

ELM: An Epidemiological Language Model via Supervised Reasoning Fine-Tuning on Global Surveillance Data

Aniket Wattamwar
Computer Science
California State University
Fullerton, CA, USA
orcid.org/0009-0001-1454-838X

Sampson Akwafuo
Computer Science
California State University
Fullerton, CA, USA
orcid.org/0000-0001-8255-4127

Abstract—The utility of general-purpose Large Language Models (LLMs) in public health is frequently limited by a lack of domain-specific investigative logic and an inability to perform structured risk assessments. This paper introduces the Epidemiological Language Model (ELM), a model that is fine-tuned specifically for the logical workflows of outbreak surveillance. We move beyond standard medical literature by targeting the World Health Organization (WHO) Disease Outbreak News (DONs) as our primary training corpus. By leveraging the WHO’s OData API, we extracted 100 reports and developed a Supervised Reasoning Fine-Tuning (SRFT) pipeline that maps raw epidemiological observations to formal risk assessments and public health response strategies. ELM serves as a specialized reasoning core designed to process real-time surveillance signals into actionable public health intelligence.

Index Terms—Large Language Models, Public Health Surveillance, Parameter-Efficient Fine-Tuning (PEFT), LoRA, Computational Epidemiology, Epidemiological Language Model, Disease Outbreak.

I. INTRODUCTION

The growth and spread of diseases have been a significant challenge. Organizations like WHO and CDC update data continuously, but the challenge to keep it structured and generate insights and the ability to make decisions is extremely difficult. A lot of disease text is available on WHO, CDC, and NCBI in multiple formats like XML, JSON, Unicode, BioC, and ASCII. Bulk downloads are available spanning manuscripts up to 100s of GB. WHO publishes articles and narratives called Disease Outbreak News (DONs). These narratives include highly important information specifically credible for epidemiology and surveillance. Every narrative talks about the situation, the description of the situation, the public health response, the risk assessment, and finally, advice.

The rise of Large Language Models (LLMs) in the past years has opened many opportunities to extract text from unorganized and unstructured data and keep it clean for use. LLMs have proven to be very effective in gathering and summarizing data. However, it is also prone to hallucination. DONs can play an important role in advancing how epidemiological data and text are perceived. There is a need to structure this data based on content, topics, outbreak and factors that are closely related

to each other. Each DON has information about a disease, its spread, action and response taken, and assessment and advice. These articles and narratives are reasoning and logic and not just facts about diseases and outbreaks. It resembles patterns and spread rates.

Public Health Surveillance is a way to understand data, analyze, and interpret at the right time. It guides us to detect outbreaks and look for patterns. By tracking patterns, the data tracked over time provides a real-time picture and measures the groups with higher risk. Effective surveillance are continuous, efficient, and practical to identify emerging threats. There are many different surveillance methods used depending on the health event. Passive surveillance, Active surveillance, behavioral, sentinel and syndromic surveillance systems are few methods. These methods form loops around healthcare professionals, public health agencies, communities to enable timely evidence based public health action. Active surveillance involves deliberate, regular outreach to collect complete data, including negative cases, and is used for rare or high-impact diseases requiring immediate public health action. Sentinel surveillance collects high-quality data from selected reporting sites to monitor trends in common diseases efficiently, though it may miss cases outside its network. Syndromic surveillance, an increasingly important approach, monitors symptom patterns from sources like emergency departments and health helplines to provide early warning signals of outbreaks before diagnoses are confirmed. Disease and Health Surveillance will be used to measure disease burden, monitor trends, risk factors, planning and implementation of resources, control and preventive measures, and as well as research [20]. Combining Generative AI with existing disease surveillance systems would increase throughput of actions and reduce time to action rate.

