

Health Facility Location in Ethiopia: Leveraging LLMs to Embed Human-AI Alignment in Algorithmic Planning

Yohai Trabelsi*

Department of Data Science and AI,
School of Computer Science,
Ariel University
Ariel, Israel
yohai.trabelsi@gmail.com

Guojun Xiong

John A. Paulson School of
Engineering and Applied Sciences,
Harvard University
Cambridge, MA, United States
xiongj1@gmail.com

Fentabil Getnet

National Data Management and
Analytics Center for Health,
Ethiopian Public Health Institute
Addis Ababa, Ethiopia
b.infen4ever@gmail.com

Stéphane Verguet

Department of Global Health and
Population, Harvard T.H. Chan
School of Public Health
Boston, MA, United States
verguet@hsph.harvard.edu

Milind Tambe

John A. Paulson School of
Engineering and Applied Sciences,
Harvard University
Cambridge, MA, United States
milind_tambe@harvard.edu

ABSTRACT

Ethiopia’s Ministry of Health is upgrading health posts to improve access to essential services, particularly in rural areas. However, limited resources necessitate careful prioritization of which facilities to upgrade in order to maximize population coverage while accounting for diverse expert and stakeholder preferences. To address this challenge, we propose the LEG framework: a hybrid approach that systematically integrates expert knowledge with optimization techniques. This framework combines a provably approximable algorithm for population coverage optimization with LLM-driven iterative refinement that incorporates human-AI alignment, ensuring solutions reflect expert qualitative guidance while preserving coverage guarantees. In collaboration with the Ethiopian Public Health Institute and Ministry of Health, we evaluate the LEG framework using human expert advice, demonstrating the framework’s effectiveness and its potential to inform equitable, data-driven health system planning.

KEYWORDS

Health Facility Location, Optimization, Human expert knowledge, Human-AI alignment, LLM

ACM Reference Format:

Yohai Trabelsi, Guojun Xiong, Fentabil Getnet, Stéphane Verguet, and Milind Tambe. 2026. Health Facility Location in Ethiopia: Leveraging LLMs to Embed Human-AI Alignment in Algorithmic Planning. In *Proc. of the 25th International Conference on Autonomous Agents and Multiagent Systems (AAMAS 2026)*, Paphos, Cyprus, May 25 – 29, 2026, IFAAMAS, 9 pages.

DISCLOSURE

This paper substantially overlaps with our submission to AAMAS 2026 [20] (same authors; accepted as a full paper). The primary

*Work conducted while the author was affiliated with Harvard University, USA.



Figure 1: Basic health post (left); Design of a comprehensive health post that provides more essential services (right). Source: Ministry of Health, Ethiopia.

additional contribution of this workshop paper lies in the collection and incorporation of human expert advice (Sections 5, and 6).

1 INTRODUCTION

Ensuring equitable access to essential health services remains a central challenge for Ethiopia’s Ministry of Health (MOH). Since the launch of the Health Extension Program (HEP) in 2003/04 [22], the country has made significant progress in expanding basic care delivery to rural populations. The HEP Optimization Roadmap (2020–2035)[15] further envisions upgrading selected health posts into comprehensive ones that can provide advanced services such as childbirth and postnatal care (see Figure 1). However, upgrading facilities is expensive, and the available public budget is severely constrained [5]. Determining which facilities to upgrade thus becomes a complex optimization problem involving limited resources, heterogeneous population needs, and diverse stakeholder opinions.

Recent studies have applied optimization and geospatial methods to identify locations where comprehensive facilities are needed [3, 5, 8]. These efforts primarily focus on maximizing population coverage under distance and capacity constraints (See Figure 2). Yet, in practice, the final allocation decisions in the stepwise construction of comprehensive health posts remain dominated by expert judgment and stakeholder negotiation, rather than by algorithmic outputs. While human expertise captures important contextual knowledge—such as terrain accessibility or local socio-political considerations—it is often expressed in natural language and difficult to

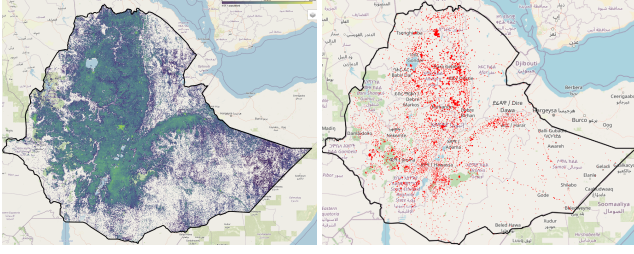


Figure 2: ([5]) Map of Ethiopia overlaid with projected population estimates for 2026 (log scale). Brighter yellow areas on the left indicate higher population densities. The map on the right shows the locations of health facilities capable of providing essential health services (red points). Many populations currently lack 2-hour access to such facilities [10].

encode into a mathematical objective function [6, 26]. It is possible that even the advisors themselves may struggle to provide their preferences as a fully coherent set of recommendations.

Recent advances in Large Language Models (LLMs) [4, 26] offer a promising avenue for bridging the gap between qualitative human judgment and quantitative optimization. LLMs can interpret and structure unformalized expert advice, enabling algorithms to incorporate contextual and domain-specific reasoning that traditional models often overlook. Nevertheless, these methods typically lack formal theoretical guarantees on performance or stability—an issue that is especially critical in high-stakes domains such as public health or infrastructure planning. In the absence of such guarantees, language-based systems risk producing allocations that appear reasonable linguistically but fail to satisfy fairness, transparency, or policy-alignment criteria required for real-world adoption.

Our LEG framework [20] addresses this gap by coupling algorithmic optimization with language-based expert reasoning. Building upon classical submodular maximization, it provides provable coverage guarantees while leveraging LLMs to interpret and iteratively incorporate human advice expressed in natural language. The LLM serves as a bridge between formal optimization and informal domain knowledge, translating verbal recommendations into structured allocation adjustments that preserve theoretical performance bounds. This integration allows the system to remain both rigorous and human-aligned.

We demonstrate the framework in collaboration with experts from the Ethiopian Public Health Institute, and from the Ministry of Health, focusing on the Somali region. Empirical results show that LLM-guided iterations significantly improve the alignment between algorithmic allocations and expert advice, while maintaining high coverage efficiency.

