

Modular Data Pipeline Architecture for Large-Scale Supply Chain Optimization: A Framework for Scalable, Configuration-Driven Input Generation

Uday Dhembare

Data Engineering Manager, Supply Chain Analytics Bellevue, WA, USA.

Received On: 01/06/2026

Revised On: 25/06/2026

Accepted On: 04/07/2026

Published On: 11/07/2026

Abstract: Large-scale supply chain optimization relies on the steady production of accurate, validated input datasets. These datasets feed the mathematical models that handle network design, capacity planning, and cost minimization. In most organizations, the data pipelines behind these models started life as monolithic scripts—single codebases that bundle business logic, configuration parameters, and execution dependencies all in one place. That works fine at small scale, but as the number of scenarios, data sources, and downstream consumers grows, these architectures turn brittle and expensive to maintain. This paper lays out a generalized framework for rebuilding supply chain input generation pipelines around modular, configuration-driven, cloud-native principles. We draw on existing work in data engineering, distributed systems, and operations research to describe the structural problems with monolithic pipeline design, spell out the requirements for a scalable replacement, and propose an architecture built on directed acyclic graph (DAG) execution, externalized configuration management, isolated multi-scenario runs, and formal version control. We evaluate the framework against quantitative benchmarks—end-to-end runtime, scenario throughput, parameter visibility, failure recovery—and show that modular pipeline architectures cut operational overhead, improve data quality, and let organizations scale their optimization capabilities across business units and planning cycles. The implications reach across manufacturing, retail, healthcare, and logistics wherever optimization models depend on complex, multi-source input generation.

Keywords: Data Engineering, Big Data, Automated Systems, ETL, Scalable Systems, Data Infrastructure, Supply Chain Optimization, Data Pipeline Architecture, Modular Systems, Configuration Management, Directed Acyclic Graph, Cloud-Native Computing, Network Design, Scenario Analysis, Distributed Computing, Large Scale Analytics.

1. Introduction

The pressure on modern supply chains to cut costs while improving service levels has pushed mathematical optimization models into widespread use for network design, inventory placement, and transportation planning [1]. These models whether built on mixed integer programming, constraint satisfaction, or machine learning all share one thing in common: they need large, accurate, consistently structured input datasets before a single optimization can run [2]. The engineering challenge of producing those inputs reliably, at scale, and across multiple planning scenarios has gotten remarkably little attention in the academic literature, even though it is one of the most practically significant bottlenecks in applied operations research [3].

In large organizations, the data pipelines that generate optimization inputs usually start as straightforward scripting efforts. Someone writes a sequence of SQL queries, file transformations, and data joins that produce the input files needed for a particular model run. It works. Then, over time, the organization's optimization capabilities mature and these pipelines accumulate complexity. New data sources get integrated, business rules get added, additional model types get introduced, downstream consumers multiply. What you

end up with is a monolithic pipeline that runs as a single indivisible unit, has dozens of configuration parameters buried directly in the code, lacks version control, and cannot easily support concurrent scenario execution [4].

The consequences are well-documented in software engineering but underexplored in supply chain contexts. A pipeline failure midway through execution means all intermediate outputs get discarded and you restart from scratch, a real cost when end-to-end runtimes run to tens of minutes or more [5]. Business analysts who want to test a new scenario must either modify the central pipeline directly (and risk breaking things for everyone else) or keep a separate copy of the entire codebase (and deal with inevitable divergence). When parameters are hardcoded in SQL, nobody outside the engineering team can review or validate the assumptions driving a model run without reading source code [6].

This paper addresses these challenges by proposing a generalized architectural framework for supply chain input generation pipelines, grounded in four principles: modularity, configuration-driven execution, multi-scenario isolation, and cloud-native deployment. We review the relevant literature, describe the structural components of the proposed

framework, and evaluate its performance against the limitations of monolithic approaches. The paper concludes with implementation considerations and future research directions.

The rest of the paper is organized as follows. Section 2 lays out the theoretical premise and reviews relevant literature. Section 3 characterizes the problems with monolithic pipeline architectures in supply chain settings. Section 4 walks through representative use cases across industries. Section 5 describes the proposed modular framework in detail. Section 6 discusses practical applications and deployment patterns. Section 7 presents impact analysis and performance benchmarks. Section 8 wraps up with contributions and directions for future work.

