

Group Marching Tree: Sampling-Based Approximately Optimal Motion Planning on GPUs

Brian Ichter
Aeronautics & Astronautics
Stanford University
Stanford, California 94305
Email: ichter@stanford.edu

Edward Schmerling
Institute for Computational &
Mathematical Engineering
Stanford University
Stanford, California 94305
Email: schmerling@stanford.edu

Marco Pavone
Aeronautics & Astronautics
Stanford University
Stanford, California 94305
Email: pavone@stanford.edu

Abstract—This paper presents a novel approach, named the Group Marching Tree (GMT*) algorithm, to planning on GPUs at rates amenable to application within control loops, allowing planning in real-world settings via repeated computation of near-optimal plans. GMT*, like the Fast Marching Tree (FMT*) algorithm, explores the state space with a “lazy” dynamic programming recursion on a set of samples to grow a tree of near-optimal paths. GMT*, however, alters the approach of FMT* with approximate dynamic programming by expanding, in parallel, the group of all active samples with cost below an increasing threshold, rather than only the minimum cost sample. This group approximation enables low-level parallelism over the sample set and removes the need for sequential data structures, while the “lazy” collision checking limits thread divergence—all contributing to a very efficient GPU implementation. While this approach incurs some suboptimality, we prove that GMT* remains asymptotically optimal up to a constant multiplicative factor. We show solutions for complex planning problems under differential constraints can be found in ~10 ms on a desktop GPU and ~30 ms on an embedded GPU, representing a significant speed up over the state of the art, with only small losses in performance. Finally, we present a scenario demonstrating the efficacy of planning within the control loop (~100 Hz) towards operating in dynamic, uncertain settings.

I. INTRODUCTION

Robotic systems are increasingly operating in real-world settings—away from the structure, repetition, and certainty of the factory floor—that require a robot to not only sense its environment and state in real time, but to react accordingly [1]. Acting in these paradigms often necessitates motion plans be computed on the basis of limited state and environmental knowledge, both of which may vary rapidly as information is gathered and the robot’s surroundings change. A major challenge in this approach is thus replanning quickly, ideally up to the bound of the control feedback loop frequency (~100 Hz), particularly for systems governed by dynamic constraints operating in complex environments.

Sampling-based motion planning has emerged as an especially successful paradigm for rapid planning in complex, high-dimensional, and unstructured environments [2], and it has been shown to extend well to planning with differential constraints [3][4]. These methods probe the state space with

a set of samples to be connected, under the supervision of a collision detection module, to form a traversable graph representation of the free state space. Sampling-based roadmap methods, such as the probabilistic roadmap algorithm (PRM) [5] and its asymptotically optimal variant PRM* [6], initially construct a graph where samples are connected to each of their near neighbors, provided the connection is collision-free. A shortest path search is performed on the resulting roadmap to yield solutions (up to the resolution constraints of the underlying graph) to the optimal planning problem [2]. As these methods are limited in their speed by the initial graph building stage, variants have been developed that simultaneously construct the graph edges while searching, e.g., Lazy PRM [7], which avoid performing any collision checks that are not required during the roadmap shortest path computation. The Fast Marching Tree algorithm (FMT*) further reduces collision checking by implementing direct dynamic programming (as opposed to full shortest path search) while constructing a tree subgraph of a disk graph defined by connection cost, increasing performance particularly in complex, high-dimensional spaces [8]. Yet even with these advances, path plan computation times are often over an order of magnitude greater than the periods of controller loops and with the slowing growth rate of CPU computational power (due primarily to limited clock frequency), it is unlikely raw CPU power will soon bridge this gap. We instead propose algorithm development for a different paradigm: parallel computing, with a particular focus on development for the interplay between algorithmic design and the many thousand core architectures of GPUs. Unfortunately, while we seek a solution inspired by the dynamic programming literature for a single pair of start/goal states (the use case relevant to control loop planning), the inherently sequential nature of dynamic programming’s minimum cost node expansion, which, e.g., FMT* is built on, complicates the necessary massive parallelization.

Statement of Contributions. In this work, we propose the use of approximate dynamic programming (ADP) methods that leverage algorithm parallelism for greater speed while incurring only a bounded degree of suboptimality. We present the Group Marching Tree (GMT*) algorithm that, like the Fast Marching Tree algorithm (FMT*) [8], performs a “lazy” dynamic programming recursion on a set of samples in the

This work was supported by a Qualcomm Innovation Fellowship and by NASA under the Space Technology Research Grants Program, Grant NNX12AQ43G. Brian Ichter was supported by the DoD NDSEG Program.

state space to grow a tree of near-optimal cost-to-arrive paths. GMT*, however, varies from the approach of FMT* with ADP by expanding the tree, in parallel, from the group of all active samples with cost below a threshold, rather than only the minimum cost sample (essentially locally relaxing the principle of optimality). This group approximation enables low-level parallelism over the sample set and removes the need for sequential data structures, allowing for massive parallelization on GPUs. The “lazy” collision checking further facilitates GPU implementation by limiting thread divergence at the lowest levels. While these approximations do introduce some suboptimality, we prove that GMT* remains asymptotically optimal up to a constant multiplicative factor and demonstrate through numerical experiments that the empirical loss is well below the theoretical bound.

We further discuss the implementation of GMT* on GPU architectures and show its application to several illustrative motion planning problems with differential constraints, for which we consider kinodynamic and nonholonomic planning. These numerical experiments show that solution trajectories can be computed in ~ 10 ms on a consumer grade GPU and ~ 30 ms on an embeddable GPU; achieving computation times two orders of magnitude faster than a state-of-the-art CPU algorithm and an order of magnitude faster than a state-of-the-art GPU algorithm, again with only small performance losses. Lastly, we demonstrate the efficacy of planning within the control loop on a simplified quadrotor in a collapsing cave environment, with state disturbances and environmental dynamism.

Related Work. Previous works have addressed planning in real-world settings through a number of methods. One approach is that of feedback motion planning, which traditionally defines a policy over the state space to allow the current state to be fed back into the controller [1]. Unfortunately, the ephemeral nature of the planning environment complicates this process; while ideally these feedback plans would always reflect the current knowledge state, they may become quickly outdated and inaccurate if updating or recomputing plans is too computationally intensive. Some methods exist to simplify this computation, such as [9], which computes a field of guiding vectors over the entire free state space, however they generally must interpolate over the state space to define the local action and assume the state space can be easily represented. The high-frequency replanning approach of [10] uses a similar approach to our own by quickly replanning full trajectories. This work leverages parallelism to generate many rapidly-exploring random tree (RRT) trajectories, selecting the trajectory with the lowest collision probability at each computation step. Our work focuses on construction of a single tree to allow planning within the loop at rates of ~ 100 Hz, rather than the 4 Hz considered in [10]. Other works, e.g., RRT^X [11], have accelerated the replanning approach by iteratively rewiring a single tree as new information becomes available. By reusing the previous tree at each time step, these methods are limited in scenarios where the environment changes drastically. Furthermore, the RRT^X tree is rooted at the goal state to enable use in scenarios with disturbances, but this limits its utility in problems with a changing goal state.

