



Journal Website:
<http://sciencebring.com/index.php/ijasr>

Copyright: Original content from this work may be used under the terms of the creative commons attributes 4.0 licence.

Research Article

Scale Mind: An Intelligent LLM-Based Framework for Scalability-Constrained Combinatorial Optimization

Submission Date: June 08, 2026, Accepted Date: July 05, 2026,

Published Date: August 19, 2026

Crossref doi: <https://doi.org/10.37547/ijasr-06-08-09>

Tran Hoang Nam

Department of Computer Science and Machine Intelligence University of Digital Engineering, Ho Chi Minh City, Vietnam

ABSTRACT

Scalability-constrained combinatorial optimization presents a fundamental computational challenge in environments where the number of decision variables, constraints, and feasible configurations grows rapidly with problem size. Conventional optimization approaches can become difficult to deploy when constraint structures are heterogeneous, dynamically changing, or expressed in forms that require substantial contextual interpretation. This paper proposes ScaleMind, an intelligent Large Language Model (LLM)-based framework designed to support scalable combinatorial optimization through natural-language constraint interpretation, structured problem representation, constraint decomposition, candidate-generation guidance, feasibility reasoning, and iterative optimization feedback. The framework is conceptually positioned around the integration of language-based reasoning with conventional combinatorial optimization mechanisms rather than treating an LLM as a standalone numerical optimizer. The methodological design emphasizes separation between semantic interpretation and mathematically verifiable optimization, thereby addressing the reliability and scalability limitations associated with unconstrained generative reasoning. The framework builds upon the scalability-oriented combinatorial LLM perspective presented by Ramamurthy, Bellamkonda, and Amanmadov (2026), while extending that conceptual direction toward a broader architecture for constraint-aware optimization. The resulting analysis indicates that ScaleMind can provide a systematic mechanism for translating complex constraint descriptions into structured optimization representations, prioritizing constraint-critical regions of the search space, and adapting optimization strategies according to problem complexity. The paper further



identifies limitations involving hallucinated constraints, computational overhead, verification requirements, and dependence on the quality of formal problem representations. The proposed framework therefore represents a research architecture for combining semantic intelligence with mathematically grounded combinatorial optimization.

KEYWORDS

Large Language Models; Combinatorial Optimization; Scalability Constraints; Constraint Reasoning; Intelligent Optimization; LLM Framework; Search-Space Reduction; Decision Intelligence; Optimization Architecture

INTRODUCTION

1.1 Background

Combinatorial optimization concerns the selection of optimal or near-optimal solutions from a discrete and potentially enormous set of alternatives. Scheduling, resource allocation, routing, assignment, configuration, and planning problems frequently exhibit this structure. As the number of variables and constraints increases, the search space can expand exponentially, making exhaustive evaluation impractical. The scalability problem is particularly significant when optimization environments contain multiple interacting constraints rather than a single objective function.

Recent advances in Large Language Models have introduced a different computational perspective. Instead of relying exclusively on predefined optimization procedures, LLMs can interpret natural-language descriptions, identify relationships among entities, organize contextual information, and generate structured reasoning sequences. Ramamurthy, Bellamkonda, and

Amanmadov (2026) position an LLM-oriented framework around combinatorial scalability constraints, indicating the emerging relevance of language models to scalable constraint-oriented computational problems.

However, an LLM should not be interpreted as a replacement for mathematically rigorous optimization. Generative models can produce plausible but invalid constraints, overlook dependencies, or generate solutions that appear reasonable while violating hard feasibility requirements. Consequently, the central research challenge is not merely how to apply an LLM to optimization, but how to construct an architecture in which semantic reasoning and formal optimization complement one another.

1.2 Problem Statement

Traditional combinatorial optimization systems generally assume that the optimization problem has already been converted into a formal representation. This assumption becomes restrictive when constraints originate from

heterogeneous sources such as policy descriptions, operational requirements, human instructions, and semi-structured documentation. Manual conversion of such information into variables, constraints, objectives, and decision rules introduces interpretation overhead and creates opportunities for inconsistency.

