
Improving Language Agents through BREW: Bootstrapping experientially-learned Environmental knoWledge

Anonymous Author(s)

Affiliation

Address

email

Abstract

1 Large Language Model (LLM)-based agents are increasingly applied to tasks re-
2 quiring structured reasoning, tool use, and environmental adaptation, such as data
3 manipulation, multistep planning, and computer-use automation. However, despite
4 their versatility, current training paradigms for model weight optimization methods,
5 like PPO and GRPO, remain relatively impractical with their high computational
6 overhead for rollout convergence. In addition, the resulting agent policies are diffi-
7 cult to interpret, adapt, or incrementally improve. To address this, we investigate
8 creating and refining structured memory of experiential learning of an agent from
9 its environment as an alternative route to agent optimization. We introduce **BREW**
10 (Bootstrapping experientially-learned Environmental knoWledge), a framework
11 for agent optimization for downstream tasks via KB construction and refinement.
12 In our formulation, we introduce an effective method for partitioning agent memory
13 for more efficient retrieval and refinement. BREW uses task graders and behavior
14 rubrics to learn insights while leveraging state-space search for ensuring robustness
15 from the noise and non-specificity in natural language. Empirical results on real
16 world, domain-grounded benchmarks – OSWorld and τ^2 Bench – show BREW
17 achieves 10 – 20% improvement in task precision, 10 – 15% reduction in API/-
18 tool calls leading to faster execution time, all while maintaining computational
19 efficiency on par with base models. Unlike prior work where memory is treated as
20 static context, we establish the KB as a modular and controllable substrate for agent
21 optimization – an explicit lever for shaping behavior in a transparent, interpretable,
22 and extensible manner.

23 1 Introduction

24 Large Language Model (LLM) based agents are rapidly being deployed for structured reasoning,
25 tool use, and autonomous interaction in real-world environments [16]. From computer-use and
26 spreadsheet automation to software engineering pipelines, these agents drive tasks such as multi-step
27 planning, data manipulation, and adaptive workflows [21, 13, 32, 2, 19]. For example, a language
28 agent might help automate a multi-step workflow like collecting data from different sources, cleaning
29 or validating it, and then uploading it onto a dedicated server, all while adjusting its plan if the
30 format or structure of the data changes unexpectedly [31, 35, 25, 3]. Yet, despite these successes, top-
31 performing agents generally score underwhelmingly on challenging real-world benchmarks—well
32 behind human experts, who routinely exceed 70% success rates [34, 4, 27, 18]. As an example,
33 consider the following scenario:

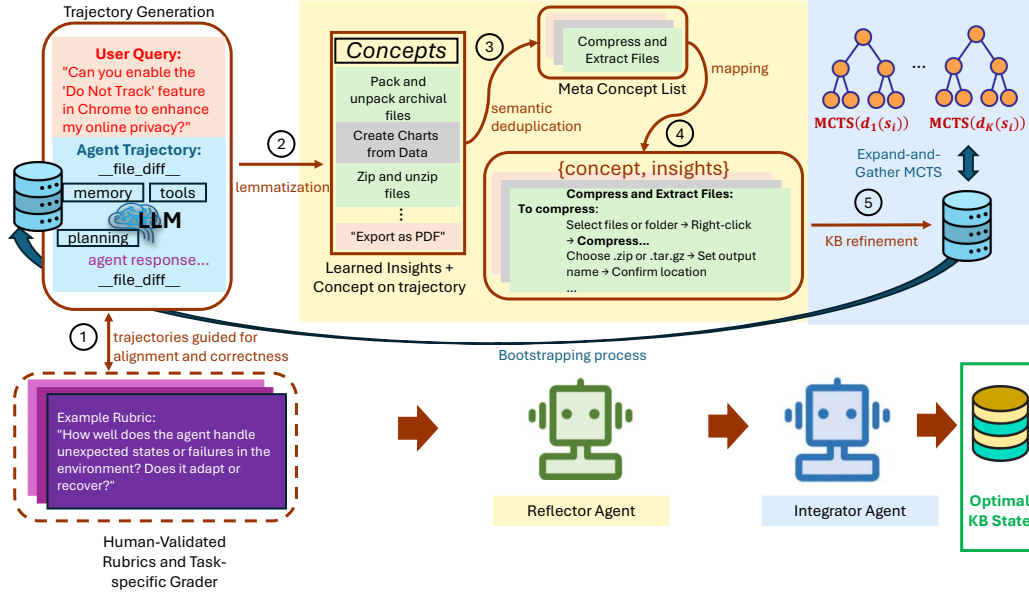


Figure 1: BREW architecture overview using examples from the OSWorld dataset. Step 1 indicates the trajectory generation process with agent alignment to human-validated rubrics and correctness using task-specific grader. Steps 2–4 indicate the Reflector Agent, which learns key concepts and corresponding insights from trajectories. Step 5 indicates the Integrator Agent, which integrates knowledge from the Reflector Agent to bootstrap the KB. We introduce Expand-and-Gather MCTS to further find the best KB configuration as the KB is iteratively refined through reward-guided optimization.

34

A computer-use agent in an Ubuntu environment tasked with automating software installation across multiple sessions. In its first encounter, it struggles through a 47-step process: opening the wrong package manager, executing redundant dependency checks, and making 23 API calls to complete what could be a 6-step workflow. When presented with a similar installation task in the next session, the agent repeats the same inefficient exploration – *as if encountering the problem for the first time*. A human user, by contrast, would likely have a recollection from internalized memory of the optimal sequence after the first attempt, recognizing the environmental patterns and tool combinations that lead to success.

35 This scenario illustrates a fundamental limitation of current language agents: despite their impressive
 36 capabilities in reasoning and tool use, they lack the ability to accumulate and apply experiential
 37 knowledge across task sessions. Each interaction begins from a blank slate, forcing agents to
 38 repeatedly explore the same action spaces and rediscover the same solutions [9]. Real-world tasks
 39 like long horizon multi-stage automation demand more than just “reactive” [33] tool loops. They
 40 require persistent & interpretable learnings from past experiences - what works, what fails and why.

41 To close this gap, recent work has explored learning agent behavior using model weight optimiza-
 42 tion [23, 22, 24], where agents are trained to maximize success across a wide variety of tool-use
 43 episodes. However, while conceptually sound, this suffers from practical limitations. First, it requires
 44 expansive exploration over large rollout spaces to converge, especially in domains where tasks are
 45 diverse, goals are sparsely defined, and intermediate feedback is noisy or delayed. Second, the
 46 resulting policies are often opaque—difficult to interpret, revise, or debug—limiting their real-world
 47 deployability. Finally, these policies are tightly coupled to the task distributions they were trained on,
 48 making it difficult to adapt or incrementally improve them when downstream requirements shift.

49 In contrast, others have explored learning of knowledge onto a memory module that remains attached
 50 to an agent. These existing memory-augmented agents can be broadly classified into either ones which
 51 (i) store only transient trajectory contexts that vanish between episodes like Mem0 [7, 29], or (ii)

embed high-level notes directly in the prompt such as MetaReflection [10] and GEPA[1]. While the latter often do not retain actionable details for future simple tasks, neither of these approach supports modular updates, fine-grained retrieval, or transparent inspection of what the agent “knows.” [28].

Leveraging learnings from both camps, we introduce BREW (**B**ootstrapping experientially-learned **E**nvironmental knowledge), a framework that incrementally constructs and refines a knowledge base (KB) a structured collection of concept-level documents in natural language, directly from an agent’s past interactions. This KB then serves as a persistent memory for the agent to retrieve knowledge in future executions to improve precision and efficiency outcomes. Our key contributions are—

- *Novel experience-driven KB construction.* We propose a technique for leveraging agent’s past interaction trajectories to generate uniquely-partitioned concept-level KB documents. This process is guided by rubrics and task-specific graders which ensures that memories are both semantically aligned with task objectives and human-interpretable.
- *State-space search for memory optimization.* We formalize the selection and update of KB entries as a state search problem and introduce an efficient reward-guided learning scheme, Expand-and-Gather Monte Carlo Tree Search (EG-MCTS), that learns to prioritize the most impactful memories for robust, multi-step reasoning.
- *State-of-the-art results.* On domain-grounded benchmarks including OSWorld and τ^2 Bench, BREW achieves significant gains of in the range of 10 – 20% towards task precision as well as 10 – 15% fewer steps leading to faster execution, while maintaining memory and compute costs comparable to base LLMs.

2 Preliminary & Related Work

Agent Learning from Demonstrations Recent work has leveraged LLMs to isolate reusable skills through interactive decomposition: one method distills sub-goals from expert trajectories into hierarchical planning and execution policies [11], and another synthesizes executable functional abstractions for advanced mathematical reasoning via program induction [14]. These approaches focus on structured skill extraction from LLM-guided interactions, yet remain reliant on static decomposition or offline synthesis. In contrast, BREW dynamically constructs and refines an experiential memory—learning necessary semantic fragments via rollout-generated insights and structured knowledge-base search (MCTS)—to support long-horizon, memory-augmented planning.

Agentic Memory The concept of providing agents with controllable memory has a rich history. [17]. Memory mechanisms are attracting more and more attention lately [20, 26, 28, 7, 30, 12]. These works focus towards storing relevant context in a structured format like graph or a tree so as to RAG over it. Despite their effectiveness these methods perform well for most cases. However, when the queries are ambiguous, requires multi-hop reasoning and long range comprehension these techniques struggle to perform the tasks [12]. In contrast to prior works BREW uses a state search to explore possible memory states. This allows BREW to select the memory state where the reward during exploration is highest making it more robust to ambiguous queries and long range comprehension. We employ MCTS [8] as a state search algorithm to explore the potential states of the memory by expanding to new and potentially different states of memory based on same interactions. We discuss the state search process more formally in Section 3.3.

3 BREW: Architecture

This section describes our proposed **B**ootstrapping experientially-learned **E**nvironmental know**W**ledge model, BREW, which constructs and iteratively refines a KB using trajectory insights guided by human-validated general-purpose agent behavior metrics, task-specific evaluation, and latent insight generation. We introduce a novel decomposition the problem of learning the optimal KB by partitioning memory as local documents associated with semantic concepts, and formulate the KB learning problem as a state space search by proposing Expand-and-Gather Monte Carlo Tree Search (EG-MCTS). Figure 1 provides an architecture overview of BREW, and Algorithm 1 describes pseudocode.

3.1 Trajectory Generation

Given the training dataset, we generate full-length trajectories, hereby referred to as rollouts, for each query using an LLM-powered agent conditioned on its associated KB. At initialization, the KB is empty, and the LLM is used with a decoding temperature of 0 to ensure deterministic behavior. Further details on training and test splits are in Experiments Section. Each rollout is evaluated using a correctness grader, which assigns a binary label: *success* or *failure* and a qualitative rubric assessment against a set of human-validated general-purpose agent behavior rubrics [6] (Step 1 in Figure 1).

3.2 Reflector and Integrator Agents

Reflector Agent: ReflAgent takes as input a rollout with its rubric and correctness labels, and outputs sentence-level insights with mapped concepts:

$$\{\text{concepts}, \text{insights}\} = \text{ReflAgent}(\{\text{rollout}, \text{eval}\}). \quad (1)$$

Examples of concept–insight pairs appear in Step 2 of Figure 1.

Concept Deduplication: Concept–insight pairs are annotated independently per rollout, often producing overlapping or paraphrased concepts. We address this via semantic clustering (Steps 3–4, Figure 1; Algorithm 1, line 3): contextual embeddings for each concept are generated using an LLM, clustered, and each insight is mapped to its cluster representative. Details appear in Algorithms 2 and 3 in Appendix A.

