

SKILLALCHEMY: Open-World Agent Skill Creation

Hengjun Wang¹, Shuyue Wei^{2*}, Boyi Liu¹, Jun Yang³, Yongxin Tong^{1*}

¹State Key Laboratory of Complex & Critical Software Environment, Beihang University, Beijing, China

²Joint SDU–NTU Centre for Artificial Intelligence Research (C-FAIR), Shandong University, Jinan, China

³School of Automation, Northwestern Polytechnical University, Xi'an, China

Abstract

Agent skills are reusable procedural artifacts that extend language agents with specialized workflows, tool conventions, and domain behaviors at inference time. However, creating reliable skills still depends largely on human authorship, model priors, or execution traces. These sources are often unavailable for unfamiliar tasks, suggesting the need to create skills from open-world materials. In this paper, we study *open-world skill creation*: given an underspecified skill brief and a source-access specification, a creator must discover behavior-relevant requirements omitted by the brief and determine how broadly each source-derived procedure is justified. We propose SKILLALCHEMY, an admission-centered framework for source-grounded skill creation. SKILLALCHEMY identifies implicit requirements through contrastive evidence, admits candidate procedures based on evidence-supported scope, and compiles the admitted content into a grammar-guided skill package. Extensive experiments across 87 SKILLSBENCH v1.1 tasks demonstrate that our SKILLALCHEMY improves pass rate over no-skill execution by 19.9pp and the strongest automated baseline by 8.6pp and is comparable to human-curated skills.

1 Introduction

Agent skills are reusable procedural artifacts that enable language agents to dynamically load and execute specialized workflows and domain-specific behaviors at inference time. In current agent ecosystems, a skill is commonly packaged as a filesystem-based artifact centered on a SKILL.md file with metadata (e.g., skill descriptions) and instructions and may bundle scripts, references, assets, or other resources that an agent can load on demand (Anthropic 2026a; OpenAI 2026a).

Equipping agents with skills enables them to perform a wide range of practical tasks beyond model priors, including code-generation workflows, data-analysis pipelines and document-processing routines (Zhou et al. 2026a; Li et al. 2026; Liu et al. 2026a). Despite their effectiveness and flexibility as deployment-time extensions, existing skills are typically created by experts, generated from model priors, or distilled from reasoning traces (Zhao et al. 2024; Ni et al. 2026; Liu et al. 2026b; Zhang et al. 2026a; Yang et al. 2026).

However, these routes rely on a distinct procedural knowledge source that is not always accessible. Expert-crafted skills demand extensive manual labor, model priors limit self-generated skills, and trace-based skills require archived execution traces. These assumptions break down most severely for

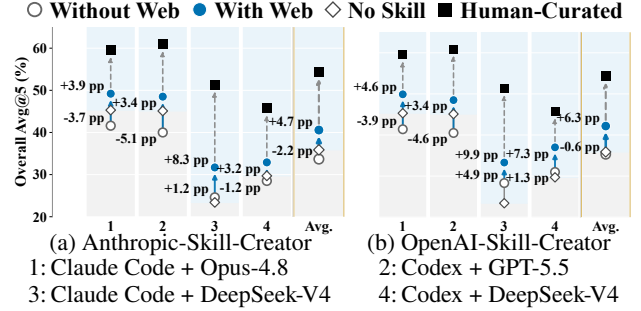


Figure 1: **Pilot Study**: Open-world source access improves skill creation but does not close the gap to human-curated skills.

unfamiliar tasks or capabilities, precisely when new custom skills are urgently required. In such cases, useful procedural knowledge may already be in open-world materials (including documentation, repositories and issue reports) yet remains largely under-exploited for reusable agent skill specifications.

Pilot Study. We examine two questions on SKILLSBENCH v1.1 (Li et al. 2026): *whether access to open-world sources improves automatic skill creation and whether such access alone closes the gap to human-curated skills?* We use two representative official skill creators released by Anthropic and OpenAI (Anthropic 2026d; OpenAI 2026a). Each creator builds skills for the same tasks with and without web access, and the resulting skills are executed by the same downstream agents (as in Figure 1). Without web access, the generated skills perform 1.4 percentage points below no-skill execution on average across four agent–model configurations. With web access, every configuration improves by 6.9 points on average over its without-web counterpart and by 5.5 points over no-skill execution. However, the stronger web-grounded creator remains about 14 points behind human-curated skills. These results show that open-world sources are beneficial but insufficient for reliable skill creation, motivating the central question of this work: how can reliable agent skills be created from open-world sources? Even with web grounding, the better web-access creator remains about 14 points behind human-curated skills, motivating our diagnosis: *open-world sources are informative but not skill-ready as they contain implicit decisions, local examples, and context-dependent practices rather than validated reusable procedures*. Specifically, converting open-world sources into reusable skills is challenging for two reasons (as illustrated in Figure 2).

*Corresponding authors: weishuyue@sdu.edu.cn, yxtong@buaa.edu.cn

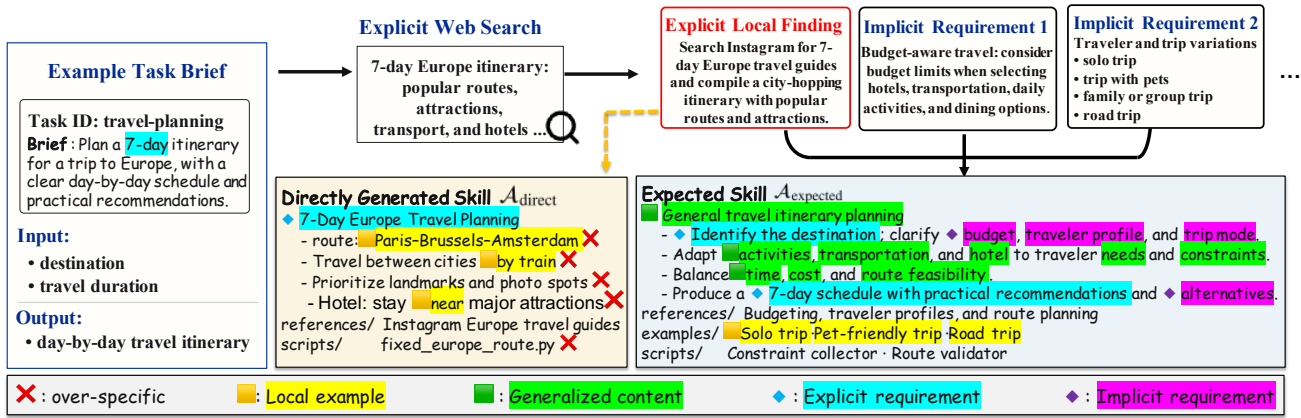


Figure 2: **Example illustrating why open-world skill creation is non-trivial:** ① task briefs underspecify implicit requirements, and ② directly adopting open-world findings without scope justification may lead to over-specific practices being mistaken for reusable instructions.

• Task briefs under-specify operational requirements.

A task brief typically states the immediate objective but leaves implicit the requirements, failure modes, and operational boundaries needed for a reusable skill. Using the brief directly as a retrieval query preserves these blind spots. Individual sources are also organized around their own subjects rather than around the complete target capability. *Thus, reliable skill creation must discover missing requirements beyond the brief before acquiring evidence.*

• Open-world findings do not justify their reusable scope.

A source-specific finding often mixes a reusable practice with local details, such as hardcoded parameters, preferred tools, fixed inputs, or environment-specific assumptions. The single occurrence is insufficient to justify promoting the entire finding into a persistent instruction. *The skill creation therefore needs to determine scope across cases, promoting consistent practices into general instructions, retaining context-bound ones as scoped examples, and excluding candidates whose support is weak or conflicting.*

We formulate the open-world skill creation as a source-grounded procedure-admission problem, *i.e.*, given an underspecified skill brief and a set of heterogeneous sources, the skill creator must first recover absent requirements from the brief and then decide whether each candidate is licensed as a reusable instruction, remains a scoped example, or should be excluded. To address these challenges, we propose SKILLALCHEMY, a framework that transforms an underspecified description of a target task or capability into an agent-usable skill. It discovers implicit requirements and admits instructions only when open-world evidence justifies their scope. SKILLALCHEMY is not merely a retrieval pipeline, but an admission-centered framework operates in three stages. (i) *Implicit Requirement Discovery* (§3.3) lifts the brief to its underlying capability, identifies omitted operational dimensions, and converts them into focused research questions. (ii) *Grounded Procedure Admission* (§3.4) aggregates relevant findings and admits a candidate as a reusable instruction only when the evidence supports both the action and its scope. (iii) *Skill Package Compilation* (§3.5) organizes admitted procedures and scoped examples into an installable skill package using the skill grammar and task-relevant exemplars.

Our main contributions are summarized as follows.

- We formulate open-world skill creation as a source-grounded procedure-admission problem, identifying two key challenges in converting an underspecified brief and heterogeneous sources into a reusable skill: implicit requirement discovery and procedure-scope justification.
- We propose SKILLALCHEMY, a skill-creation framework that turns implicit requirements in briefs into focused targets for open-world knowledge acquisition and determines whether source findings warrant general instructions, local examples, or exclusions of an installable skill package.
- We evaluate SKILLALCHEMY over 87 tasks from SKILLS-BENCH across four agent-model configurations. Our framework improves pass rate by 19.9 percentage points over no-skill execution and by 8.6 percentage points over the strongest automatic skill-creation baseline, comparable to the human-curated skills. Ablation studies further examine requirement coverage and unsupported procedure admission.

2 Related Work

Skills for Agents. Agent skills are typically treated as reusable procedural artifacts rather than ordinary prompts or atomic tool calls (Jiang et al. 2026; Zhou et al. 2026a; Vercel 2026). SkillAct (Liu et al. 2024) shows that adding reusable skill abstractions to existing prompting methods (*e.g.*, ReAct (Yao et al. 2023)) improves agent performance on interactive tasks such as ALFWorld (Shridhar et al. 2021). Skills-in-the-Wild further examines whether agents can effectively leverage skills in realistic settings, where useful skills need to be retrieved from a large and noisy collection rather than being manually curated (Liu et al. 2026a). Together, these studies establish an artifact-centric view of the skills lifecycle, in which skills can be constructed, represented, retrieved, invoked, and evaluated. Within such a lifecycle, skill construction determines what procedural knowledge is conceptualized into the repository in the first place, thereby directly shaping the utility of all downstream skill use. Aligned with this line of research, our SKILLALCHEMY studies how to produce the skill artifacts from open-world source materials.

Skill Creation. Skill creation methods can be categorized by the source of candidate skill content. (i) *Human-authored skills.* Expert-written or benchmark-provided skills (Li et al. 2026) encode human procedural knowledge and serve as strong reference artifacts. (ii) *Interaction traces* based skill creation. Voyager (Wang et al. 2023) builds an executable skill library from open-ended embodied interaction. ExpeL (Zhao et al. 2024) learns reusable lessons from task experience without updating model weights. SkillGen (Ma et al. 2026) synthesizes auditable skills from successful and failed traces and checks their net intervention effect. CoEvoSkills (Zhang et al. 2026b) iteratively evolves multi-file skill packages using surrogate verification and execution feedback. Together, these methods show how execution records can be transformed into reusable procedural knowledge, where the main evidence comes from observed attempts, failures, successes, or intervention effects. (iii) *Broad-source and scaffolded skill creation.* Official skill creation workflows provide general-purpose scaffolds for authoring and improving skills from available context (Anthropic 2026d; OpenAI 2026d). OpenSkill (Yan et al. 2026) retrieves documentation, repositories, and web resources to build transferable skills and verification anchors. SkillGenBench (Zhou et al. 2026b) benchmarks skill generation from repository- and document-grounded sources.

SKILLALCHEMY is closest to the broad-source creation line. Rather than treating retrieved or provided source content as skill content directly, it performs knowledge acquisition and evidence aggregation before skill creation, which focus complements human-authored and trace-based creation.

3 Method

3.1 Problem Formulation

We study *open-world skill creation*, where evidence needed to construct a reusable skill is not assumed to be organized as a task-complete corpus but must be identified and acquired from heterogeneous sources (*e.g.*, repositories, documents or agent experience) under a specified access policy.

The input is a tuple $(g, \mathcal{S}, \mathcal{C})$: (i) an underspecified skill brief g , namely a short natural-language description of the task or capability the skill should support, (ii) a source-access specification $\mathcal{S}$ defining permitted source types, retrieval channels, and exclusions (*e.g.*, documentation, repositories, or existing skills), and (iii) execution and packaging constraints $\mathcal{C}$ (*e.g.*, available tools and required artifact structure). A skill creator Φ maps these inputs to an installable skill package,

$$\mathcal{A} := \langle \text{SKILL}.\text{md}, \mathcal{X} \rangle = \Phi(g, \mathcal{S}, \mathcal{C}), \quad (1)$$

In this work, the skill package $\mathcal{A}$ follows the filesystem-based skill convention centered on `SKILL.md` and optionally bundlings $\mathcal{X}$ (*e.g.*, references, examples and scripts).

Source-Grounded Procedure Admission. Unlike a document summarization, which preserves what the sources state, an agent skill must specify *reusable procedures*: when an instruction applies, what the agent should do and produce, and the conditions under which it fails or falls outside scope. The aim of *open-world skill creation* is therefore not to compress or summarize a fixed source collection but to acquire evidence

materials under $\mathcal{S}$ and decide whether each candidate procedure should be admitted as a general instruction, retained as a scoped example or notes, or just be excluded from the skill.

3.2 Framework Overview

Design Rationale. Existing official general-purpose skill creation workflows typically author a skill directly from the task brief and available context (Anthropic 2026d; OpenAI 2026d). When the brief is underspecified, this process entangles three distinct decisions, *i.e.*, what omitted by the brief should be investigated, whether a source-derived candidate is supported beyond its local context, and how the admitted content should be expressed in the final artifact. SKILLALCHEMY separates these decisions explicit and resolves them sequentially. It first discovers implicit requirements and acquires evidence for them, then determines which candidate procedures are supported at which scope, and finally compiles the admitted content as an installable skill package. This separation is the central design choice of our framework.

Three-Stage Creation Process. Figure 3 illustrates the following three-stage workflow through a running example.

- *Stage 1: implicit requirement discovery* identifies behavior-relevant distinctions omitted in briefs, turns them into focused research questions, and acquires structured findings.
- *Stage 2: evidence-grounded procedure admission* aggregates findings that address the same situations, makes conflicts and missing support explicit, and distills supported content into general instructions or scoped examples.
- *Stage 3: skill package compilation* renders admitted instructions and scoped examples as an installable package under a corpus-derived skill grammar.

Writing Φ_1, Φ_2, Φ_3 for the three stages, they jointly compose the whole open-world skill creation pipeline of Eq. (1),

$$(g, \mathcal{S}) \xrightarrow{\Phi_1} (Q, F) \xrightarrow{\Phi_2} (R^g, R^e) \xrightarrow[\mathcal{C}]{\Phi_3} \mathcal{A}. \quad (2)$$

where Q is the set of research questions, F the structured findings, R^g the admitted general instructions, and R^e the retained scoped examples. Candidates admitted to neither R^g nor R^e are excluded. These intermediate outputs are preserved as explicit provenance records, whereas R^g and R^e are compiled into the skill content under $\mathcal{C}$. The following three subsections describe these creation stages subsequently.

