

Titre: Quantification de l'incertitude et fusion sémantique multimodale
pour la navigation critique des rovers spatiaux en environnement
extraterrestre
Title:

Auteur: Théo Gachet
Author:

Date: 2026

Type: Mémoire ou thèse / Dissertation or Thesis

Référence: Gachet, T. (2026). Quantification de l'incertitude et fusion sémantique
multimodale pour la navigation critique des rovers spatiaux en environnement
extraterrestre [Mémoire de maîtrise, Polytechnique Montréal]. PolyPublie.
Citation: <https://publications.polymtl.ca/76273/>

 **Document en libre accès dans PolyPublie**
Open Access document in PolyPublie

URL de PolyPublie: <https://publications.polymtl.ca/76273/>
PolyPublie URL:

**Directeurs de
recherche:** Giovanni Beltrame
Advisors:

Programme: Génie logiciel
Program:

POLYTECHNIQUE MONTRÉAL

affiliée à l'Université de Montréal

**Quantification de l'incertitude et fusion sémantique multimodale pour la
navigation critique des rovers spatiaux en environnement extraterrestre**

THÉO GACHET

Département de génie informatique et génie logiciel

Mémoire présenté en vue de l'obtention du diplôme de *Maîtrise ès sciences appliquées*
Génie informatique

Avril 2026

POLYTECHNIQUE MONTRÉAL

affiliée à l'Université de Montréal

Ce mémoire intitulé :

**Quantification de l'incertitude et fusion sémantique multimodale pour la
navigation critique des rovers spatiaux en environnement extraterrestre**

présenté par **Théo GACHET**

en vue de l'obtention du diplôme de *Maîtrise ès sciences appliquées*

a été dûment accepté par le jury d'examen constitué de :

Christopher J. PAL, président

Giovanni BELTRAME, membre et directeur de recherche

Quentin CAPPART, membre

DÉDICACE

À mes parents Céline et Nicolas, à mes grands-parents Lucette, Monique, Michel et Daniel, pour m'avoir transmis les valeurs d'exigence et de persévérance qui ont forgé mon parcours.

À Yasmine, Tom et Zoé, dont le soutien a rendu possible l'aboutissement de ces travaux.

*« L'imagination est plus importante que la connaissance.
La connaissance est limitée, alors que l'imagination englobe le monde entier,
stimulant le progrès, et donnant naissance à l'évolution. »*

– Albert Einstein.

REMERCIEMENTS

Je tiens en premier lieu à exprimer ma profonde gratitude envers mon directeur de recherche, le professeur Giovanni Beltrame, pour son encadrement, sa confiance et son soutien tout au long de mes recherches au sein du MIST Lab. Mes remerciements s'adressent également aux professeurs Christopher J. Pal et Quentin Cappart pour l'intérêt qu'ils ont porté à mes travaux en acceptant d'évaluer ce manuscrit. Je souhaite par ailleurs remercier les équipes pédagogiques et administratives de Polytechnique Montréal et de l'École des Mines de Saint-Étienne, dont l'accompagnement a rendu possible la réalisation de ce double diplôme.

Je remercie l'ensemble des chercheurs et étudiants du MIST Lab pour nos échanges scientifiques et l'environnement de recherche stimulant qu'ils contribuent à maintenir. Ma reconnaissance s'étend à mes anciens collègues de l'Agence Spatiale Européenne (ESA), dont les discussions techniques ont nourri mes réflexions sur la navigation spatiale autonome.

Enfin, je souhaite remercier ma famille pour leur soutien inconditionnel tout au long de mon parcours universitaire, et ce par-delà les continents. Leur bienveillance et leurs encouragements constants ont constitué le socle de mon épanouissement scientifique et humain.

RÉSUMÉ

L’exploration planétaire par des rovers autonomes exige une navigation hautement sécurisée au sein d’environnements extraterrestres non structurés. Si les modèles de perception multimodale basés sur l’apprentissage profond offrent une compréhension sémantique indispensable pour évaluer la traversabilité des terrains, ils souffrent d’une limitation structurelle critique : la surconfiance. Ces réseaux tendent en effet à formuler des prédictions erronées avec des probabilités extrêmement élevées, en particulier face à des données hors-distribution ou des géologies inédites, ce qui peut conduire à la perte du système robotique.

L’objectif de ce mémoire est de concevoir, d’implémenter et de valider un pipeline unifié de quantification et de calibration de l’incertitude pour la perception dense multimodale. L’enjeu est de doter le robot d’une capacité d’introspection statistique lui permettant de rejeter les prédictions ambiguës avant qu’elles n’influencent la prise de décision.

La méthodologie propose une architecture symétrique appliquée conjointement à la segmentation d’images 2D et de nuages de points 3D. Elle fusionne l’approximation bayésienne par Monte-Carlo Dropout, pour estimer l’incertitude épistémique, avec une calibration paramétrique par Temperature Scaling pour corriger la surconfiance. Pour fournir des garanties statistiques formelles, l’approche intègre la prédiction conforme : les Ensembles Prédicatifs Adaptatifs assurent une couverture marginale cible, tandis qu’une adaptation par classe de la calibration de Kandinsky garantit une fiabilité régionale sémantiquement interprétable.

L’évaluation sur des données lunaires démontre la robustesse de l’architecture. La calibration corrige drastiquement la surconfiance initiale, tandis que la prédiction conforme maintient des garanties statistiques sans sacrifier la lisibilité des prédictions. L’analyse prouve la capacité du pipeline à isoler avec justesse les défaillances locales des capteurs. En exploitant la complémentarité des modalités, le système opère un filtrage intelligent : il impose un rejet strict des éléments visuellement ambigus en 2D, tout en validant leur détection si elle s’appuie sur une topologie 3D fiable. Ce rejet sélectif améliore ainsi la précision sémantique sur l’ensemble des données conservées.

En conclusion, ce mémoire établit qu’une quantification rigoureuse de l’incertitude transmute les défaillances perceptuelles de l’apprentissage profond en signaux d’exclusion fiables. Projetés en cartes de risque, ils sont directement exploités par le nouvel algorithme probabiliste AURA* (*Adaptive Uncertainty-Aware Risk A**). En optimisant la trajectoire sur ce champ de coût continu, ce planificateur dote le système d’une navigation préventive, apportant ainsi la résilience fondamentale exigée par l’autonomie des futures missions d’exploration spatiale.

ABSTRACT

Planetary exploration by autonomous rovers demands fail-safe navigation within unstructured extraterrestrial environments. While deep learning-based multimodal perception models offer the semantic understanding required to assess terrain traversability, they suffer from a critical structural flaw: overconfidence. Indeed, these networks frequently yield erroneous predictions with high certainty, especially when encountering out-of-distribution data or novel geological formations, thereby jeopardizing the entire robotic system.

To address this vulnerability, this thesis designs, implements, and validates a unified uncertainty quantification and calibration pipeline for dense multimodal perception. The core objective is to endow the rover with a capacity for statistical introspection, enabling it to preemptively discard ambiguous predictions before they compromise downstream decision-making.

