

Supplementary Information for: Balancing Autonomy and Oversight in Language-Agent-Guided Physical Experimentation

Anderson Borch¹, Emily P.R. Avey^{1,2}, Jayden Grunde^{1,2}, Michelle A. Smeaton¹, Renae N. Gannon¹, and Steven R. Spurgeon^{1,3,4,*}

¹National Laboratory of the Rockies, Golden, CO, USA

²Materials Science and Engineering Program, Colorado School of Mines, Golden, CO, USA

³Metallurgical and Materials Engineering Department, Colorado School of Mines, Golden, CO, USA

⁴Renewable and Sustainable Energy Institute, University of Colorado–Boulder, Boulder, CO, USA

*Corresponding author: steven.spurgeon@nlr.gov

August 11, 2026

Contents

S1 System Architecture and Implementation	2
S1.1 Retrieval-Augmented Generation Pipeline	3
S1.2 Reasoning Engine and Multi-Agent Design	5
S1.3 Hardware Control and Tool Registry	5
S1.4 Microscopy Data Acquisition and Processing	6
S1.5 Software and Reproducibility	7
S2 Benchmarking Methodology and Extended Results	8
S2.1 Evaluation Suite Design	8
S2.2 Model Selection	8
S2.3 RAG System Performance Metrics	10
S2.4 Response Actionability and Safety Trade-off	12
S2.5 Procedural Automation Benchmark	12
S2.6 System Limitations	12
S3 Extended Use Case Demonstrations	14
S3.1 TEM Characterization Plan: Full Agent Response	14
S3.2 TEM Execution Summary: Full Agent Response	16
S3.3 TEM Workflow prompt	17
S3.4 PFIB Failure Analysis: Full Agent Response	18
S3.5 PFIB Failure Analysis: Cross-Section Execution	19
S3.6 PFIB Failure Analysis Workflow Prompts	20

S1 System Architecture and Implementation

The group’s instrumentation suite presents a fundamental challenge for autonomous experimentation, since it includes multiple vendors’ hardware of different configurations. The current ADAM implementation targets Thermo Fisher Scientific (TFS) instruments specifically, with the potential to be extended as additional API access is established. Each instrument exposes a distinct software interface, operates under unique physical constraints, and requires instrument specific knowledge to use safely and effectively. Building a single agentic workflow that reasons safely and correctly across all of these contexts is a qualitatively harder problem than automating any one of them in isolation, and it is the central motivation for ADAM’s architecture. Prior autonomous microscopy systems have largely targeted single instruments or narrow task classes [1–4]. ADAM is designed from the outset to operate across this full instrumentation range without requiring a purpose-built agent for each platform [5].

The ADAM framework bridges the gap between large language models (LLM) reasoning and multi-step physical hardware execution by decomposing the problem into three modular components. The majority of AI applications in materials science target either computational property prediction or post-hoc data analysis [6–8]. ADAM instead operates in real time, interacting directly with live instruments during active experiments. In order to do this safely and reliably, the framework is built around a modular human-in-the-loop (HITL) architecture that grounds the model’s reasoning in domain-specific knowledge, routes tasks to purpose-built agents, and enforces explicit human oversight before consequential operations [1, 5]. This design is a deliberate step towards full autonomy rather than a permanent ceiling. In high-consequence environments such as national laboratories, where instruments represent multi-million dollar investments and experimental errors can carry significant operational costs, deploying an unconstrained autonomous agent is not tractable. HITL architecture provides grounding and auditability required to build institutional trust in an autonomous system incrementally [1, 5]. As confidence in the system grows through validation and deployments, oversight checkpoints can be selectively relaxed or removed, moving the framework from supervised automation toward fully autonomous operation without requiring a fundamental architecture change.

The three components shown in Main Text Figure 1 are implemented as follows. The **Laboratory** houses the physical instruments and their vendor-specific control software; instrument state is continuously streamed to the AI system via the State and Image Stream, ensuring the agent reasons about live feedback rather than issuing commands blindly. The **AI System** is orchestrated using LangGraph, which manages agent state, message routing, and tool call execution across the Router, Chat, and Control agents (detailed in Section S1.2). The **Data Sources** component comprises two elements: the RAG knowledge base, built from the group’s internal documentation and updated continuously as new documents are ingested, and the instrument tool registry, which defines the validated set of API calls available to the Control Agent for each instrument (detailed in Section S1.3). Information flows between components in two modes: a passive continuous stream that monitors instrument state and synchronizes the knowledge base, and an on-demand stream that carries user prompts, retrieved document chunks, and tool execution commands. The remainder of this section describes each of these three components in the order of the data flow that connects them: the RAG pipeline, the multi-agent reasoning engine, and the instrument tool registry.

S1.1 Retrieval-Augmented Generation Pipeline

Users interact with the system through a conventional, conversational interface like standard web-based LLM platforms. However, the underlying system diverges fundamentally from standard foundation models through its integration with a custom RAG pipeline [9]. General-purpose models are trained on publicly available data and inherently lack the proprietary, instrument and group specific knowledge required for highly specific advanced materials characterizations. In a physical laboratory, this is a critical limitation. When faced with specific, proprietary workflows, ungrounded models default to generalized responses that can be vague or incorrect. Hallucinating an incorrect API call like an incorrect beam current outside the specific instruments safe operation range, or an incorrectly ordered function argument can damage fragile samples or cause harm to the instrument itself. Reducing the likelihood of hallucinations is crucial in safely deploying agentic systems in physical laboratories.

Recent benchmarking of LLM agents on real microscopy hardware has confirmed that strong performance in domain specific knowledge assessment does not reliably transfer to safe accurate instrument control [2, 3]. Because all models endpoints are normalized to a common interface, the same RAG pipeline can be evaluated against base, fine-tuned, or domain-specific models interchangeably, making the system a practical testbed for studying how domain adaptation strategies affect performance in physical laboratory settings.

To address this, the pipeline constantly ingests multi-modal data from the research group’s internal documentation, including meeting transcripts, in progress experimental results, standard operations procedures, instrument manuals, and code repositories. All data is sourced from the group’s Microsoft SharePoint and GitHub teams, authenticated with institutional credentials. All embeddings are stored within the institution’s own infrastructure to ensure data security and to guarantee proprietary information stays on site. Data privacy is a particularly important problem in national laboratory and industrial research environments, where results, protocols, and characterization data may be subject to export controls, intellectual property restrictions, or pre-publication confidentiality requirements [5]. The fully on-premises deployment ensures that no unpublished data is transmitted to external model providers. The model gateway used is approved by the lab administrators for Controlled, Unclassified Information (CUI). This architecture supports variable privacy tiers, and allows documents to be selectively included or excluded from the RAG knowledge base depending on sensitivity. Because LLMs cannot process an entire laboratory’s document history within a single context window [9], ingested data undergoes *chunking* and *embedding*. Documents are segmented into semantically similar units—a slideshow for example, is divided into individual slides—where each chunk is converted into a high dimensional vector representation that encodes its semantic meaning [9, 10].