At the same time, direct LLM-based solution generation introduces a separate reliability problem. A language model can identify potential relationships among constraints but does not inherently guarantee mathematical feasibility. The research problem addressed by this paper is therefore the design of an intelligent framework that uses LLM capabilities for semantic constraint understanding while maintaining formal validation and optimization mechanisms.

1.3 Research Relevance

The proposed ScaleMind architecture is relevant because scalability has two dimensions. The first is computational scalability, referring to the ability to process increasingly large optimization spaces. The second is semantic scalability, referring to the ability to interpret increasingly complex and heterogeneous constraint descriptions. Existing optimization architectures primarily emphasize the first dimension, whereas an LLM-enabled architecture can potentially address both.

The conceptual direction is consistent with the scalability-oriented framework of Ramamurthy, Bellamkonda, and Amanmadov (2026), but ScaleMind expands the perspective by defining

explicit functional stages between natural-language constraint acquisition and formal optimization.

1.4 Objectives

The primary objective is to develop a conceptual architecture for LLM-assisted combinatorial optimization under scalability constraints. Specific objectives are to establish a structured mechanism for constraint interpretation, transform semantic information into formal optimization components, reduce unnecessary search through constraint-aware decomposition, integrate formal feasibility verification, and provide an iterative optimization-feedback mechanism.

1.5 Scope and Significance

The framework is applicable to hypothetical domains such as production scheduling, workforce allocation, transportation routing, cloud resource assignment, and multi-resource planning. Its significance lies not in proposing another standalone optimization algorithm but in establishing an intelligent orchestration layer capable of connecting human-oriented constraint descriptions with computational optimization processes.

LITERATURE REVIEW

The provided literature consists of the work by Ramamurthy, Bellamkonda, and Amanmadov (2026), titled ScalePulse: Combinatorial Llm Framework for Scalability Constraint. This reference provides the principal conceptual



foundation for examining the relationship between LLMs, combinatorial reasoning, and scalability constraints.

The importance of this work can be understood from the intersection of two computational requirements: combinatorial problem solving and scalable constraint management. Conventional combinatorial optimization operates effectively when problem variables, feasible regions, and constraints can be explicitly formalized. However, real-world optimization environments may contain contextual information that is not initially expressed in a machine-readable optimization model. The ScalePulse perspective is therefore important because it frames LLMs as potentially useful components in addressing scalability-oriented combinatorial constraints (Ramamurthy, Bellamkonda and Amanmadov, 2026).

ScaleMind adopts this foundation but establishes a more explicitly layered architecture. Instead of allowing an LLM to directly produce an optimization answer, the proposed framework separates interpretation, representation, decomposition, candidate generation, validation, and refinement. This separation is theoretically significant because it recognizes that semantic reasoning and combinatorial optimization have different correctness criteria.

The central research gap identified from the provided reference is the need for a more explicit functional bridge between language-level reasoning and formal optimization operations. An LLM may understand that a particular resource

cannot be assigned simultaneously to incompatible tasks, but the optimization engine must ultimately represent that relationship through formal variables and constraints. ScaleMind addresses this gap conceptually by introducing a constraint translation layer and a verification loop.

The literature foundation consequently supports an architectural rather than purely algorithmic interpretation of LLM-assisted optimization. The principal theoretical positioning of ScaleMind is that an LLM should function as an intelligent reasoning and orchestration component, while deterministic optimization and validation components remain responsible for mathematical feasibility.

METHODOLOGY

3.1 Proposed Scale Mind Architecture

Scale Mind is designed as a multi-stage architecture consisting of six functional layers: constraint acquisition, semantic interpretation, formal representation, constraint-aware decomposition, optimization and candidate generation, and verification-driven refinement.

The first layer receives problem information from natural-language descriptions, structured datasets, operational rules, or mixed sources. The LLM processes this information to identify entities, decision variables, relationships, restrictions, preferences, and objectives.