Another approach to planning in changing environments is to couple low-frequency global planners with high-frequency reactive controllers that determine actions which are collision-free and optimal in a local sense. Using methods such as pre-computed trajectory libraries and funnels [12], learning [13], and potential fields [14], this approach has shown practical success in many settings, however, its focus on the local region can ignore variations in global reachability resulting from actions (e.g., not accounting for momentum or nonholonomic constraints) and their use of heuristics to inform actions may incur suboptimality (e.g., in maze-like environments). In this work GMT* is shown to be capable of computing motion plans in a tempo comparable to reactive controllers, lessening the need for consideration of only local actions. GMT* can further be used in concert with reactive controllers to better inform actions via accurate cost-to-go computations, thus potentially benefiting from properties like the robustness in [12].

A main tenet of our work is the use of approximate dynamic programming (ADP) to allow parallelism while exploring the state space. Similar search methodologies, i.e., expanding wavefronts in low-cost groups, have been used successfully for graph search to allow parallelism and complexity reduction. Dial’s algorithm [15] implements Dijkstra’s algorithm on graphs with integer weights by stepping through buckets, exactly solving for the shortest paths. The Δ -stepping algorithm [16] generalizes Dial’s algorithm to solve exactly the single source shortest path problem on non-negative real-value weighted graphs by successively relaxing edges while stepping through buckets with width Δ , however it performs extra work by revisiting edges to maintain exactness. The Group Marching Method [17] builds on the Fast Marching Method (an inspiration for FMT*) to solve the eikonal equations by advancing a group of points together in two iterations, the first forward and the second backwards to correct for instabilities. Our work employs a similar expansion strategy, but abandons any additional computation necessary to maintain exactness, instead leveraging the underlying disk graph to maintain asymptotic optimality within a constant factor.

Parallelization too has been applied successfully to motion planning by a number of researchers, finding significant algorithm accelerations. An early result in sampling-based motion planning showed that probabilistic roadmap methods are embarrassingly parallel [18], which was later extended to implementation on GPUs [19]. The focus of PRM-based approaches on entire graph construction however can be prohibitively slow even with GPUs. Common approaches to parallelization of sampling-based planning include focusing only on algorithm subroutines (such as collision checking and nearest neighbor search) [19][20], adapting serial algorithms via AND/OR-parallelism [10][21], or using load balancing and domain decomposition [22][23]. GMT*’s ADP is parallel at the sample level, meaning many of these methodologies (e.g., collision checking, domain decomposition) are applicable in implementation. Furthermore, this low-level parallelism enables massive parallelization for use on GPUs.

Organization. The remainder of this document is organized as follows. Section II describes the problem setup. Section III discusses the GMT* algorithm and proves its asymptotic

optimality up to a constant multiplicative factor. Section IV discusses its implementation on GPUs. Section V demonstrates the performance of GMT* with motion planning problems under differential constraints. Lastly, Section VI summarizes our findings and proposes directions for future work.

II. PROBLEM SETUP

In this section and the next (which contains the description and analysis of the GMT* algorithm), for ease of exposition we consider the geometric planning problem—the problem, loosely speaking, of computing the shortest free path from an initial state to a goal region where any two states can be connected by a straight line. The problem is briefly overviewed here, but a full, detailed problem formulation can be found in [8]. Let $\mathcal{X} = [0, 1]^d$ be the state space, where $d \in \mathbb{N}, d \geq 2$. Let $\mathcal{X}_{\text{obs}}$ be the obstacle space, $\mathcal{X}_{\text{free}} = \mathcal{X} \setminus \mathcal{X}_{\text{obs}}$ be the free state space, $x_{\text{init}} \in \mathcal{X}_{\text{free}}$ be the initial condition, and $\mathcal{X}_{\text{goal}} \subset \mathcal{X}_{\text{free}}$ be the goal region. A path is said to be *collision-free* if $\sigma(\tau) \in \mathcal{X}_{\text{free}}$ for all $\tau \in [0, 1]$. A path is said to be *feasible* if it is collision-free, $\sigma(0) = x_{\text{init}}$, and $\sigma(1) \in \mathcal{X}_{\text{goal}}$.

Problem 1 (Optimal path planning): Given a path planning problem $(\mathcal{X}_{\text{free}}, x_{\text{init}}, \mathcal{X}_{\text{goal}})$ and a cost measure c , find a feasible path σ^* such that $c(\sigma^*) = \min\{c(\sigma) : \sigma \text{ is feasible}\}$. If no such path exists, report failure.

For Section III we consider the cost measure $c(\sigma)$ as the arc length of σ with respect to the Euclidean metric and write $\|x - y\|$ to denote the cost of the shortest path between $x, y \in \mathcal{X}$. Though this is arguably the simplest formulation of robotic motion planning, we note the analytical distinction between (a) determining, for more general problem setups with differential constraints or alternative costs, whether a sample set in $\mathcal{X}$ admits a near-optimal trajectory as a sequence of local connections, and (b) arguing that a planning algorithm is capable of identifying such a high-quality solution given a sample set. We refer the reader to our previous works [3][4][24] for discussion on the first point, including expressions for local connection radii under both randomized and deterministic state space sampling, and abbreviate the relevant discussion (Theorems 1 and 3) in the present work. The novel theoretical contribution of this paper is establishing that GMT*, which achieves a high degree of parallelism in a single query approach unlike the FMT* and PRM* algorithms analyzed in those works, still recovers asymptotically optimal paths (up to a constant factor) under the same sampling and connection radius assumptions (Theorem 2). Our numerical experiments in Section V consider both kinodynamic and nonholonomic planning problems, specifically double integrator and Dubins airplane dynamics.

III. THE GROUP MARCHING TREE ALGORITHM

A. GMT*

We now detail the Group Marching Tree algorithm (GMT*) to be used to approximately solve the optimal path planning problem. GMT* performs a “lazy”, approximate dynamic programming recursion to grow a tree of paths in cost-to-arrive space. This amounts to iteratively attempting to expand all samples in the tree branches below a constantly increasing cost

threshold, rather than expanding only the minimum cost sample. The resulting algorithm enables simultaneous graph building and exploration of the state space, in a manner amenable to complex planning problems, such as high-dimensional, cluttered environments and differential constraints.

