



Outlook

Re: Action Needed - NVIDIA Academic Grant Program Award

From Binh Nguyen**Date** Sat 3/14/2026 1:47 PM**To** Matt Garratt <m.garratt@unsw.edu.au>

Thanks, Matt. I'll also go through their email in detail next week.

From: Matt Garratt <m.garratt@unsw.edu.au>**Sent:** Saturday, March 14, 2026 1:44:15 PM**To:** Binh Nguyen <binh.nguyen1@unsw.edu.au>**Subject:** Re: Action Needed - NVIDIA Academic Grant Program Award

Not sure will take a look tomorrow. Will approve your leave then too.

On 14 Mar 2026, at 12:38pm, Binh Nguyen <binh.nguyen1@unsw.edu.au> wrote:

That's really great news for us, Matt.

Thank you for your support throughout the process. Just wanted to check if there's anything we're required to do on our end?

Cheers,
Binh

From: Matt Garratt <m.garratt@unsw.edu.au>**Sent:** Saturday, March 14, 2026 11:24:17 AM**To:** Binh Nguyen <binh.nguyen1@unsw.edu.au>**Subject:** Fwd: Action Needed - NVIDIA Academic Grant Program Award

Hey Binh, great news!

Begin forwarded message:

From: NVIDIAAcademicGrants <nvidiaacademicgrants@nvidia.com>**Date:** 14 March 2026 at 10:13:23 am AEDT**To:** Matt Garratt <m.garratt@unsw.edu.au>**Cc:** nvidiaacademicgrants@nvidia.com**Subject:** Action Needed - NVIDIA Academic Grant Program Award

You don't often get email from nvidiaacademicgrants@nvidia.com. [Learn why this is important](#)



ACADEMIC GRANT PROGRAM

Dear Matt Garratt,

Congratulations! NVIDIA has selected your project, **VLM-Powered Voice-Commanded Autonomous Navigation with Deterministic Safety**, for inclusion in our Academic Grant Program.

In support of your project, we are donating **University of New South Wales(UNSW)** the following:

- **32K A100 GPU-Hours on Brev**
- **2x Jetson AGX Orin Dev Kit**

This grant is an unrestricted gift to support your described project and is to be used for academic purposes only.

Please be sure to acknowledge NVIDIA's support in papers, talks, posters, social, and news related to the awarded project, as follows: "Research supported by the NVIDIA Academic Grant Program using ." [PR guidelines](#) should be followed for all news pieces.

You can report these results at anytime via the [grant portal](#) by selecting "Report Results" next to the relevant application.

We encourage you to share your award on social media, tagging #NVIDIAGrant and @NVIDIAAI Dev (X) or @NVIDIAAI (LinkedIn), so we can help amplify your work.

For additional technical resources, you and your collaborators can join the [NVIDIA Developer Program](#) for free access to NIM, SDKs, self-paced training, early access programs for product releases, community forums, and unlimited viewings of [NVIDIA On-Demand](#). As an academic, you can also download industry-leading [coursework and materials](#) for teaching.

As a reminder, you have already accepted the [NVIDIA Academic Grant Program Terms and Conditions](#). We encourage you to review them and ensure you can meet all requirements before moving forward.

You can expect to receive a follow-up email from our team regarding your hardware redemption with further instructions by the end of May. This will include the necessary steps to accept your award and initiate the shipping process.

Your cloud access window is **April 1, 2026 - September 30, 2026. Extensions are not possible and any GPU hours remaining after the access window ends will be forfeited.**

An account activation email will be sent on your first day of access. Please check your spam folder if you do not see it in your inbox. Respond to that email to activate your environment. Then, add team members as collaborators to your project. You are responsible for monitoring the usage of all team members.

We have reserved 1 instance of 8xA100 80GB for the duration of your project. This node can be used by one user at a time and cannot be subdivided into smaller instances.

Usage is monitored as GPU hours, and we have provided enough GPU hours to use the provided instance continuously during the access window.

NVIDIA is excited to support your work and looks forward to the results of your research.

Thank you for your continued engagement with NVIDIA technology.