Integrator Agent: IntegAgent incrementally builds and refines KB documents $\{d(s_i)\} \in \mathcal{D}(s_i)$ during environment interaction. Instead of a centralized memory, the KB is partitioned into local documents, each tied to a meta concept. This design enables (1) efficient, context-specific retrieval; (2) modular updates with minimal interference; and (3) natural alignment with task semantics, as deduplicated meta concepts capture meaningful behavioral abstractions. Unlike prior work assuming flat memory or dialogue histories, this structure is well-suited for long-horizon, procedural tasks where behaviors cluster around discrete skills.

The KB is dynamically populated: concepts central to the dataset receive more updates, shaping memory around frequent behaviors. At each state, for meta concept k , IntegAgent updates its document d_k via

$$d_k(s_{i+1}) \leftarrow \text{IntegAgent}(k, \text{insights}_k, d_k(s_i)). \quad (2)$$

To reduce LLM variance and improve consistency, we use the Expand-and-Gather MCTS (EG-MCTS) method (Figure 2).

Formally, the KB at state s_i is the union of all concept-localized documents:

$$\mathcal{D}(s_i) = \bigcup_{k \in \mathcal{K}} \{d_k(s_i)\}, \quad (3)$$

where $\mathcal{K}$ is the set of all meta concepts and $d_k(s_i)$ is the document for concept k at state s_i .

3.3 Expand-and-Gather MCTS for Optimal KB Search

We start by creating a set of meta-concepts after deduplicating concepts extracted by ReflAgent using the first set of trajectory rollouts. We freeze this meta-concept set $\mathcal{K}$, and use it to initialize a KB with an empty document per concept $k \in \mathcal{K}$.

We model the problem of finding the optimal KB $\mathcal{D}^*$ as a search problem in the *state* space of all possible KBs $\mathcal{D}$. To simplify this state search, we model KB $\mathcal{D}$ as a collection of concept level documents. This modeling allows us to break down the larger search space into a collection of simpler document level search problems for each concept k to find the optimal document d_k^* . We then construct the optimal KB $\mathcal{D}^*$ by combining all optimal documents d_k^* for each concept k as follows:

$$\mathcal{D}^* = \bigcup_{\forall k} \{d_k^*\} \quad (4)$$

Notably, even though we are modeling document level search as independent optimization problems, each document in the KB is *not* independent of the others. For example, an agent can retrieve any

Algorithm 1 BREW: Bootstrapping Experientially-learned Environmental Knowledge

Require: Training samples $\mathcal{Q}_{\text{train}}$, eval samples $\mathcal{Q}_{\text{eval}}$, rubrics, iterations M , candidates per expansion h

Ensure: Optimized KB $\mathcal{D}^*$

Initialization

```
1:  $\mathcal{D}_0 \leftarrow \emptyset$ 
2:  $\mathcal{B} \leftarrow \text{GENERATEINSIGHTS}(\mathcal{Q}_{\text{train}}, \mathcal{D}_0, \text{rubrics})$ 
3:  $\mathcal{K} \leftarrow \text{DEDUPLICATECONCEPTS}(\mathcal{B})$  ▷ Initial concept set
4: for each  $k \in \mathcal{K}$  do
5:    $d_k^0 \leftarrow \text{INTEGAGENT}(k, \mathcal{I}_k, \emptyset)$ 
6:   Initialize  $\text{tree}_k$  with root node  $d_k^0$ 
7: end for
8:  $\mathcal{D}_{\text{current}} \leftarrow \bigcup_{k \in \mathcal{K}} \{d_k^0\}$  ▷ Initial KB
```

EG-MCTS Optimization

```
9: for  $t = 1$  to  $M$  do ▷ Parallel expansion across concepts
10:   for each  $k \in \mathcal{K}$  do
11:      $s_k \leftarrow \text{SELECTBESTNODE}(\text{tree}_k)$  ▷ UCT selection
12:      $\mathcal{D}_{\text{best}} \leftarrow \bigcup_{k' \in \mathcal{K}} \{d_{k'}^{\text{best}}\}$  ▷ Current best docs
13:      $\text{EXPANDNODE}(s_k, k, h, \mathcal{D}_{\text{current}}, \mathcal{D}_{\text{best}}, \text{tree}_k)$ 
14:   end for ▷ Update current best documents
15:   for each  $k \in \mathcal{K}$  do
16:      $d_k^{\text{best}} \leftarrow \text{best document in } \text{tree}_k$ 
17:   end for
18:    $\mathcal{D}_{\text{current}} \leftarrow \bigcup_{k \in \mathcal{K}} \{d_k^{\text{best}}\}$ 
19: end for
20: return  $\mathcal{D}_{\text{current}}$ 
```

Time Complexity: $O(|\mathcal{Q}_{\text{train}}| \cdot T_{\text{LLM}} + M \cdot |\mathcal{K}| \cdot h \cdot T_{\text{agent}})$

document in the KB during inference and this retrieval making it hard to assess the impact of changing a document in isolation. To solve this we propose Expand-and-Gather MCTS (EG-MCTS), which enables searching these disjoint state spaces concurrently using parallel MCTS explorations that are synced after each iteration. To achieve this we perform node expansions in the respective search spaces independently but condition reward calculation and insight generation on a running optimum KB state. Each iteration of EG-MCTS can be broken down two phases:

Expand Phase: During this stage, for each search tree, we pick the *best* state s^* and expand it concurrently. To perform this expansion the KB $\mathcal{D}(s^*)$ is constructed by including the *current document* $d_k(s^*)$ and the *best (oracle) documents* $\{d_i^*\}_{i \neq k}$ for all other positions. Thus, the KB at iteration t , $0 \leq t \leq E$ is defined as:

$$\mathcal{D}_t = d_t \cup d_{i:i \neq t}^* \quad (5)$$

We use this KB $\mathcal{D}(s_i)$ to generate trajectory rollouts which are consumed by the `RefAgent` to generate insights. We then use the `IntegAgent` to generate various updated variants of d_k^* e.g., $d_k(s_i), \dots, d_k(s_j)$, where $0 \leq i \leq E$ and $0 \leq j \leq E$. We then estimate a reward R for each of these newly generate states and update rewards of parent states using backpropagation.

Gather Phase: During this stage, the current best states from each document’s MCTS tree are *gathered* together and distributed to every MCTS tree for reward calculation. This is important to 1. Estimate rewards for each expanded state, and 2. Generate new insights for further node expansion.

3.4 Reward-Guided Optimization

This section describes BREW’s joint reward and loss optimization for learning an optimal KB.

Reward Objective: Each document state is rewarded based on two complementary criteria: (i) how well the current document contributes to **accurate downstream reasoning**, and (ii) how **retrievable** it is in the context of a growing KB. Formally, the total reward at time step t is defined as:

$$R_t = \lambda_{\text{corr}} \cdot R_t^{\text{corr}} + \lambda_{\text{ret}} \cdot R_t^{\text{ret}} \quad (6)$$

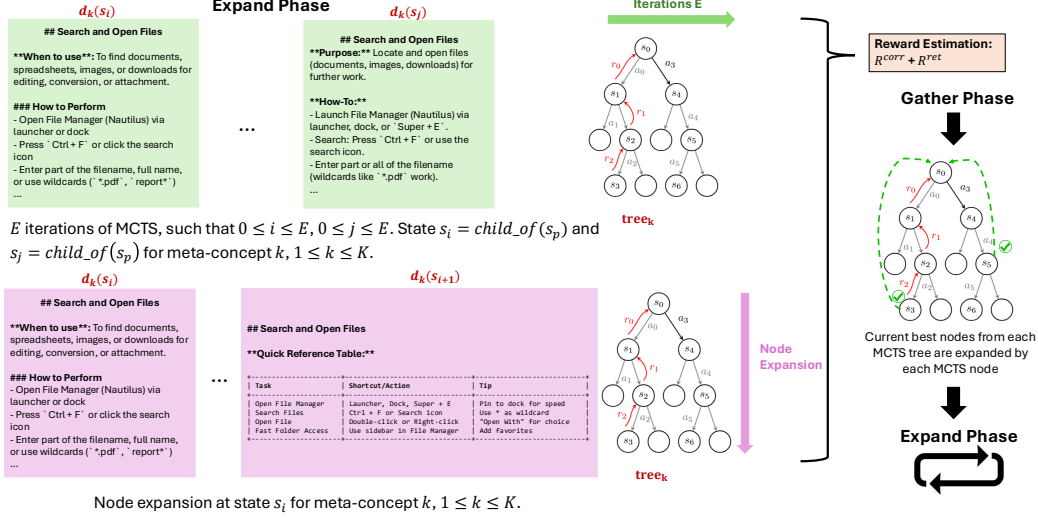


Figure 2: Illustration of BREW’s KB optimization process using Expand-and-Gather MCTS with OSWorld examples. In the **Expand Phase**, for each document k , we sample the best node from tree_k using UCT and perform node expansion. Node rewards are estimated based on correctness and retrievability. In the **Gather Phase**, the current best nodes from each tree are gathered per node, and the objective function is optimized. The process is repeated during the next iteration of KB refinement.

where R_t^{corr} is the **correctness reward**, R_t^{ret} is the **retrieval reward**, and $\lambda_{\text{corr}}, \lambda_{\text{ret}} \in [0, 1]$ are scalar weights with $\lambda_{\text{corr}} + \lambda_{\text{ret}} = 1$.

Correctness Reward: The correctness reward R_t^{corr} evaluates the accuracy of the agent’s output over a held-out query set $\mathcal{Q}$, when reasoning over the current KB $\mathcal{D}_t$. It is defined as:

$$R_t^{\text{corr}}(d_t|\mathcal{D}_t) = \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \text{Eval}_{\text{task}}(q, \text{agent} \oplus \mathcal{D}_t) \quad (7)$$

where $\text{Eval}_{\text{task}}$ is a task-specific evaluation function (e.g., question-answering accuracy, entailment correctness), and $\text{agent} \oplus \mathcal{D}_t$ denotes the agent acting over the hybrid KB.

Retrieval Reward: The retrieval reward R_t^{ret} measures how effectively the current document d_t can be retrieved from the current KB $\mathcal{D}_t$. For a held-out query set $\mathcal{Q}$, it is computed using the mean reciprocal rank (MRR):

$$R_t^{\text{ret}}(d_t|\mathcal{D}_t) = \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \text{MRR}_q(d_t, \mathcal{D}_t) \quad (8)$$

This encourages documents that are not only helpful in reasoning but also easily retrievable via the retrieval model over $\mathcal{D}_t$.

4 Experimental Setup

Datasets We evaluate BREW on three diverse benchmarks testing different aspects of interactive agent capabilities: OSWORLD for computer-use automation [27], τ^2 -Bench for tool use [5], and SPREADSHEETBENCH for data manipulation [18].

1. **OSWorld:** This benchmark tests multimodal agents on real-world computer tasks across 10 applications. We use *GTAI-7B*, a state-of-the-art computer-use agents with BREW. Tasks are evaluated using 134 custom scripts that verify final application states.
2. **τ^2 -Bench:** This benchmark evaluates conversational agents on multi-turn tool-use scenarios across *Telecom*, *Retail*, and *Airline* domains. We test o4-mini-based tool-calling agent, constructing BREW KBs for every domain.

185 3. **SpreadsheetBench:** This benchmark evaluates agents on real-world spreadsheet manipulation,
 186 spanning both cell-level and sheet-level tasks. It contains 912 authentic user instructions paired
 187 with 2,729 test cases (3 per instruction), sourced from Excel forums and blogs. Spreadsheets
 188 include diverse formats with multi-table sheets (35.7%) and non-standard tables (42.7%). We test
 189 o4-mini using a Python tool-calling agent, and enhance it with by adding an embedding based
 190 Retrieval over the BREW KB generated over a small held-out train set of 30 samples.

