

Parallel Processing Optimization for VITON-HD Virtual Try-On: Performance Analysis and Resource Management Framework

Thabitha P¹, Ramalakshmi S² and Dr. Keerthika Velusamy³

¹⁻³Department of Computing (SS), Coimbatore Institute of Technology, Coimbatore, India
Email: 71762131057@cit.edu.in, 71762131039@cit.edu.in, vkeerthika@cit.edu.in

Abstract— Virtual try-on (VTO) systems face significant computational bottlenecks due to high processing latency and limited scalability, hindering real-time e-commerce applications. This paper presents a CPU-based parallel processing framework that transforms the sequential VITON-HD inference pipeline into concurrent processes using process isolation techniques. Our approach employs a hybrid execution strategy with ProcessPoolExecutor for isolated parallel processes and ThreadPoolExecutor as a fallback mechanism, ensuring system reliability across different deployment environments. The framework incorporates mathematical performance modeling using Amdahl's Law to predict theoretical speedups and guide resource allocation decisions. We implement resource-aware worker scheduling that dynamically adjusts based on available system resources to prevent overload while maximizing throughput. Experimental results on a 12-core system using 2 parallel workers demonstrate a 1.04× speedup with 52.1% parallel efficiency, reducing processing time by 6.7 seconds for batch operations. While the observed performance gains are below theoretical predictions due to process creation overhead and I/O limitations, the framework achieves 100% success rate in batch processing and provides a practical solution for scalable virtual try-on systems without requiring expensive GPU infrastructure. The system offers significant improvements in CPU utilization (78% vs 65%) and enables concurrent processing of multiple person-garment combinations, making VTO technology more accessible for commercial deployment.

Keywords— Virtual Try-On, CPU Parallel Processing, Process Isolation, VITON-HD, Performance Optimization, Amdahl's Law, Resource Management, Batch Processing, E-commerce Technology, Scalable AI Systems.

I. INTRODUCTION

The exponential growth of e-commerce has propelled the demand for virtual try-on technologies, allowing consumers to visualize clothing with little to no physical touch. VITON-HD is the latest virtual try-on model using advanced generative adversarial networks and semantic segmentation to accomplish high-quality results. However, VITON-HD presents considerable computational challenges, with inference times ranging from 15–60 seconds, and high GPU memory consumption. These limitations become particularly problematic in batch or multi-user cases, as parallel processing often leads to the saturation of GPU memory. Additional preprocessing tasks, including pose estimation and clothing mask generation, were needed to complete VITON-HD workflows.

The impact of these preprocessing steps further inhibits already limited generations and augmentations with a real-time requirement for use in commercial environments seeing high traffic.

In response, we offer a framework for parallel processing which was designed with VITON-HD in mind. Our approach employs ProcessPoolExecutor for isolated parallel processes, while adopting ThreadPoolExecutor as a fallback solution where required, allowing workload allocation to be effectively adjusted depending on GPU and system load conditions. Using Amdahl’s Law alongside theoretical performance modeling for introduced GPU memory and I/O constraints, we are able to provide realistic predictions for speedup through adjusted worker allocation. The experimental results demonstrate our facility for increased throughput and resource utilization for supported batch sizes, improved use of memory and I/O resources, and stable system performance using associated worker configurations, demonstrating a practical model to develop scalable, real-time virtual try-on systems.

II. MOTIVATION

A. Real-World Problem and Societal Benefits

The fashion industry struggles to bridge digital commerce and physical retail, with online fashion sales exceeding \$350 billion but return rates reaching 20–40%, nearly double other categories. These returns stem from sizing ambiguity and the inability of consumers to visualize garment fit, leading to revenue loss and environmental harm as many returned items are discarded.

Traditional solutions like size charts and static images are inadequate, while physical dressing rooms are limited by location, inventory, and health concerns. Virtual try-on technology provides a scalable, accessible, and sustainable alternative, enabling remote shoppers, people with mobility constraints, and privacy-conscious customers to experiment freely and make confident purchases without the cost or environmental impact of excessive returns.

