

AGENT-BASED VECTOR SEARCH ON GPU

AKBARBEK RAKHMATULLAEV

Spring 2025 Term Report

University of Washington

Jun 15, 2025

Project Committee:

Munehiro Fukuda, Ph.D., Committee Chair

Min Chen, Ph.D., Committee Member

Wooyoung Kim, Ph.D., Committee Member

I. Introduction

Vector search is an artificial intelligence and data retrieval method that employs mathematical vectors to represent and efficiently search through complex, unstructured data. It operates by linking similar mathematical representations of data and converting queries into these same vector formats. With both queries and data represented as vectors, the search for related data involves identifying the closest matches to the query vector, a process known as nearest neighbor search. Unlike traditional search algorithms, which rely on keywords, word frequency, or word similarity, vector search utilizes the distances within the vectorized dataset to identify similarity and semantic relationships. In today's world, vector search is primarily being used in domains such as e-commerce, content discovery and recommendation systems, natural language processing (NLP), and many more.

The market offers many solutions for vector search, however, according to the latest benchmarks [1][2], one of the best-performing libraries is NGT (Neighborhood Graph and Tree) [3][4]. The NGT index combines both a graph and a tree, where the graph's vertices represent searchable objects and the tree is used to subdivide the entire vector space into smaller regions. This makes NGT a top performer compared to other solutions. Although NGT performs well with reasonably large datasets, it may face scalability issues with extremely large datasets or high-dimensional data, as it can run only on one computing node/Central Processing Unit (CPU), and does not provide Graphics Processing Unit (GPU) support for Approximate Nearest Neighbor (ANN) [5].

The Multi-Agent Spatial Simulation (MASS) [6][7] library is a parallel and distributed computing framework designed for large-scale spatial and agent-based simulations. It operates primarily with two concepts: Places and Agents. Generally, Places are the spatial locations or cells that form the simulation environment and are capable of exchanging information with any other Places. Agents are active entities that can move between Places, perform actions, and interact with each other and their environment. MASS CUDA [8] is a library designed to facilitate the execution of parallel computations using mobile agents on GPUs.

When implementing vector search on a GPU, two primary algorithmic approaches are commonly used: Hierarchical Navigable Small World (HNSW)[9] and Inverted File Index with Product Quantization (IVFPQ)[10].

HNSW (Hierarchical Navigable Small World) is a graph-based approach for approximate nearest neighbor search, designed to efficiently navigate large-scale high-dimensional spaces. It constructs a multi-layered, small-world graph structure where search operations traverse the graph using greedy strategies. HNSW provides high recall rates and fast search times, making it a popular choice for large-scale vector search applications. However, the graph-based nature of HNSW requires significant memory overhead, and updates to the index can be computationally expensive. While MASS Graph — a graph-based extension of MASS — could potentially optimize HNSW with agent-based search, its development was not mature enough for practical implementation at the time of this work. Instead, IVFPQ was chosen due to its well-established efficiency in large-scale vector search and its compatibility with MASS CUDA, allowing for immediate integration of agent-based optimizations.

IVFPQ (Inverted File System with Product Quantization) is a widely used indexing technique that balances search efficiency and memory usage. It partitions the vector space into clusters and applies product quantization to compress the data, allowing fast approximate searches while reducing storage requirements. IVFPQ is particularly effective for large-scale nearest neighbor search, as it significantly reduces computational complexity. However, its performance depends on proper tuning of parameters such as the number of clusters and quantization levels. While it provides a strong foundation for efficient search, further enhancements using agent-based approaches like MASS CUDA could improve adaptability and dynamic workload distribution across GPUs, which is a part of this project.

The idea of this project is to try to leverage the capabilities of GPUs using mobile agents by the MASS CUDA library and improve IVFPQ algorithm to perform better. A more detailed explanation of the implementation is provided in the next sections.

II. Background

The vector searching landscape is highly dynamic, with frequent advancements and discoveries. Current vector search engines, such as NGT, SPTAG [11], ANNOY [12], etc., are effective but have various limitations in terms of scalability, speed, or precision. The idea behind this project is to utilize the best state-of-the-art methods and algorithms in terms of speed and precision and scale one of the most used algorithms IVFPQ using mobile agents. As a result, there is an opportunity to develop an agile vector search engine.

While existing solutions like NGT, SPTAG, ANNOY, etc. offer effective vector search, our deliverable stands out by scaling and enhancing the IVFPQ algorithm on GPU using mobile agents.