Current LLMs available on ChatGPT (GPT-4, GPT-5.x series), Gemini Pro, and Claude (Opus) are generic and trained on the entirety of the internet information available. These LLMs are not specialized in epidemiology tasks. They do not contain the latest information as the knowledge base they were trained on. They would retrieve information via techniques like

RAG and web scraping and provide the information to the user with the possibility of hallucination. In the high-stakes environment of field epidemiology, a model must do more than recall facts, it must evaluate diagnostic uncertainty, assess regional transmission risks, and propose tactical interventions. This is the “reasoning gap” in medical AI.

This paper introduces ELM (Epidemic Language Model), a specialized LLM designed for the reasoning and the logical workflow of outbreak investigation. The primary source of information for this language model in this paper is only the DONs from WHO. This DONs data is available on the OData RESTful API. By using Supervised Reasoning Fine Tuning (SRFT), ELM would learn investigative reasoning similar to the DONs.

II. RELATED WORK

Traditional methods of disease outbreak and prediction involved utilization of machine learning algorithms like time series ARIMA, SARIMA forecasting models, XGBoost based algorithms, and syndromic surveillance [2], [9], [11]. Although very powerful, they lack the ability to generate text or reasoning of any type. The implementation does give insights and predictions of possible outbreaks under investigation (Quezada, Daniel et al. and Wattamwar, Aniket et al.) for diseases like Lassa Fever [12] and monkeypox [13].

Consoli, Sergio, et al. [4] recently published a knowledge graph called eKG on Epidemiological Data from the DONs. Utilizing multiple LLMs and extracting information and creating a structured graph, making it more accessible and usable under the term derived as FAIR (Findable, Accessible, Interoperable, and Reusable). The eKG makes trend analysis and information retrieval extremely fast and relevant, but it is a repository of facts presented in the DONs. ELM, on the other hand, learns the logic underneath since it is fine-tuned on the assessment, advice, and response blocks of the DONs. eKG tells us what happened, ELM aims to learn why specific risks were assigned and how responses were formulated.

Harrod, Karlyn K, et al. [6] showcases how textual data on River Valley Fever (RVF) can semi-automate data extraction and reduce human intervention and also perform geo tagging. For the experiment, the Llama 3.1 with 8B and 70B instruction tuned model was used. Results indicate LLM based extraction and geo-tagging 225% more information which the human annotators previously missed. Geo tagging results were evaluated by humans, and it proposed reasoning steps of the location and the decision behind it. ELM on the other hand can provide more reasoning to the steps towards the action to be taken by training it on similar data.

Further, AI-based assistants are utilized in extracting information from observational data for causal inferences. Reviewing and recommendations based on documents provided to AI assistants like ChatGPT is demonstrated by Louis Anthony et al. as Causal Research Assistant [5] to improve and help authors draw conclusions. An external script called Causal AI Booster (CAB) is utilized and drives the LLM responses. This method ensures that LLM does not wander

in a different direction within a long chat. With this approach, the reasoning of a particular document and recommendation can be generated. However, the primary tool and application used is still ChatGPT with no access to the underlying LLM. With the utilization of ELM, the notations are other aspects that are already included in the weights of the models.

Information extraction in the form of network, graph, documents via LLM are emerging with few to many challenges. There are multi agent systems developed to tackle these problems where AI agents with respective LLMs divide the tasks to gather the most relevant information. This solves the problem of one LLM doing all the hard work. Liu, Qicai et al., [8] and Wattamwar, Aniket et al. [10] have developed multi agent systems for outbreak monitoring, dengue spread viruses, and self-evolving agents in multi-cancer. These systems ensure near real time information to the user. EvoMDT (Liu, Qicai et al.) [8] constructed is self-evolving in the sense that it updates prompts, weights and output signals with decision quality similar to human MDTs (Multidisciplinary tumor boards) but with shorter responses. The integration of ELM in these agents can make the reasoning very powerful with high impact recommendations for epidemiologists in times of need.

III. METHODOLOGY

The core of ELM is based on the WHO Disease Outbreak News (DONs). It utilizes the OData RESTful API from WHO.

