review scan monitors regulatory and technological drift; full re-validation follows the triennial cadence (next: Q2 2029).

SECTION 2 · STANDARDS ALIGNMENT

Standards Alignment

This JTA was designed in alignment with the standards below governing professional certification programs and AI management systems. GAISB is a personnel credential calibrated to the practice of the AI-practitioner role; GAISB documents formal alignment between GAISB and the AI management-system standard published by ISO/IEC.

STANDARD / FRAMEWORK	RELEVANCE TO GAISB
ISO/IEC 17024:2012 — General requirements for bodies certifying persons	Mandates a current job/practice analysis as the foundation for any personnel certification examination.
ANSI/NCCA 1100-2024 — Accreditation of Certification Programs	Defines acceptable practice-analysis methodology, including subject-matter-expert panel validation, and update cadence.
Standards for Educational and Psychological Testing (AERA / APA / NCME)	Sets evidentiary standards for content validity, reliability, and fairness.
EEOC Uniform Guidelines on Employee Selection Procedures (UGESP)	Establishes job-relatedness and adverse-impact expectations for credentialing.
ISO/IEC 42001:2023 — AI Management Systems	Subject-matter referent for organizational AI management; GAISB provides personnel evidence for Clause 7.2 (Competence) and Annex A.4.6 (Human Resources). GAISB maps ~78% of GAISB task content to the ISO/IEC 42001 normative body.
NIST AI Risk Management Framework 1.0	Subject-matter referent for the Govern, Map, Measure, and Manage functions reflected in GAISB Domain 5.
EU AI Act — Regulation (EU) 2024/1689	Subject-matter referent for AI-literacy (Art. 4) and high-risk obligations (Arts. 8–15) reflected across GAISB Domains 1, 5, and 6.
OECD AI Principles (OECD/LEGAL/0449)	Subject-matter referent for the values-based principles reflected in GAISB Domain 5.

STACKING WITH ISO/IEC 42001

GAISB certifies persons; ISO/IEC 42001 certifies organizations. The two are stackable, not equivalent: an organization holding an ISO/IEC 42001 certification and employing GAISB-credentialed practitioners can evidence both organizational and personnel competence. GAISB is a private professional standards body; GAISB is not ISO-endorsed and not accredited.

SECTION 3 · METHODOLOGY

Methodology

3.1 Methodological Framework

This founding practice analysis followed a panel-validation methodology consistent with the subject-matter-expert-panel approach recognized in ANSI/NCCA 1100-2024 §8.0 for a new credential establishing its first content standard:

1. Phase 1 — CBK 1.0 finalization and domain definition.
2. Phase 2 — Task-statement generation keyed to the seven domains.
3. Phase 3 — Panel constitution from the active practitioner community.
4. Phase 4 — Task review and consensus rating by the panel.
5. Phase 5 — Retention decisions and tier assignment.
6. Phase 6 — Weight concordance review and Standards Council confirmation.

3.2 The Validation Panel

Rather than commission an anonymous external survey, GAISB validated the task model against its own founding practitioner community — the Prompt Atlas community. The active practitioners in that community (those signed in within the trailing 30 days) constitute the standing validation panel. This is a purposive expert panel, not a probability sample; its authority rests on the practitioners' direct engagement with applied AI work and on transparent, tiered consensus rather than on population inference. Panel composition and engagement are reported in Section 4 and Appendix F.

3.3 Task-Statement Generation

An initial pool of 142 candidate task statements was generated from the seven CBK 1.0 domains. Statements were edited against four rules: each begins with an action verb describing observable work; each is performable and assessable independently; each maps to exactly one CBK domain; and each is written at a Flesch-Kincaid Grade 12 reading level or below. Editing reduced the pool to 84 task statements across the seven domains for panel review.

3.4 Rating Design

The panel rated each candidate task on three dimensions — Frequency, Importance, and Criticality (anchors in Appendix A) — and consensus was expressed as an importance tier: Essential, High, or Supporting. Importance and consequence were weighted above raw frequency, consistent with ANSI/NCCA guidance that a low-frequency, high-consequence task can still be central to the role.

3.5 Retention & Tiering

A task was retained only if the panel placed it at Supporting tier or above and judged it independently assessable. Tasks the panel judged peripheral, redundant, or not generalizable across jurisdictions were removed or merged. Retained tasks and their consensus tiers appear in Appendix B.

3.6 Analysis Scope & Honesty Note

Because the validation base is a purposive expert panel, this edition reports task importance as tiered consensus and does not claim population margins of error or large-sample inferential statistics. Item-level psychometrics — difficulty, discrimination, reliability, differential item functioning, and equating — are computed on operational examination data as the item bank goes live (Sections 11 and 13). This staged approach is appropriate for a founding credential and is stated plainly so that every claim in this document can be substantiated.

SECTION 4 · PRACTITIONER COMMUNITY & VALIDATION PANEL

The Practitioner Community & Validation Panel

The empirical anchor of this practice analysis is the Prompt Atlas practitioner community — GAISB's founding community of working AI practitioners. Its active members form the standing validation panel for the GAISB task model.

203	117	21
COMMUNITY MEMBERS	ACTIVE PANEL (30-DAY)	COUNTRIES

4.1 Community Engagement

The community is genuinely active rather than nominal: members have authored 361 posts and 843 comments and have exchanged 3,219 endorsements, concentrated in a committed contributing core. This sustained engagement is what qualifies the active membership to serve as a validation panel: the panel is composed of practitioners who are demonstrably present in applied AI practice and discussion.

4.2 Geographic Distribution

Location was self-reported by 89 members and spans 21 countries. The community is anchored in Trinidad and Tobago and reaches across five world regions.

COUNTRY	MEMBERS	% OF LOCATED
Trinidad and Tobago	34	38.2%
United Kingdom	12	13.5%
United States	11	12.4%
South Africa	6	6.7%
Canada	5	5.6%
New Zealand	2	2.2%
Belgium	2	2.2%
Netherlands	2	2.2%
Thailand	2	2.2%
Australia	2	2.2%
Germany	1	1.1%

COUNTRY	MEMBERS	% OF LOCATED
France	1	1.1%
Slovenia	1	1.1%
Cyprus	1	1.1%
Finland	1	1.1%
Ireland	1	1.1%
Jamaica	1	1.1%
Singapore	1	1.1%
United Arab Emirates	1	1.1%
Malaysia	1	1.1%
Puerto Rico	1	1.1%
Total (21 countries)	89	100.0%

Self-reported location, 89 of 203 members. Percentages are of located members.