Signed by:

D5869A5B8BC84AC...

Howard Wright

VP Developer Programs

This email and any files transmitted with it may contain confidential information. If you believe you have received this email or any of its contents in error, please notify me immediately by return email and destroy this email. Do not use, disseminate, forward, print or copy any contents of an email received in error.



Thai Binh Nguyen <thethaibinh@gmail.com>

FW: Action Needed - NVIDIA Academic Grant Program Award

4 messages

Matt Garratt <m.garratt@unsw.edu.au>
To: Yu Zhou <yu.zhou8@student.unsw.edu.au>, Ryan Faulks <r.faulks@student.unsw.edu.au>, Josh Mansfield <josh.mansfield@student.unsw.edu.au>, Md Rafiqul Islam <md_rafiqul.islam@unsw.edu.au>, Udula Ranasinghe <udula_indrajith.ranasinghe@unsw.edu.au>, Gavin Kane <gavin.kane@unsw.edu.au>
Cc: "Thai Binh Nguyen (via Google Drive)" <thethaibinh@gmail.com>

6 May 2026 at 10:40

Hi Guys, we have been assigned 32,000 A100 GPU hours by NVIDIA – is this useful to any of you?

This is from a grant that Binh got but as you know he has moved to a new position

It would be great if we don't waste them. NVIDIA is sending me emails to please explain why we are not accessing the cloud.

Given its A100 and runs on NVIDIA cloud, I suspect it will support Isaac Sim.

But we have to use them by September.

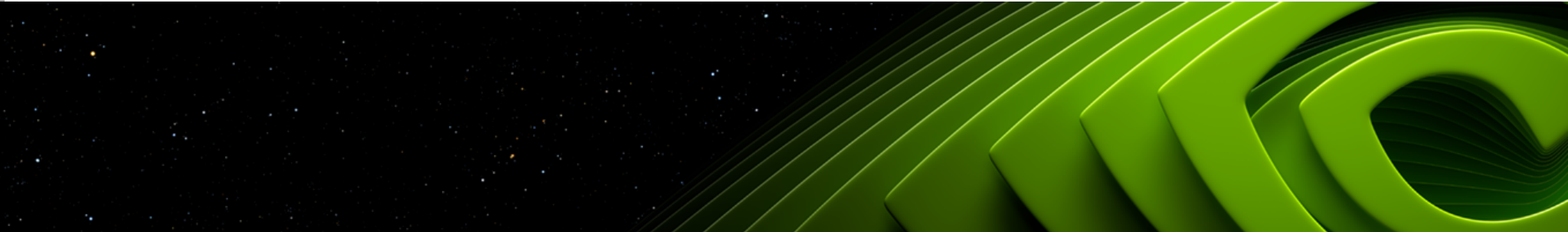
If anyone is interested, please let me know and I will add you to the project.

Kind regards,

Matt

From: NVIDIAAcademicGrants <nvidiaacademicgrants@nvidia.com>
Sent: Saturday, 14 March 2026 10:13 AM
To: Matt Garratt <m.garratt@unsw.edu.au>
Cc: nvidiaacademicgrants@nvidia.com
Subject: Action Needed - NVIDIA Academic Grant Program Award

You don't often get email from nvidiaacademicgrants@nvidia.com. [Learn why this is important](#)



ACADEMIC GRANT PROGRAM

Dear Matt Garratt,

Congratulations! NVIDIA has selected your project, VLM-Powered Voice-Commanded Autonomous Navigation with Deterministic Safety , for inclusion in our Academic Grant Program.

In support of your project, we are donating University of New South Wales(UNSW) the following:

- 32K A100 GPU-Hours on Brev
- 2x Jetson AGX Orin Dev Kit

This grant is an unrestricted gift to support your described project and is to be used for academic purposes only.

Please be sure to acknowledge NVIDIA's support in papers, talks, posters, social, and news related to the awarded project, as follows: "Research supported by the NVIDIA Academic Grant Program using ." [PR guidelines](#) should be followed for all news pieces.

You can report these results at anytime via the [grant portal](#) by selecting "Report Results" next to the relevant application.