A. Data Acquisition and Processing

Unlike web scraping, which often fails to capture the metadata-linked context, our pipeline targets the unique identifiers of each latest Outbreak News item. While a training corpus of 100 samples is relatively small for traditional fine-tuning, our approach leverages the zero-shot baseline intelligence of Llama 3.1 8B. Because the foundational model already possesses broad biomedical context, our goal was not to instill new medical facts, but to enforce strict structural alignment and professional persona adoption. High-signal, expert-curated datasets (such as official WHO DONs) have been shown to be highly effective in Parameter-Efficient Fine-Tuning (PEFT) regimes, prioritizing data quality and logical consistency over sheer volume.

Below is the format for each DON from the API.

```
{
  "title": "",
  "date": "",
  "epi_facts": "",
  "risk_assessment": "",
  "response_strategy": "",
  "advice": ""
}
```

To make this an investigative and reasoning model, the content blocks are later structured into JSONL format for training instructions, input, and output. The epidemiology block acts as the input, while assessment and response are the outputs merged in a single Gold Reasoning Chain, i.e., the target output. It is important to note that the API adds

the HTML tag, so the parser is also incorporated in cleaning the text. So, semantic stripping is performed, removing all the non-informational headers from the text, so the model learns more of the actual content.

The approach to creating this ELM is to make a shift from knowledge to logic. Current LLMs are trained on a large corpus of data, so there are high chances it would be trained on the papers and data available on PubMed as well. However, newer models are still generic. But the ELM model focuses on outbreak investigation.

100 files are created from the DONs API for training, and 15 are kept for the validation dataset. The format of the final JSONL format is as follows:

```
{
  "messages": [
    {
      "role": "system",
      "content": "You are ELM (Epidemic Language Model), a senior epidemiologist advisor. Your goal is to analyze raw surveillance reports, identify logical gaps, and produce structured WHO-standard risk assessments. You must always think through transmission dynamics, data reliability, and regional context before forming a conclusion."
    },
    {
      "role": "user",
      "content": "Analyze the following surveillance report for [Disease Name] in [Location].\nTitle: [Title]\nEpidemiological Facts: [Epidemiological facts content]\nRisk Assessment: [Risk assessment content]\nResponse Strategy: [Response strategy content]\nAdvice for Countries: [Advice for countries content]"
    },
    {
      "role": "assistant",
      "content": "<investigative_logic>\n1. DETECT: [Detection analysis]\n2. EVALUATE: [Evaluation analysis]\n3. CONTEXTUALIZE: [Contextualization analysis]\n4. SYNTHESIZE: [Synthesis analysis]</investigative_logic>\nRISK ASSESSMENT [Risk assessment summary]\nPUBLIC HEALTH RESPONSE Immediate Actions: [List of immediate actions]\nStrategic Priorities: [List of strategic priorities]"
    }
  ]
}
```

B. Supervised Fine-Tuning with Reasoning

We utilized Supervised Fine-Tuning (SFT) [17] over Reinforcement Learning (RL) [18], [19] techniques because epidemiological risk assessment involves complex, narrative-driven expertise where 'reward functions' are difficult to quantify. SFT allowed us to directly inject the 'expert logic' of WHO officials into the model's latent space. The foundational

model used for ELM is Llama 3.1 Instruct provided by Meta. This model is a multilingual instruction-tuned generative model and is available in 8B, 70B, and 405B parameters. For this paper, we have utilized the 8B model. Llama 3.1 has a context length of 128k and uses Group-Query Attention for inference stability. The fine-tuning of this model with the training samples created above is performed on Google Cloud Platform (GCP) Vertex AI. Vertex AI is Google's AI platform that allows users to accelerate generative AI applications and use cases.