4.3 Regional Distribution

REGION	MEMBERS	% OF LOCATED
Caribbean	36	40.4%
Europe	22	24.7%
North America	16	18.0%
Asia-Pacific	8	9.0%
Africa	6	6.7%
Middle East	1	1.1%
Total	89	100.0%

4.4 Why a Community Panel Is a Legitimate Basis

For a credential's founding content standard, a purposive panel of engaged practitioners is a recognized and defensible basis under ANSI/NCCA 1100-2024 and ISO/IEC 17024, which require a current, documented practice analysis validated by subject-matter experts. What matters at this stage is that the raters are genuinely practitioners, that the task model is

comprehensive across the domain, and that the linkage from task to item is enforced. Statistical generalization to the global practitioner population is not claimed at this stage and is built as the candidate population and item bank grow (Section 14).

CONCENTRATION & ITS TREATMENT

The community is concentrated in Trinidad and Tobago (38% of located members), with the remainder distributed across 20 further countries. Concentration is disclosed rather than obscured; as the panel base broadens, geographic and role diversity are tracked and reported with each revalidation, and any task whose consensus depends on a single region is flagged for re-review.

SECTION 5 · COMMON BODY OF KNOWLEDGE

Common Body of Knowledge — Domains & Sub-Domains

The Common Body of Knowledge for the GAISB credential was published by GAISB as CBK Edition 1.0 with seven domains and pre-published weights. This practice analysis tested whether each domain and its tasks were validated by the panel and whether the published weights were concordant with panel consensus. All seven domains were validated; weighting concordance is documented in Section 7.

DOMAIN	SUB-DOMAINS
D1. AI Strategy & the Modern AI Economy (10%)	1.1 The AI economic shift · 1.2 Strategic AI thesis development · 1.3 Competitive AI mapping · 1.4 Use-case identification and prioritization · 1.5 Capability and resource gap analysis
D2. Foundations of Generative AI (12%)	2.1 Model families and capabilities · 2.2 Cost and economics of generative AI · 2.3 Vendor and model selection · 2.4 Risk classification of AI use cases · 2.5 Capability-to-outcome translation
D3. Prompt Engineering & LLM System Design (18%)	3.1 Production prompt design · 3.2 RAG patterns and grounded outputs · 3.3 Structured outputs and tool use · 3.4 Prompt evaluation and testing · 3.5 Prompt safety and injection defense · 3.6 Cross-model portability and swap-readiness · 3.7 Production monitoring
D4. AI Agents & Agentic Workflows (18%)	4.1 Agent architecture and patterns · 4.2 Tool integration and action spaces · 4.3 Human-in-the-loop design · 4.4 Agent control boundaries · 4.5 ROI and economic modeling for agents · 4.6 Agent testing and adversarial robustness · 4.7 Agent monitoring and operations
D5. Ethics, Data & Responsible AI (12%)	5.1 Bias and fairness audits · 5.2 Governance approval workflows · 5.3 Responsible-AI principles in design · 5.4 Data provenance, licensing, consent · 5.5 Transparency and disclosure · 5.6 Lifecycle responsible-AI checkpoints
D6. AI Strategy, Transformation & Business Innovation (15%)	6.1 90-day transformation planning · 6.2 Quick-win identification and sequencing · 6.3 ROI projection and validation · 6.4 Executive and board communication · 6.5 Change management for AI adoption · 6.6 Cross-functional team coordination · 6.7 OKR and metric alignment
D7. Innovation & Applied AI Foundations (15%)	7.1 AI opportunity portfolio scoring · 7.2 Business case development · 7.3 Prototype planning and execution · 7.4 Experimentation frameworks · 7.5 PoC-to-production transition · 7.6 Innovation pipeline operations · 7.7 Lessons-learned dissemination

SECTION 6 · TASK INVENTORY & PANEL CONSENSUS

Task Inventory & Panel Consensus

6.1 Consensus Tiers

Each retained task carries a panel consensus tier reflecting how central it is to competent practice:

TIER	MEANING
ESSENTIAL	Core to competence; high importance and high consequence. Defines the credential.
HIGH	Important and regularly performed; strongly expected of a competent holder.
SUPPORTING	Valuable and situational; retained but weighted below Essential and High tasks.

6.2 Tier Distribution

Of 84 candidate tasks, 67 were retained: 25 Essential, 32 High, and 10 Supporting. The 17 removed tasks fell below the Supporting threshold, failed cross-jurisdictional generalization, or were merged as redundant.

6.3 Panel Priority Tasks (Essential tier)

The tasks the panel placed at the top of the Essential tier concentrate in Prompt Engineering (D3) and Agentic Workflows (D4) — the practical anchor of current AI-practitioner work.

DOMAIN	TASK STATEMENT	TIER
D3	Author production-grade prompts with system context, tools, and constraints.	ESSENTIAL
D3	Design retrieval-augmented generation (RAG) patterns for grounded outputs.	ESSENTIAL
D4	Design multi-step agent workflows with clear step boundaries.	ESSENTIAL
D4	Embed human-in-the-loop checkpoints at high-risk steps.	ESSENTIAL
D2	Build LLM selection matrices comparing performance, cost, and risk.	ESSENTIAL
D3	Define structured output schemas (JSON, function calling, tool use).	ESSENTIAL
D3	Apply prompt safety techniques (injection defense, output filtering).	ESSENTIAL
D4	Integrate tools, APIs, and data sources into agent action spaces.	ESSENTIAL
D1	Identify high-value AI use cases prioritized by business outcome.	ESSENTIAL
D2	Risk-classify AI use cases by impact, sensitivity, and reversibility.	ESSENTIAL
D3	Build prompt evaluation harnesses with test sets and scoring rubrics.	ESSENTIAL

DOMAIN	TASK STATEMENT	TIER
D3	Monitor production prompt performance, regression, and drift.	ESSENTIAL

SECTION 7 · DOMAIN WEIGHTING & PANEL CONCORDANCE

Domain Weighting & Panel Concordance

7.1 Method

The published CBK 1.0 domain weights were tested for concordance with panel consensus. For each domain, the panel's task-tier profile (the count of Essential, High, and Supporting tasks) was compared against the published weight. A weight was confirmed when the domain's task coverage and tier profile were consistent with its published share — that is, higher-weighted domains carry more validated tasks overall, with the highest-weighted domains additionally carrying the greatest Essential-tier density.