191 **Baselines** We compare BREW against two widely used experiential memory approaches, *Cognee*¹
 192 and *Agent-Mem* [30], both of which serve as established baselines for AI memory evaluation. Cognee
 193 is an open-source AI memory engine that employs a graph-plus-vector memory architecture through
 194 an Extract–Connect–Learn pipeline, enabling agents to construct cross-document and cross-context
 195 connections entirely from previously available trajectories. In contrast, Agent-Mem provides a
 196 scalable memory layer for dynamically extracting and retrieving information from conversational data,
 197 with enhanced variants incorporating graph-based memory representations. While Cognee primarily
 198 emphasizes cross-document relational reasoning, Agent-Mem focuses on scalable personalization for
 199 conversational agents.

200 **Other Experimental Configs:** For all experiments, we use GPT-4.1-2025-04-14 as the base
 201 LLM with expansion width $e = 3$, max depth $k = 3$, and balanced reward weights $\lambda_{\text{corr}} = \lambda_{\text{ret}} = 0.5$.
 202 During MCTS node selection, we use the UCT [15] for balancing exploration and exploitation Full
 203 experimental details are provided in the Appendix.

204 5 Analysis & Discussion

205 In this section, we present findings from our evaluation of BREW. For more details on qualitative
 206 insights and discussion you may refer to the supplementary material.

207 **Variations with State Search Strategy** BREW performs a search across possible KB states using
 208 MCTS. We compare different state search strategies to determine the relative trade-offs:

- 209 1. *Iterative Refinement:* In this strategy we generate one version of each document to generate an
 210 initial KB, followed by a round of evals. We then use the aggregator agent to refine the documents
 211 over the newly learned insights. We repeat this step multiple times up to a maximum number of
 212 refinements. Note that in contrast to MCTS, in this strategy we *do not* perform node expansions
 213 and rather explore a path in the search tree.
- 214 2. *Greedy Search:* In this strategy we greedily pick the best state during each node expansion and
 215 only explore the sub-tree within it. This is in contrast to MCTS where, we explore different states
 216 using the UCT algorithm that balances exploration and exploitation.

217 Table 1 presents how MCTS achieves consistent performance gains across all benchmarks. These
 218 represent 1-5% improvements over alternative search strategies across tasks. Iterative refinement’s
 219 poor performance reveals core limitations in the integrator agent feedback incorporation- which can
 220 be attributed to inherent stochasticity in LLMs. This makes state exploration especially important for
 221 textual optimization tasks like ours. We present a detailed analysis on how varying MCTS parameters
 222 result in different final states in appendix.

223 5.1 Trends across Sub-Tasks

224 **BREW learns recipes from sub-trajectories in OSWorld.** Figure 3 shows that BREW(BREW)
 225 improves success rates in 5 out of 10 OSWorld categories, achieving absolute gains of 4–16%
 226 while maintaining performance parity in the remaining categories (Chrome, Gimp, LibreOffice
 227 Calc, LibreOffice Impress, OS). The largest improvements appear in text-processing applications
 228 (LibreOffice Writer: 14% $\rightarrow$ 24%, Thunderbird: 38% $\rightarrow$ 54%) and multimedia tools (VLC:
 229 20% $\rightarrow$ 27%), with moderate gains in multi-application and development environments. Even in
 230 settings with limited improvements in task correctness, BREW consistently reduces execution length
 231 by 14–23 steps, highlighting more efficient planning.

¹github.com/topoteretes/cognee

Method	OSWorld GTA1-7B	τ^2 Bench o4-mini	SpreadsheetBench o4-mini
Baseline	44.20	56.63	44.30
Cognee	46.70	57.71	42.10
Agent-Mem	43.83	52.69	42.00
BREW-Iterative	46.13	57.34	42.98
BREW-Greedy	45.55	59.14	45.94
BREW-MCTS	47.56	59.14	46.80

Table 1: Comparison of models under different evaluation setups, including Baseline model and BREW augmented model. We report task success rate for OSWorld, ratio of independent tasks that succeeded for τ^2 Bench, and the 1st test case pass rate for SpreadsheetBench.

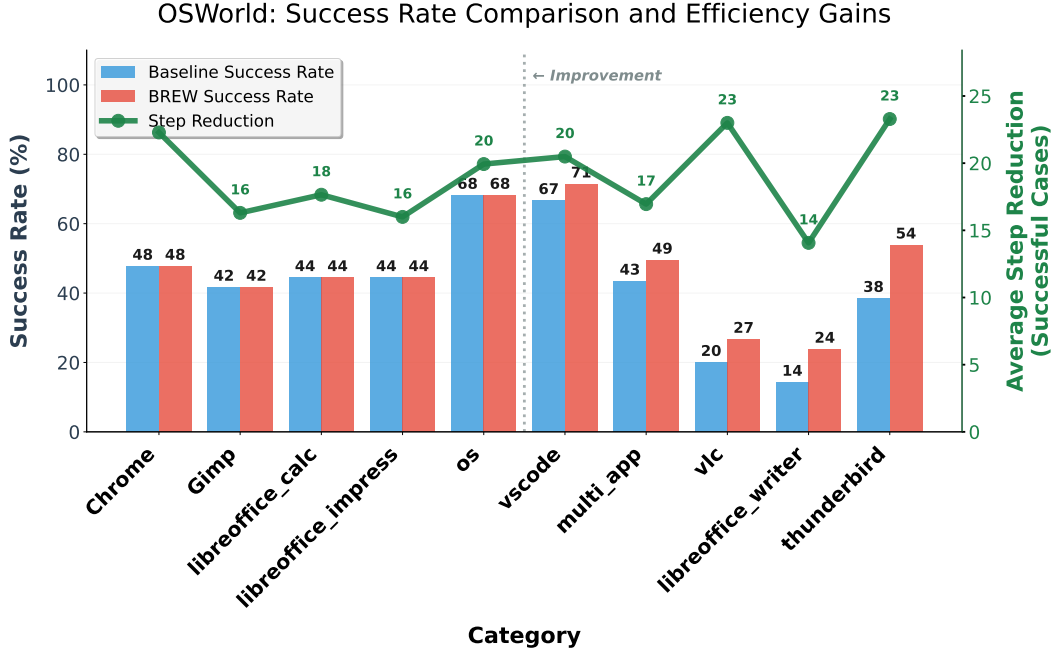


Figure 3: The bar plot represents the category-wise success rate over various tasks in the OSWorld dataset over the GTA1-agent, whereas the line plot demonstrates the reduction in the number of steps for the successful cases. Note that even in scenarios where the KB doesn’t help increase the success rate, it significantly reduces the number of steps needed to succeed.

232 This pattern suggests that BREW’s architectural enhancements are particularly effective for tasks
233 requiring complex sequential reasoning and inter-application coordination, while preserving baseline
234 robustness in domains constrained by intrinsic task complexity.

235 A qualitative analysis of the knowledge bases (KBs) constructed by BREW further supports this
236 finding. We observe that BREW captures and represents sub-trajectory characteristics in *natural*
237 *language*, including application shortcuts, standard operating procedures, and strategies for localizing
238 UI elements. Since many UI tasks share common sub-trajectories, this representation facilitates
239 knowledge transfer across tasks within the same application. Moreover, BREW substantially reduces
240 reliance on granular UI interactions: while the baseline GTA1 model executes approximately 19,000
241 clicks and 17,821 keyboard actions, BREW significantly decreases this interaction complexity.

242 **BREW learns aggressive resolution strategies for τ^2 –Bench** To evaluate robustness of BREW,
243 we analyzed the distribution of failure modes across the τ^2 –retail dataset, focusing on four key error
244 categories: *Wrong Argument*, *Wrong Info*, *Wrong Decision*, and *Partially Resolve*. Figure 4 presents
245 a comparative chart for the baseline, BREW, Cognee and Agent-Mem[30].

Overall, BREW demonstrated consistent improvements across most error types compared to the baseline and competing approaches. Specifically, BREW showed a **notable reduction in “Wrong Argument” and “Wrong Decision” errors**, indicating that it was better at capturing logical dependencies in retail dialogues and making accurate decisions.

Interestingly, *Partially Resolve* errors were slightly higher for BREW than for Cognee, likely because BREW attempted more aggressive resolution strategies that occasionally failed to fully satisfy user queries. Cognee appears to capture *richer factual details* given its relatively lower *Wrong Info* errors, whereas Agent-Mem excels in *tracking conversation state and decision accuracy*, as reflected in its reduced *Wrong Decision* failures.

Improvements in Task Efficiency We observe that overall, BREW enables agents to come to a correct response quicker.

OSworld. Figure 3 demonstrates that BREW enables GTA1 to complete tasks more efficiently. Compared to the baseline GTA1 model’s average of ~ 75 steps, the BREW-augmented model completes tasks 14% faster with an average of ~ 64 steps. Analyzing performance by outcome reveals that while step counts remain unchanged for failed cases, successful completions show a substantial 39% (rel.) reduction in execution steps, indicating improved planning efficiency for achievable tasks.

τ^2 *Bench*. Similarly, BREW reduces average conversation turns from 29.47 to 28.43 (-3.5%), while maintaining consistent step reductions across categories. Step reductions average 1.7 steps for Retail and Telecom, but 3.1 steps for Airline, indicating greater efficiency gains in complex domains. Qualitative analysis seconds these numbers showing how knowledge base integration enables more direct task completion paths and improved planning quality, though multi-turn interactions remain necessary for complex sub-tasks.

SpreadsheetBench. While we observe a slight increase in the number of turns across the entire benchmark suite ($4.5 \rightarrow 5.4$) in the case of the baseline versus BREW, an interesting pattern emerges in more than 82% of the cases the baseline and the BREW appended agent performs similarly with similar turn consumption. BREW leads to an improvement in 12% of the cases where the KB is able to address gaps in the baseline technique to enable the agent to go exploring further leading to positive outcomes with an average of 1 step increase in the interactions.

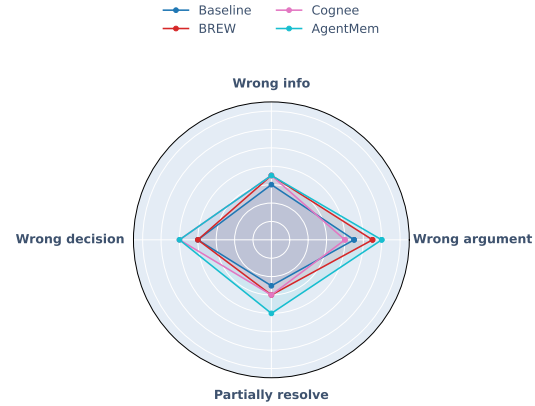


Figure 4: Distribution of errors in τ^2 Bench Retail

6 Conclusions

In this work, we explored an alternative approach to agent optimization by focusing on experiential knowledge retention rather than direct model fine-tuning. We introduced BREW, a framework that aims to construct and refine a structured, interpretable knowledge base from past agent interactions. By decomposing agent memory into concept-level documents and applying a state-search optimization strategy, BREW provides a modular and transparent substrate for memory formation. Our evaluations across OSWorld and τ^2 Bench benchmarks suggest that such structured memory can support measurable improvements in task success and efficiency, while maintaining manageable computational costs. Although the observed gains are promising, we recognize that BREW’s effectiveness is influenced by the quality and coverage of its training data. Future work could explore more adaptive and domain-general memory refinement techniques, as well as tighter integrations with ongoing agent planning. Ultimately, we hope this study encourages further investigation into more interpretable, memory-driven approaches to language agent development—especially in real-world environments where long-term consistency and adaptability are essential.