The second layer performs semantic constraint interpretation. For example, a statement such as

“Machine A cannot process Product X and Product Y simultaneously” must be interpreted as a resource-capacity constraint rather than treated as ordinary descriptive text. The LLM identifies the semantic relationship, while subsequent modules convert the interpretation into a formal representation.

3.2 Formal Constraint Representation

The central mathematical representation can be expressed as:

$$\left[\min_{\{x\}} F(x) \right]$$

subject to

$$\left[g_i(x) \leq 0, \quad i=1, \dots, m \right]$$

and

$$\left[h_j(x)=0, \quad j=1, \dots, p \right]$$

where (x) represents the combinatorial decision variables, $(F(x))$ represents the objective function, $(g_i(x))$ represents inequality constraints, and $(h_j(x))$ represents equality constraints.

ScaleMind does not permit the LLM to determine feasibility solely through generated text. Instead, interpreted constraints are converted into a structured intermediate representation and submitted to a formal validation layer. This architecture provides a distinction between semantic correctness and mathematical correctness.

3.3 Constraint Classification

The framework classifies constraints into hard constraints, soft constraints, resource constraints, temporal constraints, dependency constraints, and preference-based conditions. Hard constraints define feasibility and cannot normally be violated. Soft constraints represent desirable characteristics that can be incorporated into an objective or penalty function.

For example, in workforce scheduling, “an employee cannot be assigned to two shifts occurring simultaneously” is a hard constraint, whereas “employees should preferably receive their requested shifts” can be represented as a soft preference.

The classification process reduces ambiguity and provides the optimization engine with information about constraint priority. The LLM therefore operates primarily as a semantic classifier and reasoning coordinator rather than as an unrestricted solution generator.

3.4 Constraint-Aware Search-Space Reduction

A major methodological component of ScaleMind is constraint-aware search-space reduction. If a

combinatorial problem contains (n) binary decision variables, its theoretical solution space may contain up to (2^n) configurations. Evaluating every configuration becomes increasingly impractical as (n) grows.

ScaleMind uses interpreted constraints to identify incompatible assignments before complete candidate solutions are constructed. Suppose three resources (R_1, R_2, R_3) are assigned to four tasks under capacity and temporal restrictions. Rather than generating every possible assignment, the framework can eliminate configurations that violate known constraints during candidate construction.

The resulting principle is:

[
 $S' \subseteq S$
]

where (S) represents the original search space and (S') represents the constraint-filtered space. The effectiveness of the framework depends on whether the reduction removes infeasible alternatives without incorrectly eliminating feasible solutions.

3.5 LLM-Guided Candidate Generation

Following search-space reduction, the LLM can assist in prioritizing candidate regions. Its role is to identify promising combinations based on contextual relationships, constraint interactions, and optimization objectives. Candidate solutions

are subsequently evaluated by a formal optimization or scoring mechanism.

This hybrid design avoids a fundamental weakness of purely generative optimization: textual plausibility is not equivalent to mathematical optimality. ScaleMind instead establishes a controlled relationship in which LLM reasoning proposes or prioritizes candidate structures, while computational procedures determine objective values and constraint satisfaction.

3.6 Verification and Feedback Loop

Verification represents the principal reliability mechanism. Every generated candidate is evaluated against formal constraints. If a candidate violates a hard constraint, the violation is converted into structured feedback.

The optimization loop can therefore be represented as:

[
 $\text{Input} \rightarrow \text{Interpretation} \rightarrow \text{Formalization}$
 $\rightarrow \text{Candidate Generation} \rightarrow \text{Verification}$
 $\rightarrow \text{Refinement}$
]

The feedback mechanism enables iterative correction. For example, if a generated production schedule exceeds machine capacity, the verifier identifies the violated capacity

constraint and returns the information to the reasoning layer. The LLM can then revise the candidate rather than generating an entirely unrelated solution.

