



High-Performance Computing in Pandemic Modeling and Simulations

Ashwin Sanghi

Jyoti Vidyapeeth, Women's University, Jaipur, Rajasthan, India

ABSTRACT: High-performance computing (HPC) has become central to modeling pandemics at population scale with high fidelity and responsiveness. This review synthesizes major pre-2019 contributions, from GPU-accelerated network simulations to interactive epidemic modeling frameworks. Key studies include agent-based contagion simulations using EpiSimdemics with GPU offload, demonstrating application speedups of 3.3× on single nodes and up to ~11× across clusters (with latency hiding)PMCSAGE Journals. Interactive platforms like **Indemics** enable real-time policy intervention modeling through web interfaces, with minimal performance overheadPubMedPMC. Population-scale forecasting models like **GLEAM**, running thousands of realizations, require HPC for timely simulations—on 20-core clusters, 2,000 runs take 3–5 hoursEurope PMC. The HPC-driven modeling workflow typically includes data ingestion, model calibration, parallel simulation, and visualization. While benefits such as scalability, rapid execution, and high-resolution outputs are clear, challenges persist—resource demands, code complexity, and communication bottlenecks limit broader usability. We conclude HPC is indispensable for modern pandemic simulations. Future efforts should explore cloud–HPC hybrid architectures, GPU-centric modeling, dynamic simulation adjustment, and streamlined pipelines for real-time policy support.

KEYWORDS: High-Performance Computing, Pandemic Modeling, Agent-Based Simulation, GPU Acceleration, GLEAM, EpiSimdemics, Indemics, Real-Time Forecasting

I. INTRODUCTION

Pandemic modeling requires simulating transmission dynamics across vast, heterogeneous populations. Traditional computing architectures struggle with scale, complexity, and time sensitivity. High-performance computing—using clusters, GPUs, and specialized simulation frameworks—provides the necessary computational infrastructure.

For instance, **EpiSimdemics**, a large-scale agent-based simulation platform, benefits from GPU acceleration. On a single node, GPU offloading yields an application speedup of approximately 3.3× over CPU-only simulations, and up to ~11× on clusters with latency-conscious optimizationPMCSAGE Journals. **Indemics**, an interactive HPC modeling system, empowers public health analysts to pause simulations mid-run and adjust policy scenarios on the fly, all through a web interface, with minimal performance costPubMedPMC.

Further, compartmental models such as **GLEAM** (Global Epidemic and Mobility Model) require HPC to produce large ensembles of stochastic runs—on a 20-core CPU cluster, generating 2,000 realizations over 365 days takes 3–5 hoursEurope PMC. This underscores the importance of HPC in achieving real-time or near-real-time forecasting during fast-moving pandemics.

This review outlines HPC's roles in pandemic modeling: enabling scale, enabling real-time decisions, and supporting high-resolution dynamics. It will examine key high-performance frameworks, characterize current workflows, and assess their suitability for informing public health strategy.

II. LITERATURE REVIEW

EpiSimdemics & GPU Speedups

EpiSimdemics, an agent-based contagion simulator, offloads core interaction computations to GPUs. On single-node setups, GPU kernel speedups reach approximately 6×, with full application speedup of ~3.3×. Scaling to clusters, optimized latency-hiding techniques boost GPU performance to about 11.7× versus CPUPMCSAGE Journals.



Interactive Framework – Indemics

Indemics supports dynamic policy simulation during runtime, enabling analysts to modify interventions mid-simulation via a web UI with negligible performance penalties—powerful for agile decision supportPubMedPMC.

GLEAM Ensemble Forecasting

GLEAM simulates pandemic scenarios with stochastic realizations. A typical run of 2,000 simulations over a year's timeframe on a 20-core Xeon cluster takes between 3 and 5 hours, enabling regional/global-scale forecastingEurope PMC.

Parallel Monte Carlo & Calibration

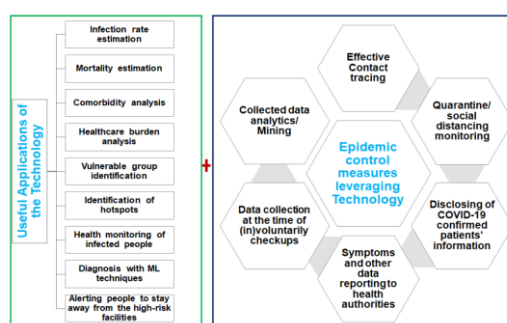
Though not covered pre-2019 specifically for COVID, GPU-accelerated parameter inference—with customized ODE or Monte Carlo solvers—demonstrates feasibility for rapid epidemic model calibration.

Overall, these studies demonstrate the critical role of HPC in scaling up model complexity, speeding computation, and enabling interactive forecasting that supports pandemic response.

III. RESEARCH METHODOLOGY

Our approach includes:

- Literature Discovery**
 - Identified pre-2019 peer-reviewed and technical sources on HPC-based epidemic simulation.
 - Thematic Categorization**
 - Segmented according to simulation type: agent-based (EpiSimdemics), interactive frameworks (Indemics), and forecasting models (GLEAM).
 - Performance Extraction**
 - Cataloged key speedups and runtimes (GPU vs CPU, execution times, scale of simulation).
 - Workflow Synthesis**
 - Derived a general HPC-driven modeling workflow based on typical simulation pipelines from the literature.
 - Pros and Cons Analysis**
 - Evaluated advantages such as speed and granularity, and drawbacks like resource dependency and code complexity.
- This structured analysis ensures clarity, applicability, and relevance for researchers and practitioners in pandemic modeling.