2. Premise

2.1. Supply Chain Optimization and Data Dependencies

The academic literature on supply chain optimization is vast, spanning network design [1], vehicle routing [7], inventory optimization, and demand forecasting. What much of this literature conveniently abstracts away is the data engineering infrastructure you actually need to operationalize these models at scale. Take a mixed integer programming model for network design: it might require hundreds of input parameters like facility capacities, transit times, cost coefficients, demand forecasts, operational constraints, each coming from different systems, updating at different frequencies, and needing different transformation logic before the model can consume them [2].

Sanders and Swink [3] made an observation that resonates with practitioners: the gap between analytical capability and operational deployment in supply chain contexts often comes down not to model sophistication but to data infrastructure limitations. Organizations that pour resources into optimization science frequently discover that the real bottleneck sits upstream, in the processes that assemble, validate, and deliver the inputs those models need. That observation is what motivates treating the pipeline architecture itself as a first-class engineering problem.

2.2. Data Pipeline Architecture in the Literature

Software engineering and data systems research has built up a rich vocabulary for pipeline architecture over the years. The Directed Acyclic Graph (DAG) as an organizing principle for data workflows is well-established at this point, with systems like Apache Airflow, Luigi, and Prefect offering production-grade implementations [5]. DAG-based execution makes sure each processing step runs only after its upstream dependencies have completed, which opens the door to parallel execution of independent steps and gives you a clear picture of data lineage.

Kleppmann [4] articulated what is really a fundamental tension in data pipeline design: simple pipelines are easy to build and understand but hard to extend, while flexible pipelines take more investment up front but handle evolving requirements more gracefully. This tension hits supply chain contexts especially hard, because planning cycles introduce

periodic changes to model structure, data sources, and business rules that need to be accommodated without disrupting ongoing operations.

The idea of configuration-driven execution that is separating what a pipeline does from the parameters that control how it does it is a well-worn software engineering practice [6]. When you pull parameters out into a configuration layer, non-technical stakeholders can review and tweak pipeline behavior without touching source code, you reduce the risk of silent errors from hardcoded values, and you can trigger pipeline runs programmatically with different parameter sets.

2.3. Cloud-Native Data Engineering

The rise of cloud-native data engineering platforms has opened up architectural possibilities for supply chain pipelines that simply did not exist a decade ago [8]. Container-based compute gives you reproducible, isolated execution environments that scale on demand. Orchestration services handle workflow management, retry logic, and monitoring without dedicated infrastructure. Object storage provides cheap, durable storage for intermediate and final outputs, with access controls that support multi-tenant usage.

Armbrust et al. [8] argued that cloud-native architectures fundamentally change the economics of data engineering by separating compute from storage and enabling fine-grained resource allocation. For supply chain pipelines, the practical upshot is significant: you can run multiple scenarios concurrently without fighting over resources, store complete pipeline outputs for audit and reproducibility, and plug pipeline execution into broader automated workflows through standard APIs.

3. Problems

3.1. The Monolithic Pipeline Anti-Pattern

The most common architecture for supply chain input generation is what I will call the monolithic sequential pipeline: a single script or job that runs a fixed sequence of data transformations and produces all required outputs in one go. This pattern makes perfect sense early on, when you have a small number of inputs, a limited set of downstream consumers, and infrequent scenario changes. The trouble is that these conditions rarely persist. As they change, the monolithic pattern introduces structural limitations that compound over time.

Parameter Opacity. In a monolithic pipeline, business logic and configuration parameters live together in the same codebase. Utilization percentages, cost coefficients, date ranges, topology rules, scenario flags, they are all embedded directly in the transformation logic rather than pulled out into a configuration layer. The practical result is that anyone who needs to understand or validate the assumptions behind a model run either needs access to the source code or has to rely on documentation that may be incomplete or out of date. When parameters are wrong as they frequently are during transitions between planning cycles, the errors may not surface until model runs have completed and outputs have

been reviewed, wasting considerable computational time and delaying planning decisions.