Critically, the pipeline implements a linked-chunk system in which each chunk maintains explicit references to adjacent and semantically related chunks, enabling the model to follow a thread of information across document boundaries when a single chunk provides insufficient context. This is a deliberate design decision to address the limitation of a single matched chunk capturing only a fragment of a procedural workflow. In complex laboratory settings, questions require pulling context from multiple sources. In this system, chunks can be linked by several relationship types, including parent document membership, document type, semantic content similarity, and positional proximity within a structured source. When the top chunk retrieved is insufficient for a fully grounded response, the system traverses these links to pull in additional related content before formulating a full answer. At query time, incoming

prompts are embedded using the same model and matched against the database via semantic similarity, retrieving the most contextually relevant chunks to inject into the agents reasoning context before it formulates a response. Furthermore, the pipeline is fully multi-modal, capable of processing not only plain text and source code, but also complex visual content including schematics, figures, images, and images embedded in documents or presentations. This is particularly relevant for materials science where critical procedural knowledge is frequently encoded in instrument diagram, images or visual workflows [7, 11]. In electron microscopy specifically, multi-modal reasoning remains in its infancy, with most ML models operating on single image types with deterministic preprocessing pipelines [6, 7]. Few systems are able to expand on and reason jointly across acquisition images, EDS maps, diffraction patterns and associated metadata within a single step [4, 11]. Models that can reason across modalities are therefore not just convenient, but necessary for capturing the full context of characterization sessions where decisions may depend equally on visual information and a constraint buried in a written procedure.

A final design feature of the RAG system is its transparency. In addition to the embedded vector database, a parallel document store retains the original, human-readable source files. Each database chunk carries a pointer to its parent document, and whenever the model draws a chunk for its response, it provides a direct download link to the full source document. This allows the user to independently verify any information the model provides at any time, helping to build trust in a system that operates sensitive equipment and can make consequential decisions. It also enables natural follow-up interaction where a user can retrieve a specific procedural document the model references and ask for further clarification.

S1.2 Reasoning Engine and Multi-Agent Design

The reasoning core of the system, represented by Main Text Figure 1 **b** is structured as a multi-agent hierarchy rather than a single model [3, 12]. This separation of responsibilities is a deliberate decision. When a user submits a natural language query, a lightweight central Router Agent evaluates the intent and directs it to the appropriate downstream agent. Conversational requests such as explaining a concept, discussing strategy, or interpreting a result is handled by a dedicated Chat Agent. Physical execution commands are delegated exclusively to the Control Agent. By design, the Chat Agent has no access to the instrument API, so the Control Agent is the only agent that can issue commands that affect the microscope. This ensures that a conversational exchange doesn’t inadvertently trigger a physical hardware operation [3].

Both the Chat Agent and the Control Agent maintain independent, active connections to the RAG database, allowing each to retrieve relevant context before responding. This distinction is important, because the Chat Agent RAG access allows it to function as a substantive scientific collaborator capable of discussing instrument parameters, historical experimental results, and lab specific procedures [13].

The Control Agent was selected through internal benchmarking across multiple frontier models, with selection criteria focused specifically on tool-calling accuracy and reliability under physical constraints. The full benchmarking methodology and model comparison results are detailed in Section S2. A key architectural feature of the reasoning layer is the asynchronous decoupling of the LLM inference from the hardware execution. Rather than requiring the agent to complete a full reasoning cycle before the instrument can act, the system dispatches tool calls as they are resolved and stream instrument state back to the agent. This means that manual operator interventions are immediately registered by the agent without interrupting the active reasoning loop, allowing genuine collaborative operation between the human operator and the automated system [5].

S1.3 Hardware Control and Tool Registry

To execute commands on physical instruments, ADAM interfaces with laboratory hardware (Thermo Fisher Helios 5 Laser Hydra and Spectra 200 TEM), via an abstracted tool registry. A persistent challenge in the transition toward autonomous experimentation is the relative immaturity of commercial instrument APIs [1, 2, 5]. Programmable interfaces like Thermo Fisher’s AutoScript were designed for scripted, human-authored automation rather than dynamic interaction with AI agents. They are instrument specific, functionally incomplete relative to the full capabilities of manual operation, and sensitive to argument ordering and hardware state dependencies in ways that are not clearly documented. Allowing a model to generate raw API calls without guardrails introduces significant risk where even a minor mistake can cause significant setbacks [3].

To overcome this, the Control Agent does not interact with the hardware API directly. Instead, it is provided with a curated, markdown-formatted registry of available functions, like `set_beam_current()`, `auto_focus()`, and `grab_frame()`. Each function is accompanied by a natural language description of its purpose, required arguments, and known constraints. The underlying Python implementations are never exposed to the model. This abstraction enforces safe, pre-validated operational boundaries and eliminates the risk of the model calling made-up functions. Importantly, this design also makes the system instrument-agnostic. By simply swapping out the API call registry, the same agent architecture can be deployed to any instrument with an API [5]. This has been demonstrated across two instruments

presented in this work.

Before consequential tool calls are dispatched to the instrument, the HITL safeguard engages. The system pauses execution and presents the proposed action to the operator for review. The checkpoint engages at semantic decision boundaries that have potentially dangerous consequences if done incorrectly such as ion beam milling or stage repositioning that would lose the current region of interest. This timing is intentional as reversible setup operations like beam parameter adjustment may proceed with lighter oversight, but the system will not cross an irreversible threshold without explicit confirmation. At each checkpoint, the operator has four response options: **Continue**, authorizing the proposed action and proceeding; **Reject**, canceling the current step while keeping the session active; **Different Direction**, providing natural language guidance to redirect the plan before it proceeds; and **Revert to Best**, rolling back to the most recently approved checkpoint state. This final operation is particularly important in multi-step workflow where the agent may have pursued a valid but suboptimal path, because it allows the operator to recover a known good state without restarting entirely. The operator may additionally authorize the agent to iterate autonomously for a defined number of steps between checkpoints, even for irreversible decisions, if needed. This prevents the oversight mechanism from becoming a bottleneck in routine operations.

Beyond planned HITL checkpoints, the Control Agent also handles unplanned API failures autonomously. When a tool call returns an error, due to incorrect parameter ordering, unsatisfied hardware dependence, or instrument state conflicts, the agent reads the returned error, reasons about its cause, and reformulates the call [14]. This is a meaningful distinction from HITL oversight, which manages anticipated uncertainty at semantic decision boundaries, while autonomous error recovery ensures the system degrades gracefully rather than catastrophically when real hardware behaves unexpectedly.

Together, these three components enable ADAM to securely translate conversational scientific intent into verifiable hardware execution. Having established the safety, grounding, and control mechanisms of this architecture, we next demonstrate its real-world efficacy through two distinct experimental deployments, a procedural automation task on the PFIB and an open ended exploratory analysis on the TEM.