3.3 Implicit Requirement Discovery

Operational Framing. A task brief may describe a requested instance, whereas a reusable skill must operate over a family of instances. Given a brief g and source-access specification $\mathcal{S}$, Φ_1 maps g to an operational frame $L_g = \text{Frame}(g)$. The frame treats brief-specific values as candidate operational factors and exposes the procedural decisions left unspecified for producing the requested output, handling failures, and verifying the result. For example, a brief requesting “a 7-day family trip to Europe” contains instance-specific values such as `7-day`, `family trip`, and `Europe`. The operational frame instead represents them as candidate factors such as `duration`, `traveler type`, and `destination`, rather than hard-coding them

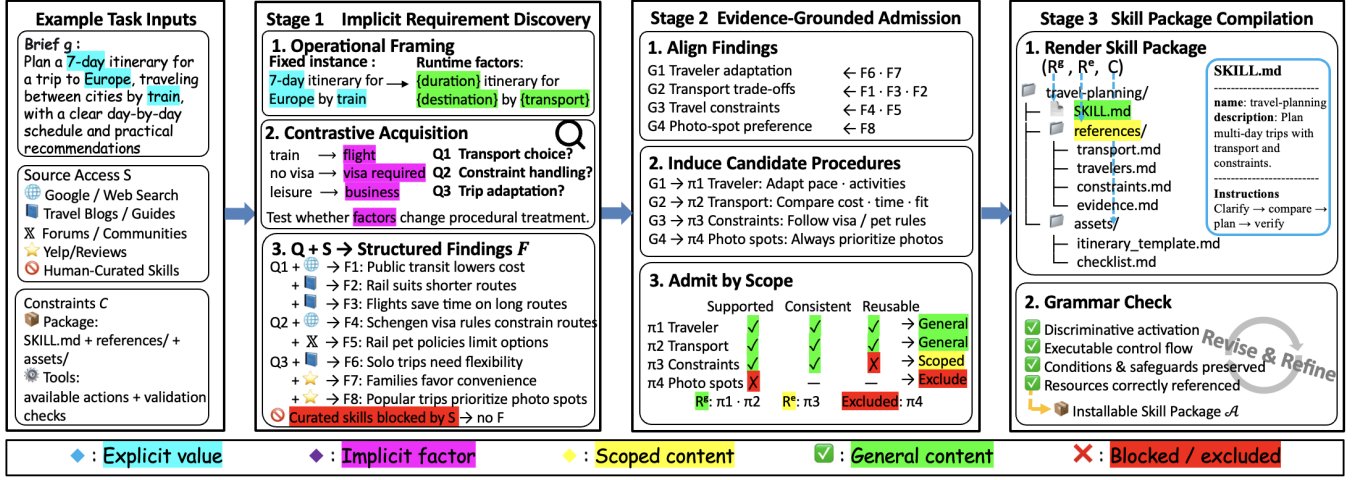


Figure 3: **Overview of SKILLALCHEMY**: Φ_1 converts an underspecified brief into focused questions and structured findings, Φ_2 induces candidate procedures and admits them via the evidence-supported scope, and Φ_3 compiles the admitted content into an agent skill package.

into a reusable procedure. The frame therefore proposes operational factors for investigation rather than treating them as requirements and a factor becomes an implicit operational factor only when acquired evidence shows that varying it changes procedural behavior.

Contrastive Evidence Acquisition. Directly searching with the original brief tends to retrieve repeated evidence about the same local instance, whereas open-ended requirement expansion may introduce many conditions unrelated to skill behavior. Thus, SKILLALCHEMY constructs paired acquisition targets $h = \langle d, x, x' \rangle$, where x and x' are matched task contexts that differ along one candidate operational factor d . Here, a source-stated applicability boundary along d counts as evidence that the corresponding component has different applicability across the matched contexts while missing evidence for either context alone does not establish such a difference. A target may vary a declared factor to test whether a procedure transfers beyond the seed instance, or an omitted factor to test whether the brief lacks a behavior-changing condition. Specifically, (i) a *substitution probe* varies a subject, method, or tool within the same capability family. For example, a substitution probe may replace a family traveler with a solo traveler while keeping the destination and duration fixed. (ii) A *boundary probe* introduces an omitted precondition, failure, or operating constraint. (iii) A *neighbor probe* compares the target with a sibling capability under the same output interface. Each target is converted into a focused question asking whether the matched contexts require different treatment in any component $k \in \mathcal{K} = \{c, a, r, v\}$, where c, a, r , and v denote conditions, action, recovery, and verification, respectively. These components characterize procedural behavior rather than just prescribing the layout of the SKILL.md. For findings F_h acquired for targets h , let

$$K_F(h) = \{k \in \mathcal{K} \mid F_h \models \text{Treat}_k(x) \neq \text{Treat}_k(x')\}. \quad (3)$$

Only when $K_F(h) \neq \emptyset$ does SKILLALCHEMY record an implicit requirement to condition procedural behavior on d , together with the affected components and any evidence-stated boundary. Evidence supporting the same treatment

across non-equivalent contexts is retained as cross-context invariance evidence. Conflicting findings remain explicit, whereas insufficient evidence leaves the contrast unresolved. Finally, Φ_1 returns questions Q and structured findings F , which remain source-grounded observations rather than executable procedures. Next, Stage 2 Φ_2 determines whether these findings are justified to be general instructions, scoped examples, or exclusions.

3.4 Evidence-Grounded Procedure Admission

Decision-Aligned Induction. Since the findings F remain tied to local source contexts, Φ_2 first groups findings by the procedural decision they inform, rather than by their source topic or document,

$$\{G_1, \dots, G_m\} \leftarrow \text{Align}(F), \quad \pi_j \leftarrow \text{Induce}(G_j). \quad (4)$$

Each Stage-1 finding retains the acquisition target that produced it, its operating context, and the treatment reported by the source. Using this information, $\text{Align}(\cdot)$ places findings about the same decision point, such as route selection, accommodation choice, or failure handling, into one group G_j , while preserving their recorded conditions. The induced candidate is represented as $\pi_j = \langle c_j, a_j, r_j, v_j, P_j \rangle$, where c_j records its applicability conditions. a_j, r_j , and v_j denote its action, recovery, and verification components and P_j maps each populated component to the findings that support it. Induction includes only content directly supported by findings in G_j . It may canonicalize synonymous source terms, but does not generalize named entities or fixed choices beyond their observed contexts unless cross-context evidence supports doing so. Candidates without supporting evidence are left unspecified. When different treatments are supported under distinct recorded conditions, the candidate retains them as separate conditional cases, and incompatible treatments under matched conditions remain unresolved conflicts.

Scope-Aware Admission. For each candidate procedure π_j , SKILLALCHEMY constructs an admission record as follows,

$$\mathcal{D}(\pi_j) = \langle F_j^+, F_j^-, \sigma_j \rangle,$$

where σ_j is the widest applicability scope justified by the current evidence, F_j^+ contains findings supporting the components specified in π_j within σ_j , and F_j^- contains findings prescribing incompatible treatment under overlapping operating conditions. c_j is part of the procedure, whereas σ_j records how broadly the available evidence licenses that procedure. A candidate is *supported* if every component specified in π_j is backed by evidence under the conditions for which it is claimed. It is *consistent* if no unresolved conflicting finding applies under matched conditions within σ_j . When the evidence supports the candidate only under a narrower scope, SKILLALCHEMY restricts σ_j to that scope before admission. A supported and consistent candidate is *reusable* only when the procedure is justified beyond a single source-local case: either an eligible source under $\mathcal{S}$ explicitly states broader applicability, or compatible evidence supports the same treatment across non-equivalent contexts identified by Φ_1 . Let S_j , C_j , and U_j denote whether π_j is supported, consistent, and reusable, respectively. The admission decision is,

$$\text{Admit}(\pi_j) = \begin{cases} \text{GENERAL}, & S_j \wedge C_j \wedge U_j, \\ \text{SCOPED}, & S_j \wedge C_j \wedge \neg U_j, \\ \text{EXCLUDE}, & \text{otherwise.} \end{cases} \quad (5)$$

The general candidates form R^g and are compiled as reusable instructions. Scoped candidates form R^e and are retained as context-bound examples. Excluded candidates remain in the audit record and are not passed to Φ_3 as the skill content.

3.5 Skill Package Compilation

Package compilation does not create new procedures but maps the admitted procedures into an executable skill artifact. Given admitted instructions R^g , scoped examples R^e , and constraints $\mathcal{C}$, Φ_3 renders them, without changing admitted scope, as a standard skill package containing `SKILL.md` and optional bundled resources $\mathcal{X}$,

$$\mathcal{A} = \text{Render}(R^g, R^e; \mathcal{C}) = \langle \text{SKILL.md}, \mathcal{X} \rangle. \quad (6)$$

Specifically, SKILLALCHEMY writes the skill name and a description of what the skill does and when it should be used to the YAML frontmatter of `SKILL.md`. It then renders the admitted procedures, together with their applicability conditions and safeguards, as executable instructions in its body. Supporting content not needed in the initially loaded `SKILL.md` context is externalized as optional bundled resources $\mathcal{X}$ and referenced from `SKILL.md`, e.g., detail notes and scoped examples in `references/*`, executable routines in `scripts/*`, and templates or static resources in `assets/*`.

Finally, SKILLALCHEMY uses a corpus-derived skill grammar $\mathcal{G}_{\text{skill}}$, distilled from numerous public qualified skills, to guide package organization. The grammar supplies recurrent presentation patterns for descriptions, executable sequences or conditional structures, applicability conditions, safeguards, and progressive disclosure through package-relative references. It affects how admitted content is rendered, but does not add new procedures or broaden the admitted scope. (*The complete skill grammar is provided in supplementary materials.*)

4 Experiments

4.1 Experimental Setup

Evaluation benchmark. We follow SKILLSBENCH v1.1 (Li et al. 2026) and evaluate all 87 tasks across 8 domains. We report avg@5 pass rate, computed as the mean verified success rate over 5 independent runs for each task. Domain scores are averaged over tasks of each domain, while the overall score is averaged over all tasks.

Compared baselines. We compare SKILLALCHEMY against six baselines, covering a no-skill setting, a human-authored reference, and four automated skill-construction methods. More implementation and configuration details are provided in the supplementary material.

- **No Skill.** The agent uses only task descriptions and visible context, without installed skills or task-specific procedural guidance.
- **Human-Curated Skill** (Li et al. 2026). The agent uses the original human-authored task-specific skills released with SKILLSBENCH v1.1 without any modification.
- **Anthropic Skill-Creator** (Anthropic 2026d). It an official skill creator released by Anthropic, which drafts skill instructions, evaluates them on representative prompts, and iteratively revises the skill based on evaluation feedback.
- **OpenAI Skill-Creator** (OpenAI 2026d). It an official skill creator from OpenAI, which scaffolds a modular skill package, adds relevant resources, and validates the package before evaluation.
- **OpenSkill** (Yan et al. 2026). It retrieves open-world knowledge and verification anchors, synthesizes a skill, and refines it against the self-constructed virtual tasks.
- **MUSE-Autoskill** (Lin et al. 2026). It distills task-solving experience into reusable procedures, validation steps, and common failure modes for reliable downstream execution.

Configurations. To provide a fair comparison, we also enable web access for skill creators from Anthropic and OpenAI, as with other baselines. We evaluate four configurations spanning two agent runtimes and three models: Claude Code (Anthropic 2026b) with DeepSeek-V4-Pro (DeepSeek-AI 2026) and Claude Opus 4.8 (Anthropic 2026c), and Codex (OpenAI 2026b) with DeepSeek-V4-Pro and GPT-5.5 (OpenAI 2026c). More implementation details about the evaluation protocol are provided in supplementary materials.

4.2 Main Results

Table 1 reports the complete overall performance and the domain-level results across all four agent-model configurations.

Overall Performance. SKILLALCHEMY achieves the highest overall performance in three of the four agent-model configurations. This exceeds no-skill execution by 19.9 percentage points and the strongest automated baseline, MUSE-Autoskill, by 8.6 percentage points. At the aggregate level, SKILLALCHEMY reaches an observed avg@5 of 55.8%, 1.5 percentage points above the Human-Curated Skill. We also report 95% confidence intervals for both conditions in §4.3.

Agent	Model	Skill Setting	Overall ($n = 87$)	Δ (pp)	Soft. ($n = 16$)	Office ($n = 14$)	Sci. ($n = 14$)	Media ($n = 5$)	Cyber ($n = 7$)	Fin. ($n = 9$)	Ind. ($n = 14$)	Math ($n = 8$)
Claude Code	DeepSeek-V4-Pro	No Skill	23.4	—	22.5	32.9	18.6	20.0	14.3	26.7	25.7	20.0
Claude Code	DeepSeek-V4-Pro	Anthropic Skill-Creator	31.7	+8.3	35.0	35.7	32.9	44.0	22.9	13.3	34.3	32.5
Claude Code	DeepSeek-V4-Pro	OpenAI Skill-Creator	33.3	+9.9	33.8	41.4	44.3	32.0	22.9	17.8	28.6	35.0
Claude Code	DeepSeek-V4-Pro	Human-Curated Skill	51.3	+27.9	43.8	47.1	72.9	80.0	42.9	40.0	44.3	50.0
Claude Code	DeepSeek-V4-Pro	OpenSkill	42.3	+18.9	35.0	38.6	61.4	68.0	37.1	28.9	37.1	42.5
Claude Code	DeepSeek-V4-Pro	MUSE-Autoskill	43.2	+19.8	37.5	38.6	61.4	64.0	40.0	31.1	37.1	45.0
Claude Code	DeepSeek-V4-Pro	SKILLALCHEMY	54.7	+31.3	53.8	54.3	65.7	72.0	57.1	48.9	44.3	50.0
Claude Code	Opus 4.8	No Skill	45.3	—	56.3	55.7	57.1	32.0	57.1	15.6	32.9	37.5
Claude Code	Opus 4.8	Anthropic Skill-Creator	49.2	+3.9	56.3	55.7	62.9	56.0	57.1	17.8	42.9	35.0
Claude Code	Opus 4.8	OpenAI Skill-Creator	49.9	+4.6	55.0	57.1	65.7	52.0	54.3	17.8	45.7	37.5
Claude Code	Opus 4.8	Human-Curated Skill	59.5	+14.2	57.5	68.6	75.7	80.0	60.0	51.1	45.7	40.0
Claude Code	Opus 4.8	OpenSkill	51.5	+6.2	51.3	60.0	68.6	72.0	51.4	44.4	40.0	22.5
Claude Code	Opus 4.8	MUSE-Autoskill	53.3	+8.0	51.3	64.3	68.6	76.0	57.1	44.4	41.4	25.0
Claude Code	Opus 4.8	SKILLALCHEMY	60.9	+15.6	63.8	71.4	81.4	64.0	57.1	48.9	47.1	40.0
Codex	DeepSeek-V4-Pro	No Skill	29.7	—	42.5	35.7	28.6	8.0	22.9	28.9	21.4	30.0
Codex	DeepSeek-V4-Pro	Anthropic Skill-Creator	32.9	+3.2	45.0	38.6	24.3	20.0	40.0	20.0	30.0	35.0
Codex	DeepSeek-V4-Pro	OpenAI Skill-Creator	37.0	+7.3	42.5	45.7	42.9	12.0	42.9	24.4	30.0	37.5
Codex	DeepSeek-V4-Pro	Human-Curated Skill	45.7	+16.0	52.5	52.9	41.4	52.0	62.9	28.9	32.9	50.0
Codex	DeepSeek-V4-Pro	OpenSkill	40.7	+11.0	47.5	45.7	37.1	48.0	54.3	26.7	30.0	42.5
Codex	DeepSeek-V4-Pro	MUSE-Autoskill	40.2	+10.5	43.8	48.6	35.7	48.0	60.0	24.4	28.6	42.5
Codex	DeepSeek-V4-Pro	SKILLALCHEMY	43.9	+14.2	46.2	48.6	47.1	36.0	48.6	42.2	34.3	45.0
Codex	GPT-5.5	No Skill	45.1	—	57.5	58.6	48.6	20.0	51.4	11.1	34.3	57.5
Codex	GPT-5.5	Anthropic Skill-Creator	48.5	+3.4	58.8	60.0	51.4	20.0	51.4	17.8	48.6	52.5
Codex	GPT-5.5	OpenAI Skill-Creator	48.5	+3.4	60.0	62.9	48.6	28.0	51.4	20.0	42.9	52.5
Codex	GPT-5.5	Human-Curated Skill	60.9	+15.8	60.0	62.9	75.7	60.0	65.7	33.3	62.9	57.5
Codex	GPT-5.5	OpenSkill	49.4	+4.3	50.0	48.6	64.3	44.0	57.1	28.9	47.1	47.5
Codex	GPT-5.5	MUSE-Autoskill	52.0	+6.9	53.8	70.0	54.3	24.0	65.7	17.8	45.7	67.5
Codex	GPT-5.5	SKILLALCHEMY	63.7	+18.6	66.3	70.0	74.3	52.0	71.4	37.8	57.1	70.0

Table 1: **Main results on SkillsBench.** n is the task number and $\Delta = p_{\text{skill}} - p_{\text{no-skill}}$ reports the difference from no-skill in percentage points.