1.1 Process description

In Figure 3, we present the LEG framework at a high level. The process begins with the problem inputs, an initial greedy allocation, and a list of advice sentences. These inputs are provided to a large language model (e.g., Gemini, ChatGPT), which iteratively refines the allocation strategy. The model proposes an allocation that is optimized using a guided greedy procedure parameterized by (α, β) .

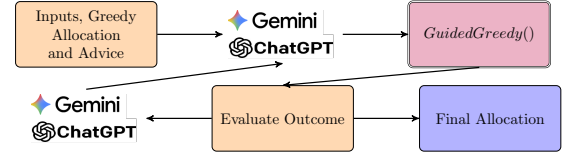


Figure 3: Overview of the proposed method.

The resulting outcome is then evaluated, and the evaluation feedback is transformed into a new prompt for further refinement. This iterative loop continues until the maximum number of iterations is reached, producing the final optimized allocation.

1.2 Contributions

In collaboration with the Ethiopian Public Health Institute and Ministry of Health, this work implements the unified LEG framework for upgrading health facilities in Ethiopia that jointly optimizes population coverage and alignment with expert guidance. Our main contributions, beyond the framework presented in [20], are as follows:

- Implementing the LEG framework [20] in realistic settings.
- Constructing abstract and concrete advice sentences and refining them through analysis by Ethiopian experts.
- *Real-world validation in Ethiopia.* Evaluating the LEG framework and the human expert advice sentences through numeric experiments in Ethiopia’s Somali region, highlighting their potential for real-world deployment.

2 RELATED WORKS

In this section we focus on applicative related works, that emphasize real world applications and aspects. More related work on relevant approaches and their theoretical foundations can be found in [20].

2.1 Health facility location

Many studies have addressed health facility location problems (see, e.g., the survey by [2]). [21] examined the allocation of emergency health facilities while incorporating expert preferences into their model. [7] examined the placement of health facilities in the Philippines, explicitly considering trade-offs between equity and efficiency in their optimization approach. [18] optimized clinic locations in Kuala Langat, Malaysia, to maximize population coverage within 3–5 km. Other studies, such as [3, 5], focused on optimizing health facility locations in Ethiopia. However, most of these studies do not account for alignment with human expert preferences.

2.2 Human-AI alignment

Recent work has focused on aligning AI decisions with human expertise (see [11, 19] for surveys) and aligning LLMs to better suit human-oriented tasks and expectations (see [23] for a survey). Our work contributes to the literature on aligning algorithmic decision-making with human preferences (e.g., [1, 14, 26]). However, whereas most existing approaches rely on the explicit construction of a numerical reward function, which is not always practical, our method

enables LLMs to directly guide the allocation process, bypassing the need for an intermediate reward-function construction step.

2.3 Human Agent Interaction

[17] proposed a multiagent system for training incident commanders in large-scale disasters, observing that in some scenarios, following human advice actually degraded system performance. To address this, [16] introduced a method for resolving conflicts between humans and agents that improved performance even when human advice conflicted with agents’ preferred decisions. More recently, [27] surveys approaches in which LLM-based human-agent systems incorporate human-provided information, feedback, or control. Building on this line of work, we demonstrate how LLMs can facilitate alignment with human advice while maintaining system base utility.

3 SETTINGS AND PROBLEM FORMULATION

3.1 Settings

Let $r \in \mathbb{N}_{>0}$ denote the number of districts. Let the ground set be $V = T_1 \uplus \dots \uplus T_r$, where each T_i represents the set of candidate grid cells in district i , and $\uplus$ denotes a disjoint union (See the example in Section 1.1). For simplicity, we assume that all grid cells in these districts are candidates for locating a facility.¹ A **grid-cell allocation** is any subset $S \subseteq V$.

For any grid-cell allocation $S \subseteq V$, we define the *district allocation* $h(S) \in \mathbb{N}^r$ by

$$h_i(S) = |S \cap T_i|, \quad \text{for } i = 1, \dots, r,$$

so that $h_i(S)$ represents the number of facilities placed in district i . We denote by b the total budget, i.e., the number of health posts to be upgraded in a given year.

Coverage function. Let $\text{covered}(S)$ be the set of grid cells served by the facilities in S . The total coverage associated with S is defined as

$$f(S) = \sum_{c \in \text{covered}(S)} w_c, \quad (1)$$

where w_c denotes the population of grid cell c . This monotone submodular function quantifies the cumulative population covered by the selected facilities.

Advice alignment function. Because the notion of “alignment with expert advice” is inherently qualitative, defining a corresponding numerical objective is nontrivial. We represent this alignment through a function $g : 2^V \rightarrow \mathbb{R}$ that assigns a score to each grid-cell allocation, reflecting how well the allocation adheres to the provided human or LLM-generated advice. (Here, 2^V denotes the set of all possible subsets of V .)

For evaluation, we approximate this alignment function using a set of advice sentences A . For each advice statement $a \in A$, we define an auxiliary scoring function $g_a : 2^V \rightarrow \mathbb{R}$ that measures the extent to which the advice a is satisfied. An LLM is used to define

these component functions and aggregate them into an overall alignment metric:

$$g_{\text{eval}}(S) = \sum_{a \in A} g_a(h(S)). \quad (2)$$

This construction serves as a proxy for human evaluation, enabling language-based advice to be incorporated into a quantitative optimization framework.

3.2 Problem Formulation

The multi-objective problem. Given the functions f and g , and a total budget b , we define the following multi-objective formulation that jointly optimizes *population coverage* and *alignment with expert advice*:

$$\max_{S \subseteq V, |S|=b} \{f(S), g(S)\}. \quad (3)$$

This formulation captures the fundamental trade-off between quantitative performance (coverage) and qualitative consistency (alignment).

The formulation in (3) defines an idealized bi-objective optimization problem that simultaneously maximizes population coverage $f(S)$ and alignment with expert advice $g(S)$. However, since the problem is intractable (there is no natural way to scalarize the two objectives) it is often necessary to provide guarantees with respect to the more measurable coverage objective. To make the problem operational, we introduce a constrained bi-objective relaxation that preserves theoretical guarantees on coverage while allowing flexible adaptation to both coverage and advice alignment.