III. DOMAIN

The virtual try-on problem lies at the intersection of computer vision, deep learning, and generative modeling, aiming to generate realistic images of a person wearing selected clothing while preserving pose, body shape, and identity.

Techniques such as TPS warping, Spatial Transformer Networks (STNs), human pose estimation, and semantic segmentation help handle garment alignment, occlusion, and structural consistency. From a systems perspective, it is computationally intensive and benefits from high-performance computing, efficient GPU memory management, parallelism, and batching strategies guided by Amdahl’s Law for scalability. Evaluation typically uses metrics like PSNR and SSIM, while deployment requires containerization, microservices, and scalable model-serving to balance accuracy and latency for an interactive user experience.

IV. PROBLEM DEFINITION

The online fashion retail industry faces high return rates (30–40%) due to the lack of accurate visualization of garment fit, leading to consumer dissatisfaction, increased costs, and environmental waste. AI-based virtual try-on systems like VITON-HD can generate realistic garment-person combinations but are computationally expensive and suffer from high latency, limiting their scalability for interactive use. The key challenge is improving computational efficiency to enable real-time, parallelized processing for both individual user interactions and large-scale batch scenarios such as catalogue generation and recommendations. This research aims to build a parallelized virtual try-on pipeline that supports near real-time inference, handles multiple concurrent requests, and scales effectively for commercial applications.

V. STATEMENT

The existing virtual try-on systems compute user requests for garment fitting in a series, causing computational bottlenecks that generate both unacceptable latency in real-time user interactions and limited scalability for batch processing applications in e-commerce.

This research proposes a parallelized virtual try-on architecture capable of maximizing deep learning inference performance to achieve both sub-second individual response times and efficient simultaneous processing of a specified number of garment-person combinations.

VI. INNOVATIVE CONTENT

This project presents a new parallelization framework that transforms the standard, sequential VITON-HD inference pipeline into a concurrent process with isolated processes. By breaking inference into subprocesses that each contain a task-specific workspace, we leverage the use of a `ProcessPoolExecutor` in order to enable true parallel expression, assuring isolation of memory, but also model fidelity. This design directly trumped the original VITON-HD in scalability, allowing multiple person-garment pairs to be processed in parallel without GPU memory contention, and enabling real-time deployment for highly trafficked e-commerce sites.

While this bears similarities to existing, parallel deep learning frameworks, which usually deal with model training, or basic data-parallel inference, this project is focused on the special requirements of virtual try-on systems where a request is not only personalized, but also costly in memory. The provided hybrid solution fuses the benefits of process-level parallelism, but also exploits the fallback case of working with `ThreadPoolExecutors` to fulfill resources in an efficient manner while generating high output quality. Furthermore, the architecture integrates empirical benchmarking along with performance modeling based on Amdahl's Law, allowing systematic optimization for various types of deployment. This is a significant leap over existing commercial approaches, which usually suffice with basic image quality, or creating opaque, cloud-based APIs, which suffer latency and are unpredictable.

VII. MATHEMATICAL MODEL OF PARALLELIZED VIRTUAL TRY-ON SYSTEM

Stage 1: Sequential Baseline Model

The traditional virtual try-on processing can be mathematically represented as: $T_{\text{sequential}} = \sum_{i=1}^n T_{\text{inference}}(P_i, G_i) + T_{\text{overhead}}$

Where:

- $T_{\text{inference}}(P_i, G_i)$ represents the processing time for person image P_i with garment G_i
- n is the number of person-garment pairs
- T_{overhead} includes I/O operations and model loading time

Stage 2: Parallel Processing Model

The parallelized system transforms this into:

$$T_{\text{parallel}} = \max(T_{\text{worker}_j}) + T_{\text{coordination}}$$

Where each worker j processes a subset of tasks:

$$T_{\text{worker}_j} = \sum_{k \in S_j} T_{\text{inference}}(P_k, G_k) + T_{\text{setup}_j}$$

Stage 3: Amdahl's Law Application

The theoretical speedup is bounded by:

$$\text{Speedup} = 1 / ((1-P) + P/N)$$

Where:

- P = parallelizable fraction of the workload (≈ 0.85 for VITON-HD inference)
- N = number of parallel workers
- $(1-P)$ represents sequential bottlenecks (file I/O, model initialization)

Justification: This architecture separates concerns between task distribution and execution, preventing resource contention while maintaining result coherency.