This architecture is conceptually aligned with the scalability-oriented LLM optimization direction established by Ramamurthy, Bellamkonda, and Amanmadov (2026), while emphasizing verification as a necessary component of reliable deployment.

3.7 Complexity-Aware Orchestration

ScaleMind dynamically determines the appropriate degree of LLM involvement based on problem complexity. Small problems may require only semantic translation and conventional optimization. Larger problems may benefit from decomposition into subproblems.

For a problem (P), decomposition can be expressed as:

$$\begin{aligned} &[\\ &P=\{P_1,P_2,\ldots,P_k\} \\ &] \end{aligned}$$

where each (P_i) represents a constrained subproblem. The LLM can identify relationships among subproblems and determine which constraints are local and which are global. Local constraints can be solved independently where appropriate, while global constraints are maintained through coordination.

This decomposition mechanism is particularly important for scalability because it prevents the LLM context window and optimization engine from simultaneously processing every component of a large problem.

3.8 Evaluation Framework

The proposed framework should be evaluated through four dimensions: solution feasibility, objective quality, computational scalability, and semantic interpretation accuracy.

Feasibility measures the percentage of generated solutions satisfying all hard constraints. Objective quality compares generated solutions against a baseline objective or known optimum where available. Scalability examines computational behavior as variables and constraints increase. Semantic interpretation accuracy evaluates whether natural-language requirements are correctly transformed into formal constraints.

A useful evaluation criterion can be expressed as:

$$\begin{aligned} &[\\ &Q=\alpha F+\beta O+\gamma S+\delta I \\ &] \end{aligned}$$

where (F) represents feasibility, (O) objective quality, (S) scalability, and (I) interpretation accuracy, while the coefficients represent application-specific priorities.

RESULTS

The methodological analysis indicates that the principal contribution of ScaleMind is architectural integration rather than the introduction of a new standalone combinatorial algorithm. The framework establishes a separation between semantic reasoning and mathematical verification, producing a system in which the LLM is responsible for interpreting and organizing information while formal computational mechanisms remain responsible for feasibility evaluation. This separation is important because scalability-oriented LLM frameworks require mechanisms that prevent semantic flexibility from undermining optimization correctness (Ramamurthy, Bellamkonda and Amanmadov, 2026).

A first finding concerns constraint interpretation. Natural-language requirements can contain implicit relationships that are difficult to encode manually. ScaleMind addresses this problem by introducing an intermediate semantic layer. The expected result is a reduction in the conceptual distance between human-defined operational requirements and machine-executable optimization models. For example, a statement involving employee availability, task incompatibility, and maximum working hours can be decomposed into multiple formal constraint categories rather than being represented as an undifferentiated text instruction.

A second finding concerns search-space management. Constraint-aware decomposition provides a mechanism for reducing the number of candidates that require complete evaluation. This does not necessarily change the theoretical

worst-case complexity of a combinatorial problem, but it can improve practical scalability by preventing obviously infeasible candidates from reaching expensive evaluation stages. The effectiveness of this mechanism is therefore dependent on the quality and correctness of constraint identification.

A third finding concerns hybrid reasoning. ScaleMind demonstrates conceptually that LLM reasoning and conventional optimization should be assigned complementary responsibilities. The LLM can prioritize candidate structures and recognize semantic relationships, whereas deterministic optimization modules can calculate objective values and verify feasibility. This architecture is more defensible than relying on generated textual answers as optimization outputs.

A fourth finding concerns verification-driven refinement. The feedback loop transforms optimization from a one-pass generation process into an iterative procedure. When a candidate violates a constraint, the violation can be represented explicitly and returned to the reasoning stage. This provides a mechanism for controlled correction and potentially reduces repeated generation of similar invalid solutions.

However, the framework also reveals several limitations. Semantic interpretation errors may propagate into the formal optimization model. A correctly implemented optimizer cannot compensate for an incorrectly interpreted requirement. Furthermore, LLM inference introduces computational overhead, which may

offset gains obtained through search-space reduction in smaller problems. The architecture therefore appears most appropriate for complex, heterogeneous optimization environments in which semantic processing itself represents a significant component of the computational challenge.