References

- [1] Lakshya A. Agrawal, Shangyin Tan, Dilara Soylu, Noah Ziem, Rishi Khare, Krista Opsahl-Ong, Arnav Singhvi, Herumb Shandilya, Michael J. Ryan, Meng Jiang, Christopher Potts, Koushik Sen, Alexandros G. Dimakis, Ion Stoica, Dan Klein, Matei Zaharia, and Omar Khattab. Gepa: Reflective prompt evolution can outperform reinforcement learning. *arXiv preprint arXiv:2507.19457*, July 2025.
- [2] Anthropic. Introducing computer use, a new Claude 3.5 Sonnet, and Claude 3.5 Haiku, October 2024. URL <https://www.anthropic.com/news/3-5-models-and-computer-use>. Accessed: 2025.
- [3] Yasharth Bajpai, Bhavya Chopra, Param Biyani, Cagri Aslan, Dustin Coleman, Sumit Gulwani, Chris Parnin, Arjun Radhakrishna, and Gustavo Soares. Let’s fix this together: Conversational debugging with github copilot. In *2024 IEEE Symposium on Visual Languages and Human-Centric Computing (VL/HCC)*, pages 1–12, 2024. doi: 10.1109/VL/HCC60511.2024.00011.
- [4] Victor Barres, Honghua Dong, Soham Ray, Xujie Si, and Karthik Narasimhan. τ^2 -bench: Evaluating conversational agents in a dual-control environment, 2025. URL <https://arxiv.org/abs/2506.07982>.
- [5] Victor Barres, Honghua Dong, Soham Ray, Xujie Si, and Karthik Narasimhan. τ^2 -bench: Evaluating conversational agents in a dual-control environment, 2025. URL <https://arxiv.org/abs/2506.07982>.
- [6] Param Biyani, Yasharth Bajpai, Arjun Radhakrishna, Gustavo Soares, and Sumit Gulwani. Rubicon: Rubric-based evaluation of domain-specific human ai conversations. In *Proceedings of the 1st ACM International Conference on AI-Powered Software*, AIware 2024, page 161–169, New York, NY, USA, 2024. Association for Computing Machinery. ISBN 9798400706851. doi: 10.1145/3664646.3664778. URL <https://doi.org/10.1145/3664646.3664778>.
- [7] Prateek Chhikara, Dev Khant, Saket Aryan, Taranjeet Singh, and Deshraj Yadav. Mem0: Building production-ready ai agents with scalable long-term memory. *arXiv preprint arXiv:2504.19413*, 2025.
- [8] Rémi Coulom. Efficient selectivity and backup operators in monte-carlo tree search. In *Proceedings of the 5th International Conference on Computers and Games (CG 2006)*, pages 72–83. Springer, 2006. doi: 10.1007/978-3-540-75538-8_7.
- [9] Lutfi Eren Erdogan, Nicholas Lee, Sehoon Kim, Suhong Moon, Hiroki Furuta, Gopala Anumanchipalli, Kurt Keutzer, and Amir Gholami. Plan-and-act: Improving planning of agents for long-horizon tasks. *The Forty-Second International Conference on Machine Learning*, 2025.
- [10] Priyanshu Gupta, Shashank Kirtania, Ananya Singha, Sumit Gulwani, Arjun Radhakrishna, Gustavo Soares, and Sherry Shi. MetaReflection: Learning instructions for language agents using past reflections. In Yaser Al-Onaizan, Mohit Bansal, and Yun-Nung Chen, editors, *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pages 8369–8385, Miami, Florida, USA, November 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.emnlp-main.477. URL <https://aclanthology.org/2024.emnlp-main.477/>.
- [11] Maryam Hashemzadeh, Elias Stengel-Eskin, Sarath Chandar, and Marc-Alexandre Cote. Sub-goal distillation: A method to improve small language agents, 2024. URL <https://arxiv.org/abs/2405.02749>.
- [12] Yuanzhe Hu, Yu Wang, and Julian McAuley. Evaluating memory in llm agents via incremental multi-turn interactions, 2025. URL <https://arxiv.org/abs/2507.05257>.
- [13] Carlos E Jimenez, John Yang, Alexander Wettig, Shunyu Yao, Kexin Pei, Ofir Press, and Karthik R Narasimhan. SWE-bench: Can language models resolve real-world github issues? In *The Twelfth International Conference on Learning Representations*, 2024. URL <https://openreview.net/forum?id=VTF8yNQm66>.

- [14] Zaid Khan, Elias Stengel-Eskin, Archiki Prasad, Jaemin Cho, and Mohit Bansal. Executable functional abstractions: Inferring generative programs for advanced math problems. 2025.
- [15] Levente Kocsis and Csaba Szepesvári. Bandit based monte-carlo planning. In Johannes Fürnkranz, Tobias Scheffer, and Myra Spiliopoulou, editors, *Machine Learning: ECML 2006*, pages 282–293, Berlin, Heidelberg, 2006. Springer Berlin Heidelberg. ISBN 978-3-540-46056-5.
- [16] Xinzhe Li. A review of prominent paradigms for llm-based agents: Tool use, planning (including rag), and feedback learning. In *Proceedings of the 31st International Conference on Computational Linguistics (COLING)*, pages 9760–9779, Abu Dhabi, UAE, 2025. Association for Computational Linguistics. URL <https://aclanthology.org/2025.coling-main.652>.
- [17] Michael L. Littman. An optimization-based categorization of reinforcement learning environments. 1993. URL <https://api.semanticscholar.org/CorpusID:17988064>.
- [18] Zeyao Ma, Bohan Zhang, Jing Zhang, Jifan Yu, Xiaokang Zhang, Xiaohan Zhang, Sijia Luo, Xi Wang, and Jie Tang. Spreadsheetbench: Towards challenging real world spreadsheet manipulation. *Advances in Neural Information Processing Systems*, 37:94871–94908, 2024.
- [19] OpenAI. Introducing Operator, January 2025. URL <https://openai.com/index/introducing-operator/>. Accessed: 2025.
- [20] Charles Packer, Sarah Wooders, Kevin Lin, Vivian Fang, Shishir G. Patil, Ion Stoica, and Joseph E. Gonzalez. Memgpt: Towards llms as operating systems, 2024. URL <https://arxiv.org/abs/2310.08560>.
- [21] Yujia Qin, Yining Ye, Junjie Fang, Haoming Wang, Shihao Liang, Shizuo Tian, Junda Zhang, Jiahao Li, Yunxin Li, Shijue Huang, et al. Ui-tars: Pioneering automated gui interaction with native agents. *arXiv preprint arXiv:2501.12326*, 2025.
- [22] Rafael Rafailov, Archit Sharma, Eric Mitchell, Stefano Ermon, Christopher D. Manning, and Chelsea Finn. Direct preference optimization: Your language model is secretly a reward model, 2024. URL <https://arxiv.org/abs/2305.18290>.
- [23] John Schulman, Filip Wolski, Prafulla Dhariwal, Alec Radford, and Oleg Klimov. Proximal policy optimization algorithms. In *Proceedings of the 34th International Conference on Machine Learning (ICML 2017)*, 2017. URL <https://arxiv.org/abs/1707.06347>.
- [24] Zhihong Shao, Peiyi Wang, Qihao Zhu, Runxin Xu, Junxiao Song, Xiao Bi, Haowei Zhang, Mingchuan Zhang, Y. K. Li, Y. Wu, and Daya Guo. Deepseekmath: Pushing the limits of mathematical reasoning in open language models, 2024. URL <https://arxiv.org/abs/2402.03300>.
- [25] Noah Shinn, Federico Cassano, Ashwin Gopinath, Karthik Narasimhan, and Shunyu Yao. Reflexion: Language agents with verbal reinforcement learning. In *Proceedings of the 37th Conference on Neural Information Processing Systems (NeurIPS 2023)*, New Orleans, LA, USA, 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/hash/1b44b878bb782e6954cd888628510e90-Abstract-Conference.html.
- [26] Yu Wang, Chi Han, Tongtong Wu, Xiaoxin He, Wangchunshu Zhou, Nafis Sadeq, Xiusi Chen, Zexue He, Wei Wang, Gholamreza Haffari, Heng Ji, and Julian McAuley. Towards lifespan cognitive systems, 2025. URL <https://arxiv.org/abs/2409.13265>.
- [27] Tianbao Xie, Danyang Zhang, Jixuan Chen, Xiaochuan Li, Siheng Zhao, Ruisheng Cao, Toh Jing Hua, Zhoujun Cheng, Dongchan Shin, Fangyu Lei, Yitao Liu, Yiheng Xu, Shuyan Zhou, Silvio Savarese, Caiming Xiong, Victor Zhong, and Tao Yu. Osworld: Benchmarking multimodal agents for open-ended tasks in real computer environments. In A. Globerson, L. Mackey, D. Belgrave, A. Fan, U. Paquet, J. Tomczak, and C. Zhang, editors, *Advances in Neural Information Processing Systems*, volume 37, pages 52040–52094. Curran Associates, Inc., 2024. URL https://proceedings.neurips.cc/paper_files/paper/2024/file/5d413e48f84dc61244b6be550f1cd8f5-Paper-Datasets_and_Benchmarks_Track.pdf.

- 396 [28] Ran Xu, Yuchen Zhuang, Yue Yu, Haoyu Wang, Wenqi Shi, and Carl Yang. Rag in the wild: On
397 the (in)effectiveness of llms with mixture-of-knowledge retrieval augmentation. *arXiv preprint*
398 *arXiv:2507.20059*, 2025.
- 399 [29] Wujiang Xu, Zujie Liang, Kai Mei, Hang Gao, Juntao Tan, and Yongfeng Zhang. A-mem:
400 Agentic memory for llm agents. *arXiv preprint arXiv:2502.12110*, 2025.
- 401 [30] Wujiang Xu, Kai Mei, Hang Gao, Juntao Tan, Zujie Liang, and Yongfeng Zhang. A-mem:
402 Agentic memory for llm agents, 2025. URL <https://arxiv.org/abs/2502.12110>.
- 403 [31] Hui Yang, Sifu Yue, and Yunzhong He. Auto-gpt for online decision making: Benchmarks and
404 additional opinions. *arXiv preprint arXiv:2306.02224*, 2023. doi: 10.48550/arXiv.2306.02224.
405 URL <https://doi.org/10.48550/arXiv.2306.02224>.
- 406 [32] John Yang, Carlos E Jimenez, Alexander Wettig, Kilian Lieret, Shunyu Yao, Karthik
407 Narasimhan, and Ofir Press. Swe-agent: Agent-computer interfaces enable automated software
408 engineering. *Advances in Neural Information Processing Systems*, 37:50528–50652, 2024.
- 409 [33] Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan
410 Cao. React: Synergizing reasoning and acting in language models. In *Proceedings of the*
411 *11th International Conference on Learning Representations (ICLR 2023)*, 2023. URL https://openreview.net/forum?id=WE_vluYUL-X.
412
- 413 [34] Shunyu Yao, Noah Shinn, Pedram Razavi, and Karthik Narasimhan. τ -bench: A bench-
414 mark for tool-agent-user interaction in real-world domains. In *NeurIPS (Workshops)*, 2024.
415 State-of-the-art agents (e.g. GPT-4o) succeed on <50
- 416 [35] Yuyan Zhou, Liang Song, Bingning Wang, and Weipeng Chen. Metagpt: Merging large
417 language models using model exclusive task arithmetic. In *Proceedings of the 2024 Conference*
418 *on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1711–1724, Miami,
419 Florida, USA, 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.
420 emnlp-main.102.