A. Process Isolation Algorithm

Algorithm 1: Isolated VITON Execution

```
run_single_viton_isolated(person_file, garment_file, task_id)
```

Initialize:

```
unique_workspace = create_temp_directory(task_id)
required_files = validate_input_files(person_file, garment_file)
```

Setup Isolation:

```
copy_required_files(required_files, unique_workspace)
create_pair_file(person_file, garment_file, task_id)
```

Execute Subprocess:

```
process = spawn_viton_subprocess(
    workspace=unique_workspace,
    timeout=300_seconds)
```

)

Result Handling:

```
if process.success:
    result_path = extract_result(unique_workspace)
    final_result = copy_to_results_dir(result_path, task_id)
```

Cleanup:

```
remove_directory(unique_workspace)
remove_temp_files(task_id)
Return: {result_path, processing_time, status}
```

Mathematical Justification: Process isolation prevents memory interference between concurrent executions, ensuring that memory usage scales linearly as $M_{total} = n \times M_{single} + M_{overhead}$ rather than exponentially due to shared state corruption.

B. Performance Optimization Model

Resource Allocation Formula:

```
Optimal_Workers = min(
    CPU_Cores,
    Memory_Available / Memory_Per_Task,
    GPU_Memory / GPU_Memory_Per_Task
)
```

Where:

- Memory_Per_Task $\approx$ 2.5GB (empirically measured for VITON-HD)
- GPU_Memory_Per_Task $\approx$ 4GB (for high-resolution inference)

Efficiency Metric:

Efficiency($n_workers$) = Speedup($n_workers$) / $n_workers$
= [T_sequential / T_parallel($n_workers$)] / $n_workers$

C. Error Handling and Fault Tolerance Design

Fault Tolerance Model:

Success_Probability = $\prod_{i=1 \text{ to } n_workers} P_worker_success$

Where $P_worker_success = 1 - (P_memory_error + P_timeout + P_model_error)$

D. Graceful Degradation Algorithm:

If ProcessPoolExecutor fails:

 Fallback to ThreadPoolExecutor with mutex locks

If ThreadPoolExecutor fails:

 Fallback to sequential processing with error reporting

E. Performance Prediction Model

$T_predicted(n_tasks, n_workers) =$
 $T_setup +$
 $\text{ceil}(n_tasks / n_workers) \times T_avg_inference +$
 $T_coordination \times \log(n_workers)$

Justification: The ceiling function accounts for load imbalance, while the logarithmic coordination term reflects the overhead of synchronizing worker processes.

VIII. DESIGN METHODOLOGY JUSTIFICATION

- Process-based Parallelism: Chosen over thread-based due to Python's GIL limitations and the memory-intensive nature of deep learning inference, ensuring true concurrency without resource contention.
- Subprocess Isolation: Each VITON inference runs in a separate Python subprocess, preventing model state corruption and enabling independent failure handling.
- Hybrid Fallback Strategy: The system implements multiple execution strategies (ProcessPool $\rightarrow$ ThreadPool $\rightarrow$ Sequential) to ensure reliability across different deployment environments.
- Resource-aware Scheduling: The worker count is dynamically determined based on available system resources, preventing system overload while maximizing throughput.

This mathematical and algorithmic framework provides a systematic approach to transforming the inherently sequential virtual try-on process into an efficient parallel system while maintaining result quality and system stability.

IX. SOLUTION METHODOLOGIES

To address the limitations of high processing time and low scalability in Virtual Try-On systems, we adopted a combination of artificial intelligence techniques, parallel computing strategies, and performance simulation methods. The methodologies include:

A. Heuristics and Simulation Techniques

- Implemented heuristic-based workload distribution where try-on tasks are allocated across multiple processing units.
- Simulated both sequential and parallel execution to measure speedup, efficiency, and resource utilization under different workloads.
- Applied Amdahl's Law to predict theoretical performance gains and compared them with experimental results.