The α - β guarantee problem. To provide theoretical guarantees while maintaining flexibility in alignment, we introduce two control parameters $\alpha, \beta \in [0, 1]$. Let OPT_b denote the optimal allocation of size b that maximizes the coverage function f . We then define the constrained optimization problem as:

$$\begin{aligned} \max_{S \subseteq V, |S|=b} \{f(S), g(S)\}, \\ \text{s.t. } f(S) \geq (1 - e^{-\alpha\beta}) f(OPT_b). \end{aligned} \quad (4)$$

The constraint ensures that the final allocation retains at least a fraction $(1 - e^{-\alpha\beta})$ of the optimal coverage value. By adjusting α and β , one can explicitly control the balance between maintaining strong theoretical coverage guarantees and achieving closer alignment with expert guidance. The parameter α defines the proportion of selections made by the greedy approach relative to the LLM-based approach, while β controls the greedy approach’s flexibility to deviate from the LLM allocation. The detailed operational meaning of these parameters is discussed further in Section 4.

REMARK 1 (CHALLENGES OF MULTI-OBJECTIVE OPTIMIZATION). *Classical scalarization methods—such as weighted sums or Pareto front exploration [9, 12]—require predefined numerical weights for each objective, which are often unavailable when stakeholder preferences are expressed in natural language. Moreover, these approaches treat objectives as static, whereas expert opinions may evolve iteratively during the planning process. Our formulation differs by explicitly integrating language-based feedback within the optimization loop, thereby allowing the trade-off between coverage and alignment*

¹This assumption simplifies exposition; our framework remains valid even if only a subset of cells within each district are eligible, as the optimization operates on any predefined candidate set.

to be dynamically adjusted while preserving provable performance guarantees.(see [20] for details on the provable guarantees.)

4 PROPOSED METHOD

This section details the LEG framework; its guarantees are in [20].

4.1 Algorithm Overview

Our framework integrates optimization algorithms with LLM-driven reasoning to balance two competing objectives: maximizing population coverage and aligning with qualitative expert advice as in Eq. (3). The overall process proceeds iteratively across five stages (Algorithm 1), alternating between algorithmic updates that ensure theoretical guarantees and language-based refinements that incorporate human guidance. At a high level, it contains 5 steps, i.e., Step 1: A classical greedy algorithm produces an initial coverage-maximizing (low level) grid-cell allocation. Step 2: Given the grid-cell allocation, an LLM produces a (high-level) district allocation using expert advice expressed in natural language. Step 3: A constrained greedy procedure refines the district level allocation into a grid-cell allocation while maintaining an α - β coverage guarantee. Step 4: The LLM receives structured feedback to improve alignment. Step 5: The prompts themselves are iteratively optimized through a form of “prompt gradient descent.” We detail each step below.

Step 1: Initial grid-cell Allocation (Greedy in line 3 of Algorithm 1). We first allocate facilities to grid-cells using the standard greedy algorithm of Nemhauser et al. [13], which provides a $(1 - 1/e)$ -approximation for maximizing monotone submodular functions. The resulting allocation of facilities to grid-cells, S_0 , serves as the initial baseline. From this grid-cell allocation, we extract the district allocation $h(S_0)$, which encodes how many facilities are assigned to each district (excluding specific cell details), and will be used to guide the subsequent district-level updates.

Step 2: LLM-Powered district allocation (line 6 of Algorithm 1). Next, we refine the district allocation $d = h(S)$ using an LLM that processes both quantitative and qualitative contextual information. The model considers the current and previous district allocations, budget, population statistics, and a set of expert advice sentences A. To maintain stability, we restrict the LLM to modify the districts of at most two facilities per iteration.

This step produces a revised district allocation that respects both expert intent and practical feasibility.

Step 3: GuidedGreedy (line 7 of Algorithm 1; Algorithm 2). We consider the district allocation d , as a vector of per-district budget and perform a guided greedy selection (Algorithm 2) to determine specific facility locations. In line 3 of Algorithm 2, we initialize the set S as an empty set. In the loop on lines 4–12 of Algorithm 2, the algorithm greedily adds grid-cells to the set S . On lines 5–6, we compute the cells with the maximum marginal gains, either ignoring or considering the district allocation d , respectively. The condition in line 8 ensures that at least $\lceil \alpha b \rceil$ cells attain a marginal gain of at least β times the maximum marginal gains computed with no restrictions (i.e., ignoring d). A running example demonstrating Algorithm 2 is given in Table 1.

Returning to Algorithm 1, after obtaining the set S_i , we compute in line 7, the values $f(S_i)$ and per-district contributions $f(S_i \cap T_j)$

for all $j \leq r$. These values constitute the quantitative feedback signals, Δf , and Δh which are then used to guide the next LLM iteration.

Table 1: Running example of Algorithm 2, inputs: $\alpha = 0.25$, $\beta = 0.5$, $d = \{1 : 3, 2 : 0\}$; columns: *Iter* = algorithm iteration, $|S|$ = size of grid-cell allocation at iteration start, d = per district budgets at iteration start, $c(L5)$ and $c_d(L6)$ = values in lines 5 and 6, *L8 Check* = condition outcome (line 8), *Action* = selected action (line 9 or 12).

Iter	$ S $	d	$c(L5)$	$c_d(L6)$	L8 Check	Action
1	0	$\{1 : 3, 2 : 0\}$	10	4	$0 > 0?$ No $4 \geq 0.5 \cdot 10?$ No	Pick $c(10; T_2)$
2	1	$\{1 : 3, 2 : 0\}$	8	4	$1 > 0?$ Yes	Pick $c_d(4; T_1)$
3	2	$\{1 : 2, 2 : 0\}$	8	3	$2 > 0?$ Yes	Pick $c_d(3; T_1)$

Step 4: Verbal Reinforcement via scalar, and vectors’ Comparison. (lines 8-9 of Algorithm 1). We quantify the change in both coverage and allocation between iterations by computing:

$$\Delta f = f(S_{i+1}) - f(S_i), \quad \Delta h = h(S_{i+1}) - h(S_i). \quad (5)$$

These differences form the basis of a new prompt to the LLM, requesting verbal feedback to improve future alignment. Specifically, the *Feedback* in the prompt instructs the LLM to (i) describe the observed differences and (ii) propose adjustments that enhance both alignment with expert advice and overall coverage.