S1.4 Microscopy Data Acquisition and Processing

PFIB Instrument. PFIB-SEM experiments were conducted on a Thermo Fisher Scientific Helios 5 Laser Hydra and Thermo Fisher Scientific Helios 5 PFIB CXe. Image library SEM images for the emulator were acquired on a standard tin on carbon specimen at an accelerating voltage of 5 kV and beam currents of 0.025, 0.2, 0.4, and 13 nA using an ETD detector and 4 mm working distance (WD). Focus (working distance) swept a range of 2–5 mm and stigmation swept a range of -0.99 to 0.99. Survey SEM imaging on the PN Diodes sample was performed at an accelerating voltage of 2 kV and current of 0.4 nA using an ETD detector at a stage tilt of 0° (normal to the SEM). The fabricated PN diodes sample consisted of Cr₂O₃:N sputtered films on epitaxial Ga₂O₃ (001) with Au/Ni contacts. The PFIB regular cross-section was executed using a Xenon plasma ion source at 30 kV and 60 nA with the stage tilted to 52° to orient the sample surface normal to the ion column. Post-mill verification images were acquired with the SEM at 52° stage tilt, at the same imaging conditions used for survey images.

TEM Instrument. TEM experiments were conducted on the Thermo Fisher Scientific Spectra 200 at 200 kV in STEM mode using HAADF and BF STEM detectors for morphological and structural

analysis and an EDS detector for elemental mapping. The sample consisted of CeTaN_x nanoparticles, which were dispersed in acetone and drop cast onto a TEM grid with a carbon support film. The open-ended characterization demonstration was performed on the Thermo Fisher AutoScript TEM emulator, which reproduces the instrument software interface and state using a library of real session images; scripts were then deployed on the Spectra 200 STEM.

Image Processing. During the PFIB spatial survey, image sharpness was assessed after each acquisition using a Tenengrad gradient sharpness metric. The sharpness score and acquired image were passed to the Control Agent after each acquisition. No additional image preprocessing was applied.

S1.5 Software and Reproducibility

The agentic workflow is implemented in LangGraph (v0.2+) and LangChain Core (v0.3+), with Chainlit (v1.3+) providing the conversational user interface, as described in Table S1. The system is compiled as a structured state graph with three runtime execution paths: conversational, hardware control, and code generation. Hardware commands are dispatched to the instrument via the Thermo Fisher AutoScript Python API over port 7520 through a dedicated subprocess bridge. Failed tool calls are retried autonomously up to N iterations (configurable via `ADAM_MAX_RETRIES`) before escalating to the operator.

All foundation model endpoints are normalized to an OpenAI-compatible interface via an on-premises LiteLLM instance; non-native endpoints are adapted using a FastAPI converter. All endpoints are approved for Controlled Unclassified Information (CUI) by institutional administrators. API routing, latency, and token usage are monitored via an enterprise LangSmith instance.

RAG embeddings use `sentence-transformers/all-MiniLM-L6-v2` (384-dimensional) for text, CLIP ViT-B/32 (512-dimensional) for visual content, and a `MicroNet InceptionV4` model pretrained on microscopy images (1536-dimensional) for image-specific queries. All vectors are stored in a PostgreSQL database with the `pgvector` extension. Retrieval uses cosine similarity, returning the top $k = 6$ chunks (configurable via `ADAM_RAG_TOP_K`). Source documents are staged in an on-premises VAST S3-compatible object store; no proprietary data is transmitted to external services. Benchmarking visualizations were generated using Matplotlib and Plotly.

Table S1: Key software dependencies.

Component	Package	Version
Workflow orchestration	LangGraph	v0.2+
Agent framework	LangChain Core	v0.3+
User interface	Chainlit	v1.3+
Model gateway	LiteLLM	—
API adapter	FastAPI	—
Evaluation logging	LangSmith	enterprise
Text embedding	sentence-transformers/all-MiniLM-L6-v2	—
Visual embedding	CLIP ViT-B/32	—
Microscopy embedding	MicroNet InceptionV4	—
Vector database	PostgreSQL + pgvector	—
Instrument control	Thermo Fisher AutoScript	—
Visualization	Matplotlib / Plotly	—

S2 Benchmarking Methodology and Extended Results

S2.1 Evaluation Suite Design

These qualitative use cases demonstrate the frameworks operational flexibility, but deploying autonomous agents to safely orchestrate physical instruments require empirical validation. Standard AI performance benchmarks fail to capture the physical risks and domain specific constraints required for autonomous microscopy. Because no standardized benchmark currently exists for evaluating LLMs in real world materials characterization tasks, we constructed a custom evaluation suite consisting of 20 expert curated Question & Answer pairs (QA pairs). These prompts were specifically designed to reflect the nuanced challenges of materials science labs. Each prompt is associated with a human expert generated correct answer, to aid the model in producing responses that are correct and formatted correctly in the context of our group’s experiments. In this setting, a useful response to a scientist may differ vastly from a standard AI interaction. A generic LLM often defaults to vague answers or dangerous guesswork, whereas an effective autonomous agent must provide highly specific, accurate instructions. Therefore, we anchored our quantitative evaluation on metrics that directly impact materials science like literature grounding, analytical actionability, hardware safety and procedural efficiency.

S2.2 Model Selection

Selecting a reasoning model suited to laboratory use requires navigating a similar trade-off between conversational helpfulness and factual safety. We evaluated four frontier models — Gemini 3.1 Pro (Google, closed-source API), Nemotron Super (NVIDIA, open-weight, self-hostable), Claude Opus 4.7 (Anthropic, closed-source API), and GPT-5.4 (OpenAI, closed-source API) — across both axes; Supplementary Section S2). Both axes were derived from automated LLM-as-judge evaluation using LangSmith. Helpfulness was scored directly using the LangSmith built-in criteria evaluator; factual safety was computed as $1 - \text{fabrication risk}$, where fabrication risk was assigned by a custom LLM judge prompt that assessed whether each response introduced claims not grounded in the retrieved context or the model’s verifiable knowledge. This evaluation was not designed to identify a single best model; the goal was to characterize the tradeoff space along dimensions directly relevant to laboratory deployment. The four models differ in ways that extend beyond benchmark scores: Nemotron Super, as an open-weight model, can be self-hosted within a facility network without routing queries to an external API — a meaningful consideration for institutions with data sensitivity or export control requirements. The closed-source models vary substantially in cost per token and inference latency, both of which constrain the feasibility of closed-loop acquisition workflows where the agent may issue dozens of sequential calls per session. A full comparison of latency, cost, and grounding behavior across models is provided in Supplementary Section S2.3. Gemini 3.1 Pro achieved the highest factual safety score ($\sim 57\%$) at the cost of the lowest helpfulness ($\sim 56\%$). In a setting where model outputs translate directly into hardware commands, this trade-off favors conservatism because an agent that gives a terse or cautious response is preferable to one that confidently generates a physically incorrect API call [1]. Gemini 3.1 Pro was therefore selected as the Control Agent model, while higher-helpfulness models remain available for the Chat Agent where conversational fluency is more valuable. These rankings reflect a snapshot in time of four common models; our goal is only to emphasize the design principle, not endorse any specific model.