Overall, the findings support the proposition that scalability should be addressed simultaneously at the computational and semantic levels. ScaleMind provides an architecture for this integration, but empirical benchmarking remains necessary before claims about performance superiority can be established.

DISCUSSION

The findings suggest that LLM-based combinatorial optimization should be understood as a hybrid systems problem rather than a simple application of generative AI to mathematical optimization. The conceptual contribution of ScaleMind is the explicit assignment of responsibilities across system layers. LLMs provide semantic flexibility, contextual interpretation, decomposition guidance, and candidate prioritization, while formal optimization components provide deterministic feasibility and objective evaluation.

This distinction has important theoretical implications. The ScalePulse framework establishes a foundation for considering LLMs in scalable combinatorial constraint environments (Ramamurthy, Bellamkonda and Amanmadov, 2026). ScaleMind extends that direction by

treating constraint interpretation as a distinct computational stage. This creates an intermediate representation through which human-oriented requirements can be transformed into machine-oriented optimization structures.

The framework also exposes an important trade-off between reasoning flexibility and computational determinism. Greater LLM involvement may improve the ability to interpret ambiguous or heterogeneous requirements, but it simultaneously introduces uncertainty into the optimization pipeline. Conversely, strict formalization improves reliability but requires all relevant requirements to be known and encoded beforehand. ScaleMind attempts to balance these characteristics by using the LLM before and around formal optimization rather than allowing it to bypass formal verification.

Another implication concerns scalability. Increasing the number of variables does not automatically make an LLM-based architecture scalable. In fact, large optimization problems can create additional context-management and inference costs. Consequently, decomposition and constraint filtering become essential. The architecture should selectively expose only relevant problem components to the reasoning layer. This suggests a future research direction involving adaptive context allocation in which the LLM processes only constraints that are relevant to the current subproblem.

The principal limitation is that the proposed framework is architectural and requires

empirical validation. No experimental dataset or benchmark results were provided in the source material, so quantitative claims regarding speed, accuracy, or optimality cannot responsibly be established. Evaluation should therefore include controlled combinatorial benchmarks, increasing constraint densities, multiple problem sizes, and comparisons with conventional optimization baselines.

A second limitation concerns semantic hallucination. An LLM may infer a constraint that was never specified or omit a constraint that was explicitly provided. Such errors can have greater consequences than ordinary textual inaccuracies because they alter the feasible region. A robust implementation must therefore preserve traceability from each formal constraint back to its originating input statement.

A third limitation concerns objective ambiguity. Real-world optimization frequently involves competing objectives such as cost, time, quality, and resource utilization. The LLM can help interpret stakeholder preferences, but objective weighting must ultimately be made explicit. Otherwise, the system may produce a semantically plausible solution without a clearly defensible optimization criterion.

Despite these limitations, the architecture has practical potential in scheduling, logistics, resource planning, and configuration problems where optimization requirements are frequently communicated through natural language. Its strongest contribution is therefore not the replacement of established optimization

algorithms, but the creation of an intelligent interface and orchestration mechanism between human requirements and formal combinatorial computation.

CONCLUSION

This paper proposed ScaleMind, an intelligent LLM-based framework for scalability-constrained combinatorial optimization. The framework integrates semantic constraint interpretation, formal representation, constraint-aware search-space reduction, LLM-guided candidate generation, deterministic verification, and iterative refinement.

The central contribution is the separation of language-based reasoning from mathematical feasibility. By assigning semantic interpretation to the LLM and formal validation to computational optimization components, ScaleMind provides a theoretically grounded approach to combining generative intelligence with combinatorial optimization. The architecture extends the conceptual direction of LLM-based scalability constraint reasoning identified by Ramamurthy, Bellamkonda, and Amanmadov (2026).