7.2 Concordance Table

DOMAIN	TASKS	ESSEN-TIAL	HIGH	SUPPORT-ING	WEIGHT	CONCORD-ANCE
D1. AI Strategy & the Modern AI Economy	7	2	4	1	10%	Confirmed
D2. Foundations of Generative AI	8	2	5	1	12%	Confirmed
D3. Prompt Engineering & LLM System Design	12	10	1	1	18%	Confirmed
D4. AI Agents & Agentic Workflows	12	8	3	1	18%	Confirmed
D5. Ethics, Data & Responsible AI	8	1	7	0	12%	Confirmed
D6. AI Strategy, Transformation & Business Innovation	10	1	7	2	15%	Confirmed
D7. Innovation & Applied AI Foundations	10	1	5	4	15%	Confirmed
Total	67	25	32	10	100%	—

Every domain's profile was concordant with its published weight. Weight tracks task coverage — domains carrying more validated tasks hold larger examination shares (7–8 tasks at 10–12%; 10 tasks at 15%; 12 tasks at 18%) — while Essential-tier density concentrates in and reinforces the two highest-weighted domains, D3 and D4. No domain's consensus contradicted its published weight, so the Standards Council retained the published weights as final. They are the controlling reference for examination construction.

WEIGHTING NOTE

Because this edition validates weights by panel concordance rather than by a large-sample survey, the weights are held at their published values and re-tested at each revalidation. As the item bank accumulates operational data, domain weights will additionally be checked against empirical item performance, and any material divergence will be reviewed by the Examination Development Committee.

SECTION 8 · CBK VALIDATION OUTCOMES

CBK Validation Outcomes

All seven CBK 1.0 domains were validated by the panel: each retained at least seven tasks at Supporting tier or above, and each retained at least one Essential-tier task. No domains were eliminated and none were added. Two sub-domains were rescoped following panel review.

SUB-DOMAIN	ACTION	REASON
3.6 Cross-model portability and tuning	Rescoped to 3.6 Cross-model portability, swapping, and provider-agnostic design	Panel noted swap-readiness, not tuning, is the dominant operational practice.
7.4 Experimentation frameworks	Rescoped to 7.4 Experimentation frameworks for AI features	Clarified scope versus classical product experimentation.

8.1 Final Validated CBK

#	FINAL DOMAIN	TASKS	WEIGHT
D1	AI Strategy & the Modern AI Economy	7	10%
D2	Foundations of Generative AI	8	12%
D3	Prompt Engineering & LLM System Design	12	18%
D4	AI Agents & Agentic Workflows	12	18%
D5	Ethics, Data & Responsible AI	8	12%
D6	AI Strategy, Transformation & Business Innovation	10	15%
D7	Innovation & Applied AI Foundations	10	15%
—	Total	67	100%

8.2 Domain Coherence

The panel judged the seven domains distinct and non-overlapping in the work they describe, with the expected adjacency between Prompt Engineering (D3) and Agentic Workflows (D4) — LLM system design and agent design share operational foundations without collapsing into one domain. The CBK is appropriately differentiated.

SECTION 9 · GAISB ASSESSMENT ARCHITECTURE

GAISB Assessment Architecture

GAISB uses a hybrid summative assessment: a structured knowledge examination plus the AIS-400 Capstone. All three GAISB credentials — Essentials, Practitioner, and Team Architect — carry both a written examination and an applied capstone; the credential is conferred only when both components are passed and the candidate signs the Code of Conduct at-testation.

9.1 Knowledge Examination — Design Parameters

PARAMETER	SPECIFICATION
Blueprint (Practitioner & Team Architect)	75 items, administered over a three-hour session
Essentials Form	Shorter 30-item form drawn proportionally from the same blueprint
Pretest Items	A small number of unscored pretest items are embedded per form for calibration and are not counted toward the candidate's score
Item Format	Multiple choice (single best answer); scenario-based stems
Scoring & Reporting	Scaled score 200–800; provisional pass mapped to a scaled 500 (~65% of operational items for Practitioner and Team Architect; 75% for Essentials), confirmed by standard-setting (Section 11)
Delivery	Online proctored; approved test centers on roadmap

9.2 Knowledge Exam — Item Allocation by Domain (75-item blueprint)

DOMAIN	WEIGHT	SCORED ITEMS	PRETEST
D1. AI Strategy & the Modern AI Economy	10%	8	2
D2. Foundations of Generative AI	12%	9	2
D3. Prompt Engineering & LLM System Design	18%	14	3
D4. AI Agents & Agentic Workflows	18%	13	2
D5. Ethics, Data & Responsible AI	12%	9	2
D6. AI Strategy, Transformation & Business Innovation	15%	11	2
D7. Innovation & Applied AI Foundations	15%	11	2
Total	100%	75	15

9.3 Cognitive-Level Distribution

COGNITIVE LEVEL (REVISED BLOOM'S)	% OF ITEMS	COUNT	FOCUS
Remember / Understand	20%	15	Recall concepts, definitions, regulatory text.
Apply	40%	30	Apply prompt patterns, agent designs, frameworks.
Analyze	25%	19	Compare and decompose AI design tradeoffs.
Evaluate / Create	15%	11	Judge tradeoffs; recommend architectures.
Total	100%	75	—

9.4 Capstone (AIS-400) — Design Parameters

All three credentials carry an applied capstone. GAISB AI Essentials™ and GAISB AI Practitioner™ carry a light applied capstone — an applied artifact or workflow reviewed by a single credentialed reviewer against a published rubric. GAISB AI Team Architect™ carries the intensive capstone below. Full capstone procedures are published in the Capstone Specification & Rubric (CAP-2026-01).