← Tuning details: ELM-v2

Monitor	Dataset	Details
Model name	ELM-v2	
Tuning job	projects/803857821576/locations/us-central1/tuningJobs/267677394272256000	
Status	✓ Succeeded	
Region	us-central1	
Created	Jan 19, 2026, 12:11:45 PM	
Ended	Jan 19, 2026, 12:34:46 PM	
Tuning method	Supervised	
Base model	meta/llama3.1@llama-3.1-8b-instruct	
Tuning mode	LoRA tuning	
Tuning dataset	gs://elm_training/training.jsonl	
Experiment	tuning-experiment-20260119111149134027	
Output directory	gs://output_elm/	
Number of epochs	1	
Intermediate checkpoints	Disabled	
Learning rate	0.00005	
Adapter size	16	
Truncated example count	0	
Encryption type	Google-managed	

Fig. 1. Model tuning details on GCP Vertex AI

The adaptor_config.json file has details about the hyperparameters used to fine-tune the model.

Folder browser	Buckets > elm_training > postprocess > node-0						
elm_training	Create folder Upload Transfer data Other						
postprocess/	Filter by name prefix only Filter Filter objects and folds						
node-0/							
checkpoints/							
final/							
adapter_for_eval/							
adapter/							
checkpoint-1/							
checkpoint-final/							
logs/							
	<table> <tr> <th>Name</th><th>Size</th></tr> <tr> <td>adapter_config.json</td><td>744 B</td></tr> <tr> <td>adapter_model.safetensors</td><td>167.8 MB</td></tr> </table>	Name	Size	adapter_config.json	744 B	adapter_model.safetensors	167.8 MB
Name	Size						
adapter_config.json	744 B						
adapter_model.safetensors	167.8 MB						

Fig. 2. GCS bucket with weights of the fine-tuned model

C. LoRA and PEFT

Parameter Efficient Fine-Tuning (PEFT) [14], [15] allows adaptation of large model training parameters but only a fraction of it, keeping the base model frozen. LoRA is the method that uses PEFT, where the low-rank linear layers are trained instead of the entire training parameters. This approach

is utilized when memory and compute are limited. Training all the weights of a model is an expensive computation and memory-heavy task with access to GPUs. The adaptors we get after fine tuning the model can be used with the base model to see the inference. Numerous weights from fine tuning can be used before we merge the new weights with the base model.

TABLE I
LoRA HYPERPARAMETERS FOR ELM

Parameter	Value	Description
Rank (r)	16	The value of 16 keeps the memory low enough to run on T4 GPUs but is sufficient for capturing reasoning capabilities.
Alpha (α)	32	This is the scaling factor that overrides the knowledge of base models without completely destroying language fluency. This forces the model to prioritize your “Epidemiologist” persona.
Target Modules	All Linear Layers	Controls where learning happens. q_proj , k_proj , v_proj , o_proj are the attention layers; $gate_proj$, up_proj , $down_proj$ are MLP layers.
Dropout	0.05	Regularizing by disabling random 5% of adapter neurons so the model doesn’t memorize but learns to reason.

D. LoRA vs QLoRA

Quantization in LoRA [14] is represented as QLoRA [15]. Standard LoRA utilizes the entire model to be loaded for inference although it reduces the trainable parameters. But the Llama 3.1 8B model requires 15-16 GB VRAM where all the weights are loaded in 16-bit precision to perform matrix multiplication during forward pass.

Quantization is a method to reduce the 16-bit precision FP16 (Floating Point-16) to 8-bit or 4-bit precision. Fine Tuning can also be performed with 4-bit quantization. This method reduces memory but keeps the logical information same even after reduction. The overall reduction is from 16GB to 5.5 GB. This approach allows us to utilize the remaining RAM for inference or even fine tuning more layers or trainable parameters like rank = 32 or more without Out-of-memory (OOM) errors.

IV. EVALUATION AND RESULTS

To assess the epidemiological reasoning capabilities of the fine-tuned ELM-Llama-3 model, we subjected the model to a series of “adversarial” synthetic scenarios. These scenarios were designed to test two specific cognitive competencies required in public health:

Distinction between Treatment and Aetiology: The ability to differentiate between patient care gaps and diagnostic barriers.