The proposed methodology relies on a symmetrical architecture applied concurrently to 2D image and 3D point cloud segmentation. It merges Bayesian approximation via Monte-Carlo Dropout to estimate epistemic uncertainty, with global parametric calibration via Temperature Scaling to mitigate network overconfidence. Furthermore, to provide formal statistical guarantees, the framework incorporates conformal prediction: Adaptive Predictive Sets ensure target marginal coverage, while a class-wise adaptation of Kandinsky calibration yields a semantically interpretable regional reliability.

Extensive evaluation on a lunar dataset highlights the framework’s robustness. The calibration phase drastically rectifies the models’ initial overconfidence, while conformal prediction upholds strict statistical guarantees without degrading prediction interpretability. Analytical results further demonstrate the pipeline’s capability to accurately pinpoint localized sensor anomalies. By exploiting the complementarity of both modalities, the system performs an intelligent filtering process: it rejects visually ambiguous features in 2D imagery while validating their detection whenever backed by reliable 3D topology. Consequently, this selective rejection mechanism enhances the semantic accuracy of the retained data.

Ultimately, this research demonstrates that rigorous uncertainty quantification can translate the perceptual limitations of deep learning into reliable exclusion signals. Projected as risk maps, these signals are directly leveraged by AURA* (*Adaptive Uncertainty-Aware Risk A**), a novel probabilistic path-planning algorithm developed within this scope. By optimizing trajectories over a continuous cost field, AURA* fosters preventive navigation, thus delivering the resilience required to achieve full autonomy in future space exploration missions.

TABLE DES MATIÈRES

DÉDICACE ET REMERCIEMENTS	iii
RÉSUMÉ	v
ABSTRACT	vi
LISTE DES TABLEAUX	ix
LISTE DES FIGURES	x
LISTE DES SIGLES ET ABRÉVIATIONS	xi
CHAPITRE 1 INTRODUCTION	1
CHAPITRE 2 REVUE DE LITTÉRATURE	6
2.1 Navigation et traversabilité en milieu planétaire	6
2.2 Perception multimodale pour l’environnement non-structuré	8
2.3 Quantification de l’incertitude en apprentissage profond	9
2.4 De la perception à la planification robuste	11
2.5 Synthèse et positionnement de la thèse	13
CHAPITRE 3 PIPELINE DE PERCEPTION MULTIMODALE	14
3.1 Le jeu de données LuSNAR	14
3.2 Synchronisation et calibration multimodale des données	16
3.3 Segmentation des images pour la perception embarquée	20
3.4 Segmentation des points LiDAR pour la modélisation de terrain	23
CHAPITRE 4 CALIBRATION ET QUANTIFICATION DE L’INCERTITUDE	27
4.1 Introduction et motivation	27
4.2 Contributions et cadre unifié du pipeline d’incertitude	28
4.3 Monte-Carlo Dropout : Réduction de la surconfiance par échantillonnage	30
4.4 Temperature Scaling : Calibration paramétrique des probabilités	33
4.5 Adaptive Prediction Sets : Ensembles prédictifs conformes	36
4.6 Kandinsky Calibration : Calibration régionale de l’incertitude	39
4.7 Architecture complète du pipeline d’incertitude : schémas et pseudocodes	41
4.8 Expérimentations et validation du pipeline d’incertitude	44
4.9 Bilan du chapitre : synthèse, limites et perspectives	62

CHAPITRE 5	NAVIGATION SOUS INCERTITUDE : L'ALGORITHME AURA*	63
5.1	Enjeux de la transition vers une commande motrice robuste	63
5.2	Formalisation des signaux d'entrée	64
5.3	Fusion conservatrice multimodale	66
5.4	Élaboration de la fonction de coût globale $J(u, v)$	67
5.5	L'algorithme AURA* : Adaptive Uncertainty-Aware Risk A*	68
5.6	Validation expérimentale : marges adaptatives et inflation statique	73
5.7	Bilan du chapitre : synthèse, limites et perspectives	86
CHAPITRE 6	CONCLUSION	87
RÉFÉRENCES		88

LISTE DES TABLEAUX

Tableau 3.1	Matrice de confusion normalisée et IoU par classe	21
Tableau 3.2	Distribution des classes sur l'ensemble des points LiDAR	24
Tableau 4.1	Impact de la calibration Kandinsky sur la modalité 2D (images). . . .	47
Tableau 4.2	Impact de la calibration Kandinsky sur la modalité 3D (LiDAR). . .	47
Tableau 4.3	Comparaison quantitative des performances pour la modalité 2D. . .	53
Tableau 4.4	Comparaison quantitative des performances pour la modalité 3D. . .	54
Tableau 5.1	Métriques cinématiques pour $\rho = 100\%$	77
Tableau 5.2	Synthèse des métriques cinématiques et sécuritaires pour $\rho = 27\%$. . .	79
Tableau 5.3	Synthèse des métriques cinématiques et sécuritaires pour $\rho = 0\%$. . .	81

LISTE DES FIGURES

Figure 3.1	Classification des scènes LuSNAR par topographie et densité.	15
Figure 3.2	Projection d'un point 3D vers le plan image via la matrice intrinsèque.	17
Figure 3.3	Structure interne d'un fichier <code>.npz</code> utilisé dans ce projet.	18
Figure 3.4	Visualisations synchronisées pour le timestamp $t = 1692601812033751296$	19
Figure 3.5	Comparaison qualitative sur une scène complète.	22
Figure 3.6	Zoom sur la région centrale pour une meilleure visibilité des erreurs.	22
Figure 3.7	Architecture utilisée pour la segmentation sémantique point par point.	23
Figure 3.8	Perte moyenne (gauche) et précision moyenne (droite) selon les époques.	25
Figure 3.9	Évolution des IoU par classe et du mIoU durant l'entraînement	25
Figure 3.10	Projection multi-échelle des prédictions 3D à quatre niveaux de zoom.	26
Figure 4.1	Pipeline unifié de calibration avec partage de paramètres.	28
Figure 4.2	Illustration d'un réseau de neurones (gauche) et du Dropout (droite).	31
Figure 4.3	Construction du tenseur d'échantillons à partir des $T_{MC} = 25$ passes.	32
Figure 4.4	Pipeline de la méthode <i>Temperature Scaling</i> (TS).	35
Figure 4.5	Pipeline de la méthode <i>Adaptive Prediction Sets</i> (APS).	38
Figure 4.6	Pipeline de la méthode <i>Kandinsky Calibration</i> (KC).	41
Figure 4.7	Pipeline d'incertitude complet et pseudocodes : calibration hors-ligne.	42
Figure 4.8	Pipeline d'incertitude complet et pseudocodes : calibration en ligne.	43
Figure 4.9	Diagrammes de fiabilité générés sur le jeu de test.	44
Figure 4.10	Analyse de sensibilité des paramètres APS sur le jeu de test.	45
Figure 4.11	Analyse de la sensibilité au bruit et de la robustesse géométrique.	49
Figure 4.12	Gain de précision par rejet d'incertitude pour la modalité 2D (images).	52
Figure 4.13	Gain de précision par rejet d'incertitude pour la modalité 3D (LiDAR).	54
Figure 4.14	Visualisation des signaux intermédiaires de l'anatomie de décision 2D.	60
Figure 4.15	Visualisation des signaux intermédiaires de l'anatomie de décision 3D.	61
Figure 5.1	Schéma décisionnel d'AURA* illustrant l'expansion d'un nœud.	72
Figure 5.2	Comparaison qualitative de la Baseline et d'AURA* pour $\rho = 100\%$	77
Figure 5.3	Comparaison qualitative de la Baseline et d'AURA* au seuil de viabilité.	78
Figure 5.4	Comparaison qualitative de la Baseline et d'AURA* pour $\rho = 0\%$	80
Figure 5.5	Évolution du taux de succès exploratoire en fonction de ρ	82
Figure 5.6	Évolution de $M_{\max}$ le long des trajectoires générées, en fonction de ρ	83
Figure 5.7	Évolution de TC le long des trajectoires générées, en fonction de ρ	84
Figure 5.8	Frontière de Pareto entre la déviation métrique et le risque épistémique.	84
Figure 5.9	Évolution relative du temps de calcul d'AURA* par rapport à la Baseline.	85