PARAMETER	SPECIFICATION
Acceptable Capstone Types	A Deployed AI System; an AI-Enabled Business or Venture; an Enterprise AI Transformation Program; or an AI Governance Deployment
Submission Components	Written narrative · evidence artifacts · recorded on-camera walkthrough
Intensive Review Panel	Four GAISB-authorized reviewers drawn from the Faculty roster and the Ethics Review Board, including at least one Standards Council member or designee
Rubric	Published six-dimension rubric (Section 10); each dimension scored 0–100
Pass Rule (Team Architect)	The 75/60 rule: the candidate passes only if the mean of all four reviewers is at least 75 and no single dimension mean falls below 60

9.5 Item & Rubric Traceability Requirement

ITEM-WRITING POLICY

Every item submitted to the GAISB knowledge-examination item bank, and every dimension of the AIS-400 Capstone rubric, must cite a specific Task ID from Appendix B and a Sub-Domain ID from Section 5. Items and rubric dimensions without traceability are rejected at Content Review. This rule preserves the empirical chain from job practice → task → assessment.

SECTION 10 · CAPSTONE RUBRIC

Capstone Rubric

The AIS-400 Capstone is evaluated against a published six-dimension rubric. Each dimension is scored on a 0–100 scale; descriptor bands anchor the scale so that scoring is criterion-referenced rather than relative. Full level anchors for each dimension appear in Appendix D.

10.1 Score Bands

BAND	SCORE	DEFINITION
Emerging	< 60	Incomplete or shows significant misunderstanding; not deliverable in a real engagement.
Developing	60–74	Partial competence; meets some requirements with notable gaps in depth, execution, or clarity.
Proficient	75–89	Meets the standard expected of a competent GAISB holder; deliverable to a real audience with minor revision.
Distinguished	90–100	Exceeds the standard; exemplary judgment, originality, and execution; usable as a reference exemplar.

10.2 Rubric Dimensions

#	DIMENSION	PRIMARY DOMAINS	ANCHORED IN
R1	Technical quality	D2, D3, D4	Appendix D.1
R2	Strategic soundness	D1, D6	Appendix D.2
R3	Responsible-AI practice	D5	Appendix D.3
R4	Applied rigor	D4, D7	Appendix D.4
R5	Professional documentation	All domains	Appendix D.5
R6	Defense & communication (on-camera walkthrough)	D1, D6	Appendix D.6

CAPSTONE PASS RULE (TEAM ARCHITECT — INTENSIVE)

The 75/60 rule governs pass/fail: the candidate passes only if the mean of all four reviewers is at least 75 and no single dimension mean falls below 60. Essentials and Practitioner light capstones are reviewed by a single credentialed reviewer against the same six-dimension rubric with published pass criteria.

10.3 Reviewer Calibration

Reviewers complete a calibration set of exemplar capstones before scoring live submissions. Inter-reviewer agreement is monitored continuously; reviewers whose agreement drops below the published threshold over a rolling window re-calibrate. Any dimension with a material scoring disagreement is routed to the panel Chair, who writes the decision record.

SECTION 11 · PASS-SCORE & STANDARD-SETTING PLAN

Pass-Score & Standard-Setting Plan

11.1 Knowledge Examination — Modified-Angoff

The knowledge-examination cut score is established using a two-round modified-Angoff procedure with a nine-member SME panel drawn from the validation panel and supplemented to maintain domain representation. Each panelist independently estimates, for every operational item, the probability that a Minimally Competent GAISB holder would answer correctly. Item-level estimates are averaged within panelist; the cut score is the mean of panelist averages, SEM-adjusted and scaled.

MINIMALLY COMPETENT GAISB HOLDER (CERTIFIED IN AI STANDARDS)

The Minimally Competent GAISB holder consistently designs, deploys, and operates AI systems — including LLM systems and agentic workflows — with sound judgment under realistic production constraints; selects models, prompts, retrieval patterns, and agent architectures appropriate to the use case; embeds responsible-AI controls (oversight, fairness, transparency, data governance) into the system as designed, not as an afterthought; recognizes when an issue exceeds their authority and escalates appropriately; and communicates the technical and economic tradeoffs of AI-system choices clearly to both technical and executive audiences.

11.2 Procedure (Knowledge Exam)

1. Round 1 — independent panelist ratings, no discussion; SEM and inter-rater agreement computed.
2. Calibration discussion — panelists review item-level disagreements; high-disagreement items receive structured re-discussion without peer pressure.
3. Round 2 — independent re-rating; final cut computed as the mean of Round 2 panelist averages.
4. Conditional SEM at the cut is applied as a one-CSEM band to set the operational pass score.
5. Cut reviewed against expected pass rate, item-difficulty distribution, and stakeholder-impact analysis.
6. Final cut approved by the Examination Development Committee and ratified by the Standards Council.

11.3 Capstone — Calibrated-Rubric Standard-Setting

The capstone standard is set by calibrating the six-dimension rubric against blind-reviewed exemplar capstones sampled to span the expected distribution. Reviewers score exemplars on the 0–100 scale; the Minimally Competent standard is fixed at the 75/60 rule (§10.2). Inter-judge concordance is measured and must meet the published threshold to certify the standard. Calibration is refreshed annually.

11.4 Reporting & Equating

Knowledge-exam raw scores are transformed to a 200–800 scaled metric with the cut mapped to 500. Across forms, common-item non-equivalent groups (CINEG) equating is used, with each operational form sharing anchor items with a prior form;

equating is performed under independent psychometric review. Capstone scores are reported per dimension and as a holistic decision (Pass / Hold for revision / Did not pass).

SECTION 12 · EXAMINATION SECURITY & ADMINISTRATION

Examination Security & Administration

This section states the operating specification for GAISB examination security and administration — the controls implemented as the item bank and testing program go live. It is the standard to which delivery is held, not a report of a program already at operational scale.

12.1 Item-Bank Security

- Item bank stored in a SOC 2-aligned credentialing platform with role-based access control and full audit logging.
- Item authors and reviewers sign per-engagement non-disclosure and conflict-of-interest agreements; access expires when the engagement ends.
- Items rotated across forms on a planned schedule; per-item exposure monitored and capped between rest periods.
- Forensic monitoring uses response-pattern, response-time, and answer-similarity analysis on every administration.