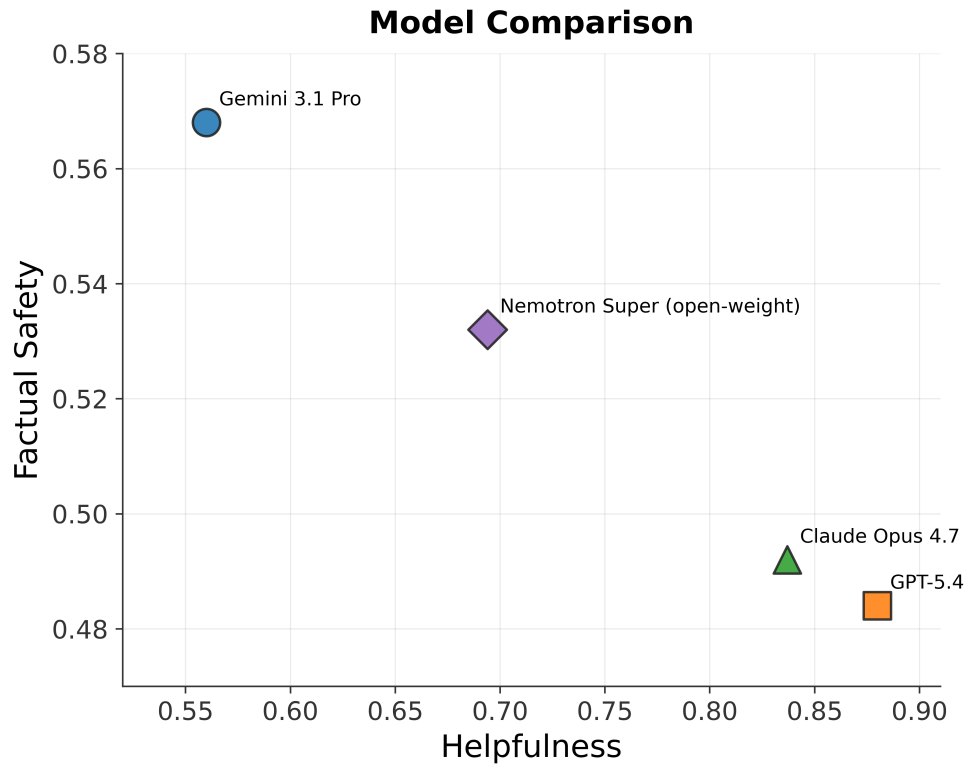


Figure S1: Factual safety score versus conversational helpfulness for four frontier models evaluated for Control Agent selection, assessed across the 20 expert-curated QA pairs. Gemini 3.1 Pro achieved the highest factual safety ($\sim 57\%$) at the cost of the lowest helpfulness ($\sim 56\%$). In a setting where model outputs translate directly into hardware commands, conservatism is preferred; Gemini 3.1 Pro was therefore selected as the Control Agent, while higher-helpfulness models are available for the Chat Agent.

S2.3 RAG System Performance Metrics

Integrating a RAG pipeline introduces measurable trade-offs in latency, operational cost, and output characteristics relative to a base model without retrieval. Figures S2–S5 characterize these trade-offs across four query categories: code generation, explanation, factual lookup, and comparison.

As shown in Figure S2, RAG integration increased latency for code queries, where retrieval of relevant documentation was slower and less precise, while reducing latency for explanation and comparison queries, where targeted context retrieval allowed the model to resolve answers more directly. Factual query latency was nearly identical across both conditions.

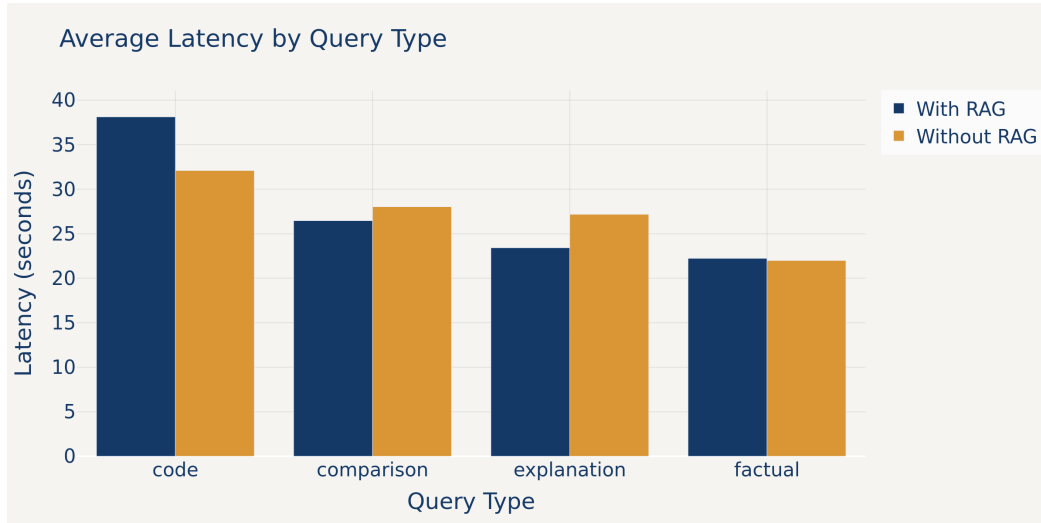


Figure S2: Average latency by query type for ADAM (RAG) versus base model without retrieval.

As shown in Figure S3, RAG consistently increased operational cost per query by a factor of 3–5 $\times$ across all categories, driven by the additional input tokens from injected retrieval context. This cost increase is expected and represents the primary operational trade-off of RAG deployment.

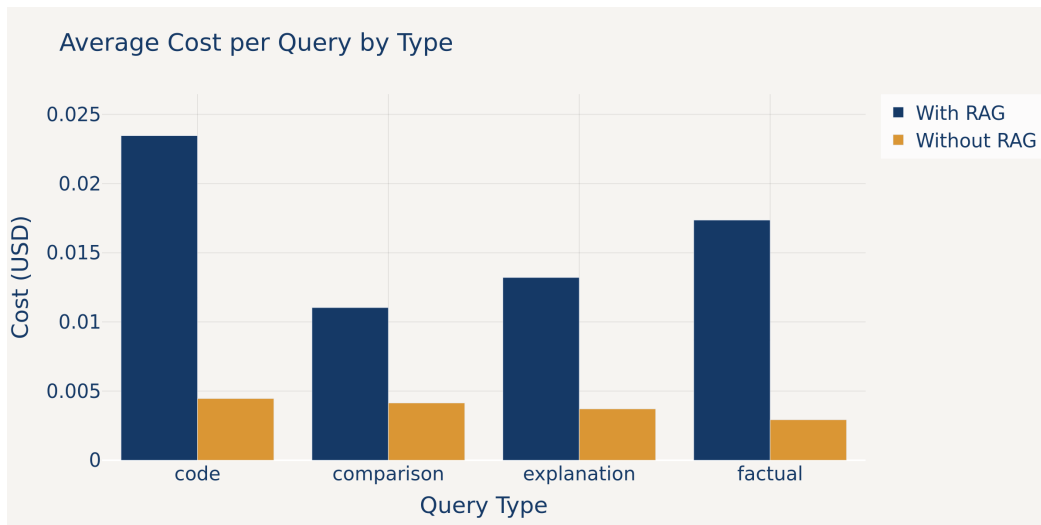


Figure S3: Average operational cost per query by type for ADAM (RAG) versus base model.