LISTE DES SIGLES ET ABRÉVIATIONS

UQ	Uncertainty Quantification
SPOC	Soil Property and Object Classification
IMU	Inertial Measurement Unit
CNN	Convolutional Neural Network
ViT	Vision Transformer
LiDAR	Light Detection and Ranging
OOD	Out-of-Distribution
BNN	Bayesian Neural Network
MC	Monte Carlo
TS	Temperature Scaling
APS	Adaptive Prediction Sets
KC	Kandinsky Calibration
RRT*	Rapidly-exploring Random Tree Star
LUSNAR	Lunar Segmentation, Navigation and Reconstruction
SLAM	Simultaneous Localization and Mapping
MLP	Multilayer Perceptron
FNN	Feedforward Neural Network
IoU	Intersection over Union
SA	Set Abstraction
FP	Feature Propagation
NLL	Negative Log-Likelihood
ECE	Expected Calibration Error
LPIPS	Learned Perceptual Image Patch Similarity
AURC	Area Under the Risk-Coverage Curve
AURA*	Adaptive Uncertainty-Aware Risk A*
STEP	Stochastic Traversability Evaluation
POMDP	Partially Observable Markov Decision Process

CHAPITRE 1 INTRODUCTION

Contexte : L'ère de l'exploration autonome

L'exploration robotique du système solaire connaît un changement de paradigme majeur, évoluant vers une science où la robotique devient le vecteur principal de découvertes en environnements extrêmes [1]. Avec le programme Artemis et les missions de retour d'échantillons martiens, les agences spatiales transitionnent d'une exploration ponctuelle vers des opérations de surface prolongées et complexes.

Dans ce contexte, la mobilité des rovers constitue le facteur limitant principal de la productivité scientifique. La contrainte fondamentale réside dans la latence de communication : la distance Terre-Mars impose un délai aller-retour de 6 à 44 minutes. Cette contrainte physique interdisant la téléopération en temps réel, les premières missions comme *Mars Pathfinder* ont requis une décomposition stricte de la navigation en quatre fonctions : désignation de but, localisation, détection de risques et sélection de trajectoire [2]. Historiquement, les opérateurs devaient adopter des cycles prudents de type *stop-and-wait*, limitant les traverses quotidiennes du rover *Sojourner* à quelques dizaines de mètres en raison de la latence du système embarqué (20 secondes pour un unique scan d'obstacles) [2]. Pour franchir ce plafond et accroître le rendement scientifique, les rovers de nouvelle génération, tels que *Perseverance*, intègrent des capacités de navigation autonome (*AutoNav*) permettant une planification continue sans intervention humaine [3].

Cependant, les systèmes de navigation actuels reposent essentiellement sur une analyse géométrique de l'environnement. À l'aide de caméras stéréoscopiques, le rover construit des cartes de disparité pour détecter les pentes abruptes et les obstacles de haute taille [4]. Bien que robuste, cette approche souffre d'un aveuglement sémantique : elle ne permet pas de distinguer un substrat rocheux stable d'une zone de sable mou si la géométrie de surface est plane. Cette limitation a déjà conduit à des situations critiques, comme l'enlèvement définitif du rover *Spirit* en 2009 dans une zone de sol meuble indétectable par la géométrie seule [5].

Pour dépasser ces limites actuelles, la prochaine génération de rovers doit impérativement intégrer une compréhension sémantique de son environnement. L'utilisation de réseaux de neurones profonds pour la segmentation de terrain offre une solution prometteuse, mais soulève une nouvelle problématique de sécurité : comment garantir la fiabilité d'une décision prise par une intelligence artificielle dans un environnement inconnu et hostile ?

Problématique : Le défi de la perception incertaine

L'intégration de l'apprentissage profond dans la chaîne de décision des rovers spatiaux représente un changement de paradigme majeur par rapport aux approches purement géométriques. En permettant une segmentation sémantique dense de l'environnement, ces modèles offrent théoriquement une capacité de distinction fine entre les différents types de sols, ce qui s'avère déterminant pour l'évitement des zones à risque non géométrique comme les sables mouvants. Cependant, cette performance accrue s'accompagne d'une opacité structurelle : les réseaux de neurones profonds opèrent comme des boîtes noires probabilistes, dont la fiabilité en conditions opérationnelles réelles reste très difficile à garantir.

Le verrou technologique central réside dans le phénomène de sur-confiance inhérent aux réseaux de neurones modernes [6]. Bien que performants sur des données similaires à leur jeu d'entraînement, ces modèles tendent à produire des prédictions erronées avec des scores de confiance extrêmement élevés lorsqu'ils sont confrontés à des données hors-distribution [7]. Ce comportement est particulièrement critique dans le contexte de l'exploration planétaire, où le rover est constamment exposé à des formations géologiques inédites ou des conditions d'éclairage non modélisées. À titre d'exemple, un réseau de neurones mal calibré pourrait classer une roche inconnue et coupante comme étant du sable fin navigable, induisant le planificateur de trajectoire en erreur et mettant en péril l'intégrité de la mission.

Dès lors, la simple maximisation de la précision de classification ne suffit plus. La problématique centrale de ce mémoire se formule ainsi :

Comment quantifier et calibrer l'incertitude des modèles de segmentation sémantique profonds afin de garantir des seuils de sécurité statistique lors de la navigation autonome d'un rover en environnement non structuré ?

La réponse à cette question impose de dépasser l'utilisation d'une seule modalité sensorielle. En effet, les caméras optiques sont vulnérables aux variations extrêmes d'illumination (ombres portées lunaires, éblouissement), tandis que les capteurs LiDAR, bien qu'insensibles à la lumière, fournissent une information éparse et dépourvue de texture [8]. La robustesse de notre système de navigation doit donc nécessairement reposer sur une stratégie de fusion multimodale, couplée à une quantification rigoureuse de l'incertitude pour détecter et rejeter les prédictions non fiables avant qu'elles n'affectent la commande du robot [9].