A description of the algorithm is given in Alg. 1, with a single iteration visualized in Fig. 1. The algorithm takes as input the planning problem $(\mathcal{X}_{\text{free}}, x_{\text{init}}, \mathcal{X}_{\text{goal}})$, a sample set $V_{\text{unexplored}}$ of n samples in $\mathcal{X}_{\text{free}}$ (at least one in $\mathcal{X}_{\text{goal}}$), a connection radius r , and a group cost threshold factor $\lambda \in (0, 1]$. Together, λ and r define the group cost threshold increment δ , equal to λr . We refer to nodes (or interchangeably, samples) as neighbors if the connection cost between them is less than the connection radius r , as defined in Theorem 1; $r = 4(1 + \eta)^{1/d} (1/d)^{1/d} (\mu(\mathcal{X}_{\text{free}})/\zeta_d)^{1/d} (\log n/n)^{1/d}$ where $\eta \geq 0$ is a tuning parameter, $\mu(\mathcal{X}_{\text{free}})$ denotes the d -dimensional Lebesgue measure of $\mathcal{X}_{\text{free}}$, and ζ_d denotes the volume of the unit ball in d -dimensional Euclidean space. This r applies for geometric planning with Euclidean cost and $V_{\text{unexplored}}$ sampled uniformly randomly from $\mathcal{X}_{\text{free}}$; smaller r may be considered if $V_{\text{unexplored}}$ is sampled with low-dispersion, deterministic sequences [24].

Algorithm 1 Group Marching Tree Algorithm

Require: connection radius r , group cost threshold factor λ ,
 set $V_{\text{unexplored}}$ of n samples in $\mathcal{X}_{\text{free}}$, at least one in $\mathcal{X}_{\text{goal}}$

- 1: Place x_{init} in V_{open} , set $i = 0$ and $\delta = \lambda r$
- 2: Initialize tree with root node x_{init}
- 3: Find nodes $\mathcal{G}$ in V_{open} with cost $\leq i\delta$
- 4: For each unexplored neighbor, x , of any node in $\mathcal{G}$:
 - 5: Find neighbor nodes y in V_{open}
 - 6: Find locally-optimal connection to x from a node in y
 - 7: If that connection is collision-free:
 - 8: Add edge to tree
 - 9: Remove x from $V_{\text{unexplored}}$ and add to V_{open}
- 10: Remove $\mathcal{G}$ from V_{open} and add to V_{closed}
- 11: Increment i
- 12: Skip to Line 3 until either:
 - a: V_{open} is empty $\Rightarrow$ return failure
 - b: A node in $\mathcal{G}$ is in $\mathcal{X}_{\text{goal}} \Rightarrow$ return min cost path to $\mathcal{X}_{\text{goal}}$

The algorithm proceeds by expanding a tree of paths outward through the state space, maintaining the samples in three sets: $V_{\text{unexplored}}$, V_{open} , and V_{closed} . $V_{\text{unexplored}}$ consists of samples not yet added to the tree. V_{open} consists of samples added to the tree and still considered for expansion; intuitively these are the samples on the tree’s outer “branches”, i.e., close to the wavefront. V_{closed} consists of samples added to the tree and no longer considered for expansion; intuitively these are the samples too far from the edge of the expanding tree to make any new connections. The algorithm begins by adding only x_{init} to V_{open} , setting V_{closed} empty, and initializing the tree of paths with x_{init} at its root (Lines 1-2). At each iteration i , all nodes in V_{open} with cost below $i\lambda r$ ($i\delta$) are placed into a set $\mathcal{G}$, denoting the group of samples to be expanded in parallel (Line 3). The amount of parallelism of this step is controlled by a tuning parameter λ , referred to as the group cost threshold factor, which represents a trade-off between parallelism and potential unconsidered optimal connections; $\lambda \rightarrow 1$ represents expanding all nodes in the open set at once, resulting in nearly a breadth first search, while

$\lambda \rightarrow 0$ represents expanding only the minimum cost nodes in a given iteration, resulting in the same final solution as FMT*. The other side of the tradeoff, the unconsidered optimal connections, arises when multiple nodes along the optimal path are considered in the same group expansion, meaning they cannot connect to each other. They may also occur when paths formed from the concatenation of several short connections fall behind the wavefront. These effects, however, are curtailed by the λ factor, the underlying disk graph's structure, and the notion that longer steps will generally be more favorable than many short steps. Even with this suboptimality, we show in Theorem 3 that asymptotic optimality within a constant multiplicative factor is maintained; furthermore, the performance loss observed in practice is studied in Section III-C and shown to be small. With the group of samples $\mathcal{G}$ in hand, all neighbors of $\mathcal{G}$ in $V_{\text{unexplored}}$ are considered for addition to the tree (Line 4). Referring to each sample in $\mathcal{G}$ as x and its neighbors in V_{open} as y , GMT* then selects the locally-optimal connection, where locally-optimal is defined as the connection with the lowest cost for the previously computed path to y concatenated with the straight line path from y to x (Line 6). Note that while this step potentially introduces suboptimal connections by *lazily* ignoring the presence of obstacles, as in FMT*, these connections become vanishingly rare as the number of samples goes to infinity, as discussed and proven in [8]. If the selected connection is collision-free, it is added to the tree and x is removed from $V_{\text{unexplored}}$ and added to V_{open} (Lines 7-9). When all samples have been considered, the samples in $\mathcal{G}$ are removed from V_{open} and added to V_{closed} , and GMT* moves to the next iteration (Lines 10-11), where an entire iteration is considered Lines 3-11. The algorithm terminates when either V_{open} is empty or a node in $\mathcal{G}$ is in $\mathcal{X}_{\text{goal}}$ (Line 12).