III. Challenges

This quarter, one of the biggest challenges was figuring out how to optimize the search part in terms of creation, distribution and proper usage of places and agents. After many iterations of different approaches, I came up with one that optimized usage of places and agents the most inside of the search part, which allowed me to save time to create places and have multiple queries at the same time. Another problem was the abundance of bugs within MASS CUDA itself and adjusting it to work as a simple importable library rather than core of application as it was before. Last but not least, bugs that occurred during the development, such as with memory allocation, and 30 minute long compilation time caused problems this quarter.

IV. Goals

The primary objective this quarter was to develop a methodology, and its implementation, for effectively utilizing MASS CUDA agents within IVFPQ, which has base implementation available from cuVS (CUDA Vector Search)[13] developed by Nvidia Corporation. A significant portion of the time was dedicated to finalize Place and Agent classes, and proper memory management. However, due to the complexity of the code, details of which will be discussed in the next section, there is still work to be done, particularly in the implementation of agents.

V. Implementation

The core idea behind our new approach is to distribute the search operation across multiple CUDA threads efficiently, leveraging MASS's ability to manage parallel computations. This approach was inspired by the work of Alex Li and Professor Fukuda[14]. Each dimension of a PQ-encoded vector is treated as a separate Place, which stores sorted values for that dimension along with their corresponding vector IDs. Agents are then dispatched to these Places to perform a binary search for the closest quantized value, ultimately contributing to an overall search result (see Figure 1). While significant progress has been made in designing the structure and implementing key components, the full implementation is still ongoing.

The implementation is structured around two custom classes: PQPlace and PQAgent. The PQPlace class extends MASS's Place class and serves as the core storage unit for each dimension's sorted PQ values and associated vector IDs represented as indices. It provides an efficient device-side method to find the closest match using binary search. The code is available in Listing 1. On the other hand, PQAgent is an agent class that can query a specific Place by executing the binary search algorithm and returning the best-matching vector ID. The code is available in Listing 2. The workflow begins by initializing MASS and creating a hashmap of Places, each for one list, and where each Place is responsible for one dimension of the PQ-encoded vectors. The quantized values from the dataset are sorted and stored in these Places to allow efficient lookups during the search phase.

Once the Places are set up, a query vector is processed by dispatching one agent per dimension. Each agent independently queries its assigned place, searching for the closest quantized value to its respective component of the query vector. This results in a set of "votes," where each Agent suggests a potential matching vector ID. These votes are then aggregated to determine the final closest vector.

While the fundamental logic and architecture are in place, the full implementation has not yet been completed due to the complexity of integrating MASS CUDA with the data structures and query workflow. One of the main challenges is to tie together kernel code available in cuVS with MASS CUDA. The current implementation is a primitive one that does not calculate how many threads, blocks, shared memory, L1 cache or global memory is required. It also lacks proper usage of LUTs (Look-up Tables). Additionally, testing and debugging parallel execution in MASS CUDA require careful verification to ensure correctness and performance gains over traditional methods. Despite these challenges, we made calculations of computational complexity of such an approach which demonstrates the feasibility of this approach and suggests that once finalized, it will provide an efficient and scalable solution for PQ-based vector search. The next steps involve refining the agent execution model, optimizing memory allocation, and validating performance improvements through extensive testing.

VI. Future Work

The objective for the upcoming quarter is to finalize agent class PQAgent, proper memory management, test and benchmark implementation, and work on a paper.