Objectifs de recherche

L’objectif de ce travail est de concevoir et de valider un cadre de perception multimodale modulaire, capable d’auto-évaluer sa fiabilité indépendamment des architectures neuronales sous-jacentes. Dans le contexte de l’exploration planétaire, où la latence de communication prohibe la téléopération temps-réel, doter le rover d’une capacité d’introspection sur sa propre perception constitue une fonction de composante critique pour la viabilité du système. Il s’agit de dépasser le paradigme actuel, où la segmentation sémantique est traitée comme une vérité terrain déterministe, pour évoluer vers une approche probabiliste où chaque prédiction est assortie d’une mesure de confiance statistiquement calibrée. Pour concrétiser cette ambition globale, notre démarche scientifique se décline en trois objectifs spécifiques interdépendants.

O1 : nous visons à développer des méthodologies de quantification de l’incertitude qui soient strictement agnostiques aux modèles de perception. L’enjeu scientifique est de traiter les sorties de segmentation comme des données brutes nécessitant un post-traitement rigoureux, indépendamment de la topologie du réseau de neurones producteur. Nous visons à implémenter et analyser un spectre de mécanismes d’évaluation, allant des approches Bayésiennes par échantillonnage aux méthodes fréquentistes récentes comme la prédiction conforme. L’objectif est de déterminer dans quelle mesure ces différentes stratégies permettent de capturer l’incertitude de la scène, indépendamment des biais architecturaux, et de valider leur compatibilité respective.

O2 : nous voulons établir une architecture de fusion multimodale asymétrique. L’exploration planétaire requiert l’exploitation conjointe de la densité texturale des caméras stéréoscopiques et de la précision géométrique des capteurs LiDAR. Ces modalités étant affectées par des sources d’incertitude fondamentalement différentes (ambiguïtés photométriques pour la vision, raréfaction spatiale et bruit thermique pour le LiDAR), nous cherchons à définir un référentiel de fusion fondé sur le principe de précaution (*worst-case fusion*). Cette approche doit permettre de résoudre les conflits perceptuels en sanctuarisant les signaux de danger validés par au moins une modalité fiable.

O3 : ce projet a pour vocation d’opérationnaliser ces métriques d’incertitude au sein de la commande cinématique du robot. Au-delà de la simple cartographie de l’incertitude, nous visons à établir un lien formel entre la confiance perceptuelle et la prise de décision motrice. L’objectif est de démontrer que la projection continue de l’incertitude épistémique sur un champ de potentiel permet au rover de contourner préventivement les zones d’ignorance, résolvant ainsi le phénomène de paralysie (*Freezing Robot Problem*) inhérent aux méthodes d’inflation géométrique statiques.

Contributions

Les travaux présentés dans ce manuscrit proposent une amélioration de la chaîne de traitement décisionnelle des rovers planétaires, en ciblant l’interface entre la perception sémantique sous incertitude et la navigation autonome. Cette avancée s’articule autour de la conception d’une architecture de quantification de l’incertitude multimodale. Nous proposons un cadre d’inférence symétrique qui applique conjointement l’approximation bayésienne et la calibration globale à deux modalités hétérogènes : l’imagerie dense 2D et les nuages de points 3D. Contrairement aux approches de l’état de l’art qui traitent l’incertitude de ces capteurs de manière isolée ou par des méthodes computationnellement lourdes, notre architecture harmonise les espaces probabilistes, ce qui permet d’atténuer la surconfiance structurelle des réseaux tout en préservant une efficacité compatible avec l’inférence en temps réel.

Deuxièmement, ce projet introduit une application novatrice du cadre de la prédiction conforme aux données de perception 3D. Alors que l’état de l’art actuel concentre ses validations sur l’imagerie 2D dense, nous étendons ces garanties statistiques aux nuages de points LiDAR caractérisés par leur éparsité et leur irrégularité. Nous formulons une méthodologie de calibration de seuils de rejet adaptatifs, offrant la garantie théorique que le taux d’erreur de segmentation reste borné par un paramètre de risque α défini par l’opérateur.

Pour exploiter pleinement cette perception qualifiée, nous introduisons l’algorithme de planification probabiliste AURA* (*Adaptive Uncertainty-Aware Risk A**), palliant ainsi les limites des planificateurs basés sur l’échantillonnage ou la recherche sur grille géométrique standard. Cette méthode substitue les rayons d’inflation statiques, générateurs de blocages virtuels, par une fonction de coût d’énergie continue. En couplant une pénalité linéaire pour les risques avérés à une barrière exponentielle pour les zones d’incertitude épistémique hors-distribution, AURA* dote le système d’une dynamique de navigation intrinsèquement prudente. Cette approche domine structurellement les heuristiques classiques au sens de Pareto, en offrant un compromis strictement supérieur entre la viabilité exploratoire et la garantie de sécurité, permettant d’optimiser la distance parcourue sans augmenter le taux de collision.

Enfin, l’ensemble de ce corpus théorique est concrétisé au sein d’une architecture logicielle soumise à une validation sur le banc d’essai lunaire LuSNAR. Cette campagne d’évaluation démontre formellement, par le biais d’une analyse stochastique de Monte-Carlo, que notre politique de fusion exploite la robustesse structurelle de l’empreinte géométrique pour neutraliser les hallucinations du modèle visuel. Le système dote ainsi la plateforme d’une fiabilité décisionnelle accrue, parvenant à réduire l’exposition au risque épistémique de plus de 90 %.

Organisation du mémoire

Afin de guider le lecteur à travers l’élaboration de ce cadre de perception et de navigation sécuritaire, ce mémoire est structuré selon une progression analytique rigoureuse, cheminant des fondations théoriques jusqu’à la validation algorithmique globale.

Le chapitre 2 dresse un panorama critique de l’état de l’art en robotique planétaire et en apprentissage profond, en analysant l’évolution des algorithmes de traversabilité depuis les modèles géométriques canoniques vers les approches sémantiques modernes. Une attention particulière y est portée aux limites inhérentes à la segmentation déterministe et à la nécessité absolue d’adopter des méthodes de quantification de l’incertitude face aux environnements imprévisibles.

Le chapitre 3 pose ensuite les fondations matérielles et logicielles de notre architecture de perception. S’appuyant sur le banc d’essai LuSNAR, il détaille les mécanismes complexes de synchronisation spatio-temporelle et de calibration extrinsèque indispensables à l’alignement des données LiDAR et optiques. Ce chapitre présente ensuite la modélisation sémantique de l’environnement, décrivant la segmentation dense des images via le modèle SegFormer et la segmentation 3D des nuages de points par PointNet++, tout en analysant la complémentarité de ces modalités pour la compréhension d’environnements planétaires non structurés.

Le cœur théorique du manuscrit réside dans le chapitre 4, consacré à la conception détaillée de notre pipeline unifié de quantification de l’incertitude. Il y est formalisé la synergie entre l’échantillonnage bayésien (*Monte-Carlo Dropout*), le lissage paramétrique (*Temperature Scaling*) et les filtres de prédiction conforme (*Adaptive Prediction Sets* et *Kandinsky Calibration*). L’évaluation microscopique des signaux y démontre comment cette architecture inédite transmute les ambiguïtés perceptuelles en cartes de rejet strictes et déterministes.