We encourage you to share your award on social media, tagging #NVIDIAGrant and @NVIDIAIDev (X) or @NVIDIAAI (LinkedIn), so we can help amplify your work.

For additional technical resources, you and your collaborators can join the [NVIDIA Developer Program](#) for free access to NIM, SDKs, self-paced training, early access programs for product releases, community forums, and unlimited viewings of [NVIDIA On-Demand](#). As an academic, you can also download industry-leading [coursework and materials](#) for teaching.

As a reminder, you have already accepted the [NVIDIA Academic Grant Program Terms and Conditions](#). We encourage you to review them and ensure you can meet all requirements before moving forward.

You can expect to receive a follow-up email from our team regarding your hardware redemption with further instructions by the end of May. This will include the necessary steps to accept your award and initiate the shipping process.

Your cloud access window is **April 1, 2026 - September 30, 2026. Extensions are not possible and any GPU hours remaining after the access window ends will be forfeited.**

An account activation email will be sent on your first day of access. Please check your spam folder if you do not see it in your inbox. Respond to that email to activate your environment. Then, add team members as collaborators to your project. You are responsible for monitoring the usage of all team members.

We have reserved 1 instance of 8xA100 80GB for the duration of your project. This node can be used by one user at a time and cannot be subdivided into smaller instances.

Usage is monitored as GPU hours, and we have provided enough GPU hours to use the provided instance continuously during the access window.

NVIDIA is excited to support your work and looks forward to the results of your research.

Thank you for your continued engagement with NVIDIA technology.

Signed by:

D5869A5B8BC84AC...

Howard Wright

VP Developer Programs

This email and any files transmitted with it may contain confidential information. If you believe you have received this email or any of its contents in error, please notify me immediately by return email and destroy this email. Do not use, disseminate, forward, print or copy any contents of an email received in error.

Thai Binh Nguyen <thethaibinh@gmail.com>

To: Matt Garratt <m.garratt@unsw.edu.au>

Cc: Yu Zhou <yu.zhou8@student.unsw.edu.au>, Ryan Faulks <r.faulks@student.unsw.edu.au>, Josh Mansfield <josh.mansfield@student.unsw.edu.au>, Md Rafiqul Islam <md_rafiqul.islam@unsw.edu.au>, Udula Ranasinghe <udula_indrajith.ranasinghe@unsw.edu.au>, Gavin Kane <gavin.kane@unsw.edu.au>

6 May 2026 at 10:46

Hi Matt,

To make it more convenient, you can log in to the NVIDIA academic grant portal and add people directly to the project. This should allow them to access the cloud resources more easily.

Kind regards,
Binh Nguyen

[Quoted text hidden]

Matt Garratt <m.garratt@unsw.edu.au>
To: Thai Binh Nguyen <thethaibinh@gmail.com>

6 May 2026 at 10:48

Indeed, that's why I sent the email.

I have already added Gavin.

[Quoted text hidden]

Gavin Kane <gavin.kane@unsw.edu.au>
To: Matt Garratt <m.garratt@unsw.edu.au>, Yu Zhou <yu.zhou8@student.unsw.edu.au>, Ryan Faulks <r.faulks@student.unsw.edu.au>, Josh Mansfield <josh.mansfield@student.unsw.edu.au>, Md Rafiqul Islam <md_rafiqul.islam@unsw.edu.au>, Udula Ranasinghe <udula_indrajith.ranasinghe@unsw.edu.au>
Cc: "Thai Binh Nguyen (via Google Drive)" <thethaibinh@gmail.com>

8 May 2026 at 10:03

G'Day Matt,

I have previously struggled to use Isaac Sim on the cloud, but the Isaac Lab should work in headless mode, and in this context there is actually something I would like to test.

Can you please add me to the project, and if we are able to meet up Monday for an update meeting, I would be interested in some feedback about the feasibility.