Execution Indivisibility. A monolithic pipeline runs as a single unit: it either finishes or it fails. When something breaks partway through, all intermediate outputs get thrown away and the whole thing has to start over from the beginning. For pipelines with runtimes of 25–30 minutes or longer, this is a significant operational cost, particularly during planning cycles when multiple reruns are common as analysts iterate on parameters and rules.

Scenario Proliferation Challenges. Supply chain planning typically involves evaluating multiple scenarios: different demand assumptions, alternative network configurations, sensitivity analyses around key parameters. With a monolithic pipeline, supporting multiple scenarios means either maintaining separate copies of the entire codebase (which leads to divergence and duplication) or running scenarios one after another (which limits throughput). When the number of scenarios grows, say from 3 to 14 in a sensitivity analysis, the operational burden scales linearly, making comprehensive scenario coverage impractical.

Version Control and Auditability Deficits. Without formal version control and code review, changes to a monolithic pipeline are hard to track, review, or reverse. When someone changes a parameter or updates a business rule, the reasoning behind that change may exist only in the head of the engineer who made it. Over time, as team composition shifts and institutional knowledge walks out the door, the pipeline becomes increasingly opaque what Kleppmann [4] aptly calls "archaeological code," where the current state can only be understood by reconstructing the history of decisions that produced it.

3.2. Quantitative Impact of Architectural Limitations

The operational consequences of monolithic architectures can be quantified along several dimensions. Consider a pipeline that produces 150+ input files from tens of millions of data points, with an end-to-end runtime of 25–30 minutes. If something breaks at the 80% mark, you lose roughly 20–24 minutes of compute time to restart. If this happens three times during a planning cycle, the total wasted time approaches 60–72 minutes, time that directly delays model inputs and, by extension, optimization results.

We can express the pipeline efficiency ratio η as:

$$\eta = T_{\text{useful}} / T_{\text{total}} = (T_{\text{total}} - T_{\text{wasted}}) / T_{\text{total}}$$

where T_{useful} is time spent on productive computation and T_{wasted} is time spent re-executing work after failures. For a monolithic pipeline with a 25% failure rate and failures occurring on average at the 70% completion mark, η comes out to roughly 0.825 meaning about 17.5% of total compute time goes to waste on redundant execution. A modular pipeline with checkpoint-based recovery brings T_{wasted} close to zero, pushing η toward 1.0.

Scenario throughput tells a similar story. The throughput capacity S of a monolithic pipeline is constrained by sequential execution:

$$S_{\text{monolithic}} = 1 / T_{\text{pipeline}}$$

A modular pipeline supporting concurrent execution achieves:

$$S_{\text{modular}} = N_{\text{concurrent}} / T_{\text{pipeline}}$$

where $N_{\text{concurrent}}$ is the number of scenarios that can run simultaneously. With $N_{\text{concurrent}} = 10$, the modular architecture delivers a 10x improvement in scenario throughput, making comprehensive sensitivity analyses feasible where they previously were not.

4. Problem Use Cases

4.1. Manufacturing: Multi-Facility Production Network Optimization

Consider a large automotive manufacturer that operates production facilities across multiple regions. Optimization models determine how production capacity is allocated, how components flow through the supply network, and when finished goods ship to distribution centers. The input generation pipeline for these models pulls data from ERP systems, supplier databases, logistics platforms, and demand forecasting engines, more than 80 distinct input sources, producing tens of millions of data points per planning cycle.

Under a monolithic architecture, the planning team faces a recurring problem: when market conditions shift rapidly as they did during semiconductor shortages and the logistics disruptions of recent years, the team needs to evaluate multiple alternative network configurations simultaneously to find the most cost-effective response. But the monolithic pipeline cannot run scenarios concurrently, and running them sequentially means planning decisions must be made before the full range of alternatives has been explored [7].

4.2. Healthcare: Pharmaceutical Distribution Network Planning

A pharmaceutical distributor manages a temperature-controlled supply chain spanning hundreds of distribution points. Optimization models govern inventory placement, transportation routing, and capacity allocation across cold chain facilities. Because pharmaceutical products are sensitive to temperature excursions, the models must incorporate detailed facility-level operational constraints like shift structures, throughput limits, processing times that change frequently as facilities are upgraded or procedures revised.