Despite higher input token counts, ADAM (RAG) produced more concise responses across all query types, with the largest reductions in explanation and comparison categories (Figure S4). This output conciseness reflects grounding in specific internal documentation rather than speculative elaboration.

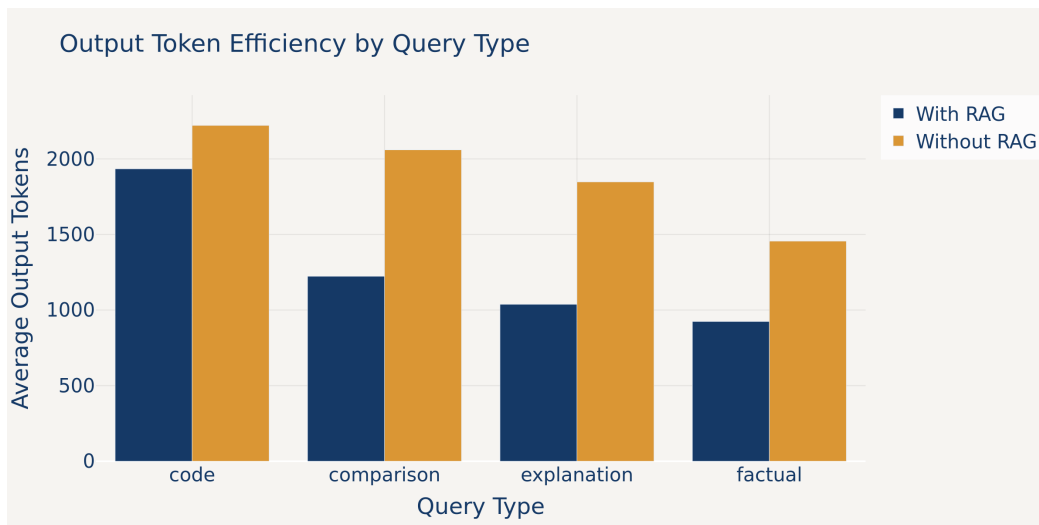


Figure S4: Average output token count by query type for ADAM (RAG) versus base model.

Factual grounding scores were higher for ADAM (RAG) across explanation, factual, and comparison queries (Figure S5). Code queries were the exception, where the base model’s pretrained knowledge outperformed retrieved documentation, suggesting that code-specific retrieval quality is a near-term improvement target for the knowledge base.

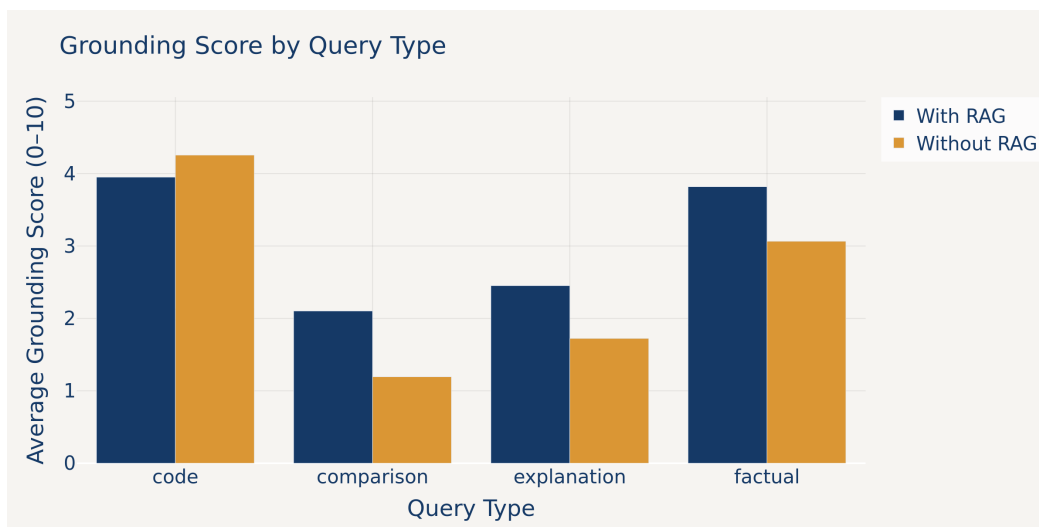


Figure S5: Average grounding score by query type for ADAM (RAG) versus base model.

Taken together, these results confirm that RAG integration improves factual reliability and output precision at the cost of higher per-query expenditure, a trade-off that is appropriate for a physical laboratory setting where response accuracy carries direct operational consequences.

S2.4 Response Actionability and Safety Trade-off

Evaluating the utility of agent responses revealed a meaningful trade of between actionability and accuracy, shown in Figure 4b. We define an *actionable* response as one that provides a specific numerical parameter an executable API tool call, or concrete step-by-step protocol. When tasked with an open-ended trouble shooting prompt, the non RAG version provides 19 actionable responses out of 50, while the RAG-enhanced model produced 13. However, the base model also generated 3 inaccurate responses, 4% more than the RAG enhanced version. The RAG-integrated agent replaced these incorrect responses with more conservative answers that acknowledged uncertainty rather than fabricating specificity. In a physical laboratory context, this is a deliberate and favorable trade-off: a response that says “consult your SOP for the exact beam current” is recoverable, while a confidently stated but incorrect beam current is not.

S2.5 Procedural Automation Benchmark

Procedural automation capabilities were evaluated on a routine PFIB image focus correction task, in which the agent was required to independently adjust working distance and stigmatism to bring the beam to a sharp, calibrated focus. This foundational step was also performed by an expert experienced operator and an inexperienced novice, with results shown in Figure 4c. The agent only required 2 user interactions, an initial natural language prompt and a final HITL confirmation. The expert required 33 interactions and the novice 55, reflecting the iterative manual adjustments of human driven focus corrections.

Image quality was assessed using the TFS integrated sharpness calculator, which reports a distance value in nanometers as a measure of image resolution. 0 is ideal, and the larger the number, the more out of focus it is. The agent achieved a final sharpness score of 6.928 nm, outperforming the novice (6.997 nm), and approaching expert level calibration (5.652). The agents total wall clock time to completion (7:20) was longer than both the expert (1:10) and novice (2:50), with the majority of this time due to LLM inference latency rather than instrument operation time. This represents a favorable trade-off in practice. By compressing the required human interaction to a single prompt and final approval, the system enables genuinely asynchronous instrument utilization. The microscope can run unattended while the researcher focuses on higher level experimental design, data interpretation, or other concurrent tasks. The ability to dispatch focus correction and routine calibration jobs during the otherwise idle periods. decouple machine uptime from human working hours, increasing the throughout of the instrument without requiring additional human hours.