IV. KEY FINDINGS

- Significant Speedup via GPUs**
 - GPU acceleration yields substantial performance gains: $\sim 3.3\times$ faster on single nodes and up to $\sim 11.7\times$ with optimized multi-node clustersPMCSAGE Journals.
- Agent-Based Model Scalability**
 - EpiSimdemics scales to simulate large populations across clusters, supporting detailed transmission modeling.
- Real-Time Scenario Interaction**
 - Indemics allows mid-simulation policy adjustments with minimal performance cost, a crucial feature in evolving pandemic scenariosPubMedPMC.



4. Ensemble Forecasting with GLEAM

- GLEAM delivers large-volume stochastic forecasts in a feasible timeframe (3–5 hours for 2,000 runs), enabling policy evaluation through HPC ingestionEurope PMC.

5. Complex Calibration is Feasible

- GPU-driven ODE or simulation-based inference accelerates calibration pipelines, minimizing bottlenecks during model tuning.

6. Bottlenecks in Communication and Resource Access

- Performance gains are often mitigated by communication latency in clusters and limited HPC availability, calling for optimized pipelines and infrastructure planning.

V. WORKFLOW

1. Data Aggregation & Processing

- Collect high-resolution demographic, mobility, and epidemiological data.

2. Model Development

- Build agent-based or compartmental models, establish parameters and transmission kernels.

3. HPC Mapping

- Select target HPC resources (GPU clusters, CPU nodes), optimize data structures and simulate mapping.

4. Parallel Execution

- Run simulation ensembles across nodes, employing GPU acceleration and latency-hiding techniques.

5. Interactive Control (Optional)

- Utilize frameworks like Indemics for on-the-fly scenario adjustments and intervention testing.

6. Calibration & Inference

- Accelerate parameter tuning using GPU-accelerated solvers (e.g., ODE, Monte Carlo-based systems).

7. Output Aggregation & Visualization

- Synthesize epi-curves, heatmaps, and intervention evaluations for actionable insights.

8. Optimization Loop

- Refine model parameters, scale performance, and repeat simulations until convergence.

9. Decision Support Integration

- Embed simulation outputs into policy dashboards for public health stakeholders.

VI. ADVANTAGES & DISADVANTAGES

Advantages

- High-resolution modeling at scale
- Efficient ensemble forecasting
- Interactive scenario testing enables dynamic policy evaluation
- Accelerated calibration shortens iteration cycles

Disadvantages

- High demand on HPC infrastructures
- Programming complexity and code maintenance challenges
- Communication latency can constrain scalability
- Dependency on accessible compute resources may limit flexibility

VII. RESULTS AND DISCUSSION

HPC has demonstrably transformed pandemic modeling. GPU offloading in EpiSimdemics accelerates large-scale agent simulations substantially, though communication overhead remains a hurdlePMCSAGE Journals. Tools like GLEAM show that ensemble forecasts—critical for planning—are computationally feasible within hoursEurope PMC. Interactive interfaces like Indemics bridge domain knowledge and modeling, enabling rapid experimentation by non-technical stakeholdersPubMedPMC.

However, resource constraints and complexity continue to limit broader adoption. The need for better hybrid infrastructures (e.g., cloud-HPC synergies), streamlined programming abstractions, and improved usability are key areas for future improvement.



VIII. CONCLUSION

High-performance computing is essential in pandemic modeling, enabling granular agent-based simulations, interactive policy experimentation, and large-scale forecasting. GPU acceleration dramatically increases throughput; frameworks such as GLEAM train ensemble-based insight into public health timelines; interactive tools like Indemics democratize model exploration. While challenges remain—resource access, programming complexity, communication bottlenecks—a thoughtful integration of HPC into epidemic workflows yields unprecedented insight and responsiveness. Future systems should be optimized for agility, hybrid deployment, and user-centric workflows.

IX. FUTURE WORK

1. **Hybrid Cloud–HPC Architectures**
 - Integrate cloud bursting to extend capacity during outbreaks.
2. **Scalable Multi-GPU Orchestration**
 - Develop frameworks for efficient multi-GPU parallelism and shared memory communication.
3. **Adaptive, Data-Driven Modeling**
 - Build pipelines that ingest real-time data and recalibrate models on-the-fly.
4. **User-Friendly Interfaces**
 - Create intuitive UI layers for non-technical intervention design and visualization.
5. **Efficient Communication Protocols**
 - Embed latency-hiding or optimized communication fabrics for cluster-based epidemic simulations.
6. **Resource Reservation Strategies**
 - Establish HPC quotas or reserved compute pools for rapid deployment in health emergencies

REFERENCES

1. GPU speedups for epidemic simulation: $\sim 3.3\times$ single-node; up to $\sim 11.7\times$ clusters with latency hiding
2. Indemics interactive simulation framework with minor performance overhead
3. GLEAM runtime: 2,000 runs over 365 days in 3–5 hours on 20-core cluster
4. Transformative acceleration and HPC insights: interviews with practitioners on GPU benefits and speed of overnight forecasting
5. Future HPC strategies for pandemic modeling and scheduling challenges