Step 5: Prompt optimization. (line 10 of Algorithm 1). Following [26], we replace conventional parameter-space gradient descent with an analogous update in the prompt space. We decompose each prompt into a static and editable part:

$$\text{Prompt}_i = P_{\text{Fix}} \parallel P_{\text{Editable},i}, \quad (6)$$

where P_{Fix} encodes the task template and $P_{\text{Editable},i}$ accumulates iteration-specific refinements. After receiving new feedback, we update

$$P_{\text{Editable},i+1} = P_{\text{Editable},i} + \text{Feedback}_{i+1}, \quad (7)$$

yielding the new prompt $\text{Prompt}_{i+1} = P_{\text{Fix}} \parallel P_{\text{Editable},i+1}$. This procedure gradually improves prompt quality and alignment consistency over time. The high level of the suggested approach is given in Algorithm 1.

REMARK 2 (A HEURISTIC IMPROVEMENT FOR $\beta = 1.0$). When $\beta = 1.0$, we use the following heuristic to improve the performance of Algorithm 1: before line 6, we run the Greedy algorithm to allocate $\lceil \alpha \cdot b \rceil$ facilities, and provide this allocation to the LLM as a part of the prompt in line 6 (the greedy allocation is serving as a minimum district allocation). We then skip line 7 and continue with line 8. This approach allows the LLM to leverage the greedy selection while preserving the same theoretical guarantee.

4.2 An example

In Figure 4, we present a toy running example of our approach. The example considers a region consisting of two districts, each containing two grid-cells. Each grid-cell represents a $1 \text{ km} \times 1 \text{ km}$ area on a map of Ethiopia, and our approach focuses on selecting

Algorithm 1 LLM-Enhanced District Resource Allocation

```

1: Input: Budget  $b$ , advice set  $A$ , balance parameters  $\alpha, \beta \in [0, 1]$ ,
   available grid-cells  $V$ , prompts  $Prompt_1, Prompt_{reflection}$ .
2: Output: Grid-cell allocation  $S_{limit}$ .
3: (Init)  $S_0 \leftarrow Greedy(b, V)$ 
4: (Init) Update  $Prompt_1$  with  $S_0, h(S_0)$  and the other inputs.
5: for  $i = 1$  to limit do
6:   (LLM)  $d \leftarrow LLM(Prompt_i)$ 
7:   (Solve)  $S_i \leftarrow GuidedGreedy(\alpha, \beta, b, V, d)$ .
8:   (Compare) Compute differences  $\Delta f, \Delta h$  (Eq. 5) and update
        $Prompt_{reflection}$ .
9:   (Verbal Feedback) Query LLM to generate verbal reflection.
10:  (Optimize Prompt) Update  $Prompt_{i+1}$ , based on Eq. (6), (7).
11: Return:  $S_{limit}$ 

```

Algorithm 2 GUIDEDGREEDY(α, β, b, V, d)

```

1: Input: Parameters  $\alpha, \beta \in [0, 1]$ ; budget  $b \in \mathbb{N}_{>0}$ ; available cells
    $V$ ; district budgets  $d \in \mathbb{N}^r$ .
2: Output: Allocation  $S$ .
3:  $S \leftarrow \emptyset$ 
4: while  $|S| < b$  do
5:    $c \leftarrow \arg \max_{v \in V} (f(S \cup \{v\}) - f(S))$ 
6:    $c_d \leftarrow \arg \max_{v \in \bigcup_{j: d_j > 0} T_j} (f(S \cup \{v\}) - f(S))$ 
7:   Let  $j_d$  be the district such that  $c_d \in T_{j_d}$ 
8:   if  $|S| > \lceil \alpha b \rceil$  or  $f(S \cup \{c_d\}) - f(S) \geq \beta \cdot (f(S \cup \{c\}) - f(S))$ 
       then
9:      $S \leftarrow S \cup \{c_d\}$ 
10:     $d_{j_d} \leftarrow d_{j_d} - 1$ 
11:   else
12:      $S \leftarrow S \cup \{c\}$ 
13: return  $S$ 

```

districts and cells for building health facilities. Given an advice sentence, we first apply the algorithm of [13] to obtain a grid-cell allocation. We then invoke the LLM to produce a district allocation, incorporating this grid-cell allocation among other factors. Based on this district allocation, we run the GuidedGreedy algorithm to compute a refined grid-cell allocation. Verbal reinforcement and prompt optimization are used to improve the prompt, and the updated prompt is supplied to the LLM in the subsequent iteration. This procedure repeats until a final grid-cell allocation is obtained.

5 HUMAN EXPERT ADVICE

5.1 Abstract advice factors

To generate the advice, we first consulted officials from the Ethiopian Public Health Institute and the Ministry of Health to identify the key abstract factors relevant to decision-making on facility placement. After establishing this set of factors, we worked with experts from the Public Health Institute to refine the advice, making it more coherent and concrete so that it can be applied to real data from the Somali region (see Section 5.2).

The advice factors identified by health officials fall into three categories: efficiency, equity, and feasibility. Below, we present the abstract factors for each category.

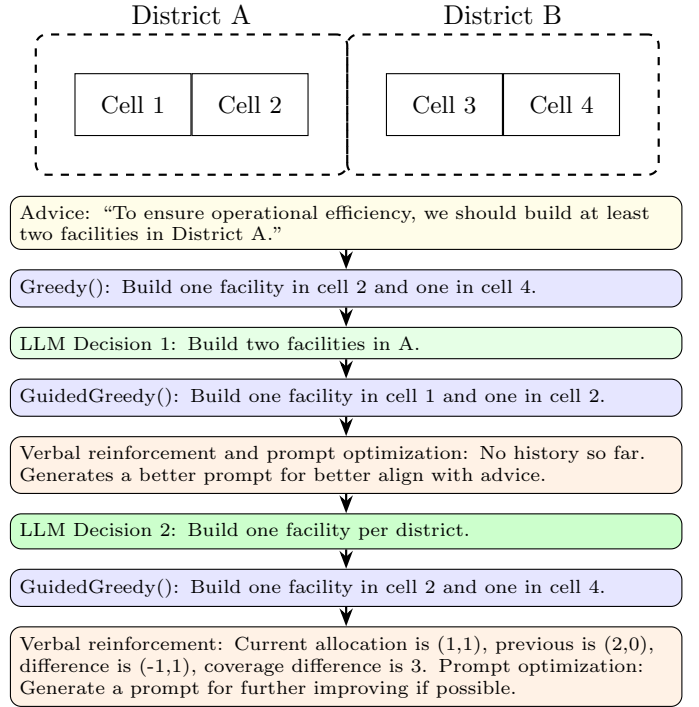


Figure 4: Toy running example with two districts and four cells; For GuidedGreedy, we use the parameters $\alpha = 0, \beta = 1$ allowing the LLM to fully determine the district allocation. Greedy() denotes the greedy algorithm of [13].