S2.6 System Limitations

Despite these results, several limitations still remain. Most significantly, RAG integration reduces but does not eliminate hallucination risk. Even with domain grounding, frontier LLMs can generate physically plausible but incorrect outputs, particularly for edge cases not well-represented in the RAG knowledge base. A similar risk is database error propagation where a source document contains outdated or incorrect information. The agent might cite this confidently despite it being incorrect. The quality of the ingested documentation is as important as the amount, and a RAG system is only as good as its sources.

In addition, the EDS interpretation in the STEM demonstration (Section S3.2) illustrates a subtle point of failure, different from hallucination. The agent’s reasoning is internally consistent but built on total X-ray counts, rather than net X-ray counts, a distinction easy to overlook without the appropriate processing pipeline. To address this gap, it is necessary for the agent to have access not only to raw data, but also properly processed (background-subtracted and deconvoluted) data, which is at present not accessible through the instrument API.

The framework also requires target instruments to expose a programmable control API. While most instruments have something like this, their maturity and completeness vary considerably across manufacturers and instrument generations. Furthermore, some capabilities are only accessible via physical controls, meaning the agent has no way of tuning those parameters, even with access to the API.

Finally, the current systems inference latency makes it intractable for experiments where response time is critical like in-situ experiments. In these contexts, the asynchronous decoupling advantage disappears and faster deterministic scripts are more appropriate. It is likely possible to significantly reduce inference latency, but it is unclear if it would be low enough for in-situ (streaming) experiments, etc. However, those kinds of optimizations are outside the scope of the paper. The demonstrations presented here encompass multiple sessions, samples, and tasks, but each scenario is only run once, in contrast to the QA and focus-correction benchmarks. These demonstrations speak to the feasibility of employing ADAM across a range of possible experiments, rather than a quantified reliability rate for checkpoint inference or failure-mode reasoning. Characterizing long-term reliability and other issues, such as reduced-latency model deployment, automated database validation, and expanded API coverage are the near-term development roadmap for this framework.

S3 Extended Use Case Demonstrations

S3.1 TEM Characterization Plan: Full Agent Response

Editorial note: The execution summary below diverges from the characterization plan in two respects. First, the agent’s ROI selection reasoning — identifying two candidate particles in the initial survey image and selecting the dominant particle as the primary analysis target — was performed but not explicitly logged in the execution summary. Second, the agent reorganized its phase structure during execution: the plan’s Step 1 (morphological survey including ROI selection) and Step 2 (focus optimization) were condensed into Phase 1, while the planned crystal structure phase was decomposed into separate HAADF (Phase 1) and BF STEM (Phase 2) acquisitions. The scientific outputs are consistent with the plan and figure; these discrepancies reflect current limitations in execution logging rather than reasoning errors. In addition, the conclusion drawn from the EDS analysis reflects a common confusion in the literature, which is the difference between total and net-counts X-ray maps; at present, the model is limited to direct viewing of raw X-ray data as opposed to background-subtracted, deconvoluted, and quantified net counts data. Improving the granularity and completeness of execution logging is an identified area for future development.

ADAM Response

Characterization Plan

To comprehensively characterize your sample on the TEM, we need to investigate its morphology, crystal structure, and chemical composition. Here is a detailed breakdown of what we should look for, how we will find it, and how we will analyze it, followed by a step-by-step experiment plan.

What to Look For, How to Find It, and How to Analyze It

1. Morphology and Size Distribution

What to look for: Particle size, shape (spherical, faceted, rods, etc.), aggregation, and overall distribution on the grid.

How to find it: Low-to-medium magnification Scanning Transmission Electron Microscopy (STEM) using High-Angle Annular Dark-Field (HAADF) or Bright-Field (BF) detectors. HAADF provides Z-contrast (heavier elements appear brighter), while BF provides phase/diffraction contrast.

How to analyze it: I will use the `InterpretImage` tool (Vision AI) to assess the field of view, count visible particles, describe their spatial arrangement, and identify Regions of Interest (ROIs) for higher magnification.

2. Crystal Structure and Defects

What to look for: Lattice fringes, crystallinity (amorphous vs. crystalline), grain boundaries, and structural defects.

How to find it: High-Resolution STEM (HR-STEM) using the HAADF detector at high magnification. Optimal focus is critical here.

How to analyze it: I will use `AnalyzeImage` to calculate the Tenengrad sharpness metric and iteratively optimize the focus. Once focused, I will acquire the image and use `InterpretImage` to check for visible lattice fringes, estimate d-spacing, and suggest possible crystal phases.

3. Chemical Composition and Distribution

What to look for: Elemental makeup, core-shell structures, elemental segregation, or impurities.

How to find it: Energy Dispersive X-ray Spectroscopy (EDS) mapping in STEM mode.

How to analyze it: I will acquire an EDS map and use `InterpretImage` to describe the elemental distribution, identifying which regions are enriched in specific elements and correlating them with the morphological features observed earlier.

Experiment Outline and Execution Plan

Acquire → Analyze → Act

Here is the step-by-step plan for how I will execute the experiment autonomously and use closed-loop reasoning to inform each subsequent step:

Step 1: Initial Survey & Morphological Analysis

Action: Ensure the microscope is in STEM mode with convergent illumination. Acquire a medium-magnification HAADF image (`FindMatchingImageTem(detector="HAADF")`).

Analysis: Pass the image to `InterpretImage` asking to describe the general morphology, particle size, and identify 1–2 promising ROIs that are well-dispersed and suitable for high-resolution imaging.

Decision: Based on the description of the ROIs, move to the best candidate for the next step.

Step 2: Focus Optimization

Action: At the chosen ROI, begin a focus sweep.

Analysis: Acquire images and use `AnalyzeImage` to measure the sharpness score, adjusting objective lens focus/defocus and comparing scores after each change.

Decision: If sharpness increases, continue in that direction; if it decreases, reverse. Stop when sharpness is maximized.

Step 3: High-Resolution Structural Analysis

Action: Once optimally focused, acquire a Bright-Field (BF) STEM image to exploit diffraction contrast for crystal domain and defect analysis, followed by a high-resolution HAADF image if lattice fringes are resolved.

Analysis: Pass the image to `InterpretImage` specifically asking to identify lattice fringes, estimate d-spacing, and describe the crystal structure or faceting of the individual particle.

Decision: Structural findings will dictate whether to examine a different particle or proceed to chemical analysis.

Step 4: Chemical Mapping

Action: Acquire an EDS spectrum/map of the same high-resolution region (`FindMatchingImageTem(detector="EDS")`).

Analysis: Pass the EDS data to `InterpretImage` to determine elemental composition and map the spatial distribution of elements across the particle.

Final Synthesis: Compile the morphological, structural, and chemical data into a comprehensive final report.

Ready to begin when you approve.