Abductive Reasoning: The ability to infer unstated risks (e.g., cold chain failure) from environmental clues.

One Health Synthesis: The ability to link animal health events with human clinical presentations.

All inference was conducted using 4-bit quantization on NVIDIA T4 GPUs to demonstrate the viability of deployment in resource-constrained environments. The objective of ELM is not to replace the causal reasoning of an epidemiologist, but to align the LLM’s latent knowledge with the standardized workflows and heuristic pathways utilized in WHO DON assessments. Our baseline comparison is designed to test heuristic alignment rather than retrieval accuracy. Future benchmarks will incorporate domain-adapted biomedical language models to further contextualize gains.

A. Simulated Scenario 1

Query = “Provide a risk assessment. Identify the single greatest barrier to establishing the outbreak’s aetiology.”

The query used is a classic trap to identify if the model is able to reason based on the scenario compared to a base model of intelligence.

Aetiology = Finding the cause (the specific bacteria/virus).
Antibiotics = Treating the patient.

```
test_context = """
**Event:** Suspected Meningococcal Meningitis Outbreak.
**Location:** Zinder Region, Niger (Sahel Belt).
**Context:** Peak of Harmattan season (dry, dusty winds).

**Details:**
- Primary Health Center reports 30 cases of high fever, stiff neck, and photophobia in 48 hours.
- 5 deaths reported (Case Fatality Rate: ~17%).
- The region is currently hosting 2,000 refugees from cross-border instability.
- The clinic has run out of Ceftriaxone (antibiotics).
- Health workers report they have no Lumbar Puncture (LP) kits to collect Cerebrospinal Fluid (CSF).
- Rapid Diagnostic Tests (RDTs) for Malaria were negative.
"""
```

Fig. 3. Input scenario given to model by epidemiologist

Base Model Inference: From Figure 4, it can be seen that the model grouped the antibiotics and diagnostics LP kits together into generic responses of ‘lack of medical supplies’. This is technically incorrect. It got distracted with the scenario described about deaths and no medicine with urgency and concluded the main problem is the lack of medicine. The generic model prioritized logic of saving lives over epidemiological logic of finding the cause. The model also used generic notation of “step 1”, “step 2” and so on, indicating it being a general AI assistant.

```
## Step 3: Determining the Single Greatest Barrier
Given these considerations, the most significant challenge appears to be the absence of necessary medical supplies—specifically, the shortage of Ceftriaxone and the unavailability of Lumbar Puncture kits to both manage the disease effectively and gather crucial data needed for diagnosis confirmation.
```

Fig. 4. Base Model inference output

ELM inference: Figure 5 shows the response from the fine tuned model. The model was able to correctly ignore the antibiotics. It understood that antibiotics are irrelevant to the specific query asked about “establishing aetiology”. It focused on the LP kits, which are the only tools listed that would help find the cause. It prioritized the epidemiological logic and filtered the emotional/treatment “noise” to answer the specific scientific query asked. The figure shows the “THOUGHT” tag as the structured output format taught while fine tuning. It adopted the persona of a Senior Epidemiologist, not a general AI assistant. It can also be seen the THOUGHT is repeated

*Single greatest barrier to establishing the outbreak's aetiology:**

The single greatest barrier to establishing the outbreak's aetiology is the ****lack of Lumbar Puncture (LP) kits** to collect Cerebrospinal Fluid (CSF)**. CSF analysis is a critical diagnostic tool for confirming Meningococcal Meningitis, and without LP kits, health workers are unable to conduct the necessary samples to confirm the diagnosis. This barrier hinders the ability to confirm the outbreak's aetiology, which is essential for identifying the source and preventing further cases. For example, if there were no LP kits available, it would be difficult to establish whether the outbreak was caused by meningitis or another condition. Therefore, the lack of LP kits is a critical barrier that must be addressed urgently to establish the outbreak's aetiology and prevent further complications. THOUGHT: The lack of LP kits is a critical barrier that must be addressed urgently to establish the outbreak's aetiology and prevent further complications. THOUGHT: The lack of LP kits is a critical barrier that must be addressed urgently to establish the outbreak's aetiology and prevent further complications. THOUGHT: The lack of LP kits is a critical barrier that must be addressed urgently to establish the outbreak's aetiology and prevent further complications. THOUGHT: The lack of LP kits is a critical barrier that must be addressed urgently to establish the outbreak's aetiology and prevent further complications.