The analysis indicates that the most significant opportunities arise when optimization problems contain both large combinatorial search spaces and heterogeneous human-readable constraints. Nevertheless, the framework should not be considered empirically validated until benchmark experiments establish its

effectiveness against conventional and hybrid optimization methods.

Future research should focus on automated constraint verification, benchmark-based evaluation, adaptive problem decomposition, uncertainty estimation, multi-objective optimization, and traceable LLM reasoning. Such developments could establish a more reliable foundation for intelligent optimization systems capable of operating across increasingly complex and scalable decision environments.

REFERENCES

1. K. Ramamurthy, N. Bellamkonda and N. Amanmadov, "ScalePulse: Combinatorial Llm Framework for Scalability Constraint," SoutheastCon 2026, Huntsville, AL, USA, 2026, pp. 1-6, doi: 10.1109/SoutheastCon63549.2026.11476075
2. Arif S, Khan NZS, Malim N, Zainudin S. Retrieval performance using different type of similarity coefficient for virtual screening. Research Journal of Applied Sciences, Engineering and Technology 2015; 9(5): 391–395. doi: 10.19026/rjaset.9.14182
3. Fujiwara T, Kwon OH, Ma KL. Supporting analysis of dimensionality reduction results with contrastive learning. arXiv 2019; arXiv:1905.03911. doi: 10.1109/TVCG.2019.293425
4. Gardiner EJ, Gillet VJ. Perspectives on knowledge discovery algorithms recently introduced in chemoinformatics: Rough set theory, association rule mining, emerging patterns, and formal concept analysis. Journal of Chemical Information and Modeling 2015; 55(9): 1781–1803. doi: 10.1021/acs.jcim.5b00198
5. Gaulton A, Hersey A, Nowotka M, et al. The ChEMBL database in 2017. Nucleic Acids Research 2017; 45(D1): D945–D954. doi: 10.1093/nar/gkw1074
6. Heikamp K, Bajorath J. Support vector machines for drug discovery. Expert Opinion on Drug Discovery 2014; 9(1): 93–104. doi: 10.1517/17460441.2014.866943
7. Li PH, Lee T, Youn HY. Dimensionality reduction with sparse locality for principal component analysis. Mathematical Problems in Engineering 2020; 2020: 1–12. doi: 10.1155/2020/972327921
8. Mahmud SMH, Chen W, Jahan H, et al. Dimensionality reduction based multi-kernel framework for drug-target interaction prediction. Chemometrics and Intelligent Laboratory Systems 2021; 212: 104270. doi: 10.1016/j.chemolab.2021.104270
9. Malavika S, Phil M, Selvam K. Reduction of dimensionality for high dimensional data using correlation measures. Available online: <http://www.ripublication.com> (accessed on 27 July 2023).
10. Steinbeck C, Han Y, Kuhn S, et al. The Chemistry Development Kit (CDK): An open-source Java library for chemo- and bioinformatics. Journal of Chemical Information Comput Sciences 2003; 43(2): 493–500. doi: 10.1021/ci025584y

11. Terol RM, Reina AR, Ziaei S, Gil D. A machine learning approach to reduce dimensional space in large datasets. IEEE Access 2020; 8: 148181–148192. doi: 10.1109/ACCESS.2020.30128362
12. M. R. Marri, S. B. Kurada, S. Gupta, S. Dey, S. Kumar and R. Chauhan, "Graph-Based Deep Learning Model for Identifying Cyber Threats in Cloud Platforms," 2025 International Conference on Computational Intelligence, Security, and Artificial Intelligence (IntelliSecAI), Al-Khobar, Saudi Arabia, 2025, pp. 1-7, doi: 10.1109/IntelliSecAI66368.2025.11472831.
13. Geo Philip, Paulson, Artificial Intelligence and Machine Learning Applications in Project Schedule Forecasting: A Predictive Framework for Time-Control in Complex Building Projects (April 16, 2026). Available at SSRN: <https://ssrn.com/abstract=6588119> or <http://dx.doi.org/10.2139/ssrn.6588119>