Domain-Level Analysis. Figure 4 reveals a clear domain-level asymmetry between SKILLALCHEMY and the human-curated, with configurations weighted equally within each domain. SKILLALCHEMY performs better in Finance and Economics, Software Engineering, and Office Tasks, whereas Media is the only domain with a substantial deficit. A task-level examination of all Media tasks identifies a recurring content gap: *the generated skills capture the overall solution procedure but less consistently preserve the specific steps and calibrated parameter choices used at failure-prone stages*. This refines the task-level variability reported in prior work by identifying execution-critical detail preservation as a plausible source of the remaining gap (Yan et al. 2026; Lin et al. 2026; Zhang et al. 2026a).

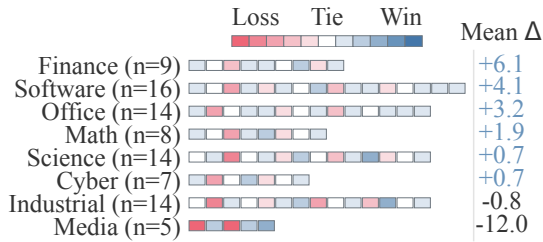


Figure 4: **Task-level comparison with Human-Curated Skills.** Each tile represents one task and is grouped by the sign of its avg@5 difference after averaging equally over the four agent-model configurations. Color intensity encodes the absolute difference, and the rightmost column reports the category mean in percentage points.

4.3 Evaluation Diagnostics

Statistical Uncertainty. Table 2 pools 1,740 binary evaluations per skill setting across 87 tasks, five runs, and four agent-model configurations. SKILLALCHEMY achieves the highest observed aggregate avg@5 at 55.8%, compared with 54.4% for the Human-Curated Skill, 47.2% for MUSE-Autoskill, and 46.0% for OpenSkill. Its 95% Wilson interval is [53.5, 58.1], versus [52.0, 56.7] for the Human-Curated Skill. These intervals quantify pooled within-setting uncertainty rather than pairwise superiority. Accordingly, the observed 1.4-point margin over the Human-Curated Skill is interpreted descriptively.

Skill Setting	Passes / Trials	avg@5	95% CI
No Skill	624 / 1740	35.9	[33.6, 38.1]
Anthropic Skill-Creator	706 / 1740	40.6	[38.3, 42.9]
OpenAI Skill-Creator	734 / 1740	42.2	[39.9, 44.5]
Human-Curated Skill	946 / 1740	54.4	[52.0, 56.7]
OpenSkill	800 / 1740	46.0	[43.6, 48.3]
MUSE-Autoskill	821 / 1740	47.2	[44.8, 49.5]
SKILLALCHEMY	971 / 1740	55.8	[53.5, 58.1]

Table 2: **Combined results on four agent-model configurations.** We combine all runs from the configurations, giving 1,740 binary outcomes per condition (over 87 tasks $\times$ 5 runs $\times$ 4 configurations). The 95% CI denotes the 95% Wilson Confidence Intervals.

Creation and Execution Cost. Table 3 reports costs for the Codex-GPT-5.5 configuration. Across all automated skill creation methods, average creation-token usage per LLM call is similar, ranging from 58.7K to 69.1K. SKILLALCHEMY requires 23.21 minutes per task, less than OpenSkill and

Skill Setting	Creation		Execution	
	Tok./Call (K)	Time/Task (min)	Tok./Run (K)	Time/Run (min)
No Skill	N/A	N/A	661	6.44
Anthropic Skill-Creator	59.4	6.47	612	5.79
OpenAI Skill-Creator	58.7	6.88	583	5.61
Human-Curated Skill	N/A	N/A	716	6.69
OpenSkill	65.9	36.37	694	6.78
MUSE-Autoskill	68.3	35.21	578	6.45
SKILLALCHEMY	69.1	23.21	709	6.39

Table 3: **Skill creation and downstream execution resource use.** Creation token is averaged per LLM call, the execution-token usage is per evaluation run and time is end-to-end per task or run.

MUSE-Autoskill (35.21–36.37 minutes) but more than the two Skill-Creator baselines (6.47–6.88 minutes). Its parallel research subagents reduce end-to-end creation latency. During downstream execution, SKILLALCHEMY uses 709K tokens and 6.39 minutes per run, with latency comparable to the other methods despite moderately higher token usage. The supplementary materials further compare skill-package length and file composition.

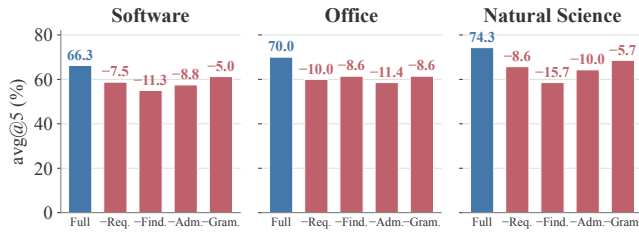


Figure 5: **Component ablation across three SKILLSBENCH domains.** Bars report avg@5 for the full method and variants without implicit requirement discovery (REQ.), structured findings (FIND.), procedure admission (ADM.), or the skill grammar (GRAM.). Labels report within-domain decrease in pp relative to the full method.

4.4 Ablation Study

We evaluate four one-component ablations on the Software Engineering, Office, and Natural Science domains of SKILLSBENCH v1.1 using Codex with GPT-5.5. The variants remove the key component of each stage, including implicit requirement discovery, structured findings, procedure admission, or grammar-guided rendering, respectively. All conditions use the same source-access scope and creation budget. For each task and condition, we create one skill and evaluate it over five independent execution runs. As shown in Figure 5, every ablation reduces avg@5 in all three domains, with observed drops of 5.0–15.7 percentage points. Removing structured findings causes the largest decrease in Software Engineering and Natural Science, while removing procedure admission has the largest effect in Office. Grammar-guided rendering yields smaller but consistent gains across the three domains. Overall, the results suggest that requirement discovery, evidence structuring, scope-aware admission, and grammar-guided artifact organization provide complementary benefits.

4.5 Robustness under Source Perturbations

Setups. Open-world sources may contain irrelevant noise, conflict procedures, or adversarially framed claims. We conduct three perturbation tests on four tasks by adding one task-specific document to the otherwise identical creation context. Specifically, (i) *Irrelevant* is topically related but does not support the target decision. (ii) *Conflict* prescribes incompatible treatment under overlapping operating conditions. (iii) *Adversarial* combines a misleading claim with directive-like language targeting the creator. All other creation and evaluation settings remain fixed. We create the skill for each (task, skill-creator, condition) tuple and run each skill five times downstream. Thus, the source-propagation unit is the generated skill package ($n = 4$ per method and perturbation) while five executions measure downstream variation. An injected payload is counted as *promoted* only when it is copied or semantically paraphrased as an affirmative runtime-facing instruction. Quotations, warnings, explicit rejections, and conflict records do not count as promotion.

Results of Perturbation. Conflict is the strongest perturbation for existing creators. Across the four baselines, 9/16 conflicting payloads are promoted, compared with 5/16 irrelevant and 4/16 adversarial payloads, while their pooled pass count decreases from 58/80 under clear evidence to 35/80 under conflict. SKILLALCHEMY does not promote any of the 12 injected payloads and retains 17–18/20 downstream passes across all conditions. The single additional failure under conflict occurs without payload promotion and is therefore as a downstream execution variation rather than evidence contamination. This experiment evaluates the containment of pre-specified undesirable claims under exposure to perturbation sources.

Skill Setting	Payload promotion ↓			Downstream passes ↑			
	Irr.	Conf.	Adv.	Clean	Irr.	Conf.	Adv.
Anthropic Skill-Creator	1/4	3/4	1/4	14/20	11/20	6/20	13/20
OpenAI Skill-Creator	1/4	2/4	1/4	13/20	12/20	8/20	12/20
OpenSkill	1/4	3/4	1/4	16/20	15/20	11/20	15/20
MUSE-Autoskill	2/4	1/4	1/4	15/20	14/20	10/20	15/20
SkillAlchemy	0/4	0/4	0/4	18/20	18/20	17/20	18/20

Table 4: **Robustness test under three-type source perturbations.** Promotion reports created skills where injected payload be a runtime-facing instruction. Passes report successful executions of all 20 runs.

4.6 Case Study: Reusable Skill Artifacts

We examine skill reusability from two complementary perspectives: whether artifacts encode operations at a reusable level, and whether a skill created for one task transfers to related tasks without revision.

Matched Artifact Audit. We audit a shared PDF-redaction operation: permanently removing sensitive text while optionally preserving an allowed fragment. Table 5 reports the shortest semantically complete instruction unit, with boilerplate compressed but scope preserved. Human-curated and generated skills often mix reusable principles with task-specific examples or execution templates. SKILLALCHEMY

Setting	Representative Instruction	Orig.	Filt.	Sem.	Δ_{sum}
No Skill	—	5/5	4/5	3/5	−3
Human-Curated	Redact sensitive data (e.g., student IDs) and insert an allowed mask.	5/5	4/5	1/5	−5
Ant. Skill-Creator	Redact rather than cover text; replace content and validate.	4/5	3/5	3/5	−2
OpenAI Skill-Creator	Select redaction, execute the edit plan, and validate.	4/5	3/5	2/5	−3
OpenSkill	Classify the edit, redact the target, and verify the result.	4/5	3/5	3/5	−2
MUSE-Autoskill	Remove text-layer content and verify deletion.	5/5	2/5	3/5	−5
SKILLALCHEMY	Bind PDF evidence to edit operations without hard-coded runtime facts.	5/5	4/5	4/5	−2

Table 5: **Artifact scope audit and frozen-skill reuse.** Representative instructions summarize the operational scope of matched PDF-redaction artifacts, while scores report unchanged reuse on the `threejs-to-obj`. Each skill is created for the original task, frozen, and evaluated on two variants. $\Delta_{\text{sum}} = (\text{Filt.} - \text{Orig.}) + (\text{Sem.} - \text{Orig.})$ reports cumulative change across the variants.

more clearly separates decision logic from runtime facts by binding the current requirement to observable PDF structure, selecting a supported operation, and rejecting unjustified execution.

Frozen-Skill Reuse. We evaluate unchanged skill reuse on the `threejs-to-obj` task family. The original task exports a `Three.js` hierarchy as an OBJ; *filtered-export* adds ancestor-aware exclusion, while *semantic-parts* assigns each geometry to its nearest semantic owner and emits per-part files with a manifest. Both variants are evaluated on held-out scenes. For each condition, the skill is created or selected only for the original task, frozen, and reused unchanged. Across five Codex-GPT-5.5 runs, SKILLALCHEMY achieves 5/5 on the original task and 4/5 on both variants, yielding the highest observed score under each requirement shift. Its cumulative degradation is only −2, compared with −3 for No Skill and between −2 and −5 for the other skill conditions. These results provide controlled evidence that SKILLALCHEMY transfers beyond its seed task while remaining robust to distinct requirement changes.

5 Conclusion

We present SKILLALCHEMY, an admission-centered framework for open-world agent skill creation. SKILLALCHEMY addresses two central challenges in open-world skill creation: recovering behavior-changing requirements omitted by underspecified briefs and restricting each source-derived procedure to its evidence-supported scope. SKILLALCHEMY operationalizes this view through implicit requirement discovery, evidence-grounded procedure admission, and scope-preserving skill package compilation. Across 87 SkillsBench v1.1 tasks and four agent-model configurations, SKILLALCHEMY improves pass rate by 19.9pp over no-skill execution and by 8.6pp over the strongest automated baseline, while reaching aggregate performance comparable to human-curated skills. Overall, these results suggest that reliable skill

creation should treat open-world knowledge as evidence to be admitted under explicit scope, rather than as instructions to be copied directly.

References

- Anthropic. 2026a. Agent Skills. <https://platform.claude.com/docs/en/agents-and-tools/agent-skills/overview>. Accessed July 8, 2026.
- Anthropic. 2026b. Claude Code. <https://github.com/anthropics/claude-code>. Accessed July 21, 2026.
- Anthropic. 2026c. Introducing Claude Opus 4.8. <https://www.anthropic.com/news/claude-opus-4-8>. Accessed July 21, 2026.
- Anthropic. 2026d. Skill Creator. <https://github.com/anthropics/skills/tree/main/skills/skill-creator>. Accessed July 21, 2026.
- DeepSeek-AI. 2026. DeepSeek V4 Preview Release. <https://api-docs.deepseek.com/news/news260424/>. Accessed July 21, 2026.
- Jiang, Y.; Li, D.; Deng, H.; Ma, B.; Wang, X.; Wang, Q.; and Yu, G. 2026. SoK: Agentic Skills – Beyond Tool Use in LLM Agents. arXiv:2602.20867.
- Li, X.; Liu, Y.; Chen, W.; You, B.; Di, Z.; He, Y.; Zheng, S.; Choe, K. W.; Sun, J.; Wang, S.; Tao, C.; Li, B.; Zhao, X.; Geng, H.; Wu, X.; Zhou, J.; Chen, X.; Xing, H.; Li, Y.; Zeng, Q.; Wang, D.; Wang, Y.; Chaim, R. B.; Jiang, P.; Shen, H.; Kong, L.; Liu, X.; Wang, R.; Liu, X.; Li, J.; Lan, X.; Lin, Y.; Ye, W.; He, J.; Li, S.; Zhang, Y.; Gao, Y.; Li, Y.; Ma, Z.; Jing, L.; Wang, T.; Li, K.; Xue, Y.; Lyu, H.; He, Y.; Tian, Y.; Wu, S.; Wang, B.; Gao, Y.; Chen, B.; Liu, L.; Cheng, S.; Bao, J.; Tong, S.; Xu, S.; Zhuo, T. Y.; Ye, T.; Qi, Q.; Li, M.; Liao, L.; Tan, Z.; Shi, C.; Tang, X.; Tankasala, S.; Yuan, B.; Qian, Y.; Tu, J.; Wang, C.; Sun, Y.; Wang, W.; Taylor, A.; Yang, Z.; Guan, C.; Dong, Z.; Zhang, X.; Dillmann, S.; Lee, H.-c.; and Song, D. 2026. SkillsBench: Benchmarking How Well Agent Skills Work Across Diverse Tasks. arXiv:2602.12670.
- Lin, H.; Li, P.; Song, J.; Jiang, F.; and Zhang, T. 2026. MUSE-Autoskill: Self-Evolving Agents via Skill Creation, Memory, Management, and Evaluation. Preprint, arXiv:2605.27366.
- Liu, A. Z.; Choi, J.; Sohn, S.; Fu, Y.; Kim, J.; Kim, D.-k.; Wang, X.; Yu, J.; and Lee, H. 2024. SkillAct: Using Skill Abstractions Improves LLM Agents. OpenReview:6LG3cIRrF4.
- Liu, Y.; Ji, J.; An, L.; Jaakkola, T.; Zhang, Y.; and Chang, S. 2026a. How Well Do Agentic Skills Work in the Wild: Benchmarking LLM Skill Usage in Realistic Settings. arXiv:2604.04323.
- Liu, Y.; Su, Z.; Xie, L.; Zhang, Y.; Zong, Q.; Guo, J.; Xie, Z.; Ji, Y.; Yim, Y.; Luo, H.; Ren, X.; Chenyu, R.; Li, H.; and Song, Y. 2026b. SkillRevise: Improving LLM-Authored Agent Skills via Trace-Conditioned Skill Revision. arXiv:2606.01139.
- Ma, Y.; Huang, Y.; Bao, H.; Zhuang, H.; Shukla, S.; Galley, M.; Zhang, X.; and Feuerriegel, S. 2026. SkillGen: Verified Inference-Time Agent Skill Synthesis. arXiv:2605.10999.
- Ni, J.; Liu, Y.; Liu, X.; Sun, Y.; Zhou, M.; Cheng, P.; Wang, D.; Zhao, E.; Jiang, X.; and Jiang, G. 2026. Trace2Skill: Distill Trajectory-Local Lessons into Transferable Agent Skills. arXiv:2603.25158.