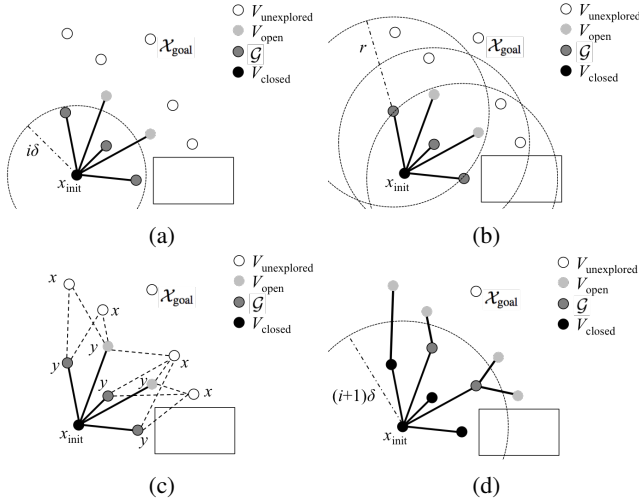


Fig. 1: One iteration step of GMT* expansion, labeled with (Fig., Line in Alg. 1). (1a, Line 3) shows the selection of the group $\mathcal{G}$. (1b, Line 4) shows the selection of the nearest neighbors of $\mathcal{G}$ in $V_{\text{unexplored}}$. (1c, Line 5) shows the candidate connections from which the locally optimal is chosen to connect the new samples. (1d, Line 3) shows the new tree after an iteration, Lines 3-11, and the new group $\mathcal{G}$.

B. GMT* Approximate Asymptotic Optimality

Our analysis of GMT* begins with the concept of *probabilistic exhaustivity* as applied in related work establishing asymptotic optimality for a range of geometric [25] and differentially constrained [3][4] batch-processing, sampling-based motion planning algorithms. Briefly, probabilistic exhaustivity is the notion that within a sufficiently large set of uniformly sampled states, a sequence of samples approximating *any* path arbitrarily well may be found. This property may be used to construct sample sequences approaching the optimal solution that are amenable for recovery by a planning algorithm.

In the subsequent analysis, we define $\text{SampleFree}(n)$ to be a function that returns n points sampled independently and identically from the uniform distribution on $\mathcal{X}_{\text{free}}$, at least one of which is in $\mathcal{X}_{\text{goal}}$. We define a path $\sigma : [0, 1] \rightarrow \mathcal{X}$ and a path $y : [0, 1] \rightarrow \mathcal{X}$ that sequentially connects the sequence of waypoints $\{y_m\}_{m=0}^M \in \mathcal{X}$ with line segments. We say the sequence of waypoints $\{y_m\}(\epsilon, r)$ -traces the path σ if the following conditions hold: (i) $\|y_m - y_{m+1}\| \leq r$ for all m , (ii) the cost of y is bounded as $c(y) \leq (1 + \epsilon)c(\sigma)$, and (iii) the distance from any point y to σ is no more than r . We formally state this in Theorem 1 (proven as Theorem IV.5 in [3]).

Theorem 1 (Probabilistic Exhaustivity): Define a planning problem $(\mathcal{X}_{\text{free}}, x_{\text{init}}, \mathcal{X}_{\text{goal}})$, a feasible path $\sigma : [0, 1] \rightarrow \mathcal{X}_{\text{free}}$, a set of samples $V = \{x_{\text{init}}\} \cup \text{SampleFree}(n)$, and $\epsilon > 0$. For a fixed n consider the event $\mathcal{A}_n$ that there exists $\{y_m\}_{m=0}^M \in V$, $y_0 = x_{\text{init}}$, $y_M \in \mathcal{X}_{\text{goal}}$, which (ϵ, r) -trace σ , where

$$r = 4(1 + \eta)^{\frac{1}{d}} \left(\frac{1}{d} \right)^{\frac{1}{d}} \left(\frac{\mu(\mathcal{X}_{\text{free}})}{\zeta_d} \right)^{\frac{1}{d}} \left(\frac{\log n}{n} \right)^{\frac{1}{d}}$$

with $\eta \geq 0$. Then, as $n \rightarrow \infty$, the probability that $\mathcal{A}_n$ does not occur (denoted by its complement $\mathcal{A}_n^c$) is asymptotically bounded as $P[\mathcal{A}_n^c] = O(n^{-\frac{\eta}{d}} \log^{-\frac{1}{d}} n)$.

We now show that the cost of the path returned by GMT* is bound within a constant multiplicative factor of the cost of any such sequence of tracing waypoints.

Theorem 2 (Bounded Suboptimality): Let $r > 0$ and suppose that the sequence of waypoints $\{y_m\}_{m=0}^M \subset \mathcal{X}_{\text{free}}$ satisfies $y_0 = x_{\text{init}}$, $y_M \in \mathcal{X}_{\text{goal}}$, $\|y_m - y_{m-1}\| \leq r$ for all $m \in \{1, \dots, M\}$ and $B(y_m, r) \subset \mathcal{X}_{\text{free}}$ for all $m \in \{0, \dots, M\}$. Let c_{GMT^*} denote the cost of the path returned by GMT* using a connection radius r and group cost threshold factor λ . Then

$$c_{\text{GMT}^*} \leq (1 + 2\lambda) \sum_{k=1}^M \|y_k - y_{k-1}\|.$$

Proof: Note that in the subsequent analysis we refer to the open set of nodes at iteration i as $V_{\text{open}}(i)$. We first assume, without loss of generality, that $\|y_m - y_{m-2}\| > r$ for all $m \in \{2, \dots, M\}$. This condition may be enforced by making a forward pass over the sequence and omitting the preceding point y_{m-1} for any y_m that violates the condition, without affecting the other lemma assumptions and only decreasing $\sum_{k=1}^M \|y_k - y_{k-1}\|$. Consider running GMT* to completion and for each y_m let $c(y_m)$ denote the cost-to-arrive of y_m in the generated tree. If y_m is not contained in any edge of the tree, we set $c(y_m) = \infty$. Let i_m denote the first iteration after

which y_m has been added to the GMT* tree, that is,

$$i_m = \min\{i \in \mathbb{N} \mid y_m \in V_{\text{open}}(i)\}.$$

Define $S_0 = 0$, $S_m = \sum_{k=1}^m \|y_k - y_{k-1}\|$, the cost of the path connecting $y_0, y_1, \dots, y_m$, and denote $\delta = \lambda r$. We show by induction that for all $m \in \{1, \dots, M\}$, one of the following two possibilities must hold:

$$c_{\text{GMT}^*} \leq S_m + m\delta, \quad (1)$$

or

$$c(y_m) \leq S_m + (m-1)\delta \quad \text{and} \quad i_m \leq \left\lceil \frac{S_{m-1}}{\delta} \right\rceil + m. \quad (2)$$

The base case $m = 1$ is trivial, since $\mathcal{G}_0$ contains only $x_{\text{init}} = y_0$, and thus the first GMT* iteration makes every collision-free connection between x_{init} and the nodes contained in $B(x_{\text{init}}, r)$, including y_1 . Then $c(y_1) = \|y_1 - y_0\| = S_1$ and $i_1 = 1 = \lceil S_0/\delta \rceil + 1$. Now suppose (1) or (2) holds for $m-1$; that means that one of the following four statements must hold.