Here the challenge of parameter management in a monolithic pipeline is especially acute. When a facility's operational parameters change, someone has to update the pipeline code. If those parameters are hardcoded in the transformation logic rather than sitting in a configuration file, the update requires a code change, a deployment, and a validation cycle, a process that can take days when the organization lacks formal version control. During that window, the optimization models run with stale parameters,

potentially generating recommendations that are physically infeasible [2].

4.3. Retail: Omnichannel Fulfillment Optimization

A large retailer operates a network of distribution centers, fulfillment centers, and retail stores. Optimization models determine how online orders are allocated across fulfillment nodes to minimize cost while meeting delivery commitments. The input generation pipeline must integrate real-time inventory data, current capacity constraints, transportation cost matrices, and demand forecasts, data sources with different update frequencies and latency requirements.

The retailer's planning team identified a need to expand their scenario analysis from 3 standard scenarios (baseline, peak demand, and disruption) to 14 scenarios capturing the network state across a two-week planning horizon on a day-by-day basis. Under a monolithic architecture, this would mean maintaining 14 separate copies of the pipeline codebase, an approach that is operationally impractical and creates serious risk of divergence between scenario-specific modifications and the main pipeline [5].

4.4. Food and Beverage: Cold Chain Logistics Planning

A food and beverage company manages a perishable goods supply chain with strict time-window constraints on transportation and storage. Optimization models determine routing, scheduling, and inventory decisions across production facilities, cold storage warehouses, and distribution points. The input pipeline must incorporate product-specific shelf life parameters, facility-specific temperature maintenance costs, and carrier-specific transit time distributions, parameters that vary considerably across product categories and shift as new products are introduced.

The company's data engineering team observed that the most frequent cause of optimization model failures is not algorithmic but data-related: incorrect or stale parameter values in the input files. In a monolithic pipeline where parameters are embedded in code, finding and fixing these errors requires source code access and understanding of the transformation logic, capabilities that tend to be concentrated in a handful of technical staff rather than distributed across the planning organization [3].

5. Solution: A Modular, Configuration-Driven Pipeline Framework

5.1. Architectural Principles

The framework I propose is organized around four architectural principles that directly address the limitations described above. None of these principles are novel in isolation; each has precedent in the software engineering and data systems literature but their combined application to supply chain input generation pipelines represents a meaningful contribution.

Principle 1: Modularity. The pipeline gets decomposed into independent modules, each responsible for generating a specific subset of output files. Modules are organized into execution layers based on their dependency relationships, with

a DAG engine working out the correct execution order automatically. Each module can run independently, which means engineers can iterate on specific parts of the pipeline without rerunning the whole thing.

Principle 2: Configuration-Driven Execution. All parameters that govern pipeline behavior like cost coefficients, date ranges, scenario identifiers, topology rules, output paths get externalized into a single configuration file. This file is human-readable, version-controlled alongside the pipeline code, and can be passed programmatically when triggering runs via API. Business stakeholders can review and validate configuration values before a run begins, which cuts down on the silent parameter errors that plague monolithic systems.

Principle 3: Multi-Scenario Isolation. Each pipeline run is scoped by a unique identifier, with all outputs like intermediate results, checkpoints, final files written to an isolated storage path. Multiple scenarios can execute concurrently without output conflicts, and experimental runs cannot overwrite production outputs. Engineers working on scenario-specific modifications can develop on separate branches while relying on the central codebase for unmodified modules.

Principle 4: Cloud-Native Deployment. The pipeline is deployed on cloud infrastructure providing containerized compute, workflow orchestration, structured logging, and infrastructure-as-code. This enables reproducible deployments, programmatic triggering from downstream workflows, cross-account execution for multi-team organizations, and automatic notification on run completion.

5.2. DAG Execution and Dependency Management

Using a directed acyclic graph to model pipeline dependencies gives you several operational advantages over sequential execution. First, independent modules in the same execution layer can run concurrently, cutting total runtime. Second, the DAG provides a formal representation of data lineage that is useful for debugging, auditing, and impact analysis. Third, it enables targeted re-execution: when a module fails or needs updating, only that module and its downstream dependents need to rerun, not the entire pipeline.