421 A Appendix

422 A.1 Details of the BREW Algorithm

423 We provide pseudocode for the core components of BREW, aligning with the stages introduced in
424 Section 3. Each algorithm plays a distinct role in constructing, organizing, or refining the knowledge
425 base over iterative interactions. GENERATEINSIGHTS (Alg. 2) produces concept-aligned insights
426 from annotated rollouts using ReflAgent. DEDUPLICATECONCEPTS (Alg. 3) clusters semantically
427 overlapping concepts into a compact meta-concept set. INTEGAGENT incrementally builds and
428 updates per-concept documents using newly generated insights. Finally, EXPANDNODE (Alg. 4)
429 performs MCTS-guided expansions to explore improved document variants, while EVALUATE (Alg. 5)
430 scores candidate KB states using correctness and retrieval-based rewards.

431 We specify the IntegAgent prompt below:

432 BREW Integrator Prompt

```
433 # Enhanced Documentation Editor Prompt
434
435 You are a meticulous documentation-level editor specializing in comprehensive
436 technical reference materials. You will be given a list of topic nodes,
437 each containing structured information that must be preserved and enhanced
438 with maximum detail retention.
439
440 ## Input Structure Analysis
441 Each node contains:
442 - **Title**: The primary topic identifier
443 - **Context**: Background information and conceptual foundation
444 - **How to Use**: Step-by-step instructions, commands, flags, parameters, and
445   implementation details
446 - **When to Use**: Specific scenarios, conditions, and decision criteria
447 - **Best Practices**: Expert recommendations, optimization techniques, and
448   common pitfalls to avoid
449
450 ## Detailed Processing Requirements
451
452 ### 1. Information Preservation (Zero Loss Policy)
453 - **Preserve every technical detail**: All command-line flags, parameter values,
454   configuration options, file paths, URLs, version numbers, and exact syntax
455 - **Maintain all examples**: Keep every code snippet, sample input/output, file
456   names, directory structures, and command sequences exactly as provided
457 - **Retain contextual nuances**: Preserve qualifying language like "typically,"
458   "usually," "in most cases," "when available," and conditional statements
459 - **Keep quantitative data**: Preserve all numbers, measurements, timeframes,
460   limits, thresholds, and statistical information
461 - **Maintain cross-references**: Keep all mentions of related tools,
462   dependencies, prerequisites, and interconnected concepts
463
464 ### 2. Enhanced Detail Extraction
465 - **Expand abbreviations**: When encountering shortened forms, expand them
466   naturally while preserving the original
467 - **Surface implicit knowledge**: Make obvious assumptions explicit (e.g., "this
468   requires root permissions," "assumes default configuration")
469 - **Clarify relationships**: Explicitly describe how different components,
470   options, or steps relate to each other
471 - **Highlight edge cases**: Emphasize special conditions, exceptions, or unusual
472   scenarios mentioned in the source
473 - **Elaborate on consequences**: When the source mentions outcomes, expand on
474   both success and failure scenarios
475
476 ### 3. Prose Transformation Guidelines
477 - **Bullet integration**: Transform each bullet point into 1-3 complete
478   sentences that naturally flow together
479
```

```

480 - **Technical precision**: Use precise technical vocabulary while maintaining
481     readability
482 - **Logical flow**: Organize information within each section to follow a logical
483     sequence (setup → execution → verification)
484 - **Contextual embedding**: Weave code snippets and technical terms seamlessly
485     into narrative sentences
486 - **Comprehensive coverage**: Ensure every sub-bullet, nested item, and
487     parenthetical note becomes part of the prose
488
489 ### 4. Structural Requirements
490 - **Heading hierarchy**: Use '# Title' for each node's main heading
491 - **Section order**: Maintain Context → How to Use → When to Use → Best
492     Practices sequence
493 - **Paragraph organization**: Create substantial paragraphs (3-6 sentences)
494     rather than brief statements
495 - **Transition quality**: Craft smooth bridges between sections and between
496     different nodes
497 - **Code formatting**: Preserve all inline code with backticks and maintain
498     proper formatting for code blocks
499
500 ### 5. Quality Assurance Checklist
501 Before finalizing, verify:
502 - [ ] Every piece of source information appears in the output
503 - [ ] All technical specifications, parameters, and examples are intact
504 - [ ] Code snippets maintain their exact syntax and formatting
505 - [ ] Prose flows naturally without choppy or fragmented sentences
506 - [ ] Each section provides comprehensive coverage of its topic area
507 - [ ] Cross-references and dependencies are clearly explained
508 - [ ] No section labels or formatting artifacts remain in the prose
509
510 ## Output Specifications
511 Generate a single, cohesive markdown document that reads as authoritative
512     technical documentation. The result should be comprehensive enough that a
513     reader could successfully implement the described tools or techniques using
514     only the information provided, without referring back to the original
515     nodes.
516
517 ---
518
519 **Input Nodes**
520 <NODES>
521 {node_list}
522 </NODES>
523
524 ---
525
526 Now, produce the aggregated markdown reference sheet with maximum detail
527     preservation and enhanced clarity.

```

Algorithm 2 GenerateInsights: Extract behavioral insights from trajectories

Require: Queries $\mathcal{Q}$, KB $\mathcal{D}$, rubrics

Ensure: Concept-insight pairs $\mathcal{B}$

```

1:  $\mathcal{B} \leftarrow \emptyset$ 
2: for each query  $q \in \mathcal{Q}$  do
3:    $\tau \leftarrow \text{LLM}(q, \mathcal{D})$  ▷ Generate trajectory
4:   label  $\leftarrow \text{GRADE}(\tau)$  ▷ Success/failure
5:    $(c, i) \leftarrow \text{REFLAGENT}(\tau, \text{rubrics}, \text{label})$ 
6:    $\mathcal{B} \leftarrow \mathcal{B} \cup \{(c, i, q)\}$  ▷ Store with source query
7: end for
8: return  $\mathcal{B}$ 

```

Algorithm 3 DeduplicateConcepts: Cluster similar concepts and map queries

Require: Concept-insight-query triples $\mathcal{B}$ **Ensure:** Meta-concepts $\mathcal{K}$ with mapped queries and insights

- 1: Extract all concepts from $\mathcal{B}$
 - 2: Embed and cluster concepts by similarity
 - 3: $\mathcal{K} \leftarrow$ cluster representatives
 - 4: **for** each $k \in \mathcal{K}$ **do**
 - 5: $\mathcal{Q}_k^{\text{train}} \leftarrow$ {training queries that contributed insights to k }
 - 6: $\mathcal{Q}_k^{\text{eval}} \leftarrow$ {held-out queries relevant to k }
 - 7: $\mathcal{I}_k \leftarrow$ {all insights mapped to concept k }
 - 8: **end for**
 - 9: **return** $\mathcal{K}$ with associated queries and insights
-

Algorithm 4 ExpandNode: Generate and evaluate new document variants

Require: Node s , concept k , candidates h , current KB $\mathcal{D}_{\text{current}}$, best docs $\mathcal{D}_{\text{best}}$, tree**Ensure:** Updated tree with new evaluated nodes

- 1: ▷ Generate new insights from concept-relevant queries
 - 2: $\mathcal{B}_{\text{new}} \leftarrow \emptyset$
 - 3: **for** query $q \in \mathcal{Q}_k^{\text{train}}$ **do**
 - 4: $\tau \leftarrow \text{LLM}(q, \mathcal{D}_{\text{current}})$
 - 5: $(c, i) \leftarrow \text{ANNOTATE}(\tau, \text{rubrics}, \cdot)$
 - 6: **if** c maps to k **then**
 - 7: $\mathcal{B}_{\text{new}} \leftarrow \mathcal{B}_{\text{new}} \cup \{i\}$
 - 8: **end if**
 - 9: **end for**
 - 10: ▷ Generate and evaluate candidate documents
 - 11: **for** $j = 1$ to h **do**
 - 12: $d_{k,j} \leftarrow \text{INTEGAGENT}(k, \mathcal{I}_k \cup \mathcal{B}_{\text{new}}, d_k^s)$
 - 13: ▷ Evaluate using hybrid KB with best docs from other concepts
 - 14: $\mathcal{D}_{\text{hybrid}} \leftarrow \{d_{k,j}\} \cup \{d_{k'} \in \mathcal{D}_{\text{best}} : k' \neq k\}$
 - 15: $R_{k,j} \leftarrow \text{EVALUATE}(d_{k,j}, \mathcal{D}_{\text{hybrid}}, \mathcal{Q}_k^{\text{eval}})$
 - 16: ▷ Add to tree and backpropagate
 - 17: Add $(d_{k,j}, R_{k,j})$ as child of s in tree
 - 18: Backpropagate $R_{k,j}$ from new node to root
 - 19: **end for**
-

Algorithm 5 Evaluate: Score document using held-out queries

Require: Document d_k , hybrid KB $\mathcal{D}_{\text{hybrid}}$, eval queries $\mathcal{Q}_k^{\text{eval}}$ **Ensure:** Reward score R

- 1: $R^{\text{corr}} \leftarrow 0$
 - 2: $R^{\text{ret}} \leftarrow 0$
 - 3: **for** each $q \in \mathcal{Q}_k^{\text{eval}}$ **do**
 - 4: $R^{\text{corr}} \leftarrow R^{\text{corr}} + \text{EVAL}(q, \text{agent} \oplus \mathcal{D}_{\text{hybrid}})$
 - 5: $R^{\text{ret}} \leftarrow R^{\text{ret}} + \text{MRR}(d_k, q, \mathcal{D}_{\text{hybrid}})$
 - 6: **end for**
 - 7: $R^{\text{corr}} \leftarrow \frac{R^{\text{corr}}}{|\mathcal{Q}_k^{\text{eval}}|}$
 - 8: $R^{\text{ret}} \leftarrow \frac{R^{\text{ret}}}{|\mathcal{Q}_k^{\text{eval}}|}$
 - 9: **return** $\lambda_{\text{corr}} \cdot R^{\text{corr}} + \lambda_{\text{ret}} \cdot R^{\text{ret}}$
-

529 **A.2 BREW Configurations**

530 **Base LLM Configuration** For all BREW algorithm steps, we use the OpenAI GPT-4.1-2025-04-14
531 model as the underlying language model. To balance exploration and stability, we set the temperature

532 to 0.7 for the IntegAgent component to encourage diversity in sampled completions, while all other
 533 calls use a temperature of 0.1 for deterministic behavior. The search process employs an expansion
 534 width of $e = 3$, a maximum search depth of $k = 3$, and a maximum of $n = 10$ iterations. Reward
 535 signals are weighted equally across correctness and retrieval relevance, with $\lambda_{corr} = \lambda_{ret} = 0.5$.

536 A.3 Baseline Methods

537 We compare BREW against two common reasoning baselines. Step-Back Prompting encourages
 538 backward reasoning by guiding the model to work from the final task objective back to the initial
 539 actions. In-Context Learning augments the input prompt with successful trajectories from related
 540 tasks, enabling the model to benefit from relevant prior examples without additional fine-tuning.

541 A.4 Benchmark Specifications

542 A.4.1 OSWorld: Computer-Use Automation

543 **Dataset Overview** OSWorld [27] comprises 369 real-world computer-use tasks spanning 10 distinct
 544 applications. The benchmark is divided into train and test sets, with the distribution of tasks across
 545 domains shown in Table 2.

546 **Agent Specifications** The UI-Tars-7B variant is a 7B-parameter multimodal transformer fine-tuned
 547 for graphical user interface understanding. It operates over an action space of PyAutoGUI commands
 548 (e.g., click, type, and key presses). The agent integrates a retrieval module that queries a task-relevant
 549 knowledge base using the user-provided description, with the top three retrieved items added to the
 550 system prompt. Inputs to the model consist of a screenshot of the active UI paired with the natural
 551 language task description.