Regards

Gavin

[Quoted text hidden]

ACADEMIC GRANTS PORTAL

To which CFP are you applying?	Application Status	Awarded
Research: Robotics & Edge AI		
CFP Subtopic	Robotics and Autonomous Vehicles	
Project Title		
VLM-Powered Voice-Commanded Autonomous Navigation with Deterministic Safety		
Interest Area	Robotics	
Description		
This project addresses a fundamental challenge in robotics: employing Vision-Language Models (VLMs) on robots while ensuring deterministic safety. VLMs offer complex scene perception but produce inherently uncertain outputs unsuitable for safety-critical actuation. We propose a novel architecture that bridges VLMs' broad reasoning with classical navigation algorithms, mapping non-deterministic high-level instructions to controlled, predictable low-level actions with strict safety guarantees.		
Institution/Organization	Website URL	
University of New South Wales Canberra	https://www.unsw.edu.au/staff/matt-garratt .https://www.unsw.edu.au/staff/matt-garratt	
First Name		
Matt	Primary Location	Australia

Last Name	Have you or are you currently collaborating with anyone at NVIDIA?
Garratt	
Email	Name of collaborator(s)
m.garratt@unsw.edu.au (mailto:m.garratt@unsw.edu.au)	
Role	Provide details on how you heard about the program
Professor	



[Privacy Policy](https://www.nvidia.com/en-us/about-nvidia/privacy-policy/) | **[Cookie Policy](https://www.nvidia.com/en-us/about-nvidia/cookie-policy/)** | **[Your Privacy Choices](https://www.nvidia.com/en-us/about-nvidia/privacy-center/)** | **[Terms of Use](https://developer.nvidia.com/legal/terms)** | **[Accessibility](https://www.nvidia.com/en-us/about-nvidia/accessibility/)** | **[Corporate Policies](https://www.nvidia.com/en-us/about-nvidia/company-policies/)** | **[Contact](https://www.nvidia.com/en-us/contact/)**

ACADEMIC GRANTS PORTAL

YOUR GRANTS

The table below lists the grants you have been awarded, if any.

If you have been awarded physical hardware, please review the details below.

ADD YOUR SHIPPING INFORMATION

To add your shipping information, click on the link under “HER Grant or Result Name” and fill out the appropriate fields. Here are some important tips for filling in the form:

- Provide a **phone number** to help our delivery partners reach you. They may require a signature at delivery.
- Ask your university's gifts office or department admin for your **ImporterID/EORI/TaxID** number. This helps us indicate to customs officials that the equipment will be used for academic purposes. Recipients in the U.S. can skip this field.

EDIT YOUR SHIPPING INFORMATION

If for any reason you need to make changes to your shipping information, you may do so as long as Grant Status = Awarded. When Grant Status = Submitted, the order has been placed through our shipping systems and can no longer be changed.

SHIPPING CONFIRMATION

Once you add your shipping information, you will receive an email from us confirming receipt of your address. You will receive another email with tracking information once your package has left the warehouse. Please note that this email may come from one of our fulfillment vendors and may not be from an NVIDIA email address.

SHIPPING TIMELINE

While we do our best to predict inventory and only award what we know we have available, sometimes there are delays while inventory is relocated, or the shipping team processes a large number of orders. We appreciate your patience during the shipping process and will do our best to get your equipment to you as quickly as possible!

QUESTIONS?

Please view our [FAQs \(https://www.nvidia.com/en-us/industries/higher-education-research/academic-grant-program/#faq\)](https://www.nvidia.com/en-us/industries/higher-education-research/academic-grant-program/#faq) before emailing us at [NVIDIAAcademicGrants@nvidia.com \(mailto:NVIDIAAcademicGrants@nvidia.com\)](mailto:NVIDIAAcademicGrants@nvidia.com) for assistance.