Le chapitre 5 opérationnalise ces métriques statistiques à travers le développement et l’évaluation de l’algorithme AURA*. La modélisation de la fonction de coût énergétique et la politique de fusion multimodale conservatrice y sont minutieusement détaillées. Par le biais d’une campagne d’échantillonnage stochastique, ce chapitre démontre la supériorité éclatante d’AURA* face aux approches par inflation statique, quantifiant formellement la résolution du compromis historique entre la sécurité absolue du système et le maintien de sa mobilité spatiale.

Enfin, le chapitre 6 offre une synthèse globale des avancées proposées. Il mène une discussion transparente et critique sur les limitations inhérentes à nos choix architecturaux, ouvrant ainsi la voie à de futures perspectives de recherche articulées autour de l’adaptation au temps de test et des processus de décision markoviens.

CHAPITRE 2 REVUE DE LITTÉRATURE

Ce chapitre positionne les travaux de recherche à l’intersection de trois domaines : la navigation robotique, la perception artificielle multimodale et la quantification de l’incertitude. L’objectif est d’analyser l’évolution des paradigmes de navigation pour démontrer que les verrous technologiques actuels ne sont plus géométriques, mais sémantiques et probabilistes.

2.1 Navigation et traversabilité en milieu planétaire

La capacité d’un rover à se déplacer de manière autonome en environnement extraterrestre repose sur son aptitude à évaluer la traversabilité, définie comme la capacité d’un système robotique à franchir un terrain donné sans compromettre son intégrité mécanique ni sa stabilité. Cette évaluation a historiquement suivi une évolution, passant de modèles strictement géométriques à des approches sémantiques et physiques intégrant la terramécanique.

2.1.1 Approches géométriques classiques : L’héritage GESTALT

Depuis les missions *Mars Exploration Rovers* en 2004, le standard industriel pour la navigation autonome est incarné par l’algorithme GESTALT (*Grid-based Estimation of Surface Traversability Applied to Local Terrain*) [10]. Ce système, ainsi que ses successeurs implémentés sur les rovers *Curiosity* et *Perseverance*, repose sur l’acquisition d’images par des caméras stéréoscopiques (*NavCams*) pour générer des nuages de points 3D locaux [11]. Ces données forment une grille topographique centrées sur le robot, où chaque cellule livre des métriques géométriques (pente, rugosité résiduelle, hauteur de marche) [4].

Les planificateurs de trajectoire locaux, tels que *Enhanced Navigation* (ENav), utilisent ces cartes de coûts pour calculer des arcs de commande minimisant le risque cinématique [12]. Bien que cette approche ait démontré sa robustesse pour l’évitement d’obstacles francs tels que les rochers de grande taille ou les falaises, elle souffre d’une limitation majeure qualifiée d’aveuglement sémantique. En ne considérant que la géométrie, ces algorithmes postulent implicitement que tout sol plat est traversable. Comme évoqué dans l’introduction, cette hypothèse a causé l’enlisement définitif du rover *Spirit* au cratère Troy en 2009 sur une croûte plane dissimulant du sable sans cohésion, indétectable par stéréoscopie [5]. Cet événement a démontré que la sécurité d’une mission ne pouvait être garantie par la seule détection de la forme du terrain, mais nécessitait une compréhension de sa matière.

2.1.2 L’apport de la terramécanique et de l’interaction roue-sol

Au-delà de la géométrie, la traversabilité dépend intrinsèquement de l’interaction physique entre la roue déformable et le régolithe [13]. Si la terramécanique, basée sur les équations de Bekker, modélise la cohésion et l’angle de friction du sol pour en déduire la capacité de traction [14], son estimation traditionnelle reste majoritairement réactive. En effet, les capteurs proprioceptifs (courant moteur, centrales inertiels, capteurs de couple) ne détectent le glissement excessif qu’une fois le robot engagé sur le terrain difficile [15]. Cette approche *a posteriori* posant un risque inacceptable pour l’autonomie à haute latence, le rover doit anticiper les propriétés mécaniques avant tout contact physique via ses capteurs extéroceptifs, motivant ainsi la transition vers l’analyse sémantique visuelle.

2.1.3 L’apport de la sémantique pour la traversabilité prédictive

L’hypothèse centrale des approches modernes de navigation est que la classe sémantique d’un terrain (texture, couleur, granulométrie) est un estimateur fiable de ses propriétés mécaniques [16, 17]. Le système SPOC (*Soil Property and Object Classification*), développé au Jet Propulsion Laboratory de la NASA pour la mission Mars 2020, a été l’un des premiers à intégrer un classifieur visuel basé sur un réseau de neurones convolutionnel [18]. Ce système segmente le terrain en classes géologiques distinctes telles que le sable, le substrat rocheux (bedrock) ou le gravier. Ces classes sémantiques sont ensuite associées à des coûts de glissement prédéfinis dans la carte de navigation globale, permettant au planificateur d’éviter les zones de sable mou même si elles apparaissent géométriquement planes.

Plus récemment, la littérature s’oriente vers l’apprentissage auto-supervisé pour affiner ces modèles. Des travaux comme ceux de Wellhausen et al. [16] proposent d’utiliser l’expérience de conduite du robot comme vérité terrain : le système apprend à associer les textures visuelles perçues par la caméra aux vibrations et au glissement mesurés par l’IMU (*Inertial Measurement Unit*) lors de la traversée. Dans cette optique, Chavez-Garcia et al. [19] démontrent que l’hybridation des données visuelles et géométriques améliore significativement la prédiction du risque par rapport à l’usage d’une modalité unique. Ces conclusions sont corroborées par De Silva et al. [20], qui confirment la nécessité d’une approche multimodale pour garantir la sécurité de navigation. Néanmoins, l’introduction de réseaux de neurones profonds dans la boucle de perception pose un nouveau problème de fiabilité. Contrairement aux mesures géométriques déterministes, la segmentation sémantique est une estimation probabiliste sujette à des erreurs de classification, notamment sur des terrains d’apparence ambiguë. La quantification rigoureuse de cette incertitude devient alors indispensable pour ne pas reproduire les erreurs du passé avec des algorithmes plus complexes.

2.2 Perception multimodale pour l’environnement non-structuré

Dépassant la simple détection d’obstacles géométriques, la navigation planétaire exige une compréhension physico-chimique des sols (ex. distinguer un régolithe compact d’un sable meuble). Pour ce faire, les systèmes modernes exploitent la segmentation sémantique dense (attribution d’une classe par pixel ou point 3D), propulsée par l’apprentissage profond. Cette section justifie nos choix technologiques en analysant l’évolution des architectures 2D et 3D.

2.2.1 Segmentation sémantique 2D : Des CNNs aux Vision Transformers

Historiquement dominants, les réseaux convolutifs (CNN) basés sur des encodeurs-décodeurs, tels que U-Net [21] ou DeepLab et ses convolutions dilatées [22], extraient efficacement les caractéristiques visuelles locales. Par nature, la convolution traite l’image hiérarchiquement, agrégeant des gradients de bas niveau pour former des concepts sémantiques abstraits.

Toutefois, le champ récepteur fini des CNN constitue une limitation intrinsèque [23]. Sur des images planétaires à haute résolution, ils peinent à capturer le contexte global [24]. Distinguer une ombre d’un cratère, par exemple, requiert d’appréhender l’illumination globale de la scène, une information souvent hors de portée d’un voisinage convolutif standard [25].