OpenAI. 2026a. Build Skills. <https://learn.chatgpt.com/docs/build-skills>. Accessed July 21, 2026.

OpenAI. 2026b. Codex CLI. <https://github.com/openai/codex>. Accessed July 21, 2026.

OpenAI. 2026c. GPT-5.5 System Card. <https://openai.com/index/gpt-5-5-system-card/>. Accessed July 21, 2026.

OpenAI. 2026d. Skill Creator. <https://github.com/openai/skills/tree/main/skills/.system/skill-creator>. Accessed July 21, 2026.

Shridhar, M.; Yuan, X.; Côté, M.-A.; Bisk, Y.; Trischler, A.; and Hausknecht, M. 2021. ALFWorld: Aligning Text and Embodied Environments for Interactive Learning. In *International Conference on Learning Representations*.

Vercel. 2026. skills.sh: The Open Agent Skills Ecosystem. <https://www.skills.sh/>. Accessed June 3, 2026.

Wang, G.; Xie, Y.; Jiang, Y.; Mandlekar, A.; Xiao, C.; Zhu, Y.; Fan, L.; and Anandkumar, A. 2023. Voyager: An Open-Ended Embodied Agent with Large Language Models. arXiv:2305.16291.

Yan, Z.; Song, D.; Zhang, H.; Liang, W.; Zhang, Y.; Dai, Y.; He, L.; Yu, P. S.; Xu, R.; Li, X.; and Sun, L. 2026. OpenSkill: Open-World Self-Evolution for LLM Agents. arXiv:2606.06741.

Yang, Y.; Li, J.; Pan, Q.; Zhan, B.; Cai, Y.; Du, L.; Zhou, J.; Chen, K.; Chen, Q.; Li, X.; Zhang, B.; and He, L. 2026. AutoSkill: Experience-Driven Lifelong Learning via Skill Self-Evolution. arXiv:2603.01145.

Yao, S.; Zhao, J.; Yu, D.; Du, N.; Shafran, I.; Narasimhan, K.; and Cao, Y. 2023. ReAct: Synergizing Reasoning and Acting in Language Models. In *International Conference on Learning Representations*.

Zhang, G.; Zhu, E.; Zhou, J.; Jia, C.; and Wang, H. 2026a. SkillEvolver: Skill Learning as a Meta-Skill. arXiv:2605.10500.

Zhang, H.; Fan, S.; Zou, H. P.; Chen, Y.; Wang, Z.; Zhou, J.; Li, C.; Huang, W.-C.; Yao, Y.; Zheng, K.; et al. 2026b. Coevoskills: Self-evolving agent skills via co-evolutionary verification. *arXiv preprint arXiv:2604.01687*.

Zhao, A.; Huang, D.; Xu, Q.; Lin, M.; Liu, Y.-J.; and Huang, G. 2024. ExpeL: LLM Agents Are Experiential Learners. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, 19632–19642.

Zhou, Y.; Shu, W.; Su, Y.; Du, W.; Fang, Y.; and Lin, X. 2026a. A Comprehensive Survey on Agent Skills: Taxonomy, Techniques, and Applications. arXiv:2605.07358.

Zhou, Y.; Zhang, Z.; Cheng, Z.; Zhang, S.; Lan, Q.; Chen, Z.; Yang, Z.; Xu, Q.; Chen, R.; Wang, H.; and Hu, S. 2026b. SkillGenBench: Benchmarking Skill Generation Pipelines for LLM Agents. arXiv:2605.18693.

This appendix provides the implementation and analysis details supporting the main paper. Appendix A documents the evaluation protocol. Appendix B analyzes generated skill packages. Appendix C describes how the skill grammar is derived and used. Appendix D.1 and Appendix D.2 present matched execution and process-level case studies. Finally, Appendix E reproduces representative `SKILL.md` files for inspection.

A Experimental Protocol

This section documents baseline reproduction, evaluation isolation, and the shared runtime configuration behind the reported experimental results.

A.1 Baseline Implementation Details

OpenSkill. We reproduce the main OpenSkill procedure described in the main paper. The method first reads the visible task information and performs a creation search D and an independent verification search D_v . It then plans and creates the skill and evaluates it with a virtual verifier. After a failed virtual test, the method determines whether the failure arises from a skill defect or a knowledge gap and refines the skill for up to three rounds. We follow the reported settings of at most four skills, three refinement rounds, three targeted searches, a pass threshold of 1.0, and at most 60 virtual-verifier tests. For each configuration, we use the corresponding main-paper model for skill creation and downstream execution.

MUSE-Autoskill. We reproduce MUSE-Autoskill following the skill-distillation procedure described in the main paper. For each task, the method distills reusable procedures, key operations, validation steps, and common errors into a task-level skill, which is then installed without further modification for downstream evaluation.

A.2 Evaluation Isolation Protocol

Isolation During Skill Creation. Our evaluation separates two stages. In the *skill-creation stage*, a skill-creation method takes a task brief together with the source materials permitted by its protocol and produces an installable skill package, *i.e.*, a `SKILL.md` artifact and its bundled resources. In the *evaluation stage*, the produced skill is loaded by a fresh downstream agent that attempts the task, and a benchmark verifier scores the resulting submission. **Evaluation-only assets are kept separate from skill creation.**

Table A.1 distinguishes the four evaluation-only asset types from the information available during skill creation. The SkillsBench runtime does not expose held-out inputs, oracle artifacts, or verifier logic to the evaluated agent. Our creation protocol excludes all four asset types from automated skill-creation methods. The Human-Curated Skill is mounted only for its evaluation condition and is never provided as a source during creation.

Evaluation-only asset	Role in evaluation and isolation
Held-Out Task inputs	Task inputs reserved for downstream evaluation. Creators receive only the visible task description and context. Kept separate by the benchmark runtime and our protocol.
Oracle/Reference Assets	Reference solutions, gold outputs, or expected artifacts used to establish task success. Isolated by the benchmark runtime and our protocol.
Verifier Implementations	Scoring Programs whose task-specific criteria map a submission to a pass/fail signal. Isolated by the benchmark runtime and our protocol.
Human-Curated Skill	Human-Authored Packages used only in the curated evaluation condition, not as creation sources.

Table A.1: **Evaluation-only assets excluded during skill creation.** The benchmark runtime isolates evaluation-time state from the downstream agent, while our protocol applies the stated exclusions to all skill-creation methods.

For Web-enabled creation, we exclude SkillsBench-related pages from retrieval. We apply this exclusion consistently across all automated skill-creation methods. This keeps Web access under a common source policy and preserves the same evaluation setup across methods.

A.3 Runtime Configuration

Model endpoints are supplied through our API layer using the exact identifiers `gpt5.5-2026.4.23`, `claude-opus-4.8`, and `deepseek-v4-pro`. The `deepseek-v4-pro` endpoint serves the preview checkpoint, which we denote as `DeepSeek-V4-Pro-Preview` in the appendix text. Agent and adapter versions are pinned by the SkillsBench runtime: `Codex v0.128.0` with `codex-acp v0.0.45`, and `Claude Code v2.1.160` with `claude-agent-acp v0.40.0`. All runs use SkillsBench commit `34256d1`. Experiments are executed on Ubuntu 22.04.5 LTS with an Intel Xeon Platinum 8360Y CPU (144 logical CPUs), 1.0 TiB of system memory, and eight NVIDIA A100-SXM4-80GB GPUs. We set the temperature to 0.2 during skill creation, while downstream execution uses the default decoding configuration of the benchmark runner.

B Generated Skill Artifacts Analysis

This section analyzes how skill-creation methods distribute executable guidance and supporting material across their generated skill packages. We compare main-file scope and package composition to determine whether the runtime-facing instructions remain focused while detailed evidence is organized in supporting resources.

Main-File Scope. The upper row of Figure A.1 shows that the median top-level `SKILL.md` length is approximately 144–157 lines across configurations, although longer task-specific packages occur. In the packaging design of `SKILLALCHEMY`, the top-level file retains the triggers, procedures, decision boundaries, expected outputs, and task bindings required during execution. Detailed evidence and scoped examples not

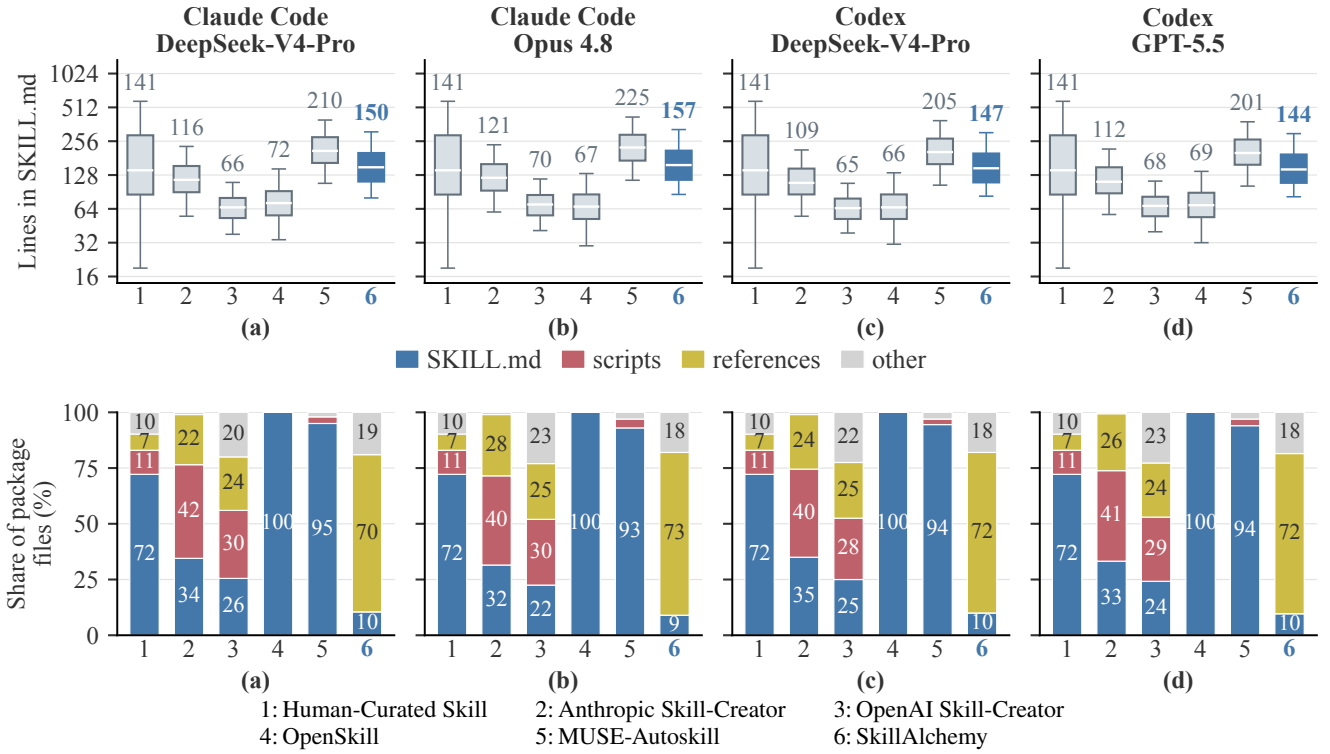


Figure A.1: **Skill package anatomy across four agent-model configurations.** Columns (a)–(d) correspond to Claude Code with DeepSeek-V4-Pro-Preview, Claude Code with Opus 4.8, Codex with DeepSeek-V4-Pro-Preview, and Codex with GPT-5.5. The upper row reports main-file length distributions, and the lower row reports the proportions of `SKILL.md`, scripts, references, and other files.

needed in the initial context are placed in package-relative reference files.

Package Organization. The lower row of Figure A.1 shows that the skill-creator baselines tend to package reusable operations as generated scripts, whereas `SKILLALCHEMY` represents them as procedures in `SKILL.md` and places supporting details in scoped references. OpenSkill and MUSE-Autoskill instead concentrate most of their user-facing files in the main `SKILL.md`. By acquiring information from open-world sources and organizing detailed supporting knowledge in reference files, `SKILLALCHEMY` preserves broad external support without placing all source-specific details in the core runtime instructions, while achieving performance comparable to human-curated skills. Complete representative `SKILL.md` files for all six skill-bearing conditions are provided in Appendix E.

C Details of Skill Grammar

This section expands the grammar-guided compilation introduced in Section 3.5 of the main paper. Section C.1 details grammar construction, Section C.2 makes its operational rules explicit, and Section C.3 explains how the grammar guides package rendering.

C.1 Skill Grammar Construction

We derive the skill grammar from public skills indexed by <https://skills.sh/>, prioritizing first-party collections from Anthropic, Vercel, Microsoft, Supabase, and Remotion. We add community-contributed skills across topic groups, remove duplicates, and retain packages with parseable metadata, nonempty instructions, and auditable source identifiers. The resulting quality-filtered set corresponds to the public qualified skills referred to in the main paper.

We extract recurrent package-level patterns instead of task-specific content. They cover specific triggers, executable and conditional procedures, applicability boundaries, safeguards, verifiable outputs, scoped examples, and progressive disclosure through package-relative references. A mechanical rubric records metadata, trigger specificity, executable steps, explicit boundaries, examples, and reference sections, and retains the top-scoring portion for deriving quality-weighted patterns. The filtered corpus identifies presentation patterns rather than certifying skill quality, and the main-paper ablation evaluates their contribution to `SKILLALCHEMY`.

C.2 Operational Skill Grammar

We use *grammar* to mean a corpus-derived operational schema, not a token-level formal grammar. Following the main paper, its role is to guide how admitted content is expressed in the final artifact. It supplies recurrent presentation patterns for

Component	Choices and enforced property
Activation metadata	<i>Choices:</i> keyword, context, explicit, hybrid, or always-on activation. <i>Purpose:</i> state what the skill does and when it should be invoked, avoiding missed or unrelated activation.
Workflow form	<i>Choices:</i> ordered steps, a decision tree, template filling, or operation cards. <i>Purpose:</i> match the structure to procedure dependencies and expose branch conditions.
Executable procedure	<i>Form:</i> condition, action, observable output, and a supported failure route. <i>Purpose:</i> replace vague advice with instructions an agent can execute and inspect.
Applicability boundary	<i>Choices:</i> source, version, capability, scope, legal, or confidence limits. <i>Purpose:</i> preserve scope and prevent unsupported generalization across tasks and contexts.
Verification form	<i>Choices:</i> a postcondition, checklist, or artifact check. <i>Purpose:</i> express an admitted verification component as a directly observable output.
Output contract	<i>Choices:</i> an analysis report, executable code, conversation, checklist, or task-dependent mixture. <i>Purpose:</i> make the output usable by its downstream consumer.
Progressive disclosure	<i>Placement:</i> core instructions in <code>SKILL.md</code> , details in <code>references/</code> , routines in <code>scripts/</code> , and static resources in <code>assets/</code> . <i>Purpose:</i> control initial context while keeping supporting material reachable.
Package organization	<i>Form:</i> complete metadata, concrete runtime instructions, and package-relative links to optional resources. <i>Purpose:</i> keep the core instructions executable while supporting progressive disclosure.