HER - Grants Results List View

2 of 2 items

HER Grant or Result N... ▾	Engagement Name ▾	Quantity ▾	Grant Type
32K A100 GPU-Hours on Brev (/academicgrantprogram/s/her-benefit/aDfVv000000QTCT)	VLM-Powered Voice- Commanded Autonomous Navigation with Deterministic Safety	32,000	Compute Credit
2x Jetson AGX Orin Dev Kit (/academicgrantprogram/s/her-benefit/aDfVv000000QTCu)	VLM-Powered Voice- Commanded Autonomous Navigation with Deterministic Safety	2	Hardware



Privacy Policy (<https://www.nvidia.com/en-us/about-nvidia/privacy-policy/>) | **Cookie Policy** (<https://www.nvidia.com/en-us/about-nvidia/cookie-policy/>) | **Your Privacy Choices** (<https://www.nvidia.com/en-us/about-nvidia/privacy-center/>) | **Terms of Use** (<https://developer.nvidia.com/legal/terms>) | **Accessibility** (<https://www.nvidia.com/en-us/about-nvidia/accessibility/>) | **Corporate Policies** (<https://www.nvidia.com/en-us/about-nvidia/company-policies/>) | **Contact** (<https://www.nvidia.com/en-us/contact/>)

[academicgrantprogram/s/results\)](#)[Logout \(https://academicgrants.nvidia.com/academicgrantprogram/s/welcome\)](https://academicgrants.nvidia.com/academicgrantprogram/s/welcome)

ACADEMIC GRANTS PORTAL

Welcome to NVIDIA's Academic Grants Program Portal! (<https://www.nvidia.com/en-us/industries/higher-education-research/academic-grant-program/>)

SUBMIT A NEW APPLICATION

You can submit a new application [here](https://academicgrants.nvidia.com/s/Application) (<https://academicgrants.nvidia.com/academicgrantprogram/s/Application>)

EDIT AN EXISTING APPLICATION

Applicants can update projects while the application is still in “applied” status. Begin by clicking on the corresponding project title. Be sure to save to lock in your changes.

ADD SHIPPING INFORMATION

Click “Grants” in the upper-right-hand menu of this portal screen and follow the instructions to provide your shipping information.

ADD PROJECT RESULTS

In the table below, click “Report Results” to the right of the project you’d like to add results for. Share updates for publications, conference talks, technical milestones, project compute speed-ups, and more. To help us understand the impact, please include a brief explanation for why the result is of particular importance. Our team uses your reported results to identify projects to highlight on our blog and social media channels.

ADD PROJECT COLLABORATORS

In the table below, click “Add Collaborator” to the right of the project you’d like to add your collaborators to. Team members can view the project and grant details as well as add project results.

QUESTIONS?

Please view our [FAQs \(https://www.nvidia.com/en-us/industries/higher-education-research/academic-grant-program/#faq\)](https://www.nvidia.com/en-us/industries/higher-education-research/academic-grant-program/#faq) before emailing us at NVIDIAAcademicGrants@nvidia.com (<mailto:NVIDIAAcademicGrants@nvidia.com>) for assistance.

Applications

Engagement ...	Institution	Status	Created Date	
VLM-Powered Voice-Commanded Autonomous Navigation with Deterministic Safety (/enterprise- projects/aDdVv0000 02tlrZKAQ),	University of New South Wales Canberra	Awarded	12/30/2025	Report Results



Privacy Policy (<https://www.nvidia.com/en-us/about-nvidia/privacy-policy/>) | **Cookie Policy** (<https://www.nvidia.com/en-us/about-nvidia/cookie-policy/>) | **Your Privacy Choices** (<https://www.nvidia.com/en-us/about-nvidia/privacy-center/>) | **Terms of Use** (<https://developer.nvidia.com/legal/terms>) | **Accessibility** (<https://www.nvidia.com/en-us/about-nvidia/accessibility/>) | **Corporate Policies** (<https://www.nvidia.com/en-us/about-nvidia/company-policies/>) | **Contact** (<https://www.nvidia.com/en-us/contact/>)

Project Title: EywaGuard: VLM-Powered Voice-Commanded Autonomous Navigation with Deterministic Safety

Principal Investigator: Professor Matt Garratt, University of New South Wales Canberra

Co-Investigator: Dr. Binh Nguyen, Research Associate, University of New South Wales Canberra