12.2 Capstone Submission Integrity

- Capstone submissions undergo automated similarity screening against the GAISB exemplar archive and public sources.
- All capstone submissions include a recorded on-camera walkthrough in which the candidate explains design choices and answers reviewer questions.
- Use of AI tools is permitted and expected in capstone work, but candidates must disclose tool use and demonstrate authorial judgment in design and decision-making.

12.3 Candidate Identity & Proctoring

- Government-issued ID verification at check-in; biometric confirmation re-checked mid-session for the knowledge exam.
- Live human proctoring for online administrations; AI-assisted behavior monitoring is a supplementary signal only, never the sole basis for adverse action.

12.4 Incident Handling

- Suspected breach handled under documented procedure; affected items quarantined; impacted forms re-equated or retired as needed.
- Score holds and investigations follow a due-process protocol with a single right of appeal to the Standards Council, consistent with the Code of Conduct.

12.5 Accommodations

Reasonable accommodations are provided for documented disabilities consistent with applicable jurisdictional law. Requests are reviewed by a credentialed accommodations specialist; psychometric integrity of accommodated administrations is monitored and reported in the annual psychometric report.

SECTION 13 · RELIABILITY, FAIRNESS & LIMITATIONS

Reliability, Fairness & Limitations

13.1 Reliability & Fairness Plan

Because this founding edition validates the task model by panel consensus, reliability and fairness evidence for the examination is generated on operational data as the item bank goes live. The following are computed and published in the annual Psychometric Report against the stated thresholds:

EVIDENCE	METHOD	THRESHOLD
Internal consistency	Cronbach's α on operational forms	≥ 0.80
Item quality	Classical difficulty & discrimination; IRT calibration as volume permits	Published bands
Differential item functioning	Mantel-Haenszel by region, role, and language	ETS A/B/C flag
Form comparability	CINEG equating with anchor items	Published tolerance
Capstone reliability	Inter-reviewer agreement, rolling window	≥ 0.75

13.2 Limitations (stated plainly)

- **Founding, panel-validated basis:** task importance is a purposive expert-panel consensus, not a large-sample probability survey. Population inferential statistics and margins of error are not claimed for this edition.
- **Community concentration:** the validation community is concentrated in Trinidad and Tobago (38% of located members) with a global tail across 20 further countries. Geographic and role diversity are tracked and broadened at each revalidation.
- **Self-reported location:** location was disclosed by 89 of 203 members; reported geography describes located members, not the full community.
- **Role data not collected:** the community roster does not capture standardized role, seniority, or industry fields; role coverage is addressed qualitatively and will be captured in the next panel cycle.
- **Examination psychometrics pending operational data:** reliability, DIF, and equating evidence is generated once candidates sit live forms (§13.1).
- **Temporal limitation:** the AI-practitioner role is evolving rapidly; the model is calibrated to 2025–2026 practice, with full re-validation on the triennial cadence (Q2 2029).

13.3 Standards Conformance

REQUIREMENT	REFERENCE	STATUS
Documented practice analysis	ISO 17024 §9.1.2; NCCA 1100 §8.1	Conforms (this report)
Validated by subject-matter experts	NCCA 1100 §8.0	Conforms (practitioner panel)
Task-to-item linkage documented	NCCA 1100 §8.4	Required by policy (§9.5)
Standard-setting plan defined	NCCA 1100 §9.0	Conforms (Section 11)
Reliability & fairness on operational data	Standards (AERA/APA/NCME)	Planned (§13.1)
Periodic review & update	ISO 17024 §9.4	Triennial (Section 14)

SECTION 14 · RECOMMENDATIONS & MAINTENANCE PLAN

Recommendations & Maintenance Plan

14.1 Recommendations

1. Adopt the validated 7-domain CBK, the examination blueprint, and the capstone rubric in this report as controlling references for the GAISB credential through the next triennial review.
2. Brief item writers and authorized training partners on the rescoped sub-domains (3.6, 7.4) and the consensus tiers.
3. Implement traceability tooling so every knowledge-exam item and every capstone rubric dimension carries a Task ID and Sub-Domain ID; reject items without traceability at Content Review.
4. Conduct knowledge-exam standard-setting (modified-Angoff) and capstone calibration before the first operational administration.
5. Capture standardized role, seniority, and industry data in the next panel cycle, and broaden the panel beyond the founding community's geographic concentration.
6. Begin generating operational reliability, DIF, and equating evidence with the first live forms, and publish it in the annual Psychometric Report.

14.2 Maintenance Plan

ACTIVITY	CADENCE	OWNER
Item-performance monitoring	Continuous	Independent Psychometric Review
Capstone inter-reviewer reliability monitoring	Continuous (rolling windows)	Capstone Review Lead
Annual Practice Review — scan & report	Annual	GAISB Standards Division
Panel expansion & role-data capture	Each panel cycle	GAISB Standards Division
Full practice-analysis re-validation	Triennial (next: Q2 2029)	GAISB Standards Division

14.3 Triggers for Off-Cycle Re-Validation

- Material change to a primary referent framework (NIST AI RMF, EU AI Act, ISO/IEC 42001) introducing new practitioner-relevant obligations.
- Sustained item-level psychometric drift over consecutive administrations.
- Notable shift in candidate-population composition from the launch baseline.
- An emergent task family judged by the panel as essential and not represented in the current inventory.

APPENDIX A

Panel Rating Instrument

A.1 Panel Eligibility

Panelists are active members of the practitioner community with substantive responsibility for, or contribution to, the design, deployment, or operation of AI systems — including LLM systems, agentic workflows, prompt engineering, AI transformation, or AI innovation work.

A.2 Rating Instructions

Each task statement describes work that may be performed by AI practitioners. For each task, panelists rate (a) how often it is performed, (b) how important it is to competent practice, and (c) how serious the consequences would be if it were performed incorrectly or omitted. The panel's aggregate judgment is expressed as a consensus tier (Essential / High / Supporting).

A.3 Rating Scales

VALUE	FREQUENCY	IMPORTANCE	CRITICALITY
1	Never	Not important	No consequence
2	Rarely (a few times/year)	Slightly important	Minor consequence
3	Occasionally (monthly)	Moderately important	Moderate consequence
4	Frequently (weekly)	Very important	Major consequence
5	Daily / continuously	Critically important	Severe / catastrophic

A.4 Tier Assignment

Consensus ratings map to tiers as follows: Essential — high importance and high consequence, performed regularly by competent practitioners; High — important and regularly performed; Supporting — valuable and situational. Tasks the panel judged peripheral, redundant, or not generalizable across jurisdictions were removed or merged.