Efficiency

- (1) Prioritize populated areas for maximizing impact/efficiency.
- (2) Prioritize zones with a low facility/population ratio.
- (3) Prioritize districts close to international borders to enhance health security through timely detection and response to cross-border epidemics.

Equity

- (1) Ensure the facilities are broadly distributed to promote geographic equity.
- (2) Give precedence to remote and underserved zones.
- (3) Prioritize districts with limited access to roads and to public transport.
- (4) Favor zones or districts with low service coverage (e.g., in-facility and post-natal coverage).
- (5) Consider fair facility distribution across districts within regions and similar population settlement distributions.

Feasibility

- (1) Do not build multiple facilities in one district.
- (2) Factor in climate risks (e.g., flood zones, drought exposure).

5.2 Concrete advice for the Somali region

While some of the advice factors were reasonably implementable, others were rather vague, or even conflicting, making their practical implementation unclear. Therefore, the abstract factors were revised to clarify their meaning in the context of the Somali region. For example, the advice “Prioritize zones with a low facility-to-population ratio” was revised to “Assign 30% of the facilities to zones with a low facility-to-population ratio.” This makes explicit the extent to which such zones should be prioritized. Nevertheless, the set of qualifying zones must still be determined by our framework. We emphasize that the provided advice sentences, are intended to be adjusted and applied for each fiscal year separately, providing flexibility to planners.

In the following lists, each advice sentence corresponds to the abstract advice factor with the same index. These concrete advice sentences serve as input to our framework. Note that our framework is agnostic to the specific advice provided and can accommodate different advice sets.

Concrete Advice: Efficiency

- (1) Promote the construction of facilities in more densely populated districts.
- (2) Assign 30% of the facilities to zones with low facility-to-population ratios.
- (3) Prioritize districts close to international borders by assigning 10% of the facilities to them.

Concrete Advice: Equity

- (1) Distribute the budget fairly across geographic areas, ensuring no more than 20% of the facilities are assigned to a single zone.
- (*) Assign 40% of the facilities to Riverine/Agro-Pastoral zones, and 60% to traditional Pastoral zones.
- (2) Assign 30% of the facilities to remote and underserved zones.
- (3) Prioritize districts with limited access to roads and to public transport.
- (4) Prioritize zones or districts with low service coverage (e.g., in-facility and postnatal) by allocating a reasonable fraction of the facilities to them.
- (5) Ensure equitable distribution of facilities across districts within zones with similar population patterns.

Concrete Advice: Feasibility

- (1) Do not build multiple facilities in one district.
- (2) Account for flood risk by assigning a fair number of facilities to districts near the Shabelle river.
- (*) Account for drought risk by allocating an appropriate share of facilities to districts that are prone to such conditions.

6 EXPERIMENTS

6.1 Experimental setup

We evaluate the proposed LEG framework using Ethiopia’s projected 2026 population data [25]. Walking accessibility is computed following [5] with the global friction map [24], assuming a maximum two-hour walking distance. The experiments focus on the Somali region, a sparsely populated area that could benefit from improved health services. We randomly selected 10 sets of six sentences from the 12 concrete advice sentences and used these in our optimization and evaluation. All iterative allocation and feedback steps were executed using Gemini-2.5-Flash as shown in Figure 3.

For Algorithm 1, we set the parameters to $b = 10$, $\alpha = 0.5$, and $\beta = 1.0$. These values were chosen empirically to balance interpretability of the resulting allocation with sufficient discriminative power to meaningfully distinguish the LEG framework from the baselines (see Section 6.2).

6.2 Evaluation baselines

To evaluate our framework, we used two baselines. The first baseline follows a purely greedy strategy and completely ignores the provided advice (*Only Coverage*). In terms of Algorithm 1, we simply return the set S_0 computed in line 3. The second baseline follows an approach similar to Algorithm 1, but selects grid-cell allocations randomly instead of greedily (*Only Alignment*), as summarized in Algorithm 3. The key difference is in line 7, where the GUIDEDRANDOM(5) subprocedure constructs the set S_i by iteratively selecting b grid cells according to the district allocation vector d . This subprocedure treats d as a per-district budget: at each iteration, it randomly selects a cell whose associated district has a positive remaining budget, then decrements that per-district’s budget, d , accordingly.

Algorithm 3 LLM-Enhanced District Resource Allocation

- 1: **Input:** Budget b , advice set A , balance parameters $\alpha, \beta \in [0, 1]$, available grid-cells V , prompts $Prompt_0, prompt_{reflection}$
 - 2: **Output:** Grid-cell allocation S_{limit} .
 - 3: (Init) $d \leftarrow LLM(Prompt_0)$
 - 4: (Init) Update $Prompt_1$ with d , and the other inputs.
 - 5: **for** $i = 1$ to limit **do**
 - 6: (LLM) $d \leftarrow LLM(Prompt_i)$
 - 7: (Solve) $S_i \leftarrow GuidedRandom(b, V, d)$
 - 8: (Compare) Compute differences Δh (Eq. 5) and update $Prompt_{reflection}$.
 - 9: (Verbal Feedback) Query LLM to generate verbal reflection.
 - 10: (Optimize Prompt) Update $Prompt_i$, based on Eq. (6), (7).
 - 11: **Return:** S_{limit}
-