Le paradigme des *Vision Transformers* (ViT) pallie ce défaut via le mécanisme d’attention [25, 26], modélisant les dépendances à longue portée dès l’encodage. Face au compromis performance-latence imposé par l’embarqué, SegFormer s’impose [27]. Son architecture hiérarchique allégée, dépourvue d’encodage positionnel complexe, résiste mieux aux variations d’échelle et perturbations visuelles, motivant son intégration dans notre pipeline 2D.

2.2.2 Segmentation sémantique 3D : Le défi de l’éparsité des données LiDAR

Complémentaire à la richesse texturale de la caméra, le LiDAR offre une géométrie précise mais complexe : un nuage de points brut (vecteurs continus, non ordonnés et de densité variable) incompatible avec les convolutions matricielles. La littérature propose trois approches de traitement.

Premièrement, la projection 2D (*Range-View* ou *Bird’s Eye View*, ex. RangeNet++ [28]) permet d’appliquer des CNN rapides. Elle induit cependant des distorsions masquant potentiellement des obstacles géométriques critiques pour le rover.

Deuxièmement, la discrétisation volumétrique (VoxelNet, Cylinder3D [29]) préserve la topologie, mais son coût mémoire croît cubiquement. Appliquée aux vastes espaces vides des environnements planétaires, elle s’avère computationnellement inefficace.

Troisièmement, les approches *point-based* traitent le nuage brut directement. L’architecture pionnière PointNet++ [30] (issue de PointNet [31]) agrège hiérarchiquement des caractéristiques via des perceptrons multicouches sur des voisinages sphériques. Invariante aux permutations et robuste aux densités variables, cette méthode offre une géométrie fidèle pour une empreinte mémoire maîtrisée, justifiant notre choix pour la branche 3D.

2.2.3 Stratégies de fusion de capteurs et limitations actuelles

Aucune modalité n’étant parfaite (sensibilité à l’illumination pour la caméra, absence de couleur pour le LiDAR), la fiabilité repose sur leur fusion. La fusion précoce (*Early Fusion*) concatène les données brutes avant classification, comme dans PointPainting [32]. Bien que riche en informations, elle reste très vulnérable aux dérives de calibration extrinsèque causées par les vibrations. Pour la sécurité critique, la fusion tardive (*Late Fusion*) offre une modularité privilégiée [8] : chaque réseau prédisant indépendamment, le système survit à l’obstruction d’un capteur (par exemple à cause de la présence de poussière sur la lentille). Cependant, fusionner par moyenne pondérée suppose que les scores de confiance reflètent fidèlement la probabilité de justesse. Or, les réseaux profonds souffrant systématiquement de sur-confiance, ces probabilités brutes demeurent inexploitable pour une décision sécuritaire sans une étape préalable de calibration.

2.3 Quantification de l’incertitude en apprentissage profond

L’intégration des réseaux profonds dans les systèmes critiques se heurte à leur incapacité structurelle à exprimer le doute. S’ils excellent en prédiction pure, ils évaluent mal leur propre fiabilité. Cette section passe en revue les méthodes permettant de quantifier cette incertitude, des approches Bayésiennes aux techniques de calibration et de prédiction conforme.

2.3.1 Le problème de la sur-confiance et taxonomie de l’incertitude

La majorité des modèles de segmentation sémantique utilisent la fonction Softmax en dernière couche pour transformer les logits en une distribution de probabilité normalisée sur les classes. Cependant, il est désormais bien établi dans la littérature que ces probabilités brutes ne reflètent pas la probabilité réelle de justesse, un phénomène connu sous le nom de désalignement ou de mauvaise calibration [6]. En raison de la minimisation de la perte d’entropie croisée durant l’entraînement, les réseaux modernes tendent vers la sur-confiance (*overconfidence*), assignant parfois des scores de probabilité proches de l’unité même lorsqu’ils produisent des prédictions erronées [33, 34].

Ce phénomène s’aggrave drastiquement face à des données hors-distribution (OOD). Lorsqu’un modèle entraîné sur des images standards est confronté à une anomalie, tel qu’un débris métallique ou une formation géologique inédite sur le sol lunaire, il projette cette entrée inconnue sur l’espace des classes familières avec une certitude souvent trompeuse, faute de mécanisme pour rejeter l’entrée [7]. Pour sécuriser la navigation, il est impératif de décomposer cette incertitude en deux composantes distinctes, formalisées par Kendall et Gal [35]. D’une part, l’incertitude aléatoire capture le bruit stochastique irréductible inhérent aux données, tel que le flou cinétique, le bruit thermique du capteur ou les réflexions spéculaires du LiDAR. D’autre part, l’incertitude épistémique, ou incertitude du modèle, mesure l’ignorance du réseau face à des données situées loin de la variété d’apprentissage. C’est cette seconde composante qui est critique pour l’exploration planétaire, car elle seule permet de détecter les anomalies nécessitant une intervention humaine ou un arrêt d’urgence.

2.3.2 Approches Bayésiennes et approximations variationnelles

Le cadre théorique le plus rigoureux pour capturer l’incertitude épistémique est celui des réseaux de neurones Bayésiens (BNN). Contrairement aux réseaux déterministes qui optimisent des poids scalaires fixes, un BNN apprend une distribution de probabilité postérieure sur chaque poids du réseau [36]. L’inférence ne produit alors plus une prédiction unique, mais une distribution prédictive marginalisée sur l’ensemble des paramètres possibles. Bien que théoriquement idéale, l’inférence exacte dans les BNNs est incalculable pour les architectures profondes modernes contenant des millions de paramètres, en raison de la complexité de l’intégration sur l’espace des poids [37].

Pour contourner cette complexité calculatoire, Gal et Ghahramani [38] ont proposé l’utilisation du *Monte Carlo Dropout* (MC-Dropout) comme une approximation variationnelle des BNNs. En maintenant l’opérateur de désactivation (*dropout*) actif lors de l’inférence et en réalisant plusieurs passes stochastiques pour une même entrée, on obtient un échantillonnage de la distribution prédictive dont la variance fournit une estimation directe de l’incertitude épistémique. De par sa simplicité d’implémentation sans aucune modification de l’architecture d’entraînement, le MC-Dropout est de plus en plus utilisé *de facto* en robotique mobile.

2.3.3 Calibration post-hoc et prédiction conforme

Plus tard, une nouvelle famille d’approches s’est développée autour de la calibration *post-hoc*. L’objectif est alors de recalibrer les sorties du modèle sur un ensemble de calibration pour qu’elles correspondent à la probabilité empirique de justesse. La méthode du *Temperature Scaling* (TS), qui consiste à diviser les logits par un scalaire optimisé, est efficace pour corriger

l’incertitude aléatoire sur des données familières mais échoue souvent à détecter les données hors-distribution, conservant une confiance élevée sur les anomalies [6].

C’est pour pallier cette lacune qu’émerge le paradigme de la prédiction conforme, théorisé par Vovk et al. [39] et popularisé récemment en vision par ordinateur par Angelopoulos et al. [40]. Contrairement aux méthodes précédentes qui fournissent une estimation ponctuelle de probabilité, la prédiction conforme construit un ensemble de prédiction qui garantit mathématiquement de contenir la vraie classe avec une probabilité spécifiée par l’utilisateur. Cette garantie de couverture marginale est libre de distribution, c’est-à-dire qu’elle reste valide quelle que soit la qualité du modèle sous-jacent ou la distribution des données, sous l’hypothèse d’échangeabilité.