1. $c_{\text{GMT}^*} \leq S_{m-1} + (m-1)\delta$,
2. $c(y_{m-1}) \leq S_{m-1} + (m-2)\delta$ and $i_{m-1} \leq \left\lceil \frac{S_{m-2}}{\delta} \right\rceil + m-1$ and
 - a. GMT* ends before considering y_m for connection ($i_m = \infty$), or
 - b. $y_{m-1} \in V_{\text{open}}$ when y_m is first considered for connection ($i_{m-1} \leq i_m - 1$), or
 - c. $y_{m-1} \notin V_{\text{open}}$ when y_m is first considered for connection ($i_{m-1} \geq i_m$).

We show that in each of these cases, (1) or (2) holds for m .

Case 1: Since $\|y_m - y_{m-1}\| \geq 0$, we have

$$c_{\text{GMT}^*} \leq S_{m-1} + (m-1)\delta \leq S_m + m\delta.$$

Case 2a: The fact that y_m goes unconsidered means that up until the time the algorithm terminates at iteration i_{term} (upon finding a goal point in group $\mathcal{G}_{i_{\text{term}}}$), y_{m-1} has never been a member of an expansion group:

$$\begin{aligned} i_{\text{term}} &\leq \max\{i_{m-1}, \lceil c(y_{m-1})/\delta \rceil\} \\ &\leq \max\{\lceil S_{m-2}/\delta \rceil + m-1, \lceil c(y_{m-1})/\delta \rceil\}. \end{aligned}$$

Then since the algorithm terminates at iteration i_{term} ,

$$c_{\text{GMT}^*} \leq i_{\text{term}}\delta \leq \max\{S_{m-2} + \delta + (m-1)\delta, c(y_{m-1}) + \delta\} \leq S_{m-1} + m\delta.$$

Case 2b: y_m must be connected to some parent when it is first considered and y_{m-1} is a candidate, so

$$c(y_m) \leq c(y_{m-1}) + \|y_{m-1} - y_m\| \leq S_m + (m-2)\delta.$$

Depending on whether y_{m-1} spends one or more steps in V_{open} , either $i_m = i_{m-1} + 1$ or $i_m = \lceil c(y_{m-1})/\delta \rceil + 1$. In the first subcase, $i_m \leq \left\lceil \frac{S_{m-2}}{\delta} \right\rceil + m$; in the second subcase, $i_m \leq \lceil S_{m-1}/\delta \rceil + m-1$.

Case 2c: When y_m is first considered for connection during iteration $i' < i_m \leq i_{m-1}$, there must be some $z \in \mathcal{G}_{i'}$ such that $\|y_m - z\| \leq r$. Then

$$\begin{aligned} c(y_m) &\leq c(z) + r \leq i'\delta + r \leq (i_{m-1} - 1)\delta + r \\ &\leq S_{m-2} + (m-1)\delta + r. \end{aligned}$$

Recalling that the y_m , by construction, are spaced so that they satisfy the property $r < \|y_m - y_{m-2}\| \leq \|y_m - y_{m-1}\| + \|y_{m-1} - y_{m-2}\|$, we have $c(y_m) \leq S_m + (m-1)\delta$. Additionally, $i_m \leq i_{m-1} \leq \left\lceil \frac{S_{m-2}}{\delta} \right\rceil + m-1$.

Thus the inductive step holds in all cases and (1) or (2) holds for all m . In particular taking $m = M$ we have $c_{\text{GMT}^*} \leq S_M + M\delta$ or $c_{\text{GMT}^*} \leq c(y_M) \leq S_M + (M-1)\delta$. Noting that $S_M \geq Mr/2 = M\delta/(2\lambda)$, we have

$$c_{\text{GMT}^*} \leq (1 + 2\lambda) \sum_{k=1}^M \|y_k - y_{k-1}\|$$

as desired. ■

Given these results, we are now in a position to prove asymptotic optimality within a suboptimality bound.

Theorem 3 (GMT* Approximate Asymptotic Optimality): Assume a δ_{obs} -robustly feasible¹ path planning problem, as defined in [3], with optimal path σ^* and cost c^* . Let c_n denote the cost returned by GMT* with n samples. Then for any $\epsilon > 0$,

$$\lim_{n \rightarrow \infty} \mathbb{P}[c_n > (1 + \epsilon)(1 + 2\lambda)c^*] = 0.$$

Proof: The proof of this theorem is conceptually identical to Theorem VI.2 in [3]; with high probability as $n \rightarrow \infty$ there exists a sequence of waypoints (Theorem 1) tracing an obstacle-clear, near-optimal path of cost $\leq (1 + \epsilon)c^*$ which GMT* recovers up to a factor $(1 + 2\lambda)$ (Theorem 2). ■

We remark that extending the results of this section to differentially-constrained system dynamics or deterministic sampling methods does not substantially change the argument in Theorem 3 (although in particular, with deterministic sampling the convergence in probability may be replaced with a standard limit). All that is required is an analogue of Theorem 1, a statement on the regularity of the dynamics that implies, with sufficient sample density, that the near-neighbor disk graph contains near-optimal paths. The only GMT*-specific analysis in Theorem 2 is a graph search bound— independent of dynamics or sampling methodology.

C. Numerical Experiments: Suboptimality Introduced

To complement the above theoretical bounds, in this subsection we examine the suboptimality incurred in practice through numerical experiments. While we will later describe implementation details and timing results in a few representative problems, this section's focus is solely on the amount of suboptimality resulting from the group expansion. Our figure

¹Briefly, we require that σ^* is a limit of paths with bounded clearance from $\mathcal{X}_{\text{obs}}$; this may be regarded as a minimum regularity assumption to guard against problems with passages of infinitesimal width that are not amenable to sampling-based motion planning.

of merit is thus only the percentage cost increase compared to FMT*, i.e., from the group expansion.

Table I lists results for two geometric planning problems over a variety of dimensions (2D to 10D), the first of which is shown in Fig. 2. This obstacle set was mapped to dimensions greater than two by expanding obstacles to fully fill the space. Fig. 2b, in particular, shows the wake-like structure of the parallel expansion. The second planning problem listed in Table I is a maze environment requiring exploration in all dimensions. For each planning problem, the same setup is run with varying λ (values of 0.2, 0.5, and 1.0) and with sufficiently high sample counts to nearly converge to the optimal. In each case, the suboptimality is significantly below the proven bound, and often below 5%. We observe the expected increase in cost error with increasing λ , and an additional increase with increasing dimension.