Abstract

This project addresses a grand challenge in autonomous robotics: how to safely deploy Vision-Language Models (VLMs) on physical robots while maintaining deterministic control over every decision that affects the robot, and more critically, the humans around it. While VLMs have demonstrated remarkable capabilities in scene understanding and high-level reasoning, their inherent uncertainty and non-deterministic outputs present fundamental safety concerns for real-world robotic deployment. We propose a novel architecture that achieves both complex environmental perception and deterministic safety guarantees by bridging VLMs' broad reasoning capabilities with the controlled execution of classical navigation algorithms. The key innovation is a mechanism that maps natural language commands and VLMs' complex, high-level decision-making to controlled, predictable low-level actions from traditional planners, enabling non-expert operators to command autonomous robots using everyday speech while maintaining strict safety guarantees. The proposed architecture provides visual and voice feedback with recommendations to the end user, while requiring voice or physical confirmation before execution, ensuring capable, convenient operation that remains strictly under human supervision. This project is built entirely around the comprehensive NVIDIA ecosystem, leveraging a hybrid cloud/edge architecture with NIM APIs, Isaac Sim, TensorRT-LLM, Isaac ROS, and DeepStream SDK. The EywaGuard framework establishes a foundation for safe, voice-commanded autonomy across diverse applications including mining, agriculture, security, and environmental monitoring.

Project Keywords

Vision-Language Models, Voice-Commanded Navigation, Safe Autonomous Systems, Edge AI, Natural Language Robotics

NVIDIA Platforms

This project is built entirely around the comprehensive NVIDIA ecosystem, leveraging a hybrid cloud/edge architecture:

Resources Request:

- 2 × NVIDIA Jetson Thor Developer Kits for edge VLM inference and robotic control
- 15,000 A100 80GB GPU hours for system integration and concurrent model inference via NVIDIA NIM™ APIs

Software – Cloud/API (Phase 1–2, Months 1–6):

- NVIDIA NIM (NVIDIA Inference Microservices)
- NIM-ready NVIDIA VLMs: Cosmos-Reason1-7B, Llama-3.1-Nemotron-Nano-VL-8B-v1, Cosmos-Nemotron-34B
- Isaac Sim: Digital-twin simulation for navigation algorithm development, cluttered environment scenarios, and automated safety validation.

Software – Edge (Phase 3, Months 7–12):

- Edge-supported NVIDIA VLMs: Cosmos-Reason1-7B, NVIDIA VILA, Cosmos-Nemotron-34B
- NVIDIA SDKs: JetPack SDK, TensorRT-LLM, Isaac ROS, DeepStream SDK

Introduction

Vision-Language Models (VLMs) and Vision-Language-Action models (VLAs) represent a paradigm shift in robotic perception and decision-making, offering unprecedented capabilities in scene understanding, natural language interaction, and high-level reasoning about the physical world. However, a critical gap exists between the impressive demonstrations of these models and their safe deployment on physical robots operating alongside humans. The fundamental challenge is this: VLMs produce outputs that are inherently uncertain and non-deterministic, yet robotic actuators require precise, predictable commands where errors can result in property damage, environmental harm, or human injury. While VLMs and VLAs have attracted significant research attention for robotic applications, the critical challenge of safe deployment, maintaining deterministic control over actions derived from inherently uncertain model outputs, remains largely unaddressed.



Figure 1: *The EywaGuard robotic platform featuring a gimbaled camera, GPS, independent suspension, and rugged 4-wheel chassis.*

This project develops **EywaGuard**, a framework for **deterministically safe, voice-commanded autonomous navigation** that integrates NVIDIA VLMs with our developing ground robot platform (Figure 1). We have already invested significant resources in the physical platform, including a custom 4-wheel chassis with independent suspension, a gimbaled camera, GPS/GNSS positioning, and a ROS2-based control software stack for fundamental locomotion. The key innovation is a novel architecture that combines the broad knowledge and reasoning capability about the physical world of a VLM with the deterministic safety of classical navigation algorithms, for example, our previous work as presented in [1], for low-level actuation and trajectory planning. A mechanism that maps natural language commands and VLMs' complex, high-level decision-making (inherently uncertain and non-deterministic) to controlled, predictable low-level actions from a traditional trajectory planner enables a new standard of intelligence and safety for industrial robotic systems.