Pour la segmentation sémantique, des méthodes récentes comme *Adaptive Prediction Sets* (APS) [41] ou la *Kandinsky Calibration* (KC) [42] étendent ce concept au niveau du pixel ou du point 3D. Elles utilisent un jeu de données de calibration pour calculer un seuil de confiance (quantile) qui assure que la couverture empirique atteint la cible définie en amont par l’utilisateur. Bien que prometteuses, ces méthodes ont été principalement validées sur des images médicales ou routières 2D. Leur application aux nuages de points LiDAR 3D, caractérisés par leur éparsité et leurs discontinuités géométriques, constitue un territoire inexploré que ce mémoire se propose d’investiguer.

2.4 De la perception à la planification robuste

La quantification de l’incertitude n’est pas une fin en soi. Pour naviguer de manière sûre, le rover doit projeter ces informations sémantiques et probabilistes dans une représentation spatiale persistante, exploitable par un planificateur pondérant le risque. Cette section examine l’intégration de cette incertitude perceptuelle dans les chaînes de décision robotiques.

2.4.1 Cartographie probabiliste et cartes de risque

La représentation canonique pour la navigation en robotique mobile demeure la grille d’occupation (*Occupancy Grid*), popularisée par les travaux fondateurs de Alberto Elfes [43]. Dans ce formalisme, l’environnement est discrétisé en cellules indépendantes, chacune stockant la probabilité postérieure d’être occupée par un obstacle physique, mise à jour récursivement via un filtre bayésien binaire. Cependant, pour la navigation en milieu naturel complexe, la simple binarité libre ou occupé s’avère insuffisante car elle ne capture pas la gradation de la traversabilité. Les approches modernes étendent ce concept vers des cartes de traversabilité sémantiques, où chaque cellule contient un vecteur de probabilités sur les classes de terrain.

Pour la représentation 3D, le formalisme OctoMap [44] utilisant des arbres octaux (*octrees*) permet de modéliser des environnements volumétriques en optimisant l’occupation mémoire.

Le défi contemporain réside dans l’intégration explicite de l’incertitude épistémique dans ces cartes. Une approche naïve consistant à moyenner les vecteurs de probabilité softmax successifs tend à lisser les pics d’incertitude temporels, masquant potentiellement les dangers. Des travaux récents en robotique de terrain, tels que ceux de Fan et al. [45], suggèrent de maintenir une couche de carte de risque distincte. Dans cette approche, la valeur de coût d’une cellule n’est pas uniquement fonction de la classe de terrain dominante, mais est pénalisée par la variance de la prédiction du modèle ou l’entropie de la distribution. Ainsi, une zone que le système de perception ne parvient pas à classer avec certitude devient artificiellement coûteuse, agissant comme une barrière virtuelle pour le planificateur.

2.4.2 Planification de trajectoire consciente du risque

Une fois la carte de risque établie, le planificateur doit générer une trajectoire optimale reliant la pose actuelle à l’objectif. Les algorithmes de recherche de chemin basés sur l’échantillonnage, tels que RRT* (*Rapidly-exploring Random Tree Star*) introduit par Karaman et al. [46], sont utilisés pour leur capacité à trouver des solutions asymptotiquement optimales dans des espaces de configuration continus. Si leur formulation classique minimise généralement la distance euclidienne ou l’énergie, ces méthodes permettent structurellement l’intégration d’une carte de risque. L’incertitude ou la probabilité de collision peuvent en effet être incorporées directement au sein de la fonction de coût comme termes de pénalité. Dans cette configuration, le coût cumulé agit comme un proxy du risque spatial, modifiant la métrique d’optimisation pour contraindre l’expansion de l’arbre de recherche vers les régions évaluées comme sûres.

Pour garantir la sécurité du rover face à une perception incertaine, il est nécessaire d’adopter une formulation de planification robuste, souvent qualifiée de *Chance-Constrained Planning* (planification sous contraintes de chances). Dans ce cadre, théorisé par Blackmore et al. [47], le problème d’optimisation est formulé de sorte que la probabilité de collision ou d’échec le long de la trajectoire reste inférieure à un seuil défini par l’utilisateur. Bien que rigoureuse, la résolution exacte de ces contraintes est souvent trop coûteuse en temps de calcul pour une replanification en ligne. Une alternative pragmatique pour l’embarqué consiste à adopter une approche consciente du risque, où le coût de traversée d’une cellule est gonflé par un facteur proportionnel à l’incertitude épistémique locale. Cette méthode force l’algorithme A* ou RRT* à contourner les zones d’ambiguïté sémantique, préférant un chemin plus long mais dont la traversabilité est certaine, une stratégie que nous implémentons dans le cadre de ce mémoire via les cartes de risque conformes.

2.5 Synthèse et positionnement de la thèse

L’analyse de la littérature scientifique menée dans ce chapitre permet de dégager un consensus clair : la sécurité des rovers planétaires ne peut plus reposer uniquement sur l’évitement d’obstacles géométriques. Elle exige désormais une compréhension sémantique fine et multimodale de l’environnement. Si certaines briques technologiques de pointe existent individuellement, leur intégration au sein d’une architecture de navigation unifiée révèle trois verrous scientifiques majeurs que ce mémoire se propose de lever.

Premièrement, il existe un fossé d’application concernant la quantification de l’incertitude. Les méthodes les plus avancées restent à ce jour confinées à l’analyse d’images 2D, principalement dans les domaines de l’imagerie médicale ou de la conduite autonome urbaine. Leur extension théorique et pratique aux données LiDAR 3D éparses et non-structurées, pourtant cruciales pour appréhender la géométrie du terrain planétaire, constitue un territoire scientifique inexploré. Pour combler ce vide, ce mémoire propose une architecture symétrique novatrice, appliquant rigoureusement ces méthodes conformes de bout en bout, tant sur l’imagerie 2D (via SegFormer) que sur les nuages de points 3D (via PointNet++).

Deuxièmement, nous identifions une dichotomie problématique dans les stratégies d’estimation de l’incertitude. D’un côté, les méthodes Bayésiennes pures offrent d’excellentes performances face aux données hors-distribution, mais s’avèrent bien trop coûteuses en mémoire et en temps de calcul pour les processeurs spatiaux. De l’autre, les méthodes de calibration *post-hoc* sont légères mais fondamentalement incapables de gérer les anomalies. La littérature manque d’une approche intermédiaire. Ce projet résout ce dilemme en concevant un pipeline hybride : il fusionne la légèreté d’une approximation variationnelle avec une calibration paramétrique globale, le tout encapsulé dans le cadre de la prédiction conforme pour offrir des garanties statistiques formelles compatibles avec les contraintes temps-réel de l’embarqué.