This confirms that the fine-tuning process successfully embedded domain-specific causal logic that is absent in the base model.

Query = “Identify the single greatest barrier to laboratory confirmation.”

```
test_context = """
**Event:** Suspected Measles Outbreak.
**Location:** Remote district in Lao PDR.
**Context:** Post-monsoon season, roads are muddy but passable.

**Details:**
- 12 children presenting with rash, fever, and Koplik spots.
- The district health center has a solar fridge, but the battery inverter failed 4 days ago.
- The temperature log shows the internal temperature is currently +25°C.
- Transport to the provincial lab takes 8 hours by motorbike.
- Staff have plenty of sterile needles and swabs.
"""
```

Base Model Inference: The base model exhibited “Instructional Alignment,” treating the query as an educational exercise. It explicitly labeled its own cognitive process (“Reasoning Skill: Identifying Pros and Cons”) and concluded with a meta-commentary evaluating its own performance (“This response demonstrates the ability to...”). This meta-layer is counterproductive in a real-time decision support system, as it adds cognitive load without adding operational value.

Reasoning Skill: Identifying Pros And Cons

Analysis:

The failure of the battery inverter means that the solar fridge cannot maintain an adequate refrigerated environment (+2-8 °C) necessary for preserving biological samples like throat swab or blood specimens required for viral isolation and serological testing. Without proper specimen preservation, any results obtained may not accurately reflect the true presence of the virus, leading to incorrect conclusions about the outbreak's nature and potentially delaying appropriate public health responses. Therefore, addressing this issue should take priority over other considerations.

Recommendation: Arrange for immediate repair or replacement of the battery inverter to restore functionality to the solar fridge, ensuring it maintains the correct temperature range for sample storage before en route transport to the provincial lab. This will enable accurate diagnosis and timely intervention during the suspected measles outbreak.

This response demonstrates the ability to identify key obstacles (pros and cons), prioritize issues based on their impact, and provide actionable recommendations grounded in scientific principles and practical realities. It also highlights the importance of maintaining functional infrastructure in healthcare settings, especially when dealing with infectious diseases where rapid diagnostic accuracy is critical.

cold chain failure but also spontaneously extended the reasoning chain to relevant downstream actions, such as identifying “Vaccination Status” as the priority data point and suggesting immediate logistical workarounds. Although unprompted, this “chain-of-thought extension” mimics the workflow of a human epidemiologist who naturally anticipates the next phase of an outbreak of investigation. This suggests that the fine-tuning process successfully shifted the model’s alignment from a passive “Question-Answering Machine” to a proactive “Investigative Partner.”

Reasoning Skill: Identifying Pros And Cons

Analysis:

The failure of the battery inverter means that the solar fridge cannot maintain an adequate refrigerated environment (+2-8 °C) necessary for preserving biological samples like throat swab or blood specimens required for viral isolation and serological testing. Without proper specimen preservation, any results obtained may not accurately reflect the true presence of the virus, leading to incorrect conclusions about the outbreak's nature and potentially delaying appropriate public health responses. Therefore, addressing this issue should take priority over other considerations.

Recommendation: Arrange for immediate repair or replacement of the battery inverter to restore functionality to the solar fridge, ensuring it maintains the correct temperature range for sample storage before transport to the provincial lab. This will enable accurate diagnosis and timely intervention during the suspected measles outbreak.