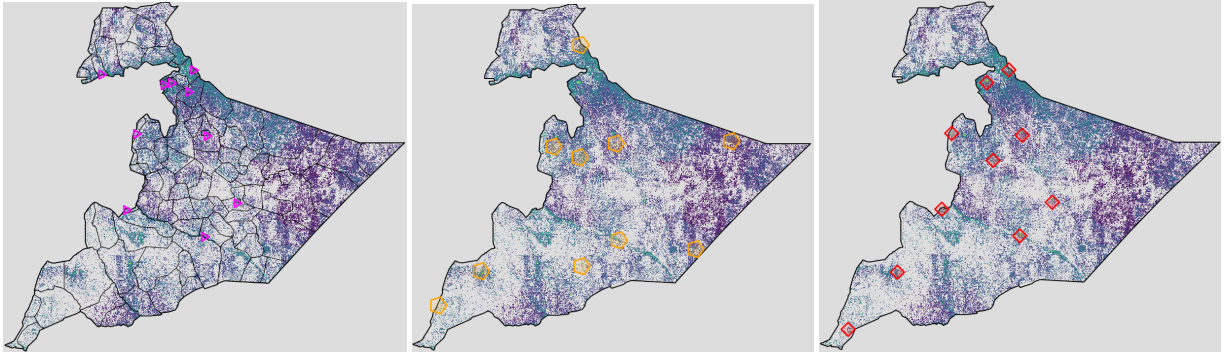


Figure 5: Only Coverage (left); Only Alignment (middle); LEG(right; $\alpha = 0.5$ and $\beta = 1.0$). Brighter areas indicate higher population density.

Algorithm 4 GUIDEDRANDOM(b, V, d)

- 1: **Input:** Budget $b \in \mathbb{N}_{>0}$; available cells V ; district budgets $d \in \mathbb{N}^r$.
 - 2: **Output:** Allocation S .
 - 3: $S \leftarrow \emptyset$
 - 4: **while** $|S| < b$ **do**
 - 5: $c_d \sim \text{Uniform}\left(\bigcup_{j:d_j>0} T_j\right)$
 - 6: $S \leftarrow S \cup \{c_d\}$
 - 7: Let j_d be the district such that $c_d \in T_{j_d}$
 - 8: $d_{j_d} \leftarrow d_{j_d} - 1$
 - 9: **return** S
-

In Figure 5, we present three maps of the Somali region showing the results of the two baseline methods and the implemented framework (LEG). Under the pure greedy approach, most facilities are concentrated in highly populated districts. In contrast, the LLM-based method distributes facilities more uniformly across the region. The LEG framework exhibits a balanced behavior, combining reasonable coverage of densely populated areas with improved allocation in remote districts.

Specifically, one concrete advisory sentence used as input to the Only Alignment and LEG frameworks is: “Distribute the budget fairly across geographic areas, ensuring that no more than 20% of the facilities are assigned to a single zone” (Equity, 1). This sentence effectively defines a constraint on district allocation. As shown in the Only Coverage figure (left), 4 out of 10 facilities (40%) are allocated to the Fafan zone (the four closely located facilities at the top), violating this constraint. In contrast, both the Only Alignment (middle) and LEG (right) approaches enforce it correctly. However, the Fafan zone has a substantial population, and adequate coverage for this area is provided exclusively by the Only Coverage baseline.

6.3 Evaluation metrics

We evaluated the LEG framework against the two baselines with respect to both coverage improvement and alignment with human expert advice.

Coverage improvement. This metric is defined as the increase in population coverage (within two hour walking distance) resulting from the grid-cell allocations.

Advice alignment. Evaluating alignment with advice is inherently challenging due to the subjective nature of such assessments. While consultation with human experts is the most natural evaluation approach, the need for rapid development cycles necessitates automated evaluation methods.

We propose an evaluation metric via LLM prompting using a powerful model (Gemini-2.5-pro), which generates a concrete evaluation function for each sentence in a predefined set of advice. For a given district allocation, this function assigns a score between 0 and 1 indicating how well the allocation aligns with the sentence, where 0 means not aligned at all and 1 means fully aligned. The overall alignment score is then computed as the average of these scores across the predefined advice set.

The prompt used for generating the scores is below.

Evaluation Prompt

Task: Generate an alignment function.

Inputs: Advice sentence: {}, sample district allocation vector: {}

Goal: The function should assess alignment with the given advice sentence and return a score between 0 and 1, where 0 indicates no alignment and 1 indicates strong alignment.

Instruction:

- (1) Allow scores in the range $[0,1]$.
- (2) Based on your training data and knowledge of Ethiopia’s Somali Region, identify the districts that best fit the groups described in the sentence.
- (3) Define a function: `advice_verifier(d)`, where `d` is the district allocation vector. The function maps the input `d` to a normalized score $0 \leq S \leq 1$.
- (4) Provide runnable Python code defining the function.

Constraints: The output must include only the function itself, with no additional text, and must be parseable by Python’s `exec` function.

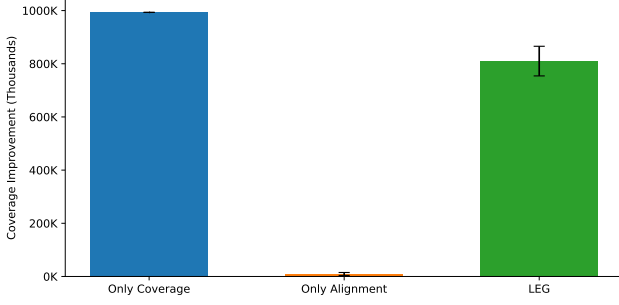


Figure 6: Comparison of coverage improvement across methods.

The function for evaluating alignment with the advice sentence: "Account for flood risk by assigning a fair number of facilities to districts near the Shabelle river" (Feasibility, (2)) operates in four steps. First, it defines which districts are considered near the river. Second, it computes the ratio of facilities assigned to districts near the river. Third, it compares this ratio to the desired threshold of 0.2. Finally, it returns the quotient of the actual ratio divided by the desired ratio. The summary of this function is given in Algorithm 5.

Algorithm 5 Evaluate the alignment with the advice of assigning fair number of facilities near the Shabelle river.