552 The GTA1-7B configuration adopts a two-agent architecture, consisting of a planner and a grounding
 553 module. The planner (GTA-1-7B) generates the high-level action sequence, while the grounding
 554 module (OpenAI O3) verifies and refines each action before execution. Knowledge retrieval is
 555 incorporated differently for each component: the planner performs a single retrieval at the start
 556 of execution, which is persisted in its prompt, whereas the grounding module performs dynamic
 557 retrievals at each verification step.

558 **Evaluation Protocol** Evaluation uses 134 task-specific scripts designed for automated verification.
 559 Success criteria include file state checks (e.g., validating .xlsx or .docx outputs), UI element
 560 validation to confirm correct interaction, and process completion checks to ensure that the intended
 561 automation sequence was executed successfully.

562 A.4.2 τ^2 -Bench: Interactive Tool Usage

563 **Dataset Overview** τ^2 -Bench [5] extends τ -Bench by introducing bidirectional tool-calling capa-
 564 bilities. The dataset covers multiple service-oriented domains, with domain-level task distributions
 565 summarized in Table 3.

566 **Domain Characteristics** The benchmark spans several domains with distinct task characteristics.
 567 The Telecom domain focuses on connectivity troubleshooting, plan modifications, and service
 568 activation workflows. The Retail domain includes order processing, return handling, and inventory
 569 queries. The Airline domain emphasizes booking modifications and policy-compliant rescheduling
 570 scenarios.

571 **Interaction Settings** Two interaction modes are defined. In Easy mode, a human proxy (imple-
 572 mented via GPT-4.1) provides detailed guidance to the agent. The knowledge base is built exclusively
 573 from Easy mode trajectories, ensuring high-quality demonstrations for learning. In Hard mode,
 574 human intervention is minimized. The knowledge base combines both Easy and Hard trajectories,
 575 testing the agent’s robustness to underspecified or noisy instructions.

576 **Evaluation Criteria** Task success is measured using domain-specific verification procedures. These
 577 include database state checks to validate final outcomes, status checks for confirming service or
 578 connection state, natural language verification to ensure correct confirmation statements appear in
 579 dialogue, and action matching to confirm that all required steps are completed. Each domain uses a
 580 tailored subset of these checks (e.g., Telecom relies primarily on status checks).

Domain	Test	Train
Calc	45	2
Chrome	44	2
Writer	21	2
Gimp	24	2
Impress	45	2
Os	22	2
Thunderbird	13	2
Multi-apps	99	2
VLC	15	2
VSCode	21	2
Total	349	20

Table 2: Test and Train samples across different domains in OSWorld.

Domain	Test	Train
Telecom	105	7
Retail	105	7
Airline	44	6
Total	254	20

Table 3: Task-wise breakdown for τ^2 -Bench with assumed 2-shot training samples per domain.

Domain Characteristics

- **Telecom:** Connectivity issues, plan management, service activation
- **Retail:** Order processing, returns, inventory queries
- **Airline:** Booking modifications, policy-compliant rescheduling

Evaluation Criteria Task success determined by:

- **Database Checks:** Final state verification
- **Status Checks:** Service/connection state validation
- **NL Checks:** Confirmation statements in dialogue
- **Action Matching:** Required action sequence completion

Note: Each domain uses specific check combinations (e.g., Telecom uses only status checks).

A.4.3 SpreadsheetBench: Real-World Spreadsheet Manipulation

Dataset Overview SpreadsheetBench [18] consists of 912 instructions collected from four major Excel forums and blogs. Each instruction is paired with spreadsheets reflecting authentic, complex user scenarios, often containing multiple tables and non-standard relational structures. The dataset totals 2,729 test cases, averaging three per instruction. A breakdown of cell-level and sheet-level manipulations is shown in Table 4.

Task Settings The benchmark defines two dimensions of evaluation:

- **Granularity:** Instructions involve either *cell-level* manipulations (specific ranges such as D2:D6) or *sheet-level* manipulations (entire tables or multi-sheet updates).
- **Evaluation:** Performance is measured using an Online Judge (OJ)-style protocol. The *soft* setting (IOI-style) awards partial credit when only some test cases are solved, while the *hard* setting (ICPC-style) requires solutions to succeed on all test cases.

Agent Configuration We evaluate

textttto4-mini using a function-calling agent connected to a single Python execution tool. The

agent translates natural language instructions into Python code for spreadsheet manipulation (e.g., modifying cells, applying formulas, restructuring tables). After each tool call, all formulas in the spreadsheet are recalculated to ensure consistency before proceeding to the next step. This setup provides a controlled environment to assess reasoning, code generation, and execution robustness across diverse spreadsheet tasks.

Granularity	Instructions	Test Cases
Cell-Level	329	986
Sheet-Level	583	1,743
Total	912	2,729

Table 4: Cell-level vs. sheet-level distribution in SpreadsheetBench.

A.5 KB Construction and Retrieval Details

Training Data Collection

- **OSWorld:** 20 successful trajectories (2 per application domain) and 10 for evals.
- τ^2 -**Bench:** 20 trajectories balanced across domains and difficulty settings and 10 for evals.
- **SpreadsheetBench:** Uniformly sample 30 trajectories for training and 10 for evaluation.

All numbers are reported on the remaining train set.

Retrieval Strategy

- **Query Formation:** For each task we take in the seed Natural Language query as the retrieval query.
- **Retrieval Count:** We take top-3 documents for all the retrieval steps
- **Integration Point:** For SPREADSHEET ENCH and **OSWorld** we insert retrievals in the system prompt augmentation. For τ^2 -bench we add perfrom retrieval after each user interaction.

B Qualitative Analysis

Exploration on MCTS parameters WE evaluate OSworld on two different MCTS parameters.

- **Increased Depth:** To increase the depth we keep maximum width of the tree as 3 and depth as 10 with max number of iterations as 25. We observe that the Knowledge base over optimizes on the train set leading to a poorer performance on test set.
- **Increased Width:** For increased width we reverse the parameters where depth is 3 and maximum width is 10 with max iterations 25. We observe many different styles of KBs are generated storing very similar information, these different styles lead to a varied performance on both eval and test set notifying the importance of state search.

We report the numbers on table ??

	Baseline	max_width=3, max_depth=3	max_width=3, max_depth=10	max_width=10, max_depth=3
OSworld	44.20	47.56	43.83	49.32

Table 5: OSworld difference in MCTS parameters

633 C Exemplar Knowledge Bases

634 C.1 Knowledge base learned for OSWorld

635 We showcase a small part of knowledge base learned through BREW . This demonstrates 3 major
636 parts on which each document is aggregated. These parts discuss when to use a piece of information,
637 why to use the information, how to use the information/tool.

```
638 ## Search and Open Files
639
640
641 **When to use**: Locating documents, spreadsheets, images, or downloads for
642 editing, conversion, or attachment.
643
644 ### How to Perform
645 - Open **File Manager (Nautilus)** from launcher or system dock
646 - Press 'Ctrl + F' or click the search icon
647 - Enter part of filename, full name, or wildcard (*.pdf', 'report*')
648 - Use right-click →**Open With** to choose the desired application
649 - Use the sidebar to navigate to **Downloads**, **Documents**, or custom folders
650
651 ### Additional Actions
652 - Right-click →**Properties** to check modification date or file type
653 - Sort results by Date, Type, or Name from the top-right dropdown
654 - Use 'F2' to rename files inline
655
656 ### Example
657 - Task: "Edit the file titled 'sales_report_march.ods'"
658   - Search for 'sales' in File Manager
659   - Confirm '.ods' type and open with LibreOffice Calc
660
661 ...
662
663 ## Insert Images
664
665 **When to use**: Adding visual elements to documents, presentations, emails, or
666 templates.
667
668 ### How to Perform
669 - Navigate to **Insert →Image →From File** (in Writer, Impress, Thunderbird)
670 - Select an image file ('.png', '.jpg', '.svg') from the file dialog
671 - Use drag handles to resize; right-click →**Wrap** or **Alignment** for layout
672
673 ### Additional Actions
674 - In GIMP: **File →Open as Layers** to insert image as a new layer
675 - Use drag-and-drop from file manager into open document windows
676 - Use **Format →Image** to apply borders, shadows, or color corrections (in
677 Writer/Impress)
678
679 ### Example
680 - Task: "Insert the logo.png image into the title slide"
681   - Open '.odp' file in Impress →Go to Slide 1 →Insert →Image →Select 'logo.
682     png'
683
684 ...
685
686 ## Export as PDF
687
688 **When to use**: Required submission format
689
690 ### How to Perform
691 - Go to **File →Export As PDF**
692 - Choose output folder (usually **Documents** or **Downloads**)
693 - Click **Save**, then confirm the exported file opens correctly
694
```

```

695   ### Additional Actions
696   - In GIMP or Impress: choose **File →Export As**, then select ‘.pdf’ from
697     format list
698   - Use **Save As** to preserve both editable and exported versions separately
699
700   ### Example
701   - Task: "Export the flyer.xcf as a PDF"
702     - Open in GIMP →File →Export As →Rename to ‘flyer.pdf’ →Click Export
703

```

704 C.2 BREW Knowledge Base for τ^2 -Bench

705 BREW enable use to learn relevant information for tau bench for across the domains in a single
706 knowledge base. This knowledge base is helpful to use relevant actions from the action pool.

```

707
708   ### Additional Actions
709
710   * Inform the user:
711     - Refunds via gift card = immediate.
712     - Refunds via other methods = -57 business days.
713
714   ### Example
715
716   * Task: "Cancel a T-shirt order placed yesterday"
717     * Validate: Status is ‘pending’
718     * Reason: "no longer needed"
719     * Confirm
720     * Execute tool call
721
722
723   # Exchange Delivered Order
724
725   **When to use**:
726   User wants to swap delivered items for a different variant (e.g., size or color)
727   .
728
729   **Why to use it**:
730   To fix sizing or option errors without needing a new purchase.
731
732   ### How to Perform
733   - Authenticate user
734   - Confirm order status is ‘delivered’
735   - Get full list of exchange items
736   > "Please ensure all items for exchange are listed. This step 'cant be repeated.
737     "
738   - Ask for refund/payment method
739   - Confirm:
740     > "'Youre exchanging item X for same product, different option. Proceed?"
741   - On confirmation:
742     ‘python
743     request_exchange(order_id="45678", item_exchanges=[...], payment_method="
744       paypal")
745     ‘‘‘
746
747   ### Additional Actions
748
749   * Mention: An email will be sent with return instructions
750   * Validate that the new variant is from the same product
751
752   ### Example
753
754   * Task: "Exchange red shirt for blue in Order #45678"
755     * Confirm all exchange items
756     * Confirm payment method for difference
757

```

```

758     * Execute tool call
759
760     ### Example
761
762     * Task: "Show me my last 2 orders"
763     * Authenticate
764     * Retrieve and present info
765
766     # Deny Unsupported Request
767
768     **When to use**:
769     User asks for an unsupported action (e.g., cancel processed order, exchange to
770     different product type, help another user).
771
772     **Why to use it**:
773     To stay compliant with platform policy.
774
775     ### How to Perform
776     - Politely reject:
777     > "Im sorry, but I 'cant process that request. 'Its outside the allowed scope.
778         "
779
780     ### Example
781
782     * Task: "Cancel a processed order"
783     * Respond with denial message
784     # Transfer to Human Agent
785
786     **When to use**:
787     User needs help outside the 'assistants permitted capabilities.
788
789     **Why to use it**:
790     To ensure user gets the right help from trained staff.
791
792     ### How to Perform
793     - Make tool call:
794     ```python
795     transfer_to_human_agents()
796     ```
797     - Then inform user:
798     > "YOU ARE BEING TRANSFERRED TO A HUMAN AGENT. PLEASE HOLD ON."
799
800     ### Example
801
802     * Task: "Delete a task"
803     * Deny deletion
804     * Transfer to human
805