APPENDIX B

Full Task Inventory (67 tasks)

All 67 retained task statements with assigned domain and panel consensus tier. Items submitted to the GAISB knowledge-examination item bank, and dimensions of the AIS-400 Capstone rubric, must cite a Task ID below.

TASK ID	TASK STATEMENT	DOMAIN	TIER
T-D1-01	Articulate where AI creates economic value across cost, revenue, and risk levers.	D1	HIGH
T-D1-02	Map competitive AI moves of industry peers and quantify exposure.	D1	HIGH
T-D1-03	Identify high-value AI use cases prioritized by business outcome.	D1	ESSENTIAL
T-D1-04	Build a 12-month AI strategic thesis tied to enterprise strategy.	D1	ESSENTIAL
T-D1-05	Conduct resource and capability gap analysis for AI adoption.	D1	HIGH
T-D1-06	Brief executives on AI economic shifts and their organizational implications.	D1	HIGH
T-D1-07	Translate enterprise objectives into AI-portfolio investment choices.	D1	SUPPORTING
T-D2-01	Differentiate among major model families (LLM, multimodal, reasoning, agentic) by use case.	D2	HIGH
T-D2-02	Build LLM selection matrices comparing performance, cost, and risk.	D2	ESSENTIAL
T-D2-03	Translate model capabilities into business outcomes for non-technical audiences.	D2	HIGH
T-D2-04	Estimate token, inference, and operating costs for AI deployments.	D2	HIGH
T-D2-05	Identify when to use foundation, fine-tuned, or purpose-built systems.	D2	HIGH
T-D2-06	Risk-classify AI use cases by impact, sensitivity, and reversibility.	D2	ESSENTIAL
T-D2-07	Define vendor evaluation criteria for foundation-model providers.	D2	HIGH
T-D2-08	Assess provider transparency and contractual control for high-impact use cases.	D2	SUPPORTING
T-D3-01	Author production-grade prompts with system context, tools, and constraints.	D3	ESSENTIAL
T-D3-02	Design retrieval-augmented generation (RAG) patterns for grounded outputs.	D3	ESSENTIAL
T-D3-03	Define structured output schemas (JSON, function calling, tool use).	D3	ESSENTIAL
T-D3-04	Build prompt evaluation harnesses with test sets and scoring rubrics.	D3	ESSENTIAL
T-D3-05	Iterate prompts using systematic testing and version control.	D3	ESSENTIAL
T-D3-06	Design few-shot and chain-of-thought patterns for complex tasks.	D3	ESSENTIAL
T-D3-07	Apply prompt safety techniques (injection defense, output filtering).	D3	ESSENTIAL

TASK ID	TASK STATEMENT	DOMAIN	TIER
T-D3-08	Tune prompts across providers for portability and swap-readiness.	D3	HIGH
T-D3-09	Monitor production prompt performance, regression, and drift.	D3	ESSENTIAL
T-D3-10	Design prompt cost controls (token budgets, caching, model routing).	D3	ESSENTIAL
T-D3-11	Translate business requirements into prompt specifications.	D3	ESSENTIAL
T-D3-12	Operate prompt change management with documented approvals.	D3	SUPPORTING
T-D4-01	Design multi-step agent workflows with clear step boundaries.	D4	ESSENTIAL
T-D4-02	Integrate tools, APIs, and data sources into agent action spaces.	D4	ESSENTIAL
T-D4-03	Embed human-in-the-loop checkpoints at high-risk steps.	D4	ESSENTIAL
T-D4-04	Define agent control boundaries (what it may and may not do).	D4	ESSENTIAL
T-D4-05	Calculate ROI for agent-driven automation.	D4	ESSENTIAL
T-D4-06	Architect agentic systems with planner-executor patterns.	D4	HIGH
T-D4-07	Design fallback and graceful-degradation behaviors.	D4	HIGH
T-D4-08	Test agent behaviors across edge cases and adversarial inputs.	D4	ESSENTIAL
T-D4-09	Monitor agent performance, success rate, and cost-per-task.	D4	ESSENTIAL
T-D4-10	Coordinate agent rollouts with operational and risk teams.	D4	HIGH
T-D4-11	Define escalation and override paths for agent operations.	D4	ESSENTIAL
T-D4-12	Decompose business processes into agent-amenable workflows.	D4	SUPPORTING
T-D5-01	Conduct AI bias and fairness audits on production systems.	D5	ESSENTIAL
T-D5-02	Build governance approval workflows with named decision-makers.	D5	HIGH
T-D5-03	Apply responsible-AI principles to design and procurement decisions.	D5	HIGH
T-D5-04	Manage data provenance, licensing, and consent across the lifecycle.	D5	HIGH
T-D5-05	Operate transparency and disclosure requirements appropriate to use case.	D5	HIGH
T-D5-06	Embed responsible-AI checkpoints in the AI development lifecycle.	D5	HIGH
T-D5-07	Coordinate ethical review of high-impact AI use cases.	D5	HIGH
T-D5-08	Translate AI regulations into operational controls within delivery work.	D5	HIGH
T-D6-01	Build 90-day AI transformation plans with quick wins and longer arcs.	D6	ESSENTIAL