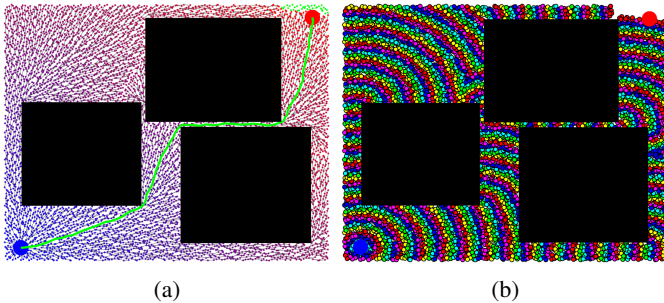


Fig. 2: Expansion of the GMT* algorithm, where Fig. 2a shows the resulting tree (colored by cost) and Fig. 2b shows individual groups (denoted by color) expanded in parallel.

Obstacle	d	Cost Error ($c_{\text{GMT}^*}/c_{\text{FMT}^*} - 1$)		
		$\lambda = 0.2$	$\lambda = 0.5$	$\lambda = 1.0$
Rectangles (Fig. 2a)	2D	0.2%	0.6%	1.8%
	3D	0.1%	0.6%	3.4%
	6D	0.4%	1.5%	2.1%
	10D	2.0%	14.8%	17.0%
Maze	3D	0.3%	1.5%	4.7%
	5D	0.9%	4.9%	7.8%

TABLE I: Suboptimality introduced by GMT* over FMT* for a range of dimensions d and group cost threshold factors λ . Results are averaged over 50 runs at $n = 5000$ samples. The variance was found to be small.

IV. GPU IMPLEMENTATION

We begin this section with a brief discussion of GPU architectures, as the ability of GMT* to exploit the computational capabilities of many-core GPUs is fundamental to this work and has driven much of the algorithm design and implementation. We particularly focus here on the CUDA enabled GPUs used in this work. CUDA C functions running on GPUs are organized in a three level thread-hierarchy. At the lowest level, threads run in groups of 32 that execute one common instruction at a time, i.e., any divergence will cause branches to execute serially. A level above this, thread groups are combined into blocks, each of which can utilize a small, low-latency shared memory and executes concurrently on the same multiprocessor, but is allowed to diverge without causing serial execution. Finally, at the highest level, blocks

are formed in grids to be dispatched to the device. A more detailed discussion can be found in [26].

We highlight here three properties of GMT* that, along with the sample-level parallelism, allow efficient application to GPU architectures. First, the use of lazy collision checking limits thread divergence at low levels by only attempting to connect new samples to the tree once per iteration. Second, the design of GMT* is such that the sample set is partitioned into $V_{\text{unexplored}}$, V_{open} , and V_{closed} , with a sample always a member of one and only one set, allowing for little overlap of memory access and easy memory representation as Boolean masks. Our work accesses these sets with thread identifiers assigned via prefix sums, a strategy described in [27]. The use of this algorithmic primitive allows fast reorganization of sparse and uneven workloads into dense uniform ones. Third, as the set of samples considered for expansion, $\mathcal{G}$, can be represented as a set of cost-thresholded buckets, there is no need for the use of serial data structures, e.g., min-heaps.

V. NUMERICAL EXPERIMENTS

A. Numerical Experiment Setup

As our goal is to show planning in changing, uncertain settings with dynamic systems, in this section we apply GMT* to the problem of planning under differential constraints with a 6D double integrator ($\ddot{x} = u$) and a Dubins airplane (Dubins car with altitude [28]). The double integrator planning problems consider a mixed time/quadratic control effort cost function, while the Dubins airplane problems consider an Euclidean distance cost function. The results were written in CUDA C and generated on an NVIDIA GeForce GTX 980 GPU on a Unix system with a 3.0 GHz CPU. We additionally provide comparison with an embeddable GPU, the NVIDIA Jetson TX1, to show many of these performance gains are similarly available for onboard computation. Our implementation of GMT* samples the state space using the deterministic, low-dispersion Halton sequence, to achieve best performance, following the discussion in [24]. This sampling, along with computation of nearest neighbor connections (edge discretization and neighbor sets), was performed offline in a precomputation phase.

B. Planning Under Differential Constraints

Through several motion planning problems, detailed in Fig. 3 and Table II, we demonstrate GMT* achieves one to two orders of magnitude speed up with relatively small performance losses compared to an implementation of FMT* on a CPU [8] and an implementation of PRM* on a GPU [6]. For each simulation, we pick a value for the connection radius r appropriate for the dynamics [3][4] and set λ to 1, which we have found allows simple implementation, maximum parallelization, and performance losses on the order of 10%. The obstacles in our simulation are represented by unions of axis-aligned bounding boxes, as commonly used for a broad phase collision checking phase [2]. This methodology can provide increasingly accurate representations of obstacle sets as more are used (e.g., as in Fig. 3b or with octree-based representations as in Fig. 3c).

The first problem (Fig. 3b) was built from point cloud data collected in [29] for an indoor office environment, with

individual environmental elements bounded by boxes. The second planning problem (Fig. 3c) represents a cave system consisting of two maze-like levels connected by three passageways. Finally, our third planning problem (Fig. 3d) represents a forest environment. To show the results extend to systems with nonlinear dynamics (and nonholonomic planning), this last simulation uses Dubins airplane dynamics rather than double integrator dynamics. Our Dubins airplane consists of a planar Dubins car augmented with a single integrator in the third dimension, with bounded control on turning rate, unbounded altitude control, and a Euclidean distance cost function [28].

As shown in Table II, in all problems a solution trajectory is found by GMT* in ~ 10 ms, a speed increase of two orders of magnitude over CPU FMT* and one order of magnitude over GPU PRM*. The algorithm also performs well on the embedded platform, only slowing by a factor of two, compared to PRM*, which slows down by approximately a factor of five. This demonstrates GMT*'s lightweight approach of building a single tree is particularly amenable to onboard computation. We also note that the cost increase incurred is less than 12% for all cases despite the high group cost threshold factor.