Table A.2: **The operational skill grammar used for package organization.** The grammar constrains organization and execution form; it does not supply task knowledge or override the evidence-based admission process.

descriptions, executable sequences or conditional structures, applicability conditions, safeguards, verifiable outputs, and progressive disclosure through package-relative references. Table A.2 makes these patterns explicit. Each component offers a small set of organization choices and guards against a recurrent skill-writing failure. `SKILLALCHEMY` selects among these choices according to the admitted content and task structure rather than imposing one fixed template.

C.3 Grammar-Guided Compilation

During package rendering, `SKILLALCHEMY` uses the grammar to select an executable and inspectable presentation for the admitted content. The selected patterns make the intended use distinguishable from nearby tasks, organize procedures as an appropriate sequence or conditional structure, preserve applicability conditions and safeguards, and place optional resources behind package-relative references. This guidance changes only organization and presentation: it cannot add pro-

cedures, broaden their scope, or alter the supporting evidence. Task-specific verification appears in the package only when it is an admitted component of a procedure; the grammar can express that component as a postcondition, checklist, or artifact check but does not invent one.

Illustrative Instantiation. Suppose the admitted content requires a workbook to be recalculated after formula edits and reopened to confirm delivered values. The grammar renders the procedure as an ordered sequence and expresses the admitted delivered-value verification as a postcondition. It places any admitted engine-specific details in a package-relative reference when they are not needed in the initially loaded instructions. The recalculation requirement, verification step, and engine details must still come from admitted evidence.

Takeaways. The grammar separates the presentation of skill-writing decisions from the admission of task-specific knowledge. Any skill-creation method can use the same patterns to make activation explicit, procedures executable, boundaries visible, outputs verifiable, and supporting resources loadable on demand. Its transferable contribution is therefore not a fixed template, but a compact compilation schema for turning admitted knowledge into an agent-usable package for downstream use.

D Case Study

D.1 Execution Case Study

This case study uses matched runs to show why outputs that satisfy most task criteria can still fail when a specific numerical requirement is missed.

Matched Comparison. `SkillsBench v1.1` contains 6 easy, 53 medium, and 28 hard tasks. We identify 33 tasks for which all seven conditions produce valid task-level and verifier-criterion results. This matched set contains 3 easy, 21 medium, and 9 hard tasks. The seven conditions are No Skill, Human-Curated Skill, Anthropic Skill-Creator, OpenAI Skill-Creator, OpenSkill, MUSE-Autoskill, and `SKILLALCHEMY`. All case-study runs use Codex with GPT-5.5. For each displayed case, we compare the same five runs with the same benchmark version and task inputs, then identify the verifier criterion that determines the pass/fail outcome. Verifier-criterion coverage supports this diagnosis and is not an additional benchmark metric. Figure A.2 shows one medium task and one hard task selected from the matched set. The pale-gold inset summarizes the task requirement highlighted by the `SKILLALCHEMY` artifact.

Case 1: Numerical conditioning. The medium task requires conditioning detector strain and searching an integer mass grid with three waveform approximants. A submission must report the expected signal-to-noise ratio (SNR) and total mass for each approximant. All seven conditions satisfy at least eight of nine verifier criteria in every matched run. Human-Curated Skill, OpenSkill, and MUSE-Autoskill pass 5/5 runs, followed by `SKILLALCHEMY` at 4/5 and No Skill at 3/5, while both skill-creator conditions remain at 0/5. The failures arise from an incorrect recovered SNR scale or best mass despite otherwise well-formed outputs. In the unsuccessful

Task: gravitational-wave-detection	
Condition detector strain and search three waveform approximants over the integer mass grid. The recovered SNR and total mass must match the expected signal.	
No Skill	Human-Curated Skill
Three runs recover the expected SNR and total mass; two runs miss because of conditioning choices. RESULT: 3/5; CRITERIA: 8-9/9	The curated PyCBC workflow recovers the expected peaks for all three approximants in every run. RESULT: 5/5; CRITERIA: 9/9
Anthropic Skill-Creator	OpenAI Skill-Creator
The bundled workflow returns SNRs near 5.5 and incorrect best masses in every run. RESULT: 0/5; CRITERIA: 8/9	The bundled script returns SNRs near 14 and misses the expected masses in every run. RESULT: 0/5; CRITERIA: 8/9
OpenSkill	MUSE-Autoskill
PSD-grid alignment and artifact-aware peak selection recover all expected outputs. RESULT: 5/5; CRITERIA: 9/9	The distilled workflow recovers the expected SNR and total mass in all five runs. RESULT: 5/5; CRITERIA: 9/9
SKILLALCHEMY: conditioning → alignment → validation	
Finding. PSD estimation, frequency-grid alignment, and edge handling jointly determine the recovered SNR scale. Validate conditioning with the recovered SNR and total mass for every approximant; a well-formed CSV is insufficient. Outcome. Four runs pass all checks; one edge-only PSD estimate changes the recovered SNR scale and best masses. RESULT: 4/5; CRITERIA: 8-9/9	

(a) Matched-filter conditioning

Task: exam-block-sequencing	
Assign 24 exam blocks to 24 ordered slots. The schedule must be feasible, internally consistent, and no more than 3% worse than the oracle objective.	
No Skill	Human-Curated Skill
All five schedules are feasible, self-consistent, and within the allowed objective gap. RESULT: 5/5; CRITERIA: 3/3	Solver and audit helpers produce five schedules that meet the objective-quality threshold. RESULT: 5/5; CRITERIA: 3/3
Anthropic Skill-Creator	OpenAI Skill-Creator
All schedules are feasible; two runs exceed the allowed objective gap. RESULT: 3/5; CRITERIA: 2-3/3	Independent metric recomputation and local improvement produce five passing schedules. RESULT: 5/5; CRITERIA: 3/3
OpenSkill	MUSE-Autoskill
All outputs are feasible; one run narrowly exceeds the allowed objective gap. RESULT: 4/5; CRITERIA: 2-3/3	All outputs are feasible; one run exceeds the allowed objective gap. RESULT: 4/5; CRITERIA: 2-3/3
SKILLALCHEMY: solve → recompute → audit	
Finding. Feasibility and self-consistent metrics do not establish solution quality; solver status and objective gap require a separate audit. Recompute the final objective, retain a feasible incumbent, and continue improvement until the required gap is met. Outcome. All five runs independently recompute the objective and remain within 3% of the oracle. RESULT: 5/5; CRITERIA: 3/3	

(b) Objective-quality audit

Figure A.2: **Case-level comparison across all seven conditions.** Results use Codex with GPT-5.5. RESULT counts successes over five matched runs, and CRITERIA gives verifier-criterion coverage. Success requires all criteria.

SKILLALCHEMY run, an edge-only power spectral density estimate similarly changes the recovered SNR scale and selected masses. This case isolates numerical conditioning rather than an artifact-format failure.

Case 2: Objective quality. The hard task assigns 24 exam blocks to 24 ordered slots and requires a feasible schedule whose objective is within 3% of the oracle value. All outputs in the matched runs are feasible and internally consistent. No Skill, Human-Curated Skill, OpenAI Skill-Creator, and SKILLALCHEMY pass 5/5 runs. Anthropic Skill-Creator passes 3/5, while OpenSkill and MUSE-Autoskill each pass 4/5 because their remaining runs exceed the allowed objective gap. Thus, feasibility and self-consistent reported metrics do not by themselves establish that the final solution is acceptable for benchmark evaluation.

Takeaways. The two cases expose a shared failure pattern:

structural validity and near-complete criterion coverage do not guarantee task success. Numerical pipelines require checks on task-determining computations, while optimization tasks require objective-quality checks beyond feasibility. Skill artifacts should therefore translate decisive task requirements into explicit, task-specific validation steps.

D.2 Implicit-Requirement Web Search

This case study examines how search framing and evidence synthesis determine whether Web access improves generated skills in practice.

Compared workflows. We compare a direct creator (OpenAI Skill-Creator), a retrieval-based method (OpenSkill), and SKILLALCHEMY, which combines implicit requirement discovery with contrastive evidence acquisition. For each workflow, we trace its questions, retrieved evidence, and final-artifact

safeguards.

Representative task. `weighted-gdp-calc` requires exports, imports, and GDP for six GCC countries over 2019–2023. The agent must add auditable formulas to an existing workbook, compute country statistics and a GDP-weighted regional mean, recalculate the file, and preserve the workbook structure and formatting.

Process comparison. The OpenAI Skill-Creator focuses on formula semantics, omitting engine compatibility and cached-value validation. OpenSkill retrieves these risks but does not connect them to preserving and delivering the workbook. SKILLALCHEMY turns interacting execution risks into safeguards on the delivered artifact.

From evidence to executable checks. Retrieved documentation helps only when it becomes checks on the submitted artifact. Function semantics do not ensure engine compatibility, recalculation, or current cached values, while separately retrieved risks do not ensure whole-workbook validation. SKILLALCHEMY connects these concerns into one execution path that preserves structure, writes compatible formulas, recalculates, reopens the saved file, and verifies values in the final delivered file.

Takeaways. Queries derived from implicit requirement discovery make Web evidence actionable by converting execution risks into artifact checks. Here, direct and broad search leave critical safeguards incomplete, whereas SKILLALCHEMY connects them to checks on the final delivered artifact.

Method	Search design, evidence, and outcome
OpenAI Skill-Creator	Framing: task-relevant spreadsheet functions. Evidence: Microsoft function documentation and openpyxl formula-writing guidance. Outcome: calculation-engine compatibility and cached-value validation remain uncovered; 0/5 runs pass.
OpenSkill	Framing: technologies for general spreadsheet construction. Evidence: workbook properties, recalculation engines, Excel functions, and IMF interfaces. Outcome: the risks are retrieved but not connected to a delivered-file requirement; 0/5 runs pass.
SKILLALCHEMY	Framing: omitted execution risks identified through implicit requirement discovery and boundary probes. Evidence: formula compatibility, recalculation, cached values, array alignment, reference stability, economic semantics, and workbook preservation. Outcome: the evidence becomes explicit final-workbook checks; 5/5 runs pass.

Table A.3: Search strategy, retrieved evidence, and downstream result on `weighted-gdp-calc`. Retrieved content is summarized from the skill-creation search records.

E Examples for SKILL.md

For all six skill-bearing conditions, we reproduce the complete runtime-facing SKILL.md for `glm-lake-mendota`, treating the aligned module from modular Human-Curated Skill and OpenSkill outputs as their SKILL.md. All numerical checks in this section use public or creator-generated artifacts; skill creation accesses no SkillsBench held-out inputs, oracle assets, or verifier implementations. In E.6, paths, dates, and targets bind the workflow to visible task context, while the candidate parameters are explicitly scoped to Lake Mendota evidence; none are transferable defaults. We preserve source line numbers and procedures and transliterate symbols for pdfLaTeX. Selected spans preserve the surrounding text while providing a reading guide. Blue marks an *explicit requirement* stated by the task, while purple marks an *implicit requirement* that is necessary for successful execution but must be discovered beyond the prompt. Orange marks a scoped *local example*, and green marks *generalized content* intended to transfer beyond the current instance. A red ✗ identifies an *over-specific* use of local evidence beyond its supported scope. Brief assessments appear immediately after the relevant source block, while neutral gray assessments describe other limitations without introducing another annotation category. The guide below summarizes these colors.

Annotation guide.

Explicit requirement

Implicit requirement

Local example

Generalized content

X Over-specific

E.1 Human-Curated Skill

Useful domain heuristics, but limited task grounding and several unscoped local rules. (92 source lines.)

SKILL.md (Human-Curated Skill)

```
1 --
2 name: glm-calibration
3 description: Calibrate GLM parameters for water temperature simulation. Use when you need to
4   adjust model parameters to minimize RMSE between simulated and observed temperatures.
5 license: MIT
6 --
7 # GLM Calibration Guide
8
9 ## Overview
10
11 GLM calibration involves adjusting physical parameters to minimize the difference between
12   simulated and observed water temperatures. The goal is typically to achieve  $RMSE < 2.0^{\circ}C$ . X
13
14   Over-specific threshold The benchmark's 2°C acceptance threshold is presented as a general definition of good GLM calibration.
15
16 ## Key Calibration Parameters
17
18 | Parameter | Section | Description | Default | Range |
19 |-----|-----|-----|-----|-----|
20 | `Kw` | `&light` | Light extinction coefficient ( $m^{-1}$ ) | 0.3 | 0.1 - 0.5 | X
21 | `coef_mix_hyp` | `&mixing` | Hypolimnetic mixing coefficient | 0.5 | 0.3 - 0.7 | X
22 | `wind_factor` | `&meteorology` | Wind speed scaling factor | 1.0 | 0.7 - 1.3 | X
23 | `lw_factor` | `&meteorology` | Longwave radiation scaling | 1.0 | 0.7 - 1.3 | X
24 | `ch` | `&meteorology` | Sensible heat transfer coefficient | 0.0013 | 0.0005 - 0.002 | X
25
26   Unscoped local defaults The fixed values and ranges are potentially useful examples, but the skill gives no lake, GLM configuration, or evidence
27   boundary for them.
28
29 ## Parameter Effects
30
31 | Parameter | Increase Effect | Decrease Effect |
32 |-----|-----|-----|
33 | `Kw` | Less light penetration, cooler deep water | More light penetration, warmer deep water |
34 | `coef_mix_hyp` | More deep mixing, weaker stratification | Less mixing, stronger stratification |
35 |
36 | `wind_factor` | More surface mixing | Less surface mixing |
37 | `lw_factor` | More heat input | Less heat input |
38 | `ch` | More sensible heat exchange | Less heat exchange |
39
40   Reusable abstraction Parameter-effect relations transfer more safely than fixed values and can support residual-driven calibration.
41
42 ## Calibration with Optimization
43
44 ```python
45 from scipy.optimize import minimize
46
47 def objective(x):
48     Kw, coef_mix_hyp, wind_factor, lw_factor, ch = x
49
50     # Modify parameters
51     params = {
52         'Kw': round(Kw, 4),
53         'coef_mix_hyp': round(coef_mix_hyp, 4),
54     }
```

SKILL.md (Human-Curated Skill) (continued)

```
44     'wind_factor': round(wind_factor, 4),
45     'lw_factor': round(lw_factor, 4),
46     'ch': round(ch, 6)
47 }
48 modify_nml('glm3.nml', params)
49
50 # Run GLM
51 subprocess.run(['glm'], capture_output=True)
52
53 # Calculate RMSE
54 rmse = calculate_rmse(sim_df, obs_df)
55 return rmse
```

Implementation gap The displayed code is not directly executable: modify_nml, calculate_rmse, sim_df, and obs_df are undefined.

```
56
57 # Initial values (defaults)
58 x0 = [0.3, 0.5, 1.0, 1.0, 0.0013]
59
60 # Run optimization
61 result = minimize(
62     objective,
63     x0,
64     method='Nelder-Mead',
65     options={'maxiter': 150}
66 )
67 ...
68
69 ## Manual Calibration Strategy
70
71 1. Start with default parameters, run GLM, calculate RMSE
72 2. Adjust one parameter at a time
73 3. If surface too warm → increase `wind_factor`
74 4. If deep water too warm → increase `Kw`
75 5. If stratification too weak → decrease `coef_mix_hyp`
76 6. Iterate until RMSE < 2.0°C ✗
77
78 ## Common Issues
79
80 | Issue | Likely Cause | Solution |
81 |-----|-----|-----|
82 | Surface too warm | Low wind mixing | Increase `wind_factor` |
83 | Deep water too warm | Too much light penetration | Increase `Kw` |
84 | Weak stratification | Too much mixing | Decrease `coef_mix_hyp` |
85 | Overall warm bias | Heat budget too high | Decrease `lw_factor` or `ch` |
86
87 ## Best Practices
88
89 - Change one parameter at a time when manually calibrating
90 - Keep parameters within physical ranges
91 - Use optimization for fine-tuning after manual adjustment
92 - Target RMSE < 2.0°C for good calibration ✗
```

Repeated local rule The same task-specific threshold is reused as both a universal stopping rule and a universal quality judgment.