TASK ID	TASK STATEMENT	DOMAIN	TIER
T-D6-02	Identify and sequence high-leverage AI quick-wins.	D6	HIGH
T-D6-03	Project ROI for AI initiatives and validate post-deployment.	D6	HIGH
T-D6-04	Develop executive and board-level communications on AI strategy and outcomes.	D6	HIGH
T-D6-05	Lead organizational change management for AI adoption.	D6	HIGH
T-D6-06	Coordinate cross-functional AI teams (product, engineering, risk, legal).	D6	HIGH
T-D6-07	Align AI investments with enterprise OKRs and strategy.	D6	HIGH
T-D6-08	Communicate AI outcomes to internal and external stakeholders.	D6	SUPPORTING
T-D6-09	Define AI program metrics and dashboards.	D6	SUPPORTING
T-D6-10	Establish AI operating cadences and governance forums for scaled delivery.	D6	HIGH
T-D7-01	Score AI opportunity portfolios across feasibility, value, and risk.	D7	ESSENTIAL
T-D7-02	Build investment-grade business cases for top AI opportunities.	D7	HIGH
T-D7-03	Plan 1-week and 30-day AI prototypes with clear success criteria.	D7	HIGH
T-D7-04	Apply experimentation frameworks to AI features and prompts.	D7	HIGH
T-D7-05	Measure AI experiment outcomes and inform go/no-go decisions.	D7	SUPPORTING
T-D7-06	Coordinate proof-of-concept to production transitions.	D7	HIGH
T-D7-07	Identify and unblock systemic barriers to AI deployment.	D7	SUPPORTING
T-D7-08	Capture and disseminate AI lessons learned across the organization.	D7	SUPPORTING
T-D7-09	Stand up an innovation pipeline with intake, scoring, and stage gates.	D7	HIGH
T-D7-10	Define reusable patterns and playbooks from delivered AI work.	D7	SUPPORTING

APPENDIX C

Sample Items with Task Traceback

The following non-operational sample items demonstrate the required traceability from validated task to assessment item. They are illustrative only and are not part of the operational item bank.

Sample Item C-1 (Apply — Domain 3)

TRACE — TASK ID T-D3-02 (DESIGN RETRIEVAL-AUGMENTED GENERATION PATTERNS FOR GROUNDED OUTPUTS)

TRACE — SUB-DOMAIN 3.2 RAG PATTERNS AND GROUNDED OUTPUTS

STEM	An organization is deploying an internal LLM assistant that must answer strictly from current company policies (revised quarterly). Hallucination from outdated training data is a stated risk. Which design choice BEST addresses both currency and groundedness?
KEY (A)	Retrieve relevant policy passages at query time from a vector index updated after each revision, and condition the response on those passages with explicit citations.
RATIONALE	Option A is the canonical RAG pattern addressing both currency (re-indexing) and groundedness (passage-conditioned generation with citations). Fine-tuning does not solve currency between cycles; a bare system-prompt instruction does not constrain grounding.

Sample Item C-2 (Analyze — Domain 4)

TRACE — TASK ID T-D4-03 (EMBED HUMAN-IN-THE-LOOP CHECKPOINTS AT HIGH-RISK STEPS)

TRACE — SUB-DOMAIN 4.3 HUMAN-IN-THE-LOOP DESIGN

STEM	A finance team is rolling out an agent that prepares and submits supplier-invoice approvals end-to-end. Which placement of human-in-the-loop checkpoints BEST balances throughput with control?
KEY (C)	Human approval required only when invoice value exceeds a defined threshold, when validation confidence is low, or when the supplier is new — straight-through processing otherwise.
RATIONALE	Risk-tiered human-in-the-loop design is the dominant pattern for agent operations balancing throughput and control; approving every step negates throughput, and after-the-fact monitoring abandons control.

Sample Item C-3 (Evaluate / Create — Domains 1 & 6)

TRACE — TASK ID T-D1-04 (12-MONTH AI STRATEGIC THESIS); T-D6-01 (90-DAY TRANSFORMATION PLAN)

TRACE — SUB-DOMAIN 1.2 STRATEGIC AI THESIS DEVELOPMENT; 6.1 90-DAY TRANSFORMATION PLANNING

STEM	A regional retailer's CEO wants to "meaningfully integrate AI within 12 months"; the CFO demands evidence of ROI within 90 days. Which 12-month thesis structure BEST serves both?
KEY (B)	A portfolio of three to five sequenced initiatives — at least one quick-win producing measurable ROI within 90 days, paired with a longer-arc transformation timed for months 9–12, all framed within an enterprise AI thesis tied to competitive positioning.
RATIONALE	Option B satisfies both the CEO's 12-month direction and the CFO's 90-day evidence requirement; a single late-delivering program yields nothing in 90 days, and a platform-first RFP inverts use-case-driven selection.

APPENDIX D

Capstone Rubric — Full Anchors

The six-dimension rubric is applied to all capstone types. Each dimension is scored 0–100 against the band definitions in §10.1 (Emerging < 60 · Developing 60–74 · Proficient 75–89 · Distinguished 90–100). Anchors below describe the Proficient (75–89) standard for each dimension; Distinguished exceeds it with quantified rigor and originality, and Developing/Emerging fall short on the italicized elements.

DIMENSION	PROFICIENT (75–89) ANCHOR
D.1 Technical quality	Production-grade prompt/agent/system design with system context, tools, and constraints; structured model comparison across performance, cost, risk, and provider transparency; testing and monitoring present and documented.
D.2 Strategic soundness	Clear strategic intent grounded in identified business outcomes; explicit economic framing across cost/revenue/risk levers; demonstrated link to enterprise strategy and a credible transformation sequence.
D.3 Responsible-AI practice	Responsible-AI principles embedded across design choices; named governance approval workflow; transparency and disclosure design appropriate to the use case; bias/fairness considerations addressed.
D.4 Applied rigor	Opportunity scoring across feasibility/value/risk; defensible business case; prototype or deployment plan with explicit success criteria, stop conditions, and a PoC-to-production pathway.
D.5 Professional documentation	Complete, well-structured written narrative and evidence artifacts; decisions traceable; deliverable to a real client or executive audience with minor revision.
D.6 Defense & communication	On-camera walkthrough clearly explains design choices and tradeoffs; candidate answers reviewer questions with authorial command of the work; communicates to both technical and executive audiences.

SCORING

Each of the six dimensions is scored 0–100 by each reviewer. For the intensive (Team Architect) capstone, the 75/60 rule governs: the mean of all four reviewers must be at least 75, and no single dimension mean may fall below 60. Full procedures are published in the Capstone Specification & Rubric (CAP-2026-01).

APPENDIX E

Comparative Positioning

Provided for stakeholder context. GAISB is positioned as a practitioner AI credential anchored in operational practice, with a hybrid examination-plus-capstone assessment. Comparison points reflect publicly available program documentation and are not endorsements.