```

806 C.3 BREW Knowledge Base for SpreadsheetBench

```

807     Header Extraction
808     1. Detecting Header Rows
809     Overview:
810     To accurately identify header rows, scan the initial region of your dataset.
811     This process is crucial for mapping column information for further
812     processing.
813
814     Approaches:
815     - Heuristic Checks:
816     - Look for rows where all cells are strings (e.g., "Name", "Date", "Region", "
817     Amount").
818     - Identify rows with distinctive formatting such as bold text or background
819     color.
820     - Example:
821

```

```

822 | Name | Date | Region | Amount | |-----|-----|-----|-----| |
823 | John | 2024-01-01 | North | 100 |
824 - Pattern Recognition:
825 - Use regex to match typical header patterns, such as column names starting with
826 uppercase letters.
827 - Score candidate rows based on the likelihood of being headers.
828 - Multi-Table Sheets:
829 - Detect gaps, empty rows, or separators indicating a new table.
830 - Assign a Table ID to each detected table for later reference.
831
832 Edge Cases:
833 - Merge multi-row headers (e.g., "Sales" over "2024", "2025" becomes "Sales 2024
834 ", "Sales 2025").
835 - Fill in missing headers by inferring from context.
836
837 2. Assigning and Validating Headers
838 Overview:
839 Once headers are detected, assign them programmatically and ensure they match
840 expected schema and data types.
841
842 Implementation:
843 - Column Naming:
844 - Set names in code, e.g., df.columns = ["Name", "Date", "Region", "Amount"].
845 - Schema Mapping:
846 - Map headers to a standardized schema, using external files or user prompts.
847 - Example:
848 - Raw header: "Amt"; Mapped header: "Amount"
849 - Quality Checks:
850 - Detect duplicate or empty headers ("Date", "Date" becomes "Date_1", "Date_2").
851 - Validate each column's expected data type.
852
853 3. Automation and Usability Enhancements
854 Overview:
855 Enhance usability and automation to streamline header extraction and user
856 interaction.
857
858 Features:
859 - Freeze Panes:
860 - Automatically freeze header rows in Excel for easier navigation.
861 - Highlighting:
862 - Use colored formatting to visually distinguish headers.
863 - Example:
864 - Yellow fill for header row.
865 - Documentation:
866 - Log extraction logic and confidence scores for each detected header.
867 - Integration:
868 - Build header extraction into ETL pipelines and record process metadata.
869
870 Block Detection
871 1. Identifying Block Boundaries
872 Overview:
873 Block detection segments data into logical units or tables.
874
875 Methods:
876 - Boundary Detection:
877 - Find empty rows, repeated labels, or formatting changes.
878 - Example:
879 | Name | Amount | |-----|-----| | John | 100 | | | <-- Empty row indicates
880 new block | Name | Amount | | Alice | 200 |
881 - Machine Learning:
882 - Train classifiers to detect block boundaries based on cell patterns.
883
884 Advanced:
885 - Detect nested blocks or hierarchies using indentation or merged cells.
886 - Identify summary blocks with keywords like "Total" or "Summary".

```

887
888 2. Processing and Tracking Blocks
889 Overview:
890 Once blocks are detected, assign IDs and enable block-level analysis.
891
892 Actions:
893 - Block ID:
894 - Assign unique IDs (e.g., Block_001, Block_002).
895 - Analysis:
896 - Perform group-by or aggregation within each block.
897 - Example:
898 - Sum "Amount" for Block_001: 100 + 150 = 250
899
900 3. Additional Block Actions
901 Overview:
902 Enable modular analysis and reporting at the block level.
903
904 Features:
905 - Summary Rows:
906 - Add computed totals/averages for each block.
907 - Export/Save:
908 - Save blocks as separate files or sheets.
909 - Example:
910 - Export Block_001 to "block1.csv"
911
912 Search for Values or Patterns
913 1. Search Execution Methods
914 Overview:
915 Efficiently locate specific values or patterns in your data.
916
917 Techniques:
918 - Manual Tools:
919 - Use Ctrl + F in Excel for quick lookups.
920 - Programmatic Search:
921 - Scan all cells using loops or vectorized code.
922 - Example:
923 - Find all instances of "North" in the "Region" column.
924 - Pattern Matching:
925 - Support exact, wildcard (*Total*), and regex (\d{4}-\d{2}-\d{2} for dates).
926
927 2. Recording and Highlighting Results
928 Overview:
929 Log and visualize search matches for user review.
930
931 Actions:
932 - Logging:
933 - Record coordinates (e.g., Sheet1, Row 3, Col "Region").
934 - Highlighting:
935 - Apply conditional formatting to search hits.
936
937 3. Advanced Search Scenarios
938 Overview:
939 Handle complex or large-scale search requirements.
940
941 Scenarios:
942 - Merged Cells:
943 - Search within merged cells or across multiple sheets.
944 - Export:
945 - Export found results for further analysis.
946 - Example:
947 - Export all rows containing "John" to "john_results.csv"
948
949 Writeback Results
950 1. Output Placement
951 Overview:

952 Choose where and how to insert results.

953

954 Options:

955 - Target Columns:

956 - Select existing or blank columns for output.

957 - Appending:

958 - Add new columns for flags, counts, or statuses.

959 - Example:

960 - Add "Approved_Flag" column next to "Status".

961

962 2. Writing and Styling Results

963 Overview:

964 Automate and style the output for visibility.

965

966 Methods:

967 - Formulas/Code:

968 - Use code (e.g., `ws.cell(row, col).value = result`) to insert results.

969 - Styling:

970 - Bold, borders, or colors for output cells.

971 - Example:

972 - Green fill for "Success", red for "Error".

973

974 3. Audit and Protection

975 Overview:

976 Maintain the integrity and traceability of results.

977

978 Measures:

979 - Lock Columns:

980 - Prevent edits to output columns.

981 - Timestamps/User Info:

982 - Add audit trail for writebacks.

983 - Example:

984 - "2024-06-01, User: admin"

985

986 Difference in State

987 1. Sheet Comparison

988 Overview:

989 Identify changes between input and output sheets.

990

991 Process:

992 - Load Sheets:

993 - Read both sheets into memory.

994 - Compare Cells:

995 - Detect differences by position and value.

996

997 2. Recording and Reporting Differences

998 Overview:

999 Log and report all detected changes.

1000

1001 Actions:

1002 - Log Mismatches:

1003 - Record cell coordinates and values.

1004 - Example:

1005 - Cell B3: "North" → "South"

1006 - Export Diff Report:

1007 - List all detected differences for review.

1008

1009 3. Visualization and Automation

1010 Overview:

1011 Make changes visible and automate validation.

1012

1013 Features:

1014 - Highlight Changes:

1015 - Color code changed cells.

1016 - Automate Checks:


```

1017 - Integrate diff comparisons into test scripts.
1018
1019 Column Selection
1020 1. Selection Criteria
1021 Overview:
1022 Choose relevant columns for analysis.
1023
1024 Methods:
1025 - Labels/Indices:
1026 - Select by name or position.
1027 - Dynamic Rules:
1028 - E.g., all numeric columns.
1029 - Assign Roles:
1030 - Example: "ID", "Date", "Metric"
1031
1032 2. Preparation and Validation
1033 Overview:
1034 Prepare columns for consistent use.
1035
1036 Actions:
1037 - Rename/Relabel:
1038 - Standardize column names.
1039 - Validate Types:
1040 - Ensure columns are of expected type.
1041 - Example:
1042 - "Date" column as datetime.
1043
1044 3. Reusability
1045 Overview:
1046 Save and reuse column selections.
1047
1048 Features:
1049 - Presets:
1050 - Save selection profiles.
1051 - Downstream Use:
1052 - Use validated columns in subsequent processes.
1053
1054 Filter Rows
1055 1. Filtering Methods
1056 Overview:
1057 Refine your dataset with filters.
1058
1059 Techniques:
1060 - Spreadsheet Tools:
1061 - Use built-in filters.
1062 - Code Logic:
1063 - Filter with code (e.g., df[df['Status'] == 'Approved']).
1064 - Multiple Criteria:
1065 - Combine conditions (AND/OR).
1066 - Example:
1067 - Status = "Approved" AND Amount > 100
1068
1069 2. Helper Columns and Complex Filters
1070 Overview:
1071 Simplify filtering using helper columns.
1072
1073 Actions:
1074 - Helper Columns:
1075 - Compute intermediate flags.
1076 - Document Logic:
1077 - Record filtering rules for audit.
1078
1079 3. Post-Filter Actions
1080 Overview:
1081 Visualize and export filtered data.

```

```

1082
1083 Features:
1084 - Highlighting:
1085 - Grey-out filtered-out rows.
1086 - Export:
1087 - Save the filtered dataset.
1088
1089 Merge Tables
1090 1. Key-Based Merging
1091 Overview:
1092 Combine tables using shared keys.
1093
1094 Techniques:
1095 - Join Operations:
1096 - Use VLOOKUP, JOIN, or code merges.
1097 - Example:
1098 - Merge "Customer_ID" from two tables.
1099 - Align Data:
1100 - Match on columns like "ID", "Name".
1101
1102 2. Stack-Based Merging
1103 Overview:
1104 Append tables when keys 'arent needed.
1105
1106 Methods:
1107 - Vertical Append:
1108 - Combine rows from similar tables.
1109 - Deduplicate:
1110 - Remove duplicate records.
1111
1112 3. Tracking and Audit
1113 Overview:
1114 Track source and unmatched records.
1115
1116 Actions:
1117 - Source Column:
1118 - Add "Source" to indicate origin.
1119 - Highlight Unmatched:
1120 - Mark or export mismatched rows.
1121
1122 Pivot or Unpivot
1123 1. Pivoting Data
1124 Overview:
1125 Summarize data using pivots.
1126
1127 Methods:
1128 - PivotTables:
1129 - Group by row/column dimensions.
1130 - Example:
1131 - Sum "Amount" by "Region".
1132 - Aggregation:
1133 - Choose SUM, AVG, COUNT, etc.
1134
1135 2. Unpivoting (Melting) Data
1136 Overview:
1137 Reshape data from wide to long format.
1138
1139 Techniques:
1140 - Melt Operations:
1141 - Convert columns into rows.
1142 - Example:
1143 -
1144 | Year | Sales_2019 | Sales_2020 | |-----|-----|-----|
1145 →
1146 | Year | Sales_Year | Value |

```

1147 - Flexible Restructuring:
1148 - Selectively unpivot non-ID columns.
1149
1150 3. Post-Pivot Actions
1151 Overview:
1152 Prepare pivoted data for export.
1153
1154 Features:
1155 - Flatten Pivot Table:
1156 - Convert back to flat for further analysis.
1157 - Reorder/Rename:
1158 - Clarify pivoted fields.
1159
1160 Map with Lookup Tables
1161 1. Mapping Techniques
1162 Overview:
1163 Standardize data using lookups.
1164
1165 Methods:
1166 - Functions:
1167 - Use VLOOKUP, merge with dictionaries.
1168 - Code-to-Label:
1169 - Example:
1170 - Code "N" →Label "North"
1171
1172 2. Application and Fallbacks
1173 Overview:
1174 Apply lookups and handle missing values.
1175
1176 Actions:
1177 - Apply Mappings:
1178 - Across selected columns.
1179 - Handle Missings:
1180 - Use defaults for missing codes.
1181
1182 3. Audit and Display
1183 Overview:
1184 Ensure mapping transparency.
1185
1186 Features:
1187 - Cache Mappings:
1188 - Store for repeated use.
1189 - Display Codes/Labels:
1190 - Show both for clarity.
1191
1192 Fill Missing Data
1193 1. Choosing Fill Methods
1194 Overview:
1195 Impute missing data appropriately.
1196
1197 Techniques:
1198 - Forward/Backward Fill:
1199 - Fill gaps with prior/next value.
1200 - Default Values:
1201 - Use fixed placeholder (e.g., 0, "Unknown").
1202 - Contextual Example:
1203 - Dates: Fill missing month with last known month.
1204
1205 2. Application and Auditing
1206 Overview:
1207 Apply fills and flag for review.
1208
1209 Actions:
1210 - Targeted Filling:
1211 - Apply to specific columns/rows.