E.2 Anthropropic Skill-Creator

Task-aware and operational, but one coordinate policy selects semantics by score. (114 source lines.)

SKILL.md (Anthropic Skill-Creator)

```
1 --
2 name: glm-lake-mendota-skill
3 description: Use this skill whenever a task asks you to run, calibrate, or evaluate the General
  Lake Model for Lake Mendota temperature simulations using glm3.nml, bcs forcing CSVs,
  field_temp_oxy.csv observations, and an RMSE target. It guides GLM execution, NetCDF temperature
  scoring, and conservative parameter iteration without hardcoding a solution.
4 --
5
6 # GLM Lake Mendota Solver Workflow
7
8 This skill helps solve Lake Mendota General Lake Model tasks where the final deliverables are:
9
10 - `/root/output/output.nc`
11 - final calibrated parameters saved in `/root/glm3.nml`
12 - vertical water temperature RMSE against `/root/field_temp_oxy.csv` below the requested threshold
13
14 Do not treat this skill as a reference solution. It is a reusable workflow, scoring harness, and
  calibration checklist. You still need to run GLM, inspect the outputs, and choose parameters based
  on the current sandbox.
15
16 ## First Moves
17
18 1. Work from `/root` unless the user explicitly stages the model elsewhere.
19 2. Confirm these inputs exist:
20   - `/usr/local/bin/glm`
21   - `/root/glm3.nml`
22   - `/root/bcs/meteo.csv`
23   - `/root/bcs/yahara.csv`
24   - `/root/bcs/pheasant.csv`
25   - `/root/bcs/outflow.csv`
26   - `/root/field_temp_oxy.csv`
27 3. Read `references/task_environment.md` for the data shape observed when this skill was created.
28 4. Read `references/sources.md` before changing model settings; it records the public
  documentation used to ground this workflow.
29 5. Run a lightweight environment check:
30
31 ```bash
32 python /path/to/this-skill/scripts/inspect_environment.py -root /root
33 ```
34
35 The checker prints paths, CSV columns, date ranges, observation depth ranges, and the key
  `glm3.nml` blocks. It does not run a simulation.
36
37 Environment requirements The task names inputs, but it does not state that their schema, coverage, depth range, and working directory must be
  validated before calibration.
38
39 ## Run GLM
40
41 GLM reads `glm3.nml` in the working simulation directory. Keep the task's relative paths intact
  unless a GLM log proves they are wrong. In this task, the configuration points at `bcs/...` CSV
  files and writes `output/output.nc`.
42
43 ```bash
44 cd /root
```

SKILL.md (Anthropic Skill-Creator) (continued)

```
43 mkdir -p output
44 /usr/local/bin/glm
45 ...
46
47 After each run, confirm ~/root/output/output.nc exists. If GLM fails, inspect the terminal log
  first; most failures are path, time-format, or malformed namelist edits.
48
49 ## Score Temperature RMSE
50
51 Score against the observation points, not just a regular grid. The bundled scorer opens the NetCDF
  output, infers the temperature variable and vertical coordinate, interpolates to each (datetime,
depth) observation, and reports RMSE.
52
53 ```bash
54 python /path/to/this-skill/scripts/score_glm_temperature.py \
55   -nc /root/output/output.nc \
56   -obs /root/field_temp_oxy.csv \
57   -json /root/output/rmse.json
58 ...
59
60 The scorer tries both direct-depth and elevation-to-depth vertical interpretations and reports the
  lower-RMSE orientation. If it cannot infer the output schema, inspect the dataset with `python -`
<<'PY' and `xarray.open_dataset(...)`, then rerun with explicit `-temp-var`, `-time-dim`, and
`-depth-dim`. ✗
```

Score-selected coordinate semantics Choosing whichever vertical interpretation yields lower RMSE can hide a coordinate error. Coordinate meaning should be established from schema or source evidence.

```
61
62 ## Calibrate Conservatively
63
64 Use short, logged candidate runs. Change a small set of physically meaningful parameters, score,
  keep the best reproducible candidate, and avoid broad blind edits that make the nml hard to trust.
65
66 Parameters worth checking first in this task type:
67
68 - Surface forcing multipliers in `&meteorology`: `wind_factor`, `sw_factor`, `lw_factor`,
  `at_factor`, `rh_factor`.
69 - Heat/light penetration in `&light`: `Kw` and, only if needed, `light_extc` / `energy_frac`.
70 - Vertical mixing in `&mixing`: `coef_wind_stir`, `coef_mix_conv`, `coef_mix_shear`,
  `coef_mix_turb`, `coef_mix_KH`, `coef_mix_hyp`, and `deep_mixing`.
71 - Inflow placement and thermal influence in `&inflow`: `inflow_factor`, `strm_hf_angle`,
  `strmbd_slope`, `strmbd_drag`.
72 - Winter behavior in `&snowice` and `&sediment` if errors cluster under ice or near the bottom.
```

Task-type examples These parameter families are concrete starting points; they are not claimed to be universal optima.

```
73
74 Use `scripts/patch_glm_nml.py` for careful single-line nml edits:
75
76 ```bash
77 python /path/to/this-skill/scripts/patch_glm_nml.py /root/glm3.nml \
78   -set meteorology.wind_factor=1.08 \
79   -set mixing.coef_wind_stir=0.28 \
80   -set light.Kw=0.34
81 ...
82
83 You can run one candidate and score it in one command:
84
85 ```bash
```

SKILL.md (Anthropic Skill-Creator) (continued)

```
86 python /path/to/this-skill/scripts/run_candidate.py \
87   -root /root \
88   -glm /usr/local/bin/glm \
89   -set meteorology.wind_factor=1.08 \
90   -set mixing.coef_wind_stir=0.28 \
91   -set light.Kw=0.34
92 ...
93
94 Keep a simple run log with parameters and RMSE. When a candidate improves RMSE, leave those
  parameters in `root/glm3.nml` and rerun GLM once from a clean `root/output` directory to verify
  the final output is reproducible.
95
96 ## Practical Calibration Pattern
97
98 1. Run and score the unmodified baseline.
99 2. Inspect residuals by time and depth if RMSE is high:
100   - warm/cold bias across most depths often points to radiation, air temperature, longwave, or
    light attenuation.
101   - surface-only errors often point to wind, heat flux, and near-surface mixing.
102   - deep-water or stratification timing errors often point to mixing, light attenuation, sediment
    heat, or inflow placement.
103   - winter-only errors often point to snow/ice and meteorological forcing.
104 3. Sweep one parameter family at a time with a small grid. Keep values near the original
    configuration unless the residual pattern justifies moving farther.
105 4. Re-run the best candidate and re-score after every accepted change.
106 5. Stop when RMSE is below the requested target and GLM can run successfully with the final
    `root/glm3.nml`.
    

Reusable calibration loop Baseline, residual diagnosis, bounded parameter-family changes, and clean reruns form a transferable process.


107
108 ## Guardrails
109
110 - Do not fetch or copy a reference answer for this exact benchmark.
111 - Do not delete files outside `root/output` during solver runs.
112 - Do not change the required simulation dates unless the task prompt changes.
113 - Do not report success from CSV point outputs alone; the expected artifact is
    `root/output/output.nc`.
114 - Preserve the final working parameters in `root/glm3.nml`.
```

E.3 OpenAI Skill-Creator

Good artifact handling, but a fixed time tolerance is promoted without validation. (101 source lines.)

SKILL.md (OpenAI Skill-Creator)

```
1 --
2 name: glm-lake-mendota-skill
3 description: Calibrate and validate a General Lake Model (GLM) Lake Mendota case using glm3.nml,
  bcs forcing CSV files, field_temp_oxy.csv observations, and output/output.nc. Use when asked to
  run GLM, tune vertical water-temperature simulation accuracy, compute RMSE against observed
  temperature profiles, or leave final GLM parameters in glm3.nml for a Lake Mendota-style benchmark
  sandbox.
4 --
5
6 # GLM Lake Mendota Calibration
7
8 ## Overview
9
```

SKILL.md (OpenAI Skill-Creator) (continued)

```
10 Use this skill to turn the Lake Mendota GLM sandbox into a reproducible calibration loop. The
    target benchmark requires `/root/output/output.nc`, a successful `/usr/local/bin/glm` run, and
    final tuned parameters saved in `/root/glm3.nml`.
11
12 Do not solve by fabricating NetCDF output or editing observations. Run GLM with physically
    plausible namelist changes, score the simulation, and keep the final parameters in `glm3.nml`.
13
14 ## Source Grounding
15
16 Before changing modeling logic, read `references/glm-mendota-notes.md`. It records the public GLM
    and glmtools sources used for this skill plus the inspected sandbox data shape.
17
18 ## Workflow
19
20 1. Confirm the sandbox files exist:
21     - `/usr/local/bin/glm`
22     - `/root/glm3.nml`
23     - `/root/bcs/meteo.csv`, `/root/bcs/yahara.csv`, `/root/bcs/pheasant.csv`,
      `/root/bcs/outflow.csv`
24     - `/root/field_temp_oxy.csv`
25 2. Inspect the current namelist. Preserve the required period and output path unless they are
    wrong:
26     - `time.start = '2009-01-01 12:00:00'`
27     - `time.stop = '2015-12-30 12:00:00'`
28     - `output.out_dir = 'output'`
29     - `output.out_fn = 'output'`
30 3. Run the baseline model from `/root`:
31     ```bash
32     cd /root
33     mkdir -p output
34     /usr/local/bin/glm
35     ...
36 4. Score the run:
37     ```bash
38     python /path/to/this-skill/scripts/score_glm_output.py \
39         -obs /root/field_temp_oxy.csv \
40         -nc /root/output/output.nc
41     ...
42 5. Tune one small group of physically meaningful parameters at a time, run GLM, and rescore. Good
    first candidates for this case are:
43     - Light attenuation: `light.Kw`
44     - Vertical mixing: `mixing.coef_mix_conv`, `mixing.coef_wind_stir`, `mixing.coef_mix_shear`,
      `mixing.coef_mix_turb`, `mixing.coef_mix_KH`, `mixing.coef_mix_hyp`
45     - Surface forcing multipliers: `meteorology.wind_factor`, `meteorology.sw_factor`,
      `meteorology.lw_factor`, `meteorology.at_factor`
46     - Drag/transfer coefficients: `meteorology.cd`, `meteorology.ce`, `meteorology.ch`
47 6. Stop when the temperature RMSE is below 2 degrees C or when further tuning is clearly
    unproductive. Leave the winning `/root/glm3.nml` in place and ensure `/root/output/output.nc` is
    from that same configuration.
```

Artifact consistency The task requires a runnable final nml, but does not explicitly say that the delivered output must be regenerated from that exact configuration.

```
48
49 ## Helper Scripts
50
51 Use scripts from this skill folder; do not copy large blocks by hand.
52
```

SKILL.md (OpenAI Skill-Creator) (continued)

```
53 ### Patch namelist scalars
54
55 ```bash
56 python scripts/patch_glm_nml.py /root/glm3.nml \
57   -set light.Kw=0.35 \
58   -set mixing.coef_wind_stir=0.30
59 ```
60
61 The patcher replaces existing scalar assignments inside a named namelist section while preserving
  the rest of the file.
62
63 ### Run and score one candidate
64
65 ```bash
66 python scripts/run_glm_candidate.py \
67   -root /root \
68   -glm /usr/local/bin/glm \
69   -set light.Kw=0.35 \
70   -set mixing.coef_wind_stir=0.30
71 ```
72
73 This backs up `glm3.nml`, applies settings, clears stale `output.nc`, runs GLM, then prints RMSE
  if output was produced.
  

Candidate lifecycle Backup, targeted patching, stale-output removal, execution, and scoring form a reusable safety pattern.


74
75 ### Coordinate-search calibration
76
77 ```bash
78 python scripts/calibrate_glm_mendota.py -root /root -glm /usr/local/bin/glm -passes 2
79   -max-runs 80
80 ```
81
82 This tries a conservative coordinate search over common GLM thermal parameters, keeps the best
  namelist, and leaves the best output in `/root/output/output.nc`. If the script stalls or GLM is
  slow, interrupt it and continue manually from the best JSON log it has printed.
82
83 ## Scoring Details
84
85 `score_glm_output.py` compares observation rows with columns `datetime`, `depth`, and `temp` to a
  GLM NetCDF temperature variable. It finds the nearest model time within 12 hours, interpolates the
  model profile to observation depth, skips missing pairs, and reports overall RMSE in degrees C. ✖
  

Fixed matching tolerance The 12-hour rule is neither task-stated nor derived from the output cadence; it may change both the match set and the
    reported RMSE.


86
87 If the NetCDF schema differs, inspect the dataset with:
88
89 ```bash
90 python - <<'PY'
91 import xarray as xr
92 ds = xr.open_dataset('/root/output/output.nc')
93 print(ds)
94 PY
95 ```
96
97 Then rerun the scorer with explicit names, for example:
98
```

SKILL.md (OpenAI Skill-Creator) (continued)

```
99 ```bash
100 python scripts/score_glm_output.py -nc /root/output/output.nc -obs /root/field_temp_oxy.csv
    -temp-var temp -depth-var z
    Unresolved coordinate meaning The schema can be overridden, but the example does not establish whether z is height or observation depth
    before using it.
101 ```
```

E.4 OpenSkill

Highly reusable, but too abstract to instantiate the current task directly. (49 source lines.)

SKILL.md (OpenSkill)

```
1 --
2 name: glm-calibration-loop
3 description: Use when a future agent must orchestrate repeatable GLM thermal calibration by
  proposing bounded parameters, editing the namelist, running GLM, scoring temperature results, and
  retaining the best reproducible configuration.
4 --
5
6 ## When To Use
7
8 Use this skill when the user needs a GLM water-temperature simulation to meet a quantitative
  agreement target against observations. It coordinates namelist editing, simulation execution,
  time-axis validation, and RMSE scoring without hard-coding final calibrated values.
9
10 Do not use this skill to create a one-shot solver, alter raw observations, or guess hidden
  evaluation behavior. It should guide a reusable, transparent calibration process.
11
12 ## Procedure
13
14 1. Establish a baseline run. Parse the namelist, validate input files, validate the GLM run target
  for the host platform, run GLM unchanged if possible, validate the saved time axis, score the
  baseline, and record diagnostics.
15
16 2. Define calibration parameters with explicit bounds and units. Candidate levers include light
  attenuation, shortwave and longwave scaling, wind scaling, heat exchange coefficients, drag,
  mixing parameters, and active ice or snow parameters.
17
18 3. Exclude high-risk fields unless justified. Morphometry, raw forcing records, and observation
  data should not be routine calibration targets.
19
20 4. Normalize parameter scales before optimization so that derivative-free searches do not
  overemphasize large numeric ranges.
21
22 5. For each candidate: patch the namelist with a parser, validate the written namelist, run GLM
  from the correct working directory with a platform-appropriate executable or interpreter, verify
  NetCDF results, validate the time axis, compute the RMSE with documented matching rules, and log
  parameters plus metrics.
23
24 6. Reject candidates that fail namelist validation, fail platform execution checks, fail
  simulation execution, lack required result variables, have an incomplete or shifted time axis,
  produce no matched observation-model pairs, or rely on undocumented coordinate assumptions.
25
26 7. Use structured search or derivative-free optimization. Practical patterns include bounded grid
  refinement, coordinate search, Nelder-Mead with transformed variables, or CMA-style search with
  normalized parameters.
27
28 8. Track the best candidate by the primary metric and keep secondary diagnostics such as bias,
  seasonal behavior, depth-specific errors, group-wise RMSE, and number of matched pairs.
29
30 9. Stop when the target metric is reached, the improvement has stalled, or the run budget is
  exhausted.
31
32 10. Re-run the best configuration from a clean state, verify that the score and execution path are
  reproducible, regenerate any supporting score summaries from the final matched table, and leave
```

SKILL.md (OpenSkill) (continued)

the namelist and model results consistent with that best run.