S3.2 TEM Execution Summary: Full Agent Response

ADAM Response

Characterization Results

Here is the comprehensive summary of the sample characterization experiment, executed autonomously as requested.

Phase 1: Morphology and Size Distribution (HAADF STEM)

Action: Switched illumination mode to Convergent (STEM mode) and acquired a High-Angle Annular Dark-Field (HAADF) image.

Analysis: The image revealed a single dominant, high-Z or relatively thick particle rather than a dispersed field of nanoparticles.

Findings:

Shape: The particle is large, faceted, and block-like/prismatic with an aspect ratio of roughly 1.3–1.6.

Features: There are small protrusions or attached fragments on the upper and left edges, suggesting secondary nucleation, agglomeration, or fracture debris.

Internal Structure: The interior contrast is mostly uniform, indicating a relatively thick and compositionally consistent particle in projection. No clear grain boundaries were visible.

Phase 2: Crystal Structure and Defects (BF STEM)

Action: Acquired a Bright-Field (BF) STEM image of the same region to look for diffraction contrast, lattice fringes, and defects.

Analysis: The BF image confirmed the blocky morphology but was dominated by mass-thickness and diffraction contrast.

Findings:

Lattice Fringes: Not visible at this magnification/field of view. The particle is likely too thick or not oriented along a low-index zone axis.

Defects: No unambiguous grain boundaries, dislocations, or stacking faults were resolved. The particle appears to be a single dominant crystal domain or composed of overlapping grains too thick to distinguish.

Recommendation: To resolve lattice fringes and confirm crystal phase, a follow-up session should tilt the particle toward a low-index zone axis for HRTEM or selected-area electron diffraction (SAED).

Phase 3: Chemical Composition (EDS Mapping)

Action: Acquired an Energy-Dispersive X-ray Spectroscopy (EDS) map of the region.

Analysis: The available EDS data for this specific region was a “total counts” map (summing all X-ray energies).

Findings:

Distribution: The total counts map was spatially uniform, showing no strong localized enrichment or segregation.

Interpretation: This suggests the particle is chemically homogeneous at this resolution, or that any secondary phases do not have a substantially different overall X-ray yield from the main body.

Proposed next acquisition: Individual elemental maps for O, Ta, Fe, N, C, and Ce to confirm elemental ratios and enable quantitative phase identification.

Summary

Phase	Technique	Key Findings	Limitations / Next Steps
1. Morphology	HAADF STEM	Single large, faceted, prismatic particle with minor attached surface fragments.	Need calibrated scale bar for absolute size.
2. Structure	BF STEM	Uniform dense contrast; no visible grain boundaries or defects.	Requires higher magnification (HRTEM) or SAED to resolve lattice fringes and phase.
3. Composition	EDS Mapping	Uniform total X-ray yield across the field of view.	Need individual elemental maps or point spectra to confirm elemental ratios.

Conclusion: The sample consists of large, faceted, likely single-crystalline blocky particles with minor surface agglomeration. Atomic-level crystallographic data would require moving to a thinner region, tilting to a zone axis, and increasing magnification significantly for HRTEM/HRSTEM imaging.

Let me know if you would like to explore a different region of the grid or perform any specific stage movements!

S3.3 TEM Workflow prompt

I want you to design and execute an experiment to characterize this sample. What things should we be looking for? Tell me what to look for, how to find that information, how to analyze it, etc. Once you have provided a summary and outline for the experiment, tell me how you would execute the experiment. Tell me how you would independently analyze the results at each step to inform the next step and each subsequent step.

S3.4 PFIB Failure Analysis: Full Agent Response

Editorial note: This session demonstrates a second, complementary oversight mechanism distinct from the semantic-reversibility HITL checkpoints described in the main text. Here, a rule-based safety gate intercepted the initial request because it required large-increment stage motion (~ 1.3 mm) while the electron beam was active and unblanked, a condition that risks sample charging and drift. The agent proposed a safe alternative (blank $\rightarrow$ move $\rightarrow$ unblank $\rightarrow$ acquire) and paused execution until operator review; the operator reviewed the flagged risk and chose to override the gate rather than adopt the suggested fix, and execution proceeded. This is distinct from the semantic reversibility checkpoints described in Section S1.3, since decisions are checked against a fixed rule, whereas the reversibility checkpoint infers whether a proposed action can be undone.

ADAM Response

Safety Gate

Safety Gate Blocked

Reason: Stage movement is requested while the electron beam is active and unblanked.

Safety Review: WARNING

Findings: The user request involves moving the stage by relatively large increments (~ 1.3 mm) while the electron beam is active and unblanked.

Recommendation: Blank the electron beam before performing stage movements to avoid sample charging, drift, or unintended exposure during transit.

Suggested safe alternative:

1. Blank the electron beam
2. Move the stage to the next diode position (Y + 1.3 mm)
3. Unblank the electron beam
4. Acquire and save the SEM image
5. Repeat the blank–move–unblank–acquire sequence for the subsequent diode

Execution has been blocked pending your approval.

Safety: WARNING — Stage movement is requested while the electron beam is active and unblanked.

Choose an action below: Do Suggested Fix / Override & Execute / Cancel

The operator selected **Override & Execute**. The agent proceeded under this explicit instruction rather than the suggested blank–move–unblank sequence.

ADAM Response

Survey and Failure Analysis

I have successfully completed the high-resolution SEM image collection sequence for the three largest diodes in your column of devices.

By following your instructions, I executed only pre-written API calls in the correct order, ensuring safety by blanking the electron beam during stage transits to prevent sample charging and drift.

1. Stored Stage Positions

The stage positions for each of the three diodes have been recorded and stored for future experiments.

2. Image Acquisition & Materials Science Analysis

All images were acquired at 850 μm HFW with a high-resolution scan setting of 3072×2048 and saved to the specified directory.

Diode 1 (Bottom-most) — diode_1.tif

Key Features: Clean, uniform, and continuous Au/Ni contact pad with sharp, well-defined boundaries. The isolation trench surrounding the active area is completely clear of debris or metal redeposition.

Defects: None detected. The semiconductor surface $\text{Cr}_2\text{O}_3/\text{Ga}_2\text{O}_3$ shows excellent morphology with no micro-cracks or pinholes.

Assessment: Pass (Working Device).

Diode 2 (Middle) — diode_2.tif

Key Features: The Au/Ni contact pad shows localized melting and structural deformation near the perimeter.

Defects: A distinct dielectric breakdown hotspot is visible at the edge of the contact pad, characterized by localized material vaporization and micro-cracking propagating into the Ga_2O_3 layer. This is a classic signature of high-field electrical failure.

Assessment: Fail (Electrical Breakdown).

Diode 3 (Top-most) — diode_3.tif

Key Features: The contact pad exhibits severe contrast variations, indicating height differences and loss of contact.

Defects: Severe contact delamination of the Au/Ni metallization from the underlying oxide layer, likely caused by thermal stress or poor adhesion. Additionally, a large metallic particle is bridging the isolation trench, creating a direct electrical short-circuit between the anode and cathode.