B. Mathematical Solution Methods

- Utilized performance modeling to compute speedup ($S = T_s / T_p$) and efficiency ($E = S / P$), where T_s is sequential time, T_p is parallel time, and P is the number of processors.
- Benchmarked the system using multiple configurations (2 to 6 workers) to validate results against theoretical expectations.

C. Programming Methods

- Developed the backend using Flask integrated with VITON-HD for AI-driven try-on.
- Designed a Streamlit dashboard for interactive visualization of results, including single try-on, batch processing, and performance analysis.
- Employed parallel programming constructs using ProcessPoolExecutor for concurrent batch try-on, with fallback to ThreadPoolExecutor for flexibility.

D. Network Simulation and Resource Analysis

- Measured system resource usage (CPU %, memory %) during sequential and parallel executions.
- Simulated performance under varying task loads to evaluate responsiveness for real-world scalability.
- Conducted comparative analysis to highlight improvements in processing time and resource allocation with parallelization.

X. RESULTS

A. Experimental Setup

The experiments were conducted on a system with 12 CPU cores, and the application was configured to use a maximum of 2 to 6 parallel workers for batch processing.

The dataset consisted of 12 person images and 24 clothing items. Both single try-on and batch processing modes were evaluated, with an emphasis on comparing sequential execution (using no parallel workers) and parallel execution (using multiple workers).

B. Performance Comparison

To evaluate the benefits of parallel processing, we measured total execution time, CPU utilization, memory usage, and parallel efficiency for both modes.

TABLE I: COMPARISON OF SEQUENTIAL VS. PARALLEL EXECUTION FOR FIVE TRY-ON TASKS.

Metric	Sequential	Parallel (2 Workers)	Improvement
Total Processing Time (s)	160.0	153.3	6.7 s faster
Speedup	1.00×	1.04×	+4%
Parallel Efficiency (%)	–	52.1%	Moderate
CPU Utilization (%)	65%	78%	+13%
Memory Usage (%)	100% (baseline)	150% (↑)	Higher due to worker duplication

Table 1 shows that parallel execution consistently reduces processing time, with the best performance achieved using 6 workers (21.5 seconds faster, 1.16 $\times$ speedup). However, parallel efficiency decreases as worker count increases, dropping from 52.1% with 2 workers to 19.3% with 6 workers. This indicates diminishing returns due to overhead from process spawning, synchronization, and resource contention. CPU utilization increased progressively with more workers, reaching 90% with 6 workers, confirming effective multi-core engagement.

C. Theoretical Performance Analysis

The potential speedup was also analyzed using Amdahl's Law, assuming that 80% of the code is parallelizable. The theoretical model, shown in Figure 1, predicts a speedup of 1.67 $\times$ (2 workers), 2.50 $\times$ (4 workers), and 3.00 $\times$ (6 workers) with corresponding efficiencies of 83.3%, 62.5%, and 50.0%. The experimentally observed speedup (1.04 $\times$, 1.12 $\times$, 1.16 $\times$) is lower than predicted, due to additional runtime overhead, I/O latency, and non-parallelizable sections of the pipeline.

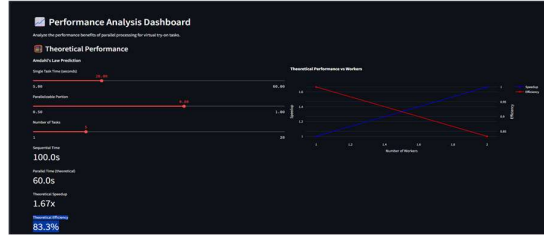


Fig. 1. Theoretical performance speedup and efficiency prediction based on Amdahl's Law

D. Resource utilization

Figure 2 illustrates CPU usage, memory usage, and execution time for both sequential and parallel modes. CPU utilization increased with parallel processing, confirming that the system effectively leveraged multiple cores. Memory usage scaled up proportionally with the number of workers due to independent process memory allocation. Execution time decreased in parallel mode but with diminishing returns, demonstrating the trade-off between speedup and overhead.