References

- [1] “ANN - Benchmarks.” [Online]. Available: <https://github.com/erikbern/ann-benchmarks?tab=readme-ov-file>
- [2] M. Aumüller, E. Bernhardsson, and A. Faithfull, “ANN-Benchmarks: A Benchmarking Tool for Approximate Nearest Neighbor Algorithms,” 2018, *arXiv*. doi: 10.48550/ARXIV.1807.05614.
- [3] M. Iwasaki, “Proximity search in metric spaces using approximate k nearest neighbor graph,”
- [4] *NGT*. Yahoo! Japan. [Online]. Available: <https://github.com/yahoojapan/NGT>
- [5] A. Andoni, P. Indyk, and I. Razenshteyn, “Approximate Nearest Neighbor Search in High Dimensions,” 2018, *arXiv*. doi: 10.48550/ARXIV.1806.09823.
- [6] J. Emau, T. Chuang, and M. Fukuda, “A multi-process library for multi-agent and spatial simulation,” in *Proceedings of 2011 IEEE Pacific Rim Conference on Communications, Computers and Signal Processing*, Victoria, BC, Canada: IEEE, Aug. 2011, pp. 369–375. doi: 10.1109/PACRIM.2011.6032921.
- [7] *MASS*. University of Washington. [Online]. Available: <https://depts.washington.edu/dslab/MASS/>
- [8] L. Kosiachenko, N. Hart, and M. Fukuda, “MASS CUDA: A General GPU Parallelization Framework for Agent-Based Models,” in *Advances in Practical Applications of Survivable Agents and Multi-Agent Systems: The PAAMS Collection*, vol. 11523, Y. Demazeau, E. Matson, J. M. Corchado, and F. De La Prieta, Eds., in Lecture Notes in Computer Science, vol. 11523. , Cham: Springer International Publishing, 2019, pp. 139–152. doi: 10.1007/978-3-030-24209-1_12.
- [9] Yu. A. Malkov and D. A. Yashunin, “Efficient and robust approximate nearest neighbor search using Hierarchical Navigable Small World graphs,” 2016, *arXiv*. doi: 10.48550/ARXIV.1603.09320.
- [10] H. Jégou, M. Douze, and C. Schmid, “Product Quantization for Nearest Neighbor Search,” *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 33, no. 1, pp. 117–128, Jan. 2011, doi: 10.1109/TPAMI.2010.57.
- [11] *SPTAG*. Microsoft. [Online]. Available: <https://github.com/microsoft/SPTAG>
- [12] E. Bernhardsson, *ANNOY*. [Online]. Available: <https://github.com/spotify/annoy>
- [13] *cuVS*. [Online]. Available: <https://github.com/rapidsai/cuvs>
- [14] A. Li and M. Fukuda, “Agent-Based Parallelization of a Multi-Dimensional Semantic Database Model,” presented at the IEEE 24th Int’l Conf. on Information Reuse and Integration for Data Science, Aug. 2023, pp. 64–69.

Appendix

A. Code

Listing 1: Place class

```
1  #include "PQPlace.h"
2  #include "DimensionMatrix.h"
3  // #include <iostream>
4
5  __device__ int PQPlace::findClosest(float queryValue) {
6      DimensionMatrix* dimMatrix = getAttribute<DimensionMatrix>(ATTRIBUTE::DIMENSION_MATRIX, 1);
7      int placeIndex = getIndex();
8
9      const auto& d_values = dimMatrix->matrix[placeIndex];
10     int numElements = dimMatrix->elementsInList;
11
12     int low = 0;
13     int high = numElements - 1;
14
15     while (low <= high) {
16         int mid = (low + high) / 2;
17         // NOTE: Maybe just adjust queryValue
18         float midVal = static_cast<float>(d_values[mid]);
19
20         if (fabsf(midVal - queryValue) < 1e-6f) {
21             return mid; // return index of vector
22         } else if (midVal < queryValue) {
23             low = mid + 1;
24         } else {
25             high = mid - 1;
26         }
27     }
28
29     // Choose the closest between low and low - 1
30     if (low == 0) return 0;
31     if (low >= numElements) return numElements - 1;
32
33     float diffLow = fabsf(static_cast<float>(d_values[low]) - queryValue);
34     float diffHigh = fabsf(static_cast<float>(d_values[low - 1]) - queryValue);
35
36     return (diffLow < diffHigh) ? low : low - 1;
37 }
38
39 __device__ void PQPlace::callMethod(int functionId, void *argument)
40 {
41     switch (functionId)
42     {
43     case FIND_CLOSEST:
44         findClosest((float) argument);
45         break;
46     default:
```

```

29     // Choose the closest between low and low - 1
30     if (low == 0) return 0;
31     if (low >= numElements) return numElements - 1;
32
33     float diffLow = fabsf(static_cast<float>(d_values[low]) - queryValue);
34     float diffHigh = fabsf(static_cast<float>(d_values[low - 1]) - queryValue);
35
36     return (diffLow < diffHigh) ? low : low - 1;
37 }
38
39 __device__ void PQPlace::callMethod(int functionId, void *argument)
40 {
41     switch (functionId)
42     {
43     case FIND_CLOSEST:
44         findClosest((float) argument);
45         break;
46     default:
47         //std::cout << "Default branch in Place::callMethod" << std::endl;
48         break;
49     }
50 }