Assessment: Fail (Delamination & Short-Circuit).

3. Defect Annotation Script (Run Outside Agent)

Below is a robust Python script using OpenCV and NumPy designed to run locally on your machine. It filters out minor surface roughness and only labels major defects (such as cracks, voids, delamination, and bridging) that would cause device failure.

[Generated script omitted here for brevity; see supplementary code listing / repository.]

Similar to the STEM demonstration (Section S3.2), the agent decided to deliver a standalone script rather than execute the annotation itself; this was inferred entirely from the user’s natural language instruction (“write a script that I can run outside of the agent”), with no explicit rule distinguishing this from a tool-registry call.

S3.5 PFIB Failure Analysis: Cross-Section Execution

Following the survey, the operator issued a follow-up instruction to execute a FIB cross-section at Diode 1, again constraining execution to direct API calls rather than agent-generated code.

ADAM Response

Cross-Section Execution

I have successfully completed the requested operations at Diode 1.

Here is a summary of the steps executed:

- 1. Safety Preparation:** Blanked the electron beam to prevent sample charging and drift during transit.
- 2. Stage Movement:** Moved the stage to the coordinates of Diode 1 (bottom-most device).
- 3. Sample Alignment:** Safely tilted the stage to 52° (making the sample surface normal to the FIB column).
- 4. Ion Beam Configuration:**
 - Turned on the ion beam
 - Set the accelerating voltage to 30 kV
 - Set the beam current to 60 nA
- 5. Patterning:** Defined a regular cross-section pattern of $10 \times 10 \times 20 \mu\text{m}$ centered at the field of view. Executed and completed the milling process.
- 6. Verification:** Unblanked the electron beam and acquired a new SEM image to verify the cross-section.

A final operator prompt requested an additional SEM image at the current stage position, confirming the completed cross-section.

ADAM Response

Verification Image

Saved Path: [filepath]

Resolution: 1536×1024

S3.6 PFIB Failure Analysis Workflow Prompts

I am on a sample that contains $\text{Cr}_2\text{O}_3/\text{Ga}_2\text{O}_3$ PN diodes with Au/Ni contacts. I want to look at the largest diodes in this column of devices, which are spaced $\sim 1.3 \text{ mm}$ apart. I am currently in the center of the bottom-most diode in this column, which is a working device. I am at the sample eucentric height and the image is in focus. I want to take high resolution SEM images at $850 \mu\text{m}$ HFW of the 3 largest diodes in this column, as I am trying to catalog which diodes have failed. Do not write your own code to execute the image collection, simply only use only API calls in the correct order. Briefly analyze each image when you're at each image location for key features and provide a summary of features at each image region, describing any defects that may be present. Save the images at [filepath]. Store the stage positions of each diode for potential future experiments. After the images are collected, please write a script that I can run outside of the agent to annotate and label any defects in each image, only with sensitivity to defects that would cause device failure. Make sure this script is robust and only labels major defects.

Can you go to diode 1 and make a $10 \times 10 \times 20 \mu\text{m}$ FIB cross-section at 30 kV 60 nA? Make sure the sample is normal to the FIB before patterning. Execute only API calls, don't write code.

Can you take and save an SEM image at the current stage position?

Supplementary References

1. Kalinin, S. V. *et al.* Automated and Autonomous Experiments in Electron and Scanning Probe Microscopy. *ACS nano* **15**, 12604–12627. doi:10.1021/acsnano.1c02104 (Aug. 24, 2021).
2. Wall, M. K., Pattison, A. J., Barnard, E. S., Ribet, S. M. & Ercius, P. TEM Agent: enhancing transmission electron microscopy (TEM) with modern AI tools. *npj Computational Materials*. doi:10.1038/s41524-026-02103-z. arXiv: 2511.08819[cond-mat.mtrl-sci] (June 10, 2026).
3. Mandal, I. *et al.* Evaluating large language model agents for automation of atomic force microscopy. *Nature Communications* **16**. Publisher: Nature Publishing Group, 9104. doi:10.1038/s41467-025-64105-7 (Oct. 14, 2025).
4. Chen, G., Yuan, W. & You, F. Bridging electron microscopy and materials analysis with an autonomous agentic platform. *Science Advances* **12**, eaed0583. doi:10.1126/sciadv.aed0583.
5. Vriza, A., Prince, M. H., Zhou, T., Chan, H. & Cherukara, M. J. Operating advanced scientific instruments with AI agents that learn on the job. *npj Computational Materials* **12**. Publisher: Nature Publishing Group, 160. doi:10.1038/s41524-026-02005-0 (Mar. 6, 2026).
6. Dehaerne, E., Dey, B., Blanco, V. & Davis, J. Scanning Electron Microscopy-based Automatic Defect Inspection for Semiconductor Manufacturing: A Systematic Review. *Journal of Micro/Nanopatterning, Materials, and Metrology* **24**. doi:10.1117/1.JMM.24.2.020901. arXiv: 2409.06833[eess.IV] (May 15, 2025).
7. Kalinin, S. V. *et al.* Machine learning for automated experimentation in scanning transmission electron microscopy. *npj Computational Materials* **9**. Publisher: Nature Publishing Group, 227. doi:10.1038/s41524-023-01142-0 (Dec. 20, 2023).
8. Guinan, G. *et al.* Revealing the hidden third dimension of point defects in two-dimensional MXenes. *Nature Communications* **17**. Publisher: Nature Publishing Group, 3473. doi:10.1038/s41467-026-71670-y (Apr. 14, 2026).
9. Lewis, P. *et al.* Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks Apr. 12, 2021. doi:10.48550/arXiv.2005.11401. arXiv: 2005.11401[cs.CL].
10. Li, J., Yuan, Y. & Zhang, Z. Enhancing LLM Factual Accuracy with RAG to Counter Hallucinations: A Case Study on Domain-Specific Queries in Private Knowledge-Bases Mar. 15, 2024. doi:10.48550/arXiv.2403.10446. arXiv: 2403.10446[cs].
11. Jamali, V., Aghazadeh, A. & Kacher, J. Thinking microscopes: agentic AI and the future of electron microscopy. *npj Computational Materials* **12**. Publisher: Nature Publishing Group, 149. doi:10.1038/s41524-026-02077-y (Apr. 10, 2026).
12. Guo, T. *et al.* Large Language Model based Multi-Agents: A Survey of Progress and Challenges Apr. 19, 2024. doi:10.48550/arXiv.2402.01680. arXiv: 2402.01680[cs.CL].
13. Lála, J. *et al.* PaperQA: Retrieval-Augmented Generative Agent for Scientific Research Dec. 14, 2023. doi:10.48550/arXiv.2312.07559. arXiv: 2312.07559[cs.CL].
14. Yao, S. *et al.* ReAct: Synergizing Reasoning and Acting in Language Models Mar. 10, 2023. doi:10.48550/arXiv.2210.03629. arXiv: 2210.03629[cs.CL].