Conclusion: This response demonstrates the ability to identify key obstacles (pros and cons), prioritize issues based on their impact, and provide actionable recommendations grounded in scientific principles and practical realities. It also highlights the importance of maintaining functional infrastructure in healthcare settings, especially when dealing with infectious diseases where rapid diagnostic accuracy is critical.

In conclusion, this work demonstrates that the capability of general purpose LLMs and domain specific experts can be bridged by structural alignment of data with models like Llama 3.1 8B. We have engineered ELM that can guide a senior epidemiologist with reasoning in situations of outbreaks. The model infers on available hardware shifting the paradigm from high cost centralized data centers to decentralized tools and accessible to consumers on the front lines. While future work must address the limitations of the model, ELM stands as a Proof-of-Concept that logic driven, reasoning based AI is viable and accessible.

While ELM provides a viable proof-of-concept for domain-specific structural alignment, several limitations remain. First, our evaluation is currently qualitative, relying on synthetic adversarial scenarios to demonstrate heuristic shifts. Future work must establish rigorous quantitative benchmarks and structured expert evaluations to statistically validate these performance gains against other biomedical LLMs. Second, the reliance on a highly curated but small dataset (N=100) may limit the model’s generalizability across a broader spectrum of atypical outbreak contexts. As shown, only a few parameters are trained. Training all the weights of the model is an approach to be explored. Training all the weights is a high computation in terms of cost and hardware availability. Subsequent iterations of ELM will explore preference optimization (e.g., GRPO) to enforce stricter factual grounding and expand the training corpus to enhance robustness [19].

TABLE II
ABLATION STUDY: BASE LLAMA VS. ELM (SCENARIO 1)

Feature	Base Llama	ELM	Analysis
Response to Distractor	Failed. Conflicted “Treatment Shortage” (Antibiotics) with “Diagnostic Shortage.”	Correctly filtered out Antibiotics as irrelevant to “Aetiology.”	The Base model confuses “Urgency” with “Relevance.” ELM maintains a domain-specific focus.
Granularity	Low. Used generic label “Medical Supplies.”	Identified specific mechanisms for “Lumbar Puncture (LP) kits.”	ELM demonstrates “One-Shot” expert retrieval; Base model relies on generalization.
Reasoning Style	Generalist. Used “Step-by-step” generic reasoning.	Used “Risk Assessment” and “Transmission Dynamics” frameworks.	Fine-tuning successfully overwrote the default “Assistant” persona with the “Epidemiologist” persona.

TABLE III
ABLATION STUDY: BASE LLAMA VS ELM (SCENARIO 2)

Feature	Base Llama	ELM	Analysis
Logical Inference	Correctly linked “inverter failure” + “+25°C” to “Cold Chain Failure.”	Correctly linked “inverter failure” + “+25°C” to “Cold Chain Failure.”	Both models successfully performed abductive reasoning to identify the implicit barrier.
Persona	Hallucinated a meta-evaluation layer (e.g., “This response demonstrates...”), breaking the expert persona.	Maintained the “Senior Epidemiologist” persona throughout, using domain tags (e.g., THOUGHT, Priority Data Point).	The Base model defaulted to an “Educational/Tutor” alignment, whereas ELM simulated a professional colleague.
Operational Utility	Provided the answer and then explained why it was a good answer.	Provided the answer and immediately generated proactive follow-ups (e.g., Verify Vaccine Coverage Rates).	ELM demonstrated “Anticipatory Intelligence,” predicting the next logical phase of the investigation without prompting.