The proposed architecture provides visual and voice feedback with recommendations to the end user, while requiring voice or physical confirmation before execution, enabling a capable, convenient system that remains strictly under human supervision. By enabling operators to command robots using natural language (e.g., "patrol the eastern perimeter and alert me if you detect any vehicles"), EywaGuard removes the barrier of specialist training and makes advanced autonomous robotics accessible to frontline workers across industries.

This work matters because it bridges advanced AI capabilities with real-world robotic deployment while addressing the critical safety gap.

Methods

Our technical approach follows a three-phase cloud-to-edge architecture (Figure 2), progressively building toward a robust, interchangeable VLM deployment strategy with deterministic safety guarantees for voice-commanded navigation.

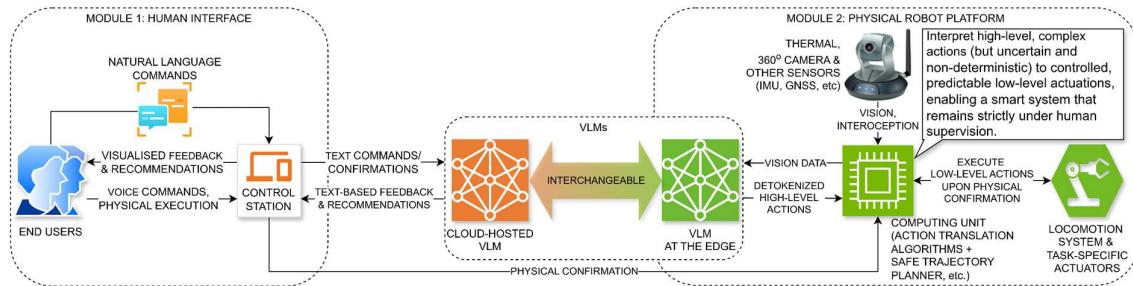


Figure 2: EywaGuard system architecture showing human interface with voice commands, interchangeable cloud/edge VLMs, and physical robot platform with action translation for safe trajectory planning.

Phase 1 – Cloud VLMs with Digital Twin Simulation (Months 1–3). The initial phase integrates NVIDIA's cloud-hosted VLMs via NIM APIs with **NVIDIA Isaac Sim** for comprehensive digital-twin simulation. Cloud-hosted models (Cosmos-Reason1-7B, Llama-3.1-Nemotron-Nano-VL-8B-v1, Cosmos-Nemotron-34B) run concurrently to process simulated sensor data and voice commands. Isaac Sim enables rapid development and automated testing of the action translation mechanism across diverse scenarios including cluttered environments, dynamic obstacles, and edge cases that would be risky or impractical to test on physical hardware. Moreover, a cross-platform software stack of the proposed system is developed to pave a platform for seamless CI/CD onto real-world hardware in the next phases.

Phase 2 – Cloud VLMs with Physical Robot (Months 4–6). The second phase transitions from simulation to physical robot deployment while maintaining cloud-based VLM inference. The EywaGuard platform streams real sensor data, including RGB, thermal, and depth imagery, to cloud-hosted models for scene understanding, voice command interpretation, and high-level decision-making. This phase validates the action translation mechanism with real-world sensor noise, environmental variability, and physical dynamics. Field trials in controlled environments verify the complete pipeline from voice command to physical execution with human confirmation, establishing baseline performance metrics for edge deployment.

Phase 3 – Physical Robot with Interchangeable Cloud and Edge VLMs (Months 7–12). The final phase implements full edge deployment using the NVIDIA Jetson Thor Developer Kit with interchangeable cloud/edge capability. Edge-optimized models (Cosmos-Reason1-7B, NVIDIA VILA, Cosmos-Nemotron-34B) are quantized and compiled using TensorRT-LLM for low latency, enabling real-time voice-commanded autonomous operation in remote, connectivity-limited environments. Isaac ROS provides GPU-accelerated SLAM and navigation, while DeepStream SDK enables concurrent multi-stream video analytics. The system is able to dynamically switch between edge inference and cloud endpoints based on connectivity and mission criticality, providing redundancy and scalability: edge handles time-critical safety decisions while cloud supports complex reasoning when available.