Strong reusable structure The skill captures a robust calibration lifecycle without hard-coding a final answer.

24

25 ## Calibration Checks

26

27 Before changing parameters, confirm that the model is reading the intended meteorological, inflow, outflow, and light files. A `bad path`, `timestamp mismatch`, `incomplete saved time axis`, or invalid run wrapper should be fixed before parameter optimization.

28

29 Use `scientifically plausible bounds`. For example, scaling factors should remain near defensible measurement uncertainty unless the data source has known bias; mixing and drag changes should be constrained to physically meaningful ranges.

30

31 Score enough `time and depth coverage` to avoid overfitting to sparse matches. If the objective improves while `coverage collapses`, treat the candidate as invalid or at least suspect.

Hidden evaluation validity Match count and time-depth coverage are not requested explicitly, but are essential to prevent a misleadingly low score.

32

33 ## Logging Pattern

34

35 For each candidate, record: candidate id, parameter values, changed namelist fields, run status, platform run method, result validation status, time-axis validation status, `RMSE`, `match count`, `time coverage`, `depth coverage`, group-wise diagnostics, and notes about warnings.

36

37 Keep logs human-readable and machine-readable where practical, but `do not encode final calibrated constants into the reusable skill` itself.

38

39 ## Pitfalls

40

41 Improving RMSE by `editing observations or forcing data is usually invalid calibration`.

42

43 Execution failures caused by an `incompatible binary or script launcher` are setup failures, not calibration evidence. Fix the run method before interpreting model quality.

44

45 `Time-axis mismatches` can change depth-band or seasonal match counts and make a score summary disagree with recomputed metrics. Fix saved-time coverage and regenerate summaries before tuning parameters.

46

47 Ice periods, stratified periods, and turnover periods can respond to different parameters; `inspect residual patterns rather than optimizing a single number blindly`.

48

49 `Result files may be stale after failed runs`. Confirm timestamps and `rerun the best candidate cleanly` before declaring success.

Under-instantiated for this task The skill never grounds its workflow in the provided paths, dates, output name, GLM schema, or Lake Mendota-specific evidence, so a solver still has to reconstruct the concrete procedure.

E.5 MUSE-Autoskill

Broad operational coverage, with an underspecified exact-time matching policy. (165 source lines.)

SKILL.md (MUSE-Autoskill)

```
1 --
2 name: glm-lake-mendota-skill
3 description: Calibrate and verify a General Lake Model run for Lake Mendota vertical
4   water temperature by running a baseline simulation, scoring NetCDF output against
5   field temperature profiles, and applying a compact physically informed parameter
6   sweep only when needed.
7 version: 1.0.0
8 tags:
9 - glm
10 - lake-modeling
11 - lake-mendota
12 - calibration
13 - netcdf
14 - water-temperature
15 - rmse
16 inputs:
17 - glm_nml_path
18 - forcing_directory
19 - field_temperature_csv
20 - output_nc_path
21 - target_start_date
22 - target_stop_date
23 - rmse_threshold
24 outputs:
25 - calibrated_glm_nml
26 - simulation_output_nc
27 - rmse_report
28 - verification_summary
29 --
30
31 # GLM Lake Mendota Calibration Skill
32
33 ## When to use
34
35 Use this skill when an agent must produce a valid General Lake Model `output.nc` for Lake Mendota
36   over a requested date range, verify simulated vertical water temperature against observed field
37   temperature profiles, and leave behind a runnable final `glm3.nml` whose RMSE is below the task
38   threshold.
39
40 This skill is intended for practical calibration work where a baseline GLM setup already exists
41   and only compact, physically informed parameter changes should be made.
42
43 ## Core principles
44
45 - Always run and score the existing setup before calibration.
46 - Fix configuration, forcing, date, and output-path problems before tuning parameters.
47 - Preserve the original namelist text as the base for candidate edits.
48 - Change only required date/output fields and a small set of calibration parameters.
49 - Score every candidate with the same RMSE routine used for the baseline.
50 - Prefer compact physically meaningful sweeps over random or broad search.
51 - Treat positive volume as the indicator of active GLM layers.
52 - Exclude inactive-layer fill values and implausible temperatures before interpolation.
53 - Compare observations to modeled depth below the moving simulated surface, not raw `z`.
```

SKILL.md (MUSE-Autoskill) (continued)

50 - Final verification must be run after the final GLM rerun so ``glm3.nml`` and ``output.nc`` correspond exactly.

Implicit execution requirements Active-layer semantics, coordinate conversion, invalid-value filtering, and final artifact consistency are necessary but absent from the task brief.

51

52 ## Recommended tools and libraries

53

54 - GLM executable: ``glm`` or ``glm3``, located with ``command -v glm`` or ``command -v glm3``

55 - Shell utilities: ``ls``, ``rm``, ``grep``, ``du``, ``find``

56 - Python 3 executable: prefer ``python3``; do not assume ``python`` exists

57 - Python libraries:

58 - ``csv.DictReader`` for observation CSV parsing

59 - ``datetime`` for timestamp parsing

60 - ``math`` for RMSE calculation

61 - ``numpy`` for finite filtering and linear interpolation

62 - ``netCDF4.Dataset`` and ``netCDF4.num2date`` for NetCDF inspection and time conversion

63 - Namelist editing:

64 - Prefer a namelist-safe parser if available

65 - Otherwise use `constrained regex replacement` that changes only the target scalar value before any inline comment

Implementation-specific fallback The tool stack is concrete and usable, although the regex fallback is less robust than a namelist-aware parser.

66

67 ## Workflow

68

69 1. Inspect the GLM setup before changing anything.

70

71 Read the provided ``glm3.nml`` from ``glm_nml_path``. List forcing CSV files in ``forcing_directory``. Preview the field temperature CSV header and first few rows from ``field_temperature_csv``. Locate the GLM executable with ``command -v glm`` or ``command -v glm3``.

72

73 Confirm the namelist points to ``target_start_date``, ``target_stop_date``, and ``output_nc_path``. If any of those required fields are wrong, update only those fields. Do not tune calibration parameters during this inspection step.

74

75 2. Run a clean baseline simulation.

76

77 From the directory containing ``glm3.nml``, remove stale GLM outputs such as ``output.nc``, ``lake.csv``, and outlet or outflow CSV files. Then execute the GLM binary from that same directory.

78

79 Treat a nonzero exit status or missing ``output.nc`` as a configuration or forcing failure. Fix that failure before attempting calibration.

80

81 3. Inspect the generated NetCDF.

82

83 Use ``python3``, ``netCDF4.Dataset``, and ``num2date`` to open ``output_nc_path``.

84

85 Verify expected variables such as ``time``, ``z``, ``V``, and ``temp``. Print or record dimensions, variable shapes, units, first timestamp, and last timestamp. Confirm the output covers the requested period.

86

87 4. Compute RMSE against field observations.

88

89 Read observations from ``field_temperature_csv`` with ``csv.DictReader``. Skip rows where depth or temperature is ``NA``. Parse timestamps with the CSV ``datetime`` format.

90

SKILL.md (MUSE-Autoskill) (continued)

91 Convert NetCDF model times to Python datetimes using ``num2date``. For each observation `timestamp` that matches a model time, extract the corresponding ``z``, ``V``, and ``temp`` arrays. ✗

Exact-time matching assumption The wording implies exact model-time matches without first verifying alignment or defining a justified tolerance, so valid observations may be silently lost.

92

93 Keep only active layers where volume is positive and ``z``, ``V``, and ``temp`` are finite. Reject GLM fill values and implausible temperatures, for example by requiring ``abs(temp) < 100``.

94

95 Convert model ``z`` coordinates to `depth below the simulated surface` using:

96

97 ``depth = surface_z - z``

98

99 Sort modeled layers by increasing depth. `Linearly interpolate simulated temperature to each observed depth` with NumPy. Count `matched observations and missing observations`. Compute:

100

101 ``RMSE = sqrt(mean((modeled_temp - observed_temp)^2))``

102

103 The RMSE report must include RMSE, `matched count, missing count, date coverage`, and any skipped-row reasons.

104

105 5. Decide whether calibration is needed.

106

107 If baseline RMSE is `below `rmse_threshold`` and date/output checks pass, keep the original parameters and finish with the baseline output.

108

109 If baseline RMSE is above threshold, run a `compact calibration sweep`. Do not use broad random search. Preserve the original namelist text as the base for every candidate and edit only targeted scalar values.

110

111 6. Sweep physically relevant `Lake Mendota GLM parameters` first.

112

113 Prioritize parameters that influence heat flux, wind mixing, and light attenuation:

114

115 - Radiation scaling: ``sw_factor`, `lw_factor``

116 - Optional atmospheric scaling: ``at_factor``

117 - Wind forcing scale: ``wind_factor``

118 - Light attenuation: ``Kw``

119 - Selected mixing terms if present, especially ``coef_wind_stir``

120

121 Use `modest ranges around the baseline`. First test individual parameter changes. Then run small combined grids over parameters that individually improve RMSE.

Correctly scoped local parameters The parameter list is explicitly tied to Lake Mendota and is therefore a local example, not a universal rule.

122

123 Wind and shortwave/longwave scaling are often high-leverage for this task, but keep all changes conservative enough to avoid destabilizing or breaking GLM.

124

125 7. `Score every candidate immediately`.

126

127 For each candidate parameter set:

128

129 - Write candidate values into ``glm3.nml``

130 - Delete stale outputs

131 - Run ``glm`` or ``glm3``

132 - Skip failed runs

133 - Score the resulting ``output.nc`` with the same RMSE routine

SKILL.md (MUSE-Autoskill) (continued)

```
134 - Track RMSE, matched count, missing count, and changed parameters
135
136 Do not stop at the first borderline pass if time allows. Continue the compact sweep to gain
margin below `rmse_threshold`.
137
138 8. Restore the best parameter set and rerun.
139
140 After the sweep, write the best parameter set into `glm3.nml`. Run GLM one final time from a
clean output state.
141
142 This final rerun is required. Do not leave `output.nc` from a candidate run that does not
exactly match the final saved namelist.
143
144 9. Independently verify final artifacts.
145
146 After the final rerun:
147
148 - Recompute RMSE from final `output_nc_path`
149 - Inspect `glm3.nml` for start date, stop date, output directory/file, and changed calibration
keys
150 - Open `output.nc` with `netCDF4.Dataset`
151 - Confirm first and last timestamps cover the requested period
152 - Confirm `output.nc` exists and has nontrivial size
153
154 The final verification summary must report final RMSE, matched observation count, missing
observation count, output path, first timestamp, last timestamp, and the small set of changed
parameters.
155
156 10. Handle common pitfalls explicitly.
157
158 - Use `python3`; `python` may be unavailable.
159 - Skip observation rows containing `NA`.
160 - Do not treat all NetCDF array entries as active layers.
161 - Identify active GLM layers by positive volume.
162 - Exclude huge inactive-layer fill values before interpolation.
163 - Bound plausible temperatures, for example `abs(temp) < 100`.
164 - Convert model `z` to depth below the moving simulated surface.
165 - Verify final artifacts only after the final rerun.
```

Remaining limitation Coverage is broad, but much of the skill remains a checklist; it gives less evidence-backed explanation for why a failure occurs or which residual should select the next parameter.

E.6 SkillAlchemy

More implicit requirements and reusable synthesis, with task-local bindings kept in scope. (254 source lines.)

SKILL.md (SkillAlchemy)

```
1 --
2 name: glm-lake-mendota
3 description: >-
4   Runs the General Lake Model (GLM 3.x) to simulate vertical lake water temperature and
5   calibrates its parameters until sim-vs-observation RMSE drops below a target (e.g. < 2 degC).
6   Sparse task anchor Only the core objective is blue. Paths, commands, and validation mechanisms are classified by where they come from, not
   recolored blue merely because they help execute the task.
7   Covers the whole closed loop: edit glm3.nml, run glm, extract the (time, depth) temperature
   field from output.nc, score it against field profiles, diagnose the residual, pick the next
```

SKILL.md (SkillAlchemy) (continued)

```
8 knob. Use for the Lake Mendota task when the visible context provides glm3.nml, output.nc,
9 field observations, a date window, and an RMSE target. For another GLM temperature task,
10 reuse only the general loop after its task-local bindings and model conventions are revalidated.
11 version: 0.1.0
12 --
13
14 # GLM Lake Temperature Simulation & Calibration
15
16 A GLM task looks like "run a model". It is really a closed calibration loop:
17
18 ```
19 edit glm3.nml -> run glm -> extract temp(time, depth) -> score vs obs -> diagnose residual -> next
20 knob
21 ```
22
23 The commands below instantiate this loop for the visible Lake Mendota task. Treat its paths, dates,
24 target, and candidate parameters as task-local bindings, not reusable defaults. For another GLM
25 setup, re-identify and validate those bindings before following the general loop.
```

Task-scoped instantiation Concrete paths, dates, targets, and candidate parameters bind the general loop to the visible task and must be re-established before reuse.

```
26
27
28 The three facts that decide whether you succeed. Internalise them before touching anything:
29
30 1. "z" in "output.nc" is bottom-up, and it is a HEIGHT, not a depth. Layer index `0` is the
31 lake bottom; index `NS[t]-1` is the surface. Reading `temp[t,0]` as "surface
32 temperature" is a likely way to fail this task. The mistake can return a plausible but incorrect
33 temperature and RMSE, with no exception, no warning, and no obvious sign that the depth
34 convention was misread.
35
36 2. GLM's exit code is not a success signal. Its `main()` ends in an unconditional `exit(0)`.
37 Non-zero reliably means failure; zero means nothing at all. Success has to be defined as:
38 "output.nc" exists and its "time" axis spans the requested window and "temp" is finite.
39
40 3. Relative paths in the nml bind to the process CWD, not to the nml's own location -- GLM's C
41 source contains no path-resolution code whatsoever. "out_dir='output'" means
42 "$PWD/output". Always "cd" to the nml's directory before running.
```

Implicit requirements beyond the prompt Coordinate semantics, unreliable exit codes, and CWD-dependent paths are not visible in the task brief, yet each can produce plausible-looking false success.

```
43
44
45 ## Activation Rules
46
47
48 Apply these instructions when:
49
50 - The visible task concerns Lake Mendota temperature simulation or vertical-profile calibration
51 - It provides a `glm3.nml`, forcing directory, observation file, expected output, and date window
52 - It specifies a calibration metric and target for the delivered GLM artifacts
53 - For another GLM setup, reuse only the loop after rebinding and validating every task-local input
54 - Post-hoc analysis of a GLM run: extracting temperature, comparing to field profiles, plotting
55 profiles
56
57
58 Do NOT trigger on:
59
60 - Generic NetCDF reading with no lake model involved → just use `netCDF4`/`xarray` directly
61 - Other lake/reservoir models (FLake, Simstrat, GOTM, MyLake, CE-QUAL-W2) -- the nml, the output
62 layout and the calibration knobs are all different. Only the "loop shape" transfers.
```

Scope boundary The transferable claim is explicitly limited to the loop structure; task paths, schemas, dates, and calibration values must be revalidated in a new setup.