```

```

1  #pragma once
2
3  #include "Place.h"
4
5  class PQPlace : public mass::Place {
6  public:
7      MASS_FUNCTION PQPlace(int index) : mass::Place(index) {}
8      MASS_FUNCTION ~PQPlace() {}
9      // callMethod to override the default behavior of the Place
10     __device__ virtual void callMethod(int functionId, void *arg = NULL);
11
12     // Attributes
13     enum ATTRIBUTE
14     {
15         DIMENSION_MATRIX
16     };
17
18     // Methods
19     enum FUNCTION
20     {
21         FIND_CLOSEST
22     };
23
24     private:
25         __device__ void findClosest(float queryValue);
26 };

```

Listing 2: Agent class

```

1  #include "PQAgent.h"
2  // #include <iostream>
3
4  __device__ void PQAgent::search() {
5      // TODO: Get Attribute and pick value by agent index
6      // Get the place where the agent resides
7      mass::Place *place = *getAttribute<mass::Place *>(mass::AgentPreDefinedAttr::RESIDE_PLACE, 1);
8      // TODO:
9      int vectorId = place->callMethod(someNumber);
10     // TODO: Aggregation
11 }
12
13 __device__ void PQAgent::callMethod(int functionId, void *argument) {
14     switch(functionId) {
15         case PQAgent::FUNCTION::SEARCH:
16             search();
17             break;
18         default:
19             printf("Default branch in Agent::callMethod");
20             break;
21     }
22 }

```

```

1  #pragma once
2
3  #include "Agent.h"
4
5  class PQAgent : public mass::Agent {
6  public:
7      MASS_FUNCTION PQAgent(int index) : mass::Agent(index) {}
8      MASS_FUNCTION ~PQAgent() {}
9      // callMethod to override the default behavior of the Agent
10     __device__ virtual void callMethod(int functionId, void *arg = NULL);
11
12     // Attributes
13     enum ATTRIBUTE
14     {
15         QUERY_VALUE
16     };
17
18     // Methods
19     enum FUNCTION
20     {
21         SEARCH
22     };
23
24 private:
25     __device__ void search();
26
27 };

```

B. How To Run a Program

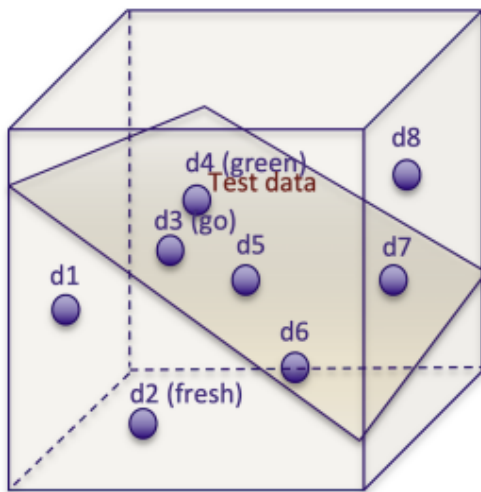
To run the Feature Extraction program, Python 3.9 or higher, along with the corresponding version of PyTorch, is required. For the IVF PQ program, the latest version of cuVS must be

installed, ensuring that all associated prerequisites, including compatibility with the appropriate version of the CUDA Toolkit, are met. Additionally, the libnpy library is necessary to convert the NumPy arrays generated by the Feature Extraction program into a C++-compatible vector data structure.

C. Figures

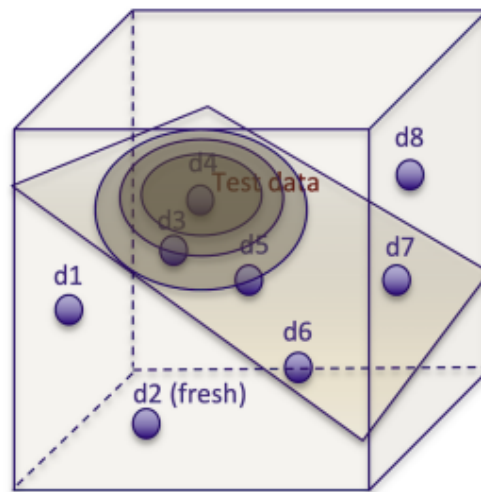
Figure 1:

a) Math-based data retrieval



Sort: d4, d3, d5, d6, d7.

b) Agent-based data retrieval



Use agent propagation and collision.