Formally, the DAG structure can be represented as $G = (V, E)$, where V is the set of processing modules and E is the set of dependency relationships. For a module v_i in V , the set of required inputs is $I(v_i) = \{outputs(v_j) : (v_j, v_i) \in E\}$, and execution order is determined by a topological sort of G . The pipeline engine verifies that all required inputs are available before starting each module and records a checkpoint on successful completion.

5.3. Checkpoint-Based Failure Recovery

One of the most important operational improvements the modular architecture provides is checkpoint-based failure recovery. After each module finishes successfully, its outputs are persisted to durable storage and a checkpoint record is

written. If a later module fails, the pipeline can pick up from the last good checkpoint instead of starting over.

The time savings are straightforward to quantify. Let T_i be the runtime of module i , and p_i the probability it fails. Under a monolithic pipeline (no checkpoints), expected wasted time is:

$$E[T_{\text{wasted_monolithic}}] = \sum_i p_i \cdot \sum_{\{j \leq i\}} T_j$$

Under a modular pipeline with checkpoints:

$$E[T_{\text{wasted_modular}}] = \sum_i p_i \cdot T_i$$

For a pipeline with 18 modules and a uniform failure probability of 5% per module, the modular architecture reduces expected wasted time by roughly 85% compared to the monolithic approach.

5.4. Configuration Management and Parameter Visibility

Pulling pipeline parameters out into a configuration file addresses one of the most operationally painful limitations of monolithic design. A well-structured configuration file gives you a single authoritative source of truth for every parameter governing a pipeline run, so stakeholders can review and validate assumptions before execution begins.

The configuration file should be organized into clearly labeled sections corresponding to different aspects of pipeline behavior: temporal parameters (date ranges, planning cycle identifiers), network parameters (topology rules, facility flags), cost parameters (utilization rates, cost coefficients), and execution parameters (output paths, scenario identifiers). This organization lets different stakeholder groups like network planners, finance analysts, data engineers review the sections relevant to their domain without needing to understand the full configuration.

6. Uses

6.1. Planned Optimization Cycles

The most straightforward application of the modular framework is in regularly scheduled optimization cycles, the periodic runs that reshape network design, update routing logic, or recalibrate cost parameters in response to changing conditions. Here the framework's value is primarily about reliability and auditability: configuration values can be reviewed and approved before a run begins, intermediate outputs can be inspected at each stage, and the complete provenance of every output file can be traced back to source data and transformation logic.

For organizations running optimization cycles monthly or quarterly, the framework also enables a more disciplined approach to parameter management. Instead of updating hardcoded values in a script and hoping you caught everything, planners work with a structured configuration file that explicitly lists all parameters and their current values. This reduces the risk of planning cycles proceeding on stale or wrong assumptions [6].

6.2. Ad-Hoc Scenario Analysis

Perhaps the biggest operational win from the modular framework is its support for ad-hoc scenario analysis, evaluating specific "what-if" questions that come up between planned optimization cycles. When a facility's capacity changes, a new transportation lane opens up, or cost parameters shift, planners need to quickly assess the impact on network performance without waiting for the next scheduled cycle.

The framework supports this through two mechanisms. First, the configuration-driven model lets planners create scenario-specific configurations that modify only the relevant parameters while inheriting everything else from a baseline. Second, the multi-scenario isolation model lets multiple scenarios run concurrently, enabling rapid comparison of alternatives. Together, these mechanisms reduce the time from conceiving a scenario to having analytical results from days (under a monolithic pipeline) to hours.

6.3. Sensitivity Analysis and Research Applications

For organizations that use optimization models to inform strategic decisions like facility location, network topology, service level commitments, sensitivity analysis is essential for understanding how recommendations change as assumptions vary. The modular framework supports systematic sensitivity analysis by enabling programmatic generation of configuration variants and concurrent execution of the resulting scenarios.

To put some numbers on it: consider a sensitivity analysis that evaluates network performance across 14 demand scenarios, each representing one day within a two-week planning window. Under a monolithic pipeline, that requires 14 sequential runs, a total of 350–420 minutes for a 25–30 minute pipeline. Under the modular framework with concurrent execution, all 14 scenarios can kick off simultaneously, with results available in 25–30 minutes regardless of how many scenarios you run. That 14x improvement in throughput makes comprehensive sensitivity analyses operationally practical for the first time.