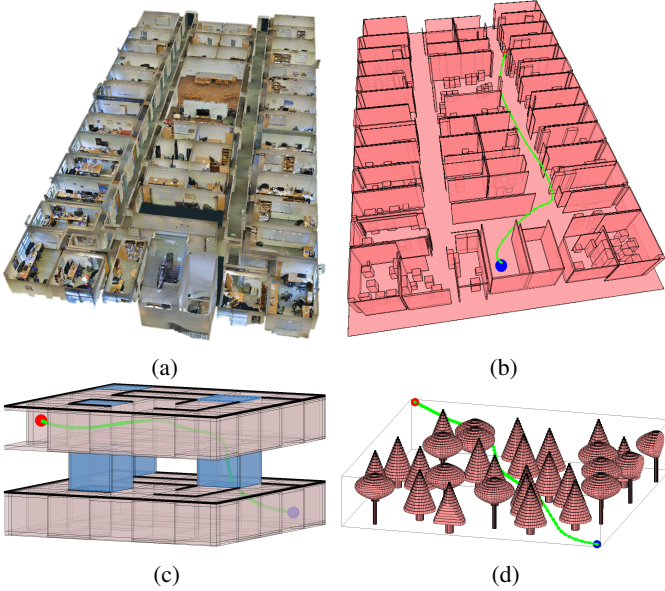


Fig. 3: The solution trajectory connecting x_{init} (blue) to x_{goal} (red) returned by GMT* is shown in green for all figures. (3a-3b) The indoor environment was generated with point cloud data from [29] with individual elements bounded by boxes. (3c) The cave environment consists of two maze-like levels, with wall outlines shown in black, connected by three passageways, shown in blue. (3d) The forest environment consists of many varying tree obstacles.

Table III further demonstrates the algorithm's scaling with sample and obstacle counts. The small increases in computation time with increasing sample count are a result of the GPU not being fully utilized at every iteration with lower sample counts, i.e., the group size may be too small to use every GPU core. The obstacle scaling too shows only slight increases in computation time with increased obstacle resolution, approximately doubling for each order of magnitude

increase, however, if obstacles and complexity of the space becomes a significant bottleneck, GMT* is amenable to space partitioning structures, e.g., k-d trees, or parallelization at the obstacle level [20].

Double Integrator, Fig. 3b				
Algorithm	Device	$c_{\text{alg}}/c_{\text{GMT}^*}$	Time (ms)	$t_{\text{alg}}/t_{\text{GMT}^*}$
FMT*	CPU	0.91	1291	129.1
PRM*	Embd. GPU	0.88	735	73.5
PRM*	GPU	0.88	158	15.8
GMT*	Embd. GPU	1	27	2.7
GMT*	GPU	1	10	1

Double Integrator, Fig. 3c				
Algorithm	Device	$c_{\text{alg}}/c_{\text{GMT}^*}$	Time (ms)	$t_{\text{alg}}/t_{\text{GMT}^*}$
FMT*	CPU	0.91	1490	99.3
PRM*	Embd. GPU	0.90	517	34.5
PRM*	GPU	0.90	140	9.3
GMT*	Embd. GPU	1	31	2.0
GMT*	GPU	1	15	1.0

Dubins Airplane, Fig. 3d				
Algorithm	Device	$c_{\text{alg}}/c_{\text{GMT}^*}$	Time (ms)	$t_{\text{alg}}/t_{\text{GMT}^*}$
FMT*	CPU	0.95	1312	218.7
PRM*	Embd. GPU	0.95	945	157.5
PRM*	GPU	0.95	96	16.0
GMT*	Embd. GPU	1	11	1.8
GMT*	GPU	1	6	1

TABLE II: Results for algorithms run with 5000 samples in the environments in Fig. 3. GPU refers to the GTX 980, while Embd. GPU refers to an embeddable Jetson TX1 GPU.

Fig.	Sample Count (n)	Obstacle Count	Time (ms)	Cost
3c	1k	300	6	552
	2k	300	8	379
	5k	300	15	290
	10k	300	21	233
3d	5k	150	7	-
	5k	500	14	-
	5k	1500	18	-
	5k	5000	26	-

TABLE III: Results for GMT* scaling with sample and obstacle counts, in terms of cost and computation time. Note the obstacle count represents varying the resolution of the obstacle representations, not varying the obstacle location or class.

C. Planning in the Loop

The numerical experiments in Section V-B have shown that it is possible to plan at rates amenable to implementation within control loops by allowing some performance loss in exchange for parallelism. We now demonstrate that this strategy is beneficial through numerical experiments for a system operating in a dynamic environment with random state disturbances. The setup mimics a collapsing cave system (Fig. 3c), which a quadrotor modeled as a double integrator must escape. A successful escape requires high performance actions to minimize time spent in the degrading cave as well as actions that account for state disturbances and variations in the environment. The collapse is modeled as randomly placed box obstacles added to the environment, with the rates representing the number of obstacles added each second. Fig. 4 shows the success rate over 50 runs of a quadrotor using a waypoint tracking controller, which tracks trajectories generated with FMT* and GMT* (replanning as quickly as possible). The results for FMT* show that, as expected, success rate decreases quickly with increased noise and environmental degradation.

Replanning with GMT*, however, shows little variation in failure rate with increased noise level. The results further show significantly higher success rates for replanning with GMT* than replanning with FMT*. Note that these planning problems may not be possible to solve for every instance, as the collapses can happen anywhere within the environment (and very quickly), potentially trapping the quadrotor. Experiment videos are available at <https://goo.gl/67RSsp>.

Environment Degradation Rate	FMT* Replan				GMT* Replan			
	0	1	2	3	0	1	2	3
0	1	0.52	0.38	0.34	1	1	1	1
1	0.8	0.54	0.3	0.22	0.94	0.98	0.96	0.98
2	0.82	0.36	0.2	0.28	0.9	0.94	0.94	0.92
4	0.78	0.24	0.18	0.14	0.9	0.88	0.76	0.88
8	0.42	0.24	0.04	0.14	0.7	0.78	0.7	0.66
16	0.14	0.06	0.06	0	0.5	0.58	0.6	0.4
	0	1	2	3	0	1	2	3
	Disturbance Levels				Disturbance Levels			

Fig. 4: Success rate of a quadrotor replanning with FMT* or GMT* in the Fig. 3c environment with varying levels of state disturbances and cave collapse rates (simulated with obstacles randomly placed in the environment).

VI. CONCLUSION

We have introduced and analyzed a novel planning algorithm, the Group Marching Tree algorithm (GMT*), that trades off parallelism for optimality in order to leverage GPU hardware. The computational speed of GMT* allows us to approach the problem of planning in real-world settings—particularly focusing on the uncertain, dynamic environments that naturally arise from active robot sensing and the uncertain, disturbed motion of systems in the field—by replanning at rates commensurate with the control loop frequency. Simulation results show planning times on the order of 10 ms (for a 6D double integrator and Dubins airplane) and demonstrate the efficacy of planning at these rates in difficult environments.