SKILL.md (SkillAlchemy) (continued)

```
55 - Water **quality**/biogeochemistry calibration (AED2: oxygen, nutrients, chlorophyll).
    Temperature
56 is calibrated first and frozen; this skill stops there.
57 - Statistical / ML lake-temperature prediction (LSTM, process-guided DL) -- no GLM binary in the
    loop
58 - Pure hydrology (streamflow, rainfall-runoff) with no lake thermal structure
59
60 --
61
62 ## Agentic Protocol
63
64 Work the loop in this order. **Do not skip Step 1 or Step 3** -- they are the two places where
65 silence is mistaken for success.
66
67 **Step 0 -- Read the model card before acting.** Read `references/sop_models.md` and match your
68 current stage to a card (H1-H9). The cards carry the evidence, the failure modes, and the exact
69 commands. Come back here for the sequencing.
70
71 **Step 1 -- Establish ground truth about the environment (never assume it).**
72 ```bash
73 ldd /usr/local/bin/glm | grep -i "not found"      # missing libnetcdf.so.X? -> troubleshooting.md
74 head -3 /root/bcs/meteo.csv /root/field_temp_oxy.csv
75 python3 scripts/nml_edit.py dump -nml /root/glm3.nml      # what is actually in the config
76 ```
77 A missing shared library is an **environment** problem, not a modelling problem. Fix it before
78 anything else, or every run will "succeed" and write nothing.
79
80 **Step 2 -- Get one baseline run to completion, from the right directory.**
81 ```bash
82 python3 scripts/run_glm.py -nml /root/glm3.nml -glm /usr/local/bin/glm
83 ```
84 This `cd`s to `/root` for you and then **verifies the artefact instead of the exit code** -- it
85 recomputes the expected record count from `start`/`stop`/`dt`/`nsave` and checks the actual `time`
86 axis against it. Do not hand-roll `glm` invocations; the CWD trap is why this script exists.
87
88 **Step 3 -- Introspect `output.nc` before you trust any recipe about it.**
89 ```bash
90 python3 scripts/extract_temp.py -nc /root/output/output.nc -header      # or: ncdump -h
91 ```
92 Confirm you see `NS(time)`, `z(time,z,lat,lon)`, `temp(time,z,lat,lon)`, and `time` with
93 `units = "hours since ...". If the names differ, your GLM is a different build -- consult the
94 version table in `references/sop_models.md` (H3) rather than guessing.
95
96 Evidence-to-protocol synthesis Concrete environment and schema findings are converted into a reusable order of operations rather than copied as
97 universal constants.
98
99
100 **Step 4 -- Score the baseline under the fixed metric.**
101 ```bash
102 python3 scripts/eval_rmse.py -nc /root/output/output.nc -obs /root/field_temp_oxy.csv -target
103 2.0
104 ```
105 This prints the pooled RMSE, the bias, and the breakdown **by depth band and by season**. The
106 breakdown is not decoration -- it is the input to Step 5. A whole-column RMSE that looks fine while
107 the hypolimnion is 3 °C off is the normal state of an uncalibrated GLM, and the aggregate hides
108 it.
109
110
111 **Step 5 -- Diagnose the residual, then pick the knob.** Use the mapping table below (full version,
```

SKILL.md (SkillAlchemy) (continued)

```
105 with the physics and the citations, in H7). Do not grid-search. Each residual pattern points
    at
106 a specific, physically-motivated parameter; move that one.
    Diagnostic generalization The skill turns residual structure into a principled next action instead of promoting one local parameter setting.
107
108 Step 6 -- Calibrate, staged and bounded.
109 ```bash
110 python3 scripts/calibrate.py -nml /root/glm3.nml -obs /root/field_temp_oxy.csv \
111     -params ch, lw_factor, sw_factor, sed_temp_mean -target 2.0 -calib-frac 0.5
112 ```
113 For this Lake Mendota task, the command uses four parameters selected by a Morris sensitivity screen
114 (Ladwig et al., HESS 2021). It iterates in a scratch directory, then writes the winning
115 parameters back into the real nml and re-runs in place -- so the delivered nml genuinely
116 reproduces the delivered `output.nc`. Runs are seconds-scale; dozens of iterations are cheap.
117
118 Step 7 -- Gate the delivery. This is the last thing you do.
119 ```bash
120 python3 scripts/verify_delivery.py -nc /root/output/output.nc -nml /root/glm3.nml \
121     -obs /root/field_temp_oxy.csv -start 2009-01-01 -stop 2015-12-30 -target 2.0
122 ```
123 Five checks: the file exists; the time axis covers the window; `temp` is finite and physical;
    the
124 final nml re-runs cleanly and reproduces that time axis; RMSE is under target. Exit 0 only if
    all
    Delivery requirements discovered The task asks for a runnable final nml, but not the full artifact, coverage, physical-validity, and clean-rerun gate
    needed to establish that claim.
125 five pass. If you did any iterating outside `/root`, this is what catches the fact that you never
126 copied the answer back.
127
128 --
129
130 ## Core Operation Models
131
132 | # | Model | Core proposition | Source |
133 |---|-----|-----|-----|
134 | H1 | Run GLM Without Losing Your Output | Relative nml paths bind to CWD, never to the nml.
    `cd` to the nml dir. `out_dir` is created with `mkdir(2)` -- one level only | GLM C source
    (`glm_main.c`, no `chdir`/`realpath` anywhere) |
135 | H2 | Never Trust GLM's Exit Code | `main()` ends in unconditional `exit(0)`. Define success
    as artefact + full time axis + finite temp | GLM `src/glm_main.c`; LakeEnsemblR issues #279, #288
    |
136 | H3 | The Lagrangian Layer Trap ! | `z` = layer-TOP height above the bottom, bottom-up.
    Index 0 = bottom, `NS[t]-1` = surface. Slice by `NS[t]`; padding is  $\geq 1e30$  and `_FillValue` is
    compile-time optional | GLM `glm_ncdf.c` + glmtools + glm-py, verified numerically on a real
    `output.nc` |
137 | H4 | Midpoint Interpolation to a Depth Grid | Temperature attaches to layer midpoints,
    not tops. Nodes `[0]+mids+[z_surf]`, values `[T0]+T+[T_last]`, linear interp, out-of-column →
    NaN not clamped | glmtools `resample_depth()` in `R/get_var.R` |
138 | H5 | The Fixed RMSE Convention | Pooled over all matched (time, depth); sim interpolated
    in depth onto obs depths; time matched at calendar-day precision; unmatched dropped; ice and
    surface not excluded | glmtools `dot_compare_to_field.R` + `calib_helpers.R`; LakeEnsemblR
    `calc_fit.R` |
139 | H6 | Staged Calibration, Not Grid Search | Screen → cap at 4-6 params → heat budget →
    light/structure (`Kw`) → bottom (`sed_temp_mean`) → mixing efficiencies last | Ladwig et al. HESS
    2021; Feldbauer et al. HESS 2025; Bruce et al. 2018 |
```

SKILL.md (SkillAlchemy) (continued)

```
140 | H7 | Residual Diagnosis → Knob Selection | Each residual pattern maps to a specific knob
    | with a physical justification. Read the residual, don't sweep the space | R03 synthesis of HESS
    | 2025 / Bruce 2018 / glmGUI 2020 |
141 | H8 | Know When To Stop |  $\geq 3$  °C = broken.  $\sim 2$  °C = passing but mediocre. 1.3-1.6 °C = good for
    | Mendota.  $\leq 1.0$  °C = suspicious | Bruce (Mendota 1.60); Ladwig (1.96 total, surface 1.30, bottom
    | 2.43) |
142 | H9 | Editing the nml Safely | Fortran namelist, `!` comments, array-length rule is `len >=
    | count`. Surgical key-value edits + round-trip verification. Never regex-rewrite wholesale | GLM
    | `glm_init.c`; AED config docs |
    | Operation models The evidence is consolidated into reusable models for execution, scoring, calibration, and safe editing.
```

```
143
144 Full cards -- with inputs, exact actions, evidence URLs, failure modes and confidence -- in
145 `references/sop_models.md`. **H3 is the one that decides the task.**
146
147 ### Residual → knob (the short version)
148
149 | What the residual looks like | Move this | Which way |
150 |--|--|--|
151 | Warm/cold bias at all depths | `sw_factor`, `lw_factor` | down if warm, up if cold |
152 | Surface/epilimnion too warm in summer | `ce` (latent heat loss), then `ch`; `wind_factor` |
    | up |
153 | Surface too cold | `ch`/`ce` down, `sw_factor` up | -- |
154 | Thermocline too deep / not enough stratification | `Kw` first; then `wind_factor`, `cd`
    | down | `Kw` up |
155 | Thermocline too shallow / won't mix down | `Kw` down; then `wind_factor`, `coef_mix_shear`
    | up | -- |
156 | Hypolimnion/bottom too cold | `sed_temp_mean` (deep zone, 3-8 °C); then `coef_mix_hyp` | up
    |
157 | Hypolimnion too warm | `sed_temp_mean` down; `coef_mix_hyp` down | down |
158 | Winter / under-ice bias | `sed_temp_mean` (dominant under ice), `lw_factor` | -- |
159 | Turnover too early/late | `wind_factor`, `coef_mix_conv` | up ⇒ earlier |
    | Residual-to-action abstraction Error patterns are mapped to physically motivated parameter families, increasing transfer beyond the original run.
```

```
160
161 Physics, directions, bounds and citations for every row: H7 in `references/sop_models.md`.
162
163 --
164
165 ## Output Style
166
167 - Lead with the metric and verdict: "Pooled RMSE is below the task-defined target; the dominant
168   bias is in the hypolimnion." Then explain.
169 - Show the command you ran and the output that matters. Not the whole log.
170 - Name the knob and the reason together: "the bottom is too warm, so `sed_temp_mean` comes down" --
171   never "let me tune some parameters".
172 - When you cite a fact, cite it the way a person would: "Ladwig's Mendota calibration got surface
173   1.30 °C but bottom 2.43 °C" -- not "according to H8 in sop_models.md".
174 - State uncertainty where the evidence is thin (the under-ice knob mapping is physics-derived, not
175   measured). Do not launder inference into fact.
176 - Banned: "let me systematically analyse...", "according to the model card...", "would you like me
177   to continue?". Answer, then stop.
178
179 --
180
181 ## Output Modes
182
```

SKILL.md (SkillAlchemy) (continued)

```
183 | Mode | Triggered when | Structure |
184 |----|-----|-----|
185 | **Diagnose** | A run finished and the RMSE is too high | Pooled RMSE + bias → depth-band &
    seasonal breakdown → the one residual pattern that dominates → the knob it implies → the next
    command |
186 | **Bootstrap** | Fresh environment, nothing has run yet | `ldd` + forcing-file survey + nml dump
    → baseline run → `ncdump -h` → baseline RMSE. No parameter changes yet |
187 | **Calibrate** | Baseline is sound, RMSE needs to come down | The parameter set + the bounds +
    why those params → the loop → the iteration log → best params, with the calibration/validation
    split reported |
188 | **Debug** | GLM "ran" but something is wrong or missing | Which of the three silent-failure
    classes it is (no file / truncated / all-NaN) → the confirming check → the fix from
    `troubleshooting.md` |
189 | **Verify** | About to declare done | The five acceptance checks, each PASS/FAIL with its reason.
    Never claim success on an unverified artefact |
190 | **Explain** | User asks how GLM or the layer scheme works | Prose, grounded in the source. Point
    to `references/` rather than pasting cards |
191
192 --
193
194 ## Boundary Rules
195
196 1. **No answer key exists in this skill**, by design.** No calibrated parameter set for Lake Mendota
197 is supplied anywhere in this pack. The literature gives *search bounds and starting points*,
    not
198 optima. The only published, *used* Mendota value citable here is  $K_w = 0.69 \text{ m}^{-1}$  (Bruce et al.
199 2018). If you find yourself pasting parameters from memory, stop -- you are inventing them.
200 2. **Never tune the metric to hit the target.** The RMSE convention (H5) is fixed. Different
201 aggregation choices can materially change RMSE on the same simulation, so the requested
202 convention is a task-local binding rather than a tunable choice. Excluding the ice season or
203 surface layer is not a result; no such exclusion exists in glmtools or LakeEnsemblR.
204 3. **Calibrate at most 4-6 parameters.** Across 73 lakes, only 2-3 of 6 were sensitive
205 (HESS 2025). Ludwig used 4 for Mendota. More free parameters buys equifinality, not accuracy.
206 4. **A parameter pinned to its bound is a diagnostic, not a victory.** It usually means the
    forcing
207 data is wrong and the parameter is absorbing the error. Investigate rather than widening the
    bound.
208 5. **RMSE < 1.0 °C whole-column should be disbelieved.** No published GLM study on a Mendota-like
209 lake reports one. Suspect a broken obs↔sim match (unit error, wrong depth convention,
    near-empty
210 match set) before celebrating. Check N.
211 6. **Report by depth band, always.** Mendota's published surface/bottom gap is 1.30 vs 2.43 °C
212 (NSE 0.97 vs 0.20). An aggregate that hides a broken hypolimnion is not a passing result.
213 7. **Scope: water temperature only, GLM only.Evidence boundary.

Local evidence remains scoped Mendota measurements, literature ranges, commands, and parameter examples are retained as scoped evidence. The reusable claim is the validation or diagnostic principle, so no over-specific marker is needed.



```
220
221 --
222
223 ## Examples
```


```

SKILL.md (SkillAlchemy) (continued)

```
224
225 `examples/worked_session.md` walks four scenarios end to end. Read it if you want to see the loop
226 before running it -- especially the two where the correct move is to distrust the result:
227
228 - Bootstrap -- cold start: `ldd` → nml dump → baseline run → `ncdump -h` → baseline RMSE, in
    that
229 order, touching no parameters until the last step.
230 - Diagnose -- a residual concentrated in the hypolimnion (+2.97 °C, worst in summer) points at
231 exactly one knob (`sed_temp_mean`), and explicitly not at `sw_factor`, which would trade away
    a
232 surface fit that is already good.
233 - The suspiciously good result -- RMSE comes back at 0.62 °C. It is not a win: `N = 37` of
    ~2800
234 observations, because the obs↔sim time match silently collapsed. A number that beats the target
235 for the wrong reason is worse than one that misses it.
236 - Delivery -- the trap people actually lose on: iterating in a scratch dir, then shipping an
    nml
237 that reproduces something other than the delivered `output.nc`. Check 4 of `verify_delivery.py`
238 is what catches it.
    Worked local cases Concrete sessions illustrate the generalized protocol without being presented as universal solutions.
239
240 ## References
241
242 | File | What it holds | Read it when |
243 |----|-----|-----|
244 | `references/sop_models.md` | The 9 operation cards (H1-H9), 8 fields each, with evidence URLs |
    Always -- Step 0 of the protocol |
245 | `references/nml_reference.md` | Every nml block + the Mendota starting config; parameter tiers,
    bounds, units; CSV forcing contracts; the literature inconsistencies | Before changing any
    parameter |
246 | `references/troubleshooting.md` | 22 failure modes: symptom → cause → confirm → fix. "First 60
    seconds" checklist | GLM misbehaves, or produces nothing |
247 | `references/sources.md` | Every URL used, what it establishes, confidence. How to re-verify this
    skill yourself | You doubt a claim here (you should) |
248 | `references/research_notes.md` | The evidence summary, the contradictions found, what stays
    uncertain | You want the reasoning, not the rule |
249 | `references/R01`-`R05` | Raw research reports (GLM source reads, literature, numerical
    verification) | Full provenance for any single claim |
250 | `scripts/` | `nml_edit` · `run_glm` · `extract_temp` · `eval_rmse` · `calibrate` ·
    `verify_delivery` | Throughout -- they encode the traps |
251
252 Three checks a skeptical solver should personally re-run: `ncdump -h output.nc` (confirm the dims
253 and vars), glmtools' `resample_depth()` (confirm the midpoint convention), GLM's `glm_main.c`
254 (confirm the unconditional `exit(0)`). Do not take this skill's word for any of them.
```