ATTRIBUTE	GAISB (THIS CREDENTIAL)	VENDOR STRATEGY COURSES	VENDOR CLOUD-AI CERTS	GOVERNANCE-ONLY CREDENTIALS
Primary scope	AI system design + deployment + operations + transformation (full lifecycle)	AI strategy concepts	Vendor-specific services	AI governance and AI risk
Empirical basis	Practitioner-panel practice analysis (203-member community, 21 countries)	Course outline	Vendor product taxonomy	Vendor-published
Assessment	Knowledge exam + capstone	Course completion	Multiple-choice exam	Multiple-choice exam
Standards alignment	ISO/IEC 42001 (Clause 7.2 + Annex A. 4.6, ~78%); NIST AI RMF; EU AI Act; OECD	Varies	Vendor	Varies
Vendor-neutral	Yes	Sometimes	No	Often

GAISB is intentionally positioned to complement, not duplicate, adjacent credentials. Holders of governance-only or vendor-specific credentials moving into AI-systems leadership are a primary candidate population.

APPENDIX F

Community Data

Real composition of the founding practitioner community that anchors this practice analysis. Figures are drawn from the community roster; location is self-reported.

F.1 Community at a Glance

MEASURE	VALUE
Community members	203
Active members (signed in, trailing 30 days) — validation panel	117
Countries represented (located members)	21
Members with self-reported location	89
Community posts authored	361
Community comments	843
Endorsements exchanged	3,219

F.2 Geographic Distribution (located members)

COUNTRY	MEMBERS	% OF LOCATED
Trinidad and Tobago	34	38.2%
United Kingdom	12	13.5%
United States	11	12.4%
South Africa	6	6.7%
Canada	5	5.6%
New Zealand	2	2.2%
Belgium	2	2.2%
Netherlands	2	2.2%
Thailand	2	2.2%
Australia	2	2.2%
Germany	1	1.1%

COUNTRY	MEMBERS	% OF LOCATED
France	1	1.1%
Slovenia	1	1.1%
Cyprus	1	1.1%
Finland	1	1.1%
Ireland	1	1.1%
Jamaica	1	1.1%
Singapore	1	1.1%
United Arab Emirates	1	1.1%
Malaysia	1	1.1%
Puerto Rico	1	1.1%
Total (21 countries)	89	100.0%

APPENDIX G

Panel Method Notes

G.1 Why a Tiered-Consensus Method

For a founding credential validating its first content standard, a tiered-consensus method from a purposive expert panel is preferred over a synthetic large-sample statistic. It is honest about the evidence available, it is defensible under ISO/IEC 17024 and ANSI/NCCA 1100 (which require SME-validated practice analysis, not a probability survey), and it does not manufacture precision the data cannot support. As the candidate population grows, empirical item statistics supplement panel consensus (§13.1).

G.2 From Ratings to Tiers

Panelists rated each task on Frequency, Importance, and Criticality. Consensus was expressed as a tier: Essential (core to competence — high importance and high consequence), High (important and regular), and Supporting (valuable and situational). Importance and consequence were weighted above raw frequency so that low-frequency, high-consequence tasks are not undervalued.

G.3 Tier Summary

TIER	TASKS	SHARE OF INVENTORY
ESSENTIAL	25	37.3%
HIGH	32	47.8%
SUPPORTING	10	14.9%
Total	67	100.0%

G.4 What Strengthens This Basis Over Time

The evidentiary strength of the practice analysis increases as: the panel base broadens beyond the founding community's geographic concentration; standardized role, seniority, and industry data are captured; and operational examination data yields item statistics, reliability, DIF, and equating evidence (§13.1). Each revalidation folds this accumulating evidence into the weighting and blueprint.

APPENDIX H

Glossary

TERM	DEFINITION
AIS-101 ... AIS-400	GAISB module identifiers for the GAISB curriculum modules plus the AIS-400 Capstone.
CBK	Common Body of Knowledge. Validated set of domains and tasks underlying the credential.
Consensus tier	The panel's aggregate importance judgment for a task: Essential, High, or Supporting.
DIF	Differential Item Functioning. Item-level fairness analysis across reference and focal sub-groups.
EDC	Examination Development Committee. GAISB body that approves examination content and blueprint.
GAISB	Global AI Standards Body. The professional standards body issuing the GAISB credential.
JTA	Job Task Analysis. The practice-analysis study described in this report.
MCC	Minimally Competent GAISB holder — the reference person used to set the cut score under modified-Angoff (see §11.2).
NCCA	National Commission for Certifying Agencies, the accreditation arm of the Institute for Credentialing Excellence.
Prompt Atlas community	GAISB's founding practitioner community, whose active members form the validation panel.
RAG	Retrieval-Augmented Generation. LLM design pattern that grounds outputs in retrieved source material.
Validation panel	The active practitioners of the community who reviewed and rated the task model.
Standards Council	GAISB governance body responsible for ratifying standards and credentials; twelve seats on four-year staggered terms.

APPENDIX I

Sources & References

This edition references published standards, regulations, and the GAISB-published Common Body of Knowledge.

REFERENCE	ISSUING BODY
Global AI Standards Body — gaisb.org	GAISB
GAISB Common Body of Knowledge, Edition 1.0 (CBK-2026-01)	GAISB
GAISB Professional Code of Conduct (COC-2026-01)	GAISB
GAISB & ISO/IEC 42001 Alignment Dossier (ISO-DSR-2026-01)	GAISB
ISO/IEC 17024:2012 — Requirements for bodies certifying persons	International Organization for Standardization
ANSI/NCCA 1100-2024 — Accreditation of Certification Programs	Institute for Credentialing Excellence
Standards for Educational and Psychological Testing	AERA, APA, NCME
EEOC Uniform Guidelines on Employee Selection Procedures	U.S. Equal Employment Opportunity Commission
ISO/IEC 42001:2023 — AI Management System	International Organization for Standardization
NIST AI Risk Management Framework 1.0	U.S. National Institute of Standards and Technology
Regulation (EU) 2024/1689 — EU AI Act	European Union
OECD AI Principles (OECD/LEGAL/0449)	Organisation for Economic Co-operation and Development

COLOPHON

The GAISB Job Task Analysis — Founding Practice Analysis · Edition 1.0 · JTA-2026-01 · Effective Q2 2026. Issued by the Global AI Standards Body (GAISB™) and ratified by the GAISB Standards Council. Public document — cite with attribution. ISO/IEC and ANSI/NCCA documents are subject to copyright held by the issuing body and are referenced under fair use.