This paper leaves several important research avenues open. Foremost, we plan to validate this approach experimentally on a platform with state and environmental sensing. We further plan to provide a more detailed theoretical analysis of GMT*, such as providing time and space complexity analysis and potentially proving tighter suboptimality bounds. We additionally plan to explore extensions to other planning paradigms, which the computational speed of GMT* may enable. To merge planning and game playing, we plan to use GMT* as a default simulation policy when many actions must be considered, such as in algorithms like Monte Carlo tree search. We also plan to show that GMT* may be used to construct policies in decision making frameworks through fast approximation of the cost-to-go. Finally, we plan to demonstrate extensions to planning with a probabilistic state belief by utilizing a backwards search in cost-to-go space. In this way we can define actions over regions of the state space, with the same computation times shown above, and select actions from criteria such as best worst-case or maximum expected performance.

REFERENCES

- [1] S. M. LaValle, “Motion planning: Wild frontiers,” *IEEE Robotics and Automation Magazine*, 2011.
- [2] S. Lavalle, *Planning Algorithms*. Cambridge University Press, 2006.
- [3] E. Schmerling, L. Janson, and M. Pavone, “Optimal Sampling-Based Motion Planning under Differential Constraints: the Driftless Case,” in *Proc. IEEE Conf. on Robotics and Automation*, 2015.
- [4] —, “Optimal Sampling-Based Motion Planning under Differential Constraints: the Drift Case with Linear Affine Dynamics,” in *Proc. IEEE Conf. on Decision and Control*, 2015.
- [5] L. E. Kavraki, P. Švestka, J. C. Latombe, and M. H. Overmars, “Probabilistic Roadmaps for Path Planning in High-Dimensional Spaces,” *IEEE Transactions on Robotics and Automation*, 1996.
- [6] S. Karaman and E. Frazzoli, “Sampling-based Algorithms for Optimal Motion Planning,” *International Journal of Robotics Research*, 2011.
- [7] R. Bohlin and L. Kavraki, “Path Planning Using Lazy PRM,” in *Proc. IEEE Conf. on Robotics and Automation*, 2000.
- [8] L. Janson, E. Schmerling, A. Clark, and M. Pavone, “Fast Marching Tree: A Fast Marching Sampling-Based Method for Optimal Motion Planning in Many Dimensions,” *International Journal of Robotics Research*, 2015.
- [9] S. R. Lindemann and S. M. LaValle, “Simple and efficient algorithms for computing smooth, collision-free feedback laws over given cell decompositions,” *International Journal of Robotics Research*, 2009.
- [10] W. Sun, S. Patil, and R. Alterovitz, “High-Frequency Replanning Under Uncertainty Using Parallel Sampling-Based Motion Planning,” *IEEE Transactions on Robotics*, 2015.
- [11] M. Otte and E. Frazzoli, “RRT^X: Asymptotically optimal single-query sampling-based motion planning with quick replanning,” *International Journal of Robotics Research*, 2015.
- [12] A. Majumdar and R. Tedrake, “Robust online motion planning with regions of finite time invariance,” in *Workshop on Algorithmic Foundations of Robotics*, 2012.
- [13] C. Richter, W. Vega-Brown, and N. Roy, “Bayesian learning for safe high-speed navigation in unknown environments,” in *International Symposium on Robotics Research*, 2015.
- [14] S. Scherer, S. Singh, L. Chamberlain, and M. Elgersma, “Flying fast and low among obstacles: Methodology and experiments,” *International Journal of Robotics Research*, 2008.
- [15] R. B. Dial, “Algorithm 360: Shortest-path forest with topological ordering [H],” *Communications of the ACM*, 1969.
- [16] U. Meyer and P. Sanders, “ Δ -stepping: a parallelizable shortest path algorithm,” *Journal of Algorithms*, 2003.
- [17] S. Kim, “An $\mathcal{O}(n)$ level set method for eikonal equations,” *SIAM Journal on Scientific Computing*, 2001.
- [18] N. M. Amato and L. K. Dale, “Probabilistic Roadmap Methods Are Embarrassingly Parallel,” in *Proc. IEEE Conf. on Robotics and Automation*, 1999.
- [19] J. Pan, C. Lauterbach, and D. Manocha, “g-Planner: Real-time Motion Planning and Global Navigation using GPUs,” in *AAAI Conf. on Artificial Intelligence*, 2010.
- [20] J. Bialkowski, S. Karaman, and E. Frazzoli, “Massively Parallelizing the RRT and the RRT*,” in *IEEE/RSJ Int. Conf. on Intelligent Robots & Systems*, 2011.
- [21] D. Devaurs, T. Siméon, and J. Cortés, “Parallelizing RRT on Large-scale Distributed-Memory Architectures,” *IEEE Transactions on Robotics*, 2013.
- [22] C. Rodriguez, J. Denny, S. A. Jacobs, S. Thomas, and N. M. Amato, “Blind RRT: A Probabilistically Complete Distributed RRT,” in *IEEE/RSJ Int. Conf. on Intelligent Robots & Systems*, 2013.
- [23] A. Fidel, S. A. Jacobs, S. Sharma, N. M. Amato, and L. Rauchwerger, “Using Load Balancing to Scalably Parallelize Sampling-Based Motion Planning Algorithms,” in *IEEE Parallel and Distributed Processing Symposium*, 2014.
- [24] L. Janson, B. Ichter, and M. Pavone, “Deterministic Sampling-Based Motion Planning: Optimality, Complexity, and Performance,” in *International Symposium on Robotics Research*, 2015.
- [25] J. A. Starek, J. V. Gomez, E. Schmerling, L. Janson, L. Moreno, and M. Pavone, “An Asymptotically-Optimal Sampling-Based Algorithm for Bi-directional Motion Planning,” in *IEEE/RSJ Int. Conf. on Intelligent Robots & Systems*, 2015.
- [26] NVIDIA, *CUDA C Programming Guide*, 2016.
- [27] D. Merrill, M. Garland, and A. Grimshaw, “Scalable GPU graph traversal,” in *ACM SIGPLAN Notices*, 2012.
- [28] H. Chitsaz and S. M. LaValle, “Time-optimal paths for a Dubins airplane,” in *Proc. IEEE Conf. on Decision and Control*, 2007.
- [29] I. Armeni, O. Sener, A. R. Zamir, H. Jiang, I. Brilakis, M. Fischer, and S. Savarese, “3D semantic parsing of large-scale indoor spaces,” in *IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, 2016.