- 1: **Input:** Grid-cell allocation S
 - 2: **Output:** Alignment score $a \in [0, 1]$
 - 3: $\text{NearRiver} \leftarrow [\text{'Gode'}, \text{'Kelafo'}, \text{'Mustahil'}, \text{'Ferfer'}, \text{'Adadle'}, \text{'East Imi'}, \text{'West Imi'}, \text{'Raso'}, \text{'Danan'}, \text{'Shilabo'}, \text{'Berocano'}, \text{'Debeweyin'}]$
 - 4: $d = h(S)$
 - 5: $q \leftarrow \frac{|\sum_{i \in \text{NearRiver}} \{d_i(S)\}|}{|S|}$ *#ratio of near river facilities*
 - 6: $q_{OPT} \leftarrow 0.2$ *#optimal ratio of near river facilities*
 - 7: **Return** $\min(1, \frac{q}{q_{OPT}})$
-

6.4 Experimental results

6.4.1 Coverage improvement. In Figure 6 we compare the performance of our framework with the *Only Coverage* and *Only Alignment* baselines. Our framework performs slightly worse than the *Only Coverage* baseline, but substantially better than the *Only Alignment* baseline.

6.4.2 Advice alignment. The results in Figure 7 indicate that the *Only LLM* method outperforms the others, while the LEG framework performs reasonably well, surpassing the *Only Greedy* method.

6.4.3 Advice alignment per advice sentence. In Figure 8, we compare the performance of the LEG framework against the baselines for each advice sentence separately. We first observe that certain sentences are straightforward to satisfy, resulting in high scores across all evaluated methods. Second, the greedy baseline achieves its best performance on sentence 1 of the efficiency set, which is expected given that this sentence prioritizes more populated areas—a domain where the greedy approach naturally excels. Finally,

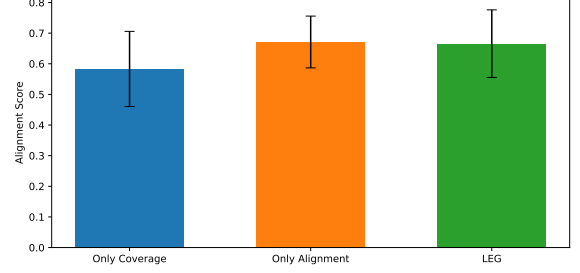


Figure 7: Comparison of Alignment Across Methods.

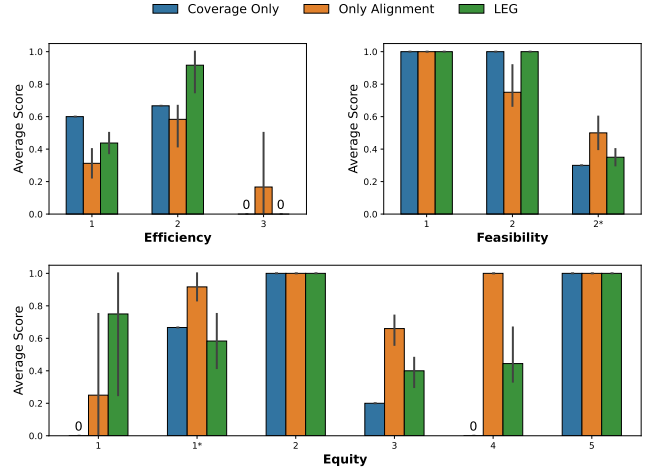


Figure 8: Comparison of Alignment Across Methods and Advice Sentences

the LEG framework significantly outperforms the Alignment-only baseline on certain sentences, such as prioritizing districts with low facility-to-population ratios (Efficiency, 2), while the Alignment-only baseline excels on others, such as prioritizing districts with low service coverage (Equity, 4). This reveals distinct strengths across different sentences.

7 CONCLUSIONS AND FUTURE WORK

In this paper, we implement the LEG framework [20] for optimizing a bi-objective problem in a complex environment. The framework employs an adaptive greedy algorithm integrated within an LLM-based refinement loop, ensuring coverage guarantees while jointly optimizing advice alignment and coverage. We evaluate the framework using human expert advice, demonstrating its effectiveness in real-world settings. Combined with existing literature, our work brings us closer to addressing pressing real-world challenges.

While the proposed evaluation approach effectively measures both population coverage and alignment with advice throughout the development cycle, we plan to complement it with human evaluation. Developing the human evaluation procedure and analyzing its results is an important next step for deployment.

ACKNOWLEDGEMENTS

We are very grateful to Kasahun Sime and Wondesen Nigatu of the Ministry of Health, Ethiopia, for their continuous support and valuable advice. The findings and conclusions in this report are those of the authors and do not represent any recommendations or official positions.