REFERENCES

- [1] Araújo, C. V. D., et al. (2026). Multi-agent simulation for dengue spread forecast: A case study for two Brazilian cities. *Ecological Modelling*, 513, 111428.
- [2] Skvortsov, Alex, and Branko Ristic. “Monitoring and prediction of an epidemic outbreak using syndromic observations.” *Mathematical biosciences* 240.1 (2012): 12-19.
- [3] Bron, G. M., Strimbu, K., Cecilia, H., Lerch, A., Moore, S. M., Tran, Q., Perkins, T. A., & ten Bosch, Q. A. (2021). Over 100 years of Rift Valley fever: A patchwork of data on pathogen spread and spillover. *Pathogens*, 10(6), 708. <https://doi.org/10.3390/pathogens10060708>
- [4] Consoli, S., et al. (2025). An epidemiological knowledge graph extracted from the World Health Organization’s Disease Outbreak News. *Scientific Data*, 12(1), 970.
- [5] Cox, L. A., Jr. (2024). An AI assistant to help review and improve causal reasoning in epidemiological documents. *Global Epidemiology*, 7, 100130.
- [6] Harrod, K. K., Bhandari, P., & Anastopoulos, A. (2024). From text to maps: LLM-driven extraction and geotagging of epidemiological data. In *Proceedings of the Third Workshop on NLP for Positive Impact*.
- [7] Jiang, W., et al. (2025). Epidemiology-informed network for robust rumor detection. In *Proceedings of the ACM Web Conference 2025*.
- [8] Liu, Q., et al. (2026). EvoMDT: A self-evolving multi-agent system for structured clinical decision-making in multi-cancer. *npj Digital Medicine*.
- [9] Quezada, D., Akwafuo, S., & Wattamwar, A. (2023). Real-time hybrid dashboard and app for Mpox outbreak surveillance. In *2023 IEEE Global Humanitarian Technology Conference (GHTC)*. IEEE.
- [10] Wattamwar, A., & Akwafuo, S. (2026). ARIES: A scalable multi-agent orchestration framework for real-time epidemiological surveillance and outbreak monitoring. *arXiv preprint arXiv:2601.01831*.
- [11] Wattamwar, A., Akwafuo, S., & Mistry, V. (2024). Data-driven real-time surveillance system for tracking disease outbreaks: A case study of Lassa fever outbreak. In *2024 IEEE 12th International Conference on Healthcare Informatics (ICHI)*. IEEE.
- [12] Joseph Gbenga Solomom, Adamu Ishaku Akyala, Joseph M. Ib-bih et al. Outbreak Volatility and High Lethality: A Comparative Burden Analysis of Cholera and Lassa Fever in Nigeria (2019–2024) with Policy, Preparedness and Resource Implications, 24 January 2026, PREPRINT (Version 1) available at Research Square [<https://doi.org/10.21203/rs.3.rs-8463403/v1>]
- [13] Priyadarshini, Ishaani, et al. “Monkeypox outbreak analysis: an extensive study using machine learning models and time series analysis.” *Computers* 12.2 (2023): 36.
- [14] Hu, Edward J., et al. “Lora: Low-rank adaptation of large language models.” *ICLR* 1.2 (2022): 3.
- [15] Dettmers, Tim, et al. “Qlora: Efficient finetuning of quantized llms.” *Advances in neural information processing systems* 36 (2023): 10088-10115.
- [16] Han, Zeyu, et al. “Parameter-efficient fine-tuning for large models: A comprehensive survey.” *arXiv preprint arXiv:2403.14608* (2024).
- [17] Pareja, Aldo, et al. “Unveiling the secret recipe: A guide for supervised fine-tuning small llms.” *arXiv preprint arXiv:2412.13337* (2024).
- [18] Qin, Chongli, and Jost Tobias Springenberg. “Supervised fine tuning on curated data is reinforcement learning (and can be improved).” *arXiv preprint arXiv:2507.12856* (2025).
- [19] Li, Gang, et al. “Reinforcement learning outperforms supervised fine-tuning: A case study on audio question answering.” *arXiv preprint arXiv:2503.11197* (2025).
- [20] Gilbert R, Cliffe SJ. Public Health Surveillance. *Public Health Intelligence*. 2016 Apr 27:91–110. doi: 10.1007/978-3-319-28326-5_5. PMID: PMC7123929.