Proposed Timeline

Month/Year	Activity + Software/Models	Resources
Mar – May 2026	Phase 1: Isaac Sim digital-twin setup; Software architecture design & development; NIM API/VLMs integration; Design simulation scenarios, automated benchmarking & testing the developed algorithm with 3 concurrent models; Voice interface development & natural language command parsing; ICRA 2027 paper preparation	7,500 A100 hours

Month/Year	Activity + Software/Models	Resources
Jun – Aug 2026	Phase 2: Deploy proposed algorithm onto Jetson Thor using Isaac ROS & DeepStream SDK for multi-vision streaming; Testing real EywaGuard hardware with cloud-based VLMs in simple, controlled environments; ICRA 2027 paper submission	7,500 A100 hours. Jetson Thor [arrival]
Sep 2026 – Feb 2027	Phase 3: VLMs deployment onto Jetson Thor using TensorRT-LLM optimization; Field trials with industrial partners or national park authorities & real-world performance benchmarking; IROS 2027 submission; Open-source release of EywaGuard framework	Jetson Thor

Projected Outcomes:

- Two conference publications (ICRA 2027, IROS 2027)
- Open-source EywaGuard framework on GitHub

Dependencies: No external grant dependencies. The EywaGuard robotic platform (chassis, electronics, basic locomotion software) is already developed and operational.

Project Support Details & Justification:

- **Cloud Hours (15,000):** Required for Phase 1 & 2 (Months 1–6) system integration and running concurrent inference of 3 VLMs simultaneously. Calculation: 24 hours × 30 days × 6 months × 3 models ≈ 13,000 hours, rounded to 15,000 to accommodate peak usage and experimentation.
- **Jetson Thor Developer Kit (2×):** Essential for Phase 2 & 3 (Months 4–12): one integrated into EywaGuard for actual deployment and field trials, the other for development. Cloud cannot replace edge inference requirement for real-time/safety-critical voice-commanded decisions in connectivity-limited environments.

Cloud Readiness

Team fully prepared for immediate development upon cloud access:

- **Cloud Infrastructure:** Utilizing stock NVIDIA VLMs (Cosmos-Reason1-7B, Llama-3.1-Nemotron-Nano-VL-8B-v1, Cosmos-Nemotron-34B) fully supported by NVIDIA AI Enterprise. Integration via NIM APIs using OpenAI-compatible Python SDK, no custom hosting required.
- **Algorithm development & Software stack:** A containerized robotics digital twin framework based on ROS2 and NVIDIA IsaacSim has been developed from our previous work [1]. Cross-platform containerization enables systematic CI/CD deployment across workstation, cloud, and NVIDIA edge hardware via NVIDIA Container Toolkit.
- **Simulation:** Isaac Sim cluttered environments have been designed and ready for critical navigation testing.

Appendix B – CV(s)

Professor Matt Garratt – Principal Investigator

Professor of Autonomous Systems, UNSW Canberra at ADFA. Theme Lead for AI in the Defence Trailblazer initiative. 20+ years experience in autonomous robotics and control engineering.

Dr. Binh Nguyen – Co-Investigator & Lead Developer

Research Associate, UNSW Canberra. PhD in Robotics & AI. 8+ years in aerial/ground robotics, embedded systems.

Selected Publications:

- Binh Nguyen, et al., "Non-Conservative Efficient Collision Checking and Depth Noise-Awareness for Trajectory Planning," *IEEE Robotics and Automation Letters*, vol. 10, no. 8, pp. 7859-7866, 2025.
- Binh Nguyen, et al., "Online State-to-State Time-Optimal Trajectory Planning for Quadrotors." *ICUAS* 2024.

References

- [1] B. Nguyen, M. Murshed, T. Choudhury, K. Keogh, G. K. Appuhamillage, and L. Nguyen, "Non-Conservative Efficient Collision Checking and Depth Noise-Awareness for Trajectory Planning," *IEEE Robotics and Automation Letters*, vol. 10, no. 8, pp. 7859-7866, 2025.