REFERENCES

- [1] Pieter Abbeel and Andrew Y Ng. 2004. Apprenticeship learning via inverse reinforcement learning. In *Proceedings of the international conference on Machine learning*, Vol. 21. 1.
- [2] Amir Ahmadi-Javid, Pardis Seyedi, and Siddhartha S Syam. 2017. A survey of healthcare facility location. *Computers & Operations Research* 79 (2017), 223–263.
- [3] Sisay Mulugeta Alemu, Gerd Weitkamp, Abera Kenay Tura, Jelle Stekelenburg, and Regien Biesma. 2025. Optimizing maternal health facility locations in Ethiopia: A geospatial approach for upgrading to emergency obstetric and newborn care. *Journal of Transport Geography* 125 (2025), 104208.
- [4] Nikhil Behari, Edwin Zhang, Yunfan Zhao, Aparna Taneja, Dheeraj Nagaraj, and Milind Tambe. 2024. A decision-language model (dLM) for dynamic restless multi-armed bandit tasks in public health. *Advances in Neural Information Processing Systems* 37 (2024), 3964–4002.
- [5] Davin Choo, Yohai Trabelsi, Fentabil Getnet, Samson Warkaye Lamma, Wondesen Nigatu, Kasahun Sime, Lisa Matay, Milind Tambe, and Stéphane Verguet. 2025. Optimizing Health Coverage in Ethiopia: A Learning-augmented Approach and Persistent Proportionality Under an Online Budget. In *Proceedings of the AAAI Conference on Artificial Intelligence (to appear)*. AISI Track, to appear; also published as arXiv:2509.00135.
- [6] Paul F Christiano, Jan Leike, Tom Brown, Miljan Martic, Shane Legg, and Dario Amodei. 2017. Deep reinforcement learning from human preferences. *Advances in neural information processing systems* 30 (2017).
- [7] Lorenzo Jaime Yu Flores, Ramon Rafael Tonato, Gabrielle Ann dela Paz, and Valerie Gilbert Ulep. 2021. Optimizing health facility location for universal health care: A case study from the Philippines. *PloS one* 16, 9 (2021), e0256821.
- [8] Abraham Haileamlak and Israel Ataro. 2023. The Ethiopian health extension program (HEP) is still relevant after 15 years of implementation although major transformation is essential to sustain its gains and relevance. *Ethiopian journal of health sciences* 33, Spec Iss 1 (2023), 1.
- [9] Conor F Hayes, Roxana Rădulescu, Eugenio Bargiacchi, Johan Källström, Matthew Macfarlane, Mathieu Reymond, Timothy Verstraeten, Luisa M Zintgraf, Richard Dazeley, Fredrik Heintz, et al. 2022. A practical guide to multi-objective reinforcement learning and planning. *Autonomous Agents and Multi-Agent Systems* 36, 1 (2022).
- [10] Nathaniel Hendrix, Samson Warkaye, Latera Tesfaye, Mesfin Agachew Woldekidan, Asrat Arja, Ryoko Sato, Solomon Tessema Memirie, Alemnesh H Mirkuzie, Fentabil Getnet, and Stéphane Verguet. 2023. Estimated travel time and staffing constraints to accessing the Ethiopian health care system: A two-step floating catchment area analysis. *Journal of Global Health* 13 (2023), 04008.
- [11] Jiaming Ji, Tianyi Qiu, Boyuan Chen, Borong Zhang, Hantao Lou, Kaile Wang, Yawen Duan, Zhonghao He, Jiayi Zhou, Zhaowei Zhang, et al. 2023. Ai alignment: A comprehensive survey. *arXiv preprint arXiv:2310.19852* (2023).
- [12] Kaisa Miettinen. 1999. *Nonlinear multiobjective optimization*. Vol. 12. Springer Science & Business Media.
- [13] George L Nemhauser, Laurence A Wolsey, and Marshall L Fisher. 1978. An analysis of approximations for maximizing submodular set functions—I. *Mathematical programming* 14, 1 (1978), 265–294.
- [14] Andrew Y Ng, Stuart Russell, et al. 2000. Algorithms for inverse reinforcement learning.. In *Seventeenth International Conference on Machine Learning*, Vol. 1. 2.
- [15] Ministry of Health [Ethiopia]. 2020. Realizing Universal Health Coverage through Primary Healthcare: A Roadmap for Optimizing the Ethiopian Health Extension Program (2020–2035).
- [16] Nathan Schurr, Janusz Marecki, and Milind Tambe. 2009. Improving adjustable autonomy strategies for time-critical domains. In *Proceedings of The 8th International Conference on Autonomous Agents and Multiagent Systems-Volume 1*. 353–360.
- [17] Nathan Schurr, Pratik Patil, Fred Pighin, and Milind Tambe. 2006. Using multiagent teams to improve the training of incident commanders. In *Proceedings of the fifth international joint conference on Autonomous agents and multiagent systems*. 1490–1497.
- [18] SS Radiah Shariff, Noor Hasnah Moin, and Mohd Omar. 2012. Location allocation modeling for healthcare facility planning in Malaysia. *Computers & Industrial Engineering* 62, 4 (2012), 1000–1010.
- [19] Hua Shen, Tiffany Knearem, Reshmi Ghosh, Kenan Alkiek, Kundan Krishna, Yachuan Liu, Ziqiao Ma, Savvas Petridis, Yi-Hao Peng, Li Qiwei, et al. 2024. Towards bidirectional human-ai alignment: A systematic review for clarifications, framework, and future directions. *arXiv preprint arXiv:2406.09264* 2406 (2024), 1–56.
- [20] Yohai Trabelsi, Guojun Xiong, Fentabil Getnet, Stéphane Verguet, and Milind Tambe. 2026. Health Facility Location in Ethiopia: Leveraging LLMs to Integrate Expert Knowledge into Algorithmic Planning. In *In Proceedings of the 2026 International Conference on Autonomous Agents and Multiagent Systems (to appear)*. also available as arXiv:2601.11479.
- [21] Hao Wang, Peng Luo, and Yimeng Wu. 2023. Research on the location decision-making method of emergency medical facilities based on WSR. *Scientific Reports* 13, 1 (2023), 18011.
- [22] Huihui Wang, Roman Tesfaye, Gandham N. V. Ramana, and Chala Tesfaye Chekagn. 2016. *Ethiopia Health Extension Program: An Institutionalized Community Approach for Universal Health Coverage*. World Bank.
- [23] Yufei Wang, Wanjun Zhong, Liangyou Li, Fei Mi, Xingshan Zeng, Wenyong Huang, Lifeng Shang, Xin Jiang, and Qun Liu. 2023. Aligning large language models with human: A survey. *arXiv preprint arXiv:2307.12966* (2023).
- [24] D. J. Weiss, A. Nelson, C. A. Vargas-Ruiz, K. Gligorić, S. Bavadekar, E. Gabrilovich, A. Bertozzi-Villa, J. Rozier, H. S. Gibson, and et al. Shekel, T. 2020. Global maps of travel time to healthcare facilities. *Nature Medicine* 26 (2020), 1835–1838.
- [25] WorldPop. 2025. WorldPop Population Counts 2015–2030: Ethiopia – Spatial Distribution of Population. <https://data.humdata.org/dataset/worldpop-population-counts-2015-2030-eth>. Constrained estimates of total population per 100m grid cell in Ethiopia for 2015–2030, Geotiff format, CC BY4.0 license.
- [26] Guojun Xiong and Milind Tambe. 2026. VORTEX: Aligning Task Utility and Human Preferences through LLM-Guided Reward Shaping. In *Proceedings of the AAAI Conference on Artificial Intelligence (to appear)*. To appear; also published as arXiv:2509.16399.
- [27] Henry Peng Zou, Wei-Chieh Huang, Yaozu Wu, Yankai Chen, Chunyu Miao, Hoang Nguyen, Yue Zhou, Weizhi Zhang, Liancheng Fang, Langzhou He, et al. 2025. A survey on large language model based human-agent systems. *Authorea Preprints* (2025).