6.4. Multi-Team and Cross-Functional Integration

In large organizations, supply chain optimization models serve multiple teams with different technical capabilities and integration needs. Data scientists may want to trigger pipeline runs programmatically as part of model training workflows. Business analysts may need to inspect intermediate outputs for data quality validation. Operations teams may need automatic notifications when new inputs are ready.

The framework's cloud-native deployment model supports these diverse patterns through standard API interfaces. Teams can trigger pipeline runs by passing configuration payloads to an orchestration API, receive notifications on completion, access outputs from isolated storage paths, and monitor execution through structured logging. This API-driven model reduces coordination overhead between teams and lets the pipeline function as a shared service rather than a centrally operated job.

7. Impact

7.1. Operational Efficiency

The most immediately measurable impact of the modular framework is the reduction in operational overhead. Eliminating separate pipeline copies for different scenarios removes the risk of logic divergence. Externalizing configuration parameters cuts the time to prepare a pipeline run from hours (when you have to find and update parameters in code) to minutes (when you can review and modify a structured configuration file).

Quantitative benchmarks from comparable pipeline modernization efforts in the data engineering literature suggest that modular architectures with checkpoint recovery reduce average failure recovery time by 70–85% versus monolithic designs [5]. For a planning organization running 50+ pipeline executions per quarter, this translates to a substantial reduction in wasted compute time and a real improvement in how quickly the planning process can move.

7.2. Data Quality and Auditability

The configuration-driven model provides a significant boost to data quality by making parameter values visible and reviewable before execution begins. When parameters live in code, errors typically get caught after model runs have completed at which point the cost includes not just fixing the parameter but rerunning the pipeline and the optimization model. When parameters sit in a configuration file, errors can be caught during pre-run review, before any compute resources have been burned.

Version control integration further strengthens data quality by creating an auditable record of every change to pipeline logic or configuration. This record lets organizations answer questions that come up regularly during post-hoc analysis: What parameters were used in that model run? When was this business rule changed, and why? What is different between the inputs from two consecutive planning cycles? Without version control, these questions are difficult or impossible to answer reliably [4].

7.3. Organizational Scalability

Perhaps the most strategically important impact is what the framework does for organizational scalability, the ability to expand optimization capabilities without a proportional increase in data engineering headcount. Under a monolithic architecture, every new scenario type, downstream consumer, or data source requires modifications to the central pipeline that must be carefully coordinated to avoid breaking existing functionality. That coordination overhead grows with the number of active use cases and can become a real bottleneck as optimization capabilities mature.

The modular framework addresses this by letting different teams develop and test modifications to specific pipeline modules independently, using the central codebase for unmodified modules and merging changes through formal code review. This is analogous to the microservices pattern in software engineering [6], where independent teams develop

and deploy individual services without needing to coordinate across the entire system.

7.4. Cross-Industry Applicability

The framework described here is not tied to any particular industry or optimization model type. The architectural principles such as modularity, configuration-driven execution, multi-scenario isolation, cloud-native deployment apply wherever complex input generation pipelines feed downstream analytical models. In manufacturing, the framework can support production planning and capacity optimization. In healthcare, pharmaceutical distribution and clinical supply chain planning. In retail, omnichannel fulfillment and inventory optimization. In food and beverage, cold chain logistics and perishable goods distribution [8].

The generalizability reflects a broader observation about supply chain data engineering: the challenges of managing complex, multi-source input generation pipelines are structural rather than domain-specific. Organizations across industries face the same fundamental tension between the simplicity of monolithic design and the flexibility that mature optimization capabilities demand. The modular framework offers a principled way to resolve that tension, adaptable to the specific requirements of different domains.

8. Conclusion

This paper has presented a generalized framework for designing and implementing modular, configuration-driven data pipelines for supply chain optimization. The framework addresses a set of structural limitations common to monolithic pipeline architectures—parameter opacity, execution indivisibility, scenario proliferation challenges, and version control deficits—through the application of four architectural principles: modularity, configuration-driven execution, multi-scenario isolation, and cloud-native deployment.

The analysis in Sections 3 and 5 shows that modular architectures with checkpoint-based recovery can reduce expected wasted compute time by roughly 85% compared to monolithic designs, while concurrent multi-scenario execution provides throughput improvements that scale linearly with available execution slots. These quantitative gains translate directly into operational benefits: faster planning cycles, more comprehensive scenario coverage, and better responsiveness to shifting business conditions.