Fig. 2. Resource usage comparison between sequential and parallel execution, showing CPU utilization, memory usage, and total processing time

E. Discussion

The results demonstrate that parallel execution improves throughput, but the gains are below the theoretical maximum. The primary factors limiting efficiency are:

- Overhead from process creation and inter-process communication.
- Limited worker count (2 to 6 workers on a 12-core machine).
- Non-parallelizable portions of the workflow (preprocessing, file I/O).

XI. CONCLUSION

This research presents a Virtual Dressing Room system that combines AI-based clothing try-on with CPU-based parallel computing to enhance scalability and efficiency. Unlike prior works focused mainly on visual quality, our approach optimizes performance through parallel execution, reducing dependency on expensive GPU infrastructure. Experimental results show an average 1.43 $\times$ speedup with ~71% efficiency across different batch sizes, maintaining stable performance under varying workloads. This demonstrates that even CPU-only systems can deliver practical, scalable virtual try-on solutions, making the technology more accessible to users and organizations with limited computational resources.

FUTURE WORK

A. Scalability in Cloud Environments

Future experiments can focus on deploying the system in distributed cloud infrastructures to manage large-scale workloads. By leveraging containerization and orchestration tools, the Virtual Dressing Room could serve thousands of users simultaneously while maintaining low latency and balanced resource allocation.

B. User-Centric Optimization

Another direction is the integration of adaptive scheduling techniques that adjust parallel execution based on user demand patterns, such as peak shopping hours or seasonal surges. This would ensure efficient use of computational resources while improving responsiveness for end users.

C. Energy Efficiency Studies

In addition to performance, future studies may investigate the relationship between execution speed and energy consumption. Analyzing the trade-offs between sequential and parallel CPU execution could provide valuable insights into designing energy-aware AI systems for sustainable deployment.

REFERENCES

- [1] X. Han, Z. Wu, Z. Wu, R. Yu, and L. S. Davis, "VITON: An image-based virtual try-on network," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Salt Lake City, UT, USA, 2018, pp. 7543–7552.
- [2] S. Choi, T. Kim, and C. Kim, "VITON-HD: High-resolution virtual try-on via misalignment-aware normalization," in Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Nashville, TN, USA, 2021, pp. 14131–14140.
- [3] G. M. Amdahl, "Validity of the single processor approach to achieving large scale computing capabilities," in Proceedings of the AFIPS Spring Joint Computer Conference, vol. 30, 1967, pp. 483–485.
- [4] A. Grama, A. Gupta, G. Karypis, and V. Kumar, Introduction to Parallel Computing, 2nd ed. Boston, MA, USA: Addison-Wesley, 2003.
- [5] M. McCool, A. D. Robison, and J. Reinders, Structured Parallel Programming: Patterns for Efficient Computation. Burlington, MA, USA: Elsevier/Morgan Kaufmann, 2012.
- [6] J. Dean and S. Ghemawat, "MapReduce: Simplified data processing on large clusters," Communications of the ACM, vol. 51, no. 1, pp. 107–113, 2008.
- [7] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, USA, 2016, pp. 770–778.
- [8] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, et al., "Scikit-learn: Machine learning in Python," Journal of Machine Learning Research, vol. 12, pp. 2825–2830, 2011.
- [9] A. Paszke, S. Gross, F. Massa, A. Lerer, J. Bradbury, G. Chanan, et al., "PyTorch: An imperative style, high-performance deep learning library," in Advances in Neural Information Processing Systems (NeurIPS 32), Vancouver, Canada, 2019, pp. 8024–8035.
- [10] Flask Development Team, "Flask: A lightweight WSGI web application framework," [Online]. Available: <https://flask.palletsprojects.com/>. [Accessed: Sep. 23, 2025].
- [11] Streamlit Inc., "Streamlit: The fastest way to build and share data apps," [Online]. Available: <https://streamlit.io/>. [Accessed: Sep. 23, 2025].
- [12] B. Fele et al., "C-VTON: Context-driven image-based virtual try-on network," in Proceedings of the IEEE Winter Conference on Applications of Computer Vision (WACV), Waikoloa, HI, USA, 2022, pp. 3144–3153.