1212 - Flag Filled Cells:
1213 - Highlight for later review.
1214
1215 3. Documentation
1216 Overview:
1217 Keep fill logic transparent.
1218
1219 Features:
1220 - Record Logic:
1221 - Document assumptions and methods.
1222 - Audit Trail:
1223 - Track all changes.
1224
1225 Flag Rows or Cells
1226 1. Defining Flag Rules
1227 Overview:
1228 Establish criteria for flagging.
1229
1230 Examples:
1231 - Simple Rule:
1232 - Flag where Amount < 0
1233 - Complex Rule:
1234 - Flag where Status = "Pending" and Amount > 1000
1235
1236 2. Applying Flags
1237 Overview:
1238 Insert flags and summarize.
1239
1240 Actions:
1241 - Flag Column:
1242 - Add "Flag" column with "Yes"/"No".
1243 - Export Flagged Rows:
1244 - Save for further inspection.
1245
1246 3. Advanced Flagging
1247 Overview:
1248 Use multiple criteria and document.
1249
1250 Features:
1251 - Multi-Criteria:
1252 - Combine several rules for granular checks.
1253 - Notes:
1254 - Document flagging rationale.
1255
1256 Sort Data
1257 1. Setting Sort Criteria
1258 Overview:
1259 Organize data for analysis.
1260
1261 Options:
1262 - Sort Columns:
1263 - By value, ascending/descending.
1264 - Multi-Level:
1265 - E.g., sort by "Region", then by "Amount".
1266
1267 2. Applying Sorts
1268 Overview:
1269 Implement sorting programmatically or manually.
1270
1271 Methods:
1272 - Spreadsheet Tools:
1273 - Built-in sort features.
1274 - Code:
1275 - E.g., df.sort_values(['Region', 'Amount'])
1276

1277 3. Post-Sort Actions
1278 Overview:
1279 Finalize sorted data.
1280
1281 Actions:
1282 - Renumber Rows:
1283 - Update indices.
1284 - Highlight Extremes:
1285 - Mark top/bottom values.
1286
1287 Validate Data
1288 1. Validation Checks
1289 Overview:
1290 Ensure data meets required standards.
1291
1292 Checks:
1293 - Type:
1294 - Ensure numeric columns contain numbers.
1295 - Range:
1296 - E.g., "Amount" > 0.
1297 - Pattern:
1298 - Date columns match YYYY-MM-DD.
1299 - Business Rule Example:
1300 - "Start Date" < "End Date"
1301
1302 2. Marking and Reporting
1303 Overview:
1304 Visualize and report errors.
1305
1306 Actions:
1307 - Highlight Invalids:
1308 - Color-code errors.
1309 - Export Summary:
1310 - Table of error counts and locations.
1311
1312 3. Integration in Workflow
1313 Overview:
1314 Make validation a routine part of processing.
1315
1316 Features:
1317 - Pre-Processing Step:
1318 - Validate before analysis.
1319 - Automation:
1320 - Integrate into data pipelines.
1321
1322 Split Sheets or Data
1323 1. Defining Split Rules
1324 Overview:
1325 Segment data for modular analysis.
1326
1327 Methods:
1328 - By Category:
1329 - E.g., split by "Region".
1330 - By Date Range:
1331 - E.g., split by year.
1332
1333 2. Exporting Segments
1334 Overview:
1335 Save segments for separate use.
1336
1337 Actions:
1338 - Export Files:
1339 - "North_Region.csv", "South_Region.csv"
1340 - Consistent Formatting:
1341 - Ensure identical columns and styling.

```

1342
1343 3. Automation and Documentation
1344 Overview:
1345 Automate splitting and track provenance.
1346
1347 Features:
1348 - Automation:
1349 - Use scripts/macros for repeated splits.
1350 - Documentation:
1351 - Record rules and export logs.
1352

```

1353 D Qualitative Analysis of BREW-Generated Knowledge Bases

1354 This section presents a comprehensive qualitative analysis of knowledge bases generated through
1355 the BREW technique applied to two distinct agent training environments: OSWorld and τ^2 Bench
1356 described in the section before. The analysis examines knowledge representation patterns, procedural
1357 sophistication, and domain-specific learning characteristics extracted from CUA agent behaviors,
1358 providing insights into the effectiveness and scope of knowledge distillation techniques across diverse
1359 task environments.

1360 D.1 Cross-Domain Knowledge Base Analysis

1361 D.1.1 Base Structure & Organization

1362 **Schema Consistency and Evolution:** Both knowledge bases demonstrate consistent structural
1363 schemas, though adapted to their respective domains. The OSWorld KB employs a four-part
1364 schema (contextual triggers, procedural steps, extended capabilities, concrete instantiation), while the
1365 τ^2 Bench KB extends this to a five-part structure, adding explicit purpose rationale (“Why to use it”).
1366 This evolution suggests that BREW adapts its extraction patterns to domain-specific requirements—
1367 conversational commerce demands explicit justification for actions due to customer interaction
1368 contexts.

1369 **Taxonomic Organization Principles:** The OSWorld KB reveals a **capability-based taxonomy**
1370 organized around computational tasks: file operations, document processing, inter-application work-
1371 flows, and data visualization. Each category represents a distinct computational domain with specific
1372 tool requirements and interaction patterns. In contrast, the τ^2 Bench KB employs a **lifecycle-based**
1373 **taxonomy** structured around transactional states: order creation, modification, fulfillment, and
1374 post-delivery operations. This organizational difference reflects fundamental domain characteristics—
1375 desktop automation focuses on tool orchestration, while conversational commerce centers on process
1376 management.

1377 **Hierarchical Task Decomposition:** Both KBs demonstrate sophisticated hierarchical reasoning, but
1378 through different decomposition strategies. OSWorld exhibits **technical decomposition**, breaking
1379 complex operations like “Create Charts from Data” into constituent technical steps (data selection,
1380 chart insertion, customization, formatting). τ^2 Bench shows **process decomposition**, structuring
1381 operations like order modification into authentication, validation, confirmation, and execution phases.
1382 This suggests BREW successfully identifies domain-appropriate decomposition strategies rather than
1383 applying uniform patterns.

1384 **Knowledge Boundary Definition:** Both KBs explicitly encode operational boundaries, but through
1385 contrasting mechanisms. OSWorld boundaries are **capability-constrained**—determined by available
1386 applications and system resources. τ^2 Bench boundaries are **policy-constrained**—explicitly defined
1387 through “Deny Unsupported Request” patterns and escalation protocols. This difference highlights
1388 how knowledge extraction adapts to domain-specific constraint types.

1389 D.1.2 Procedural Knowledge Grounding

1390 **Context-Dependent Action Selection:** Both domains demonstrate sophisticated context awareness,
1391 but grounded in different environmental factors. OSWorld exhibits **application-context sensitivity**,
1392 where identical operations (e.g., image insertion) require different procedures across LibreOffice
1393 Writer, Impress, GIMP, and Thunderbird. The agent learned application-specific affordances and

1394 interaction patterns rather than generic command sequences. τ^2 Bench demonstrates **state-context**
1395 **sensitivity**, where available actions depend on order status (pending vs. delivered), payment methods,
1396 and authentication levels. This reveals learned understanding of business process constraints and
1397 temporal operation windows.

1398 **Error Prevention and Validation Workflows:** Both KBs incorporate sophisticated error prevention
1399 mechanisms, but grounded in domain-specific failure modes. OSWorld emphasizes **technical valida-**
1400 **tion:** file integrity checks (“confirm the exported file opens correctly”), application state verification,
1401 and multi-step confirmation for irreversible operations. τ^2 Bench emphasizes **transactional valida-**
1402 **tion:** authentication cascades, confirmation dialogues with standardized templates, and explicit user
1403 consent protocols. The emergence of defensive programming practices across both domains suggests
1404 these represent fundamental principles of reliable agent behavior.

1405 **State-Dependent Decision Logic:** The procedural knowledge in both domains demonstrates sophisti-
1406 cated state machine reasoning. OSWorld exhibits **application state awareness**—understanding when
1407 applications are ready for input, when files are loaded, and when operations can be safely executed.
1408 Window management and application switching reveal learned understanding of desktop metaphors
1409 and resource constraints. τ^2 Bench demonstrates **business process state awareness**—finite state
1410 machine reasoning where order lifecycle states determine available operations. The agent learned
1411 that pending orders enable modification while delivered orders unlock return workflows, indicating
1412 internalized understanding of business logic constraints.

1413 **Security and Authentication Grounding:** While OSWorld operates in a trusted desktop environment
1414 with minimal explicit security concerns, τ^2 Bench reveals pervasive **authentication-first paradigms**.
1415 Nearly every transactional operation begins with identity verification through email, name, and zip
1416 code combinations. The KB demonstrates **graduated security reasoning:** information retrieval
1417 requires basic authentication while financial transactions trigger rigorous verification protocols. This
1418 contrast highlights how procedural knowledge adapts to domain-specific security requirements.

1419 **Cross-Application vs. Cross-Process Orchestration:** OSWorld demonstrates **technical orches-**
1420 **tration**—coordinating multiple applications (Chrome, LibreOffice suite, File Manager, GIMP) to
1421 accomplish complex workflows. The “Navigate Between Applications” section reveals learned be-
1422 haviors for window management, application switching, and resource coordination. τ^2 Bench exhibits
1423 **process orchestration**—coordinating authentication, validation, confirmation, and execution phases
1424 across different operational contexts. Both forms of orchestration require sophisticated temporal
1425 reasoning and constraint management, but applied to different environmental complexity types.

1426 **Failure Mode Internalization:** Both KBs reveal learned understanding of domain-specific failure
1427 modes. OSWorld incorporates file validation, application crash recovery suggestions, and verification
1428 steps for critical operations. τ^2 Bench includes explicit escalation protocols (“Transfer to Human
1429 Agent”), policy compliance mechanisms, and irreversibility warnings for financial operations. The
1430 consistent emergence of failure-aware procedures suggests that agents successfully internalize risk
1431 assessment and mitigation strategies during training.

1432 **Domain-Specific Communication Patterns:** The procedural knowledge reveals distinct communica-
1433 tion paradigms appropriate to each domain. OSWorld procedures are **task-oriented** with minimal
1434 user interaction—focusing on efficient command execution and verification. τ^2 Bench procedures are
1435 **dialogue-oriented** with standardized customer interaction templates, confirmation protocols, and
1436 expectation management communications. This adaptation demonstrates that BREW extracts not just
1437 procedural logic but domain-appropriate interaction modalities.

1438 The cross-domain analysis reveals that BREW successfully extracts procedural knowledge that is both
1439 **structurally consistent** (following learnable organizational patterns) and **contextually grounded**
1440 (adapted to domain-specific constraints, failure modes, and interaction requirements). This dual capa-
1441 bility suggests significant potential for knowledge transfer across related domains while maintaining
1442 appropriate domain-specific adaptations.