The use cases in Section 4 illustrate the framework's applicability across manufacturing, healthcare, retail, and food and beverage. In each context, the core challenge is the same: organizations need to generate large, accurate, consistently structured input datasets for downstream optimization models, and the pipelines responsible for this work must be flexible enough to support multiple scenario types, reliable enough to recover from failures without starting over, and transparent enough for non-technical stakeholders to review the assumptions driving model runs.

The impact analysis in Section 7 highlights three dimensions of value: operational efficiency (reduced pipeline

management overhead), data quality and auditability (improved parameter visibility and change tracking), and organizational scalability (expanding optimization capabilities without proportional increases in engineering headcount). These dimensions reinforce each other: better data quality means fewer pipeline failures, which reduces operational overhead; better organizational scalability enables more comprehensive scenario coverage, which improves the quality of optimization recommendations.

Several directions for future research emerge from this work. First, integrating real-time validation checks within the modular pipeline, verifying output consistency between modules before proceeding downstream, looks promising for further improving data quality. Second, applying machine learning to pipeline monitoring, predicting module failure likelihood based on historical execution data and alerting operators proactively could reduce operational overhead further. Third, developing standardized configuration schemas for supply chain input generation would facilitate the exchange of best practices across organizations and industries.

The broader point is that data engineering infrastructure deserves recognition as a first-class concern in supply chain optimization research and practice. The quality of optimization recommendations is bounded by the quality of the inputs they depend on, and the quality of those inputs is bounded by the architectural soundness of the pipelines that produce them. Organizations that invest in modular, configuration-driven pipeline architectures are not just improving their data engineering operations—they are expanding the frontier of what their optimization capabilities can achieve.

References

1. M. S. Daskin, *Network and Discrete Location: Models, Algorithms, and Applications*, 2nd ed. Hoboken, NJ: Wiley, 2013.
2. U. Dhembare, "Data engineering: The critical foundation for scalable supply chain optimization models," *International Journal of Science and Management Studies*, vol. 8, no. 2, pp. 45–62, 2025.
3. N. R. Sanders and M. Swink, "Integrating big data across the supply chain," *Supply Chain Management: An International Journal*, vol. 24, no. 4, pp. 1–18, 2019.
4. M. Kleppmann, *Designing Data-Intensive Applications: The Big Ideas Behind Reliable, Scalable, and Maintainable Systems*. Sebastopol, CA: O'Reilly Media, 2017.
5. M. Zaharia, R. S. Xin, P. Wendell, T. Das, M. Armbrust, A. Dave, X. Meng, J. Rosen, S. Venkataraman, M. J. Franklin, A. Ghodsi, J. Gonzalez, S. Shenker, and I. Stoica, "Apache Spark: A unified engine for big data processing," *Communications of the ACM*, vol. 59, no. 11, pp. 56–65, Nov. 2016.
6. S. Newman, *Building Microservices: Designing Fine-Grained Systems*, 2nd ed. Sebastopol, CA: O'Reilly Media, 2021.
7. U. Dhembare, "A progressive framework for supply chain optimization: From rule-based logic to advanced mathematical models," *Kevin Open Science Journal of AIML, Data Science, Robotics*, vol. 3, no. 1, pp. 12–29, 2025.
8. M. Armbrust, A. Fox, R. Griffith, A. D. Joseph, R. Katz, A. Konwinski, G. Lee, D. Patterson, A. Rabkin, I. Stoica, and M. Zaharia, "A view of cloud computing," *Communications of the ACM*, vol. 53, no. 4, pp. 50–58, Apr. 2010.
9. G. Wang, A. Gunasekaran, E. W. T. Ngai, and T. Papadopoulos, "Big data analytics in logistics and supply chain management: Certain investigations for research and applications," *International Journal of Production Economics*, vol. 176, pp. 98–110, Jun. 2016.
10. S. Tiwari, H. M. Wee, and Y. Daryanto, "Big data analytics in supply chain management between 2010 and 2016: Insights to industries," *Computers & Industrial Engineering*, vol. 115, pp. 319–330, Jan. 2018.