Troisièmement, la quantification de l’incertitude est trop souvent traitée comme une tâche de diagnostic perceptuel isolée, déconnectée de la boucle de commande du robot. Le lien opérationnel entre une métrique d’incertitude évaluée au niveau du pixel et la décision de navigation finale au niveau de la carte globale est rarement explicité. Ce mémoire construit précisément ce pont méthodologique. Il démontre comment l’incertitude est transmutée en cartes de risque dynamiques, directement exploitées par l’algorithme probabiliste AURA*, développé spécifiquement pour ces travaux. En substituant la rigidité de l’inflation géométrique classique par une optimisation continue sur les gradients de variance épistémique, AURA* dote le système d’une navigation préventive, garantissant l’intégrité physique et l’autonomie totale exigées par l’exploration spatiale.

CHAPITRE 3 PIPELINE DE PERCEPTION MULTIMODALE

L’élaboration d’une architecture de navigation autonome intégrant la quantification de l’incertitude requiert, en amont de tout développement algorithmique, un environnement de validation rigoureux. Les méthodes d’apprentissage profond et de fusion probabiliste proposées dans ce manuscrit exigent un accès synchronisé à de multiples modalités sensorielles (vision dense, géométrie LiDAR, données inertielles), ainsi qu’à une vérité terrain absolue pour la topologie sémantique. Les environnements extraterrestres réels n’offrant pas de tels jeux de données annotés à haute résolution, le recours à une simulation physiquement réaliste s’impose. Ce chapitre détaille par conséquent le socle empirique de notre étude : l’acquisition, la synchronisation temporelle et la calibration spatiale des données d’entrée qui alimenteront l’intégralité du pipeline de perception et de planification.

3.1 Le jeu de données LuSNAR

Le jeu de données LuSNAR (*Lunar Segmentation, Navigation and Reconstruction*) est un ensemble simulé conçu pour l’évaluation de pipelines complets de navigation autonome en environnement lunaire [48]. Il fournit des données synchronisées issues de capteurs standards simulés et a été conçu pour offrir un benchmark permettant la validation de modules de perception, de localisation, de reconstruction 3D et de planification. Contrairement aux benchmarks centrés sur une seule tâche, LuSNAR propose une approche multi-modalités et multi-tâches adaptée aux besoins de la perception robotique moderne.

3.1.1 Contenu du jeu de données

Chaque scène de LuSNAR est une simulation réaliste basée sur des données altimétriques lunaires, générée à l’aide d’un moteur graphique haute fidélité. Un rover virtuel, équipé de capteurs simulés, évolue dans ces environnements. Pour chaque identifiant temporel (timestamp), plusieurs modalités sont synchronisées. Des images stéréo RGB sont capturées par deux caméras, avec une carte de profondeur dense et un masque sémantique 2D projeté. Un nuage de points LiDAR simulé à 360°, bruité de manière réaliste, est généré à chaque étape, chaque point étant annoté par une classe sémantique 3D. La pose du rover est enregistrée sous la forme $(x, y, z, q_x, q_y, q_z, q_w)$, et les mesures IMU incluent accélérations et vitesses angulaires dans le repère du rover. Les modalités sont synchronisées via un timestamp unique.

3.1.2 Objectifs et structure globale

Conçu pour répondre aux défis de l’exploration autonome en terrain lunaire, LuSNAR vise à fournir un cadre réaliste, reproductible et complet pour l’entraînement et l’évaluation de modèles de segmentation sémantique (2D/3D), de localisation SLAM (*Simultaneous Localization and Mapping*) visuel ou LiDAR, de reconstruction 3D et de planification. Neuf scènes distinctes ont été simulées dans Unreal Engine, variant systématiquement selon deux axes : relief topographique et densité d’objets. Cette diversité permet de tester la robustesse des algorithmes dans des environnements à complexité croissante.

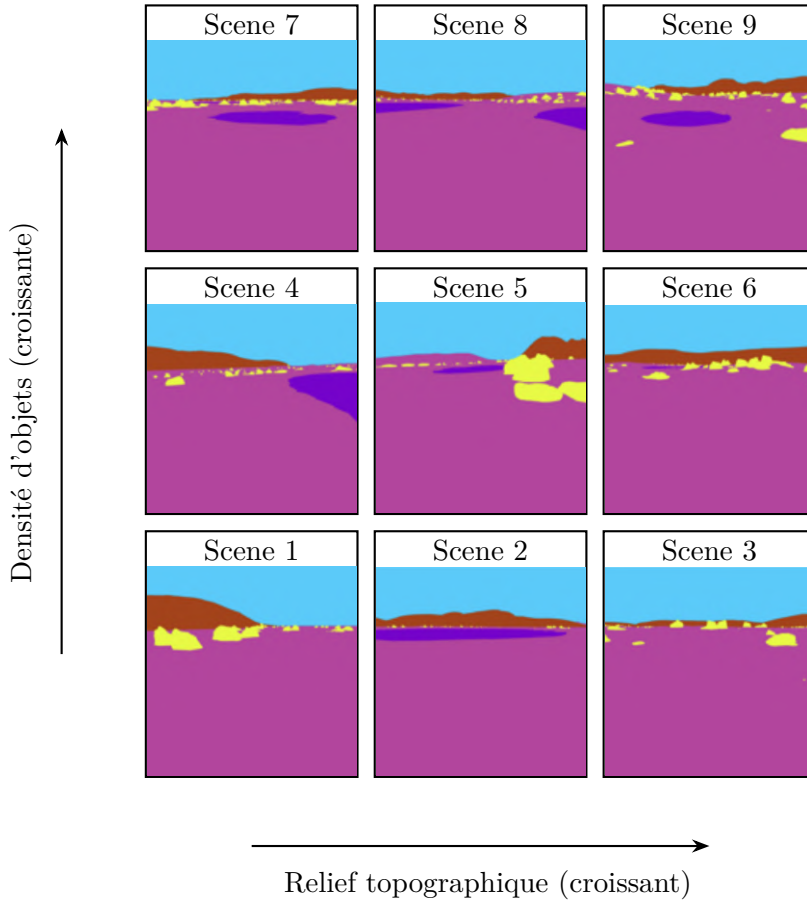


FIGURE 3.1 Classification des scènes LuSNAR par topographie et densité.

3.1.3 Validation expérimentale du jeu de données

La pertinence du jeu de données a été évaluée par ses auteurs à travers une série d’expérimentations sur trois tâches [48]. Pour la segmentation 2D, SegFormer atteint une précision moyenne (mAcc) de 94.22% et un mIoU de 89.96%, surpassant DeepLabV3+ et UNet. Pour

la localisation, des méthodes de SLAM visuel (*ORB-SLAM3*) et LiDAR (*HDL Graph SLAM*) ont été testées sur des scènes variées. Les trajectoires reconstruites sont comparées aux vérités terrain, validant la fiabilité du benchmark pour l'évaluation des algorithmes de pose. La reconstruction 3D a été évaluée via COLMAP et OpenMVS, appliquée aux paires d'images. Les résultats montrent une fidélité géométrique suffisante pour les tâches de modélisation dense sur terrains lunaires simulés.

3.2 Synchronisation et calibration multimodale des données

3.2.1 Synchronisation spatiale

Calibration extrinsèque

Les coordonnées des points LiDAR sont initialement exprimées dans le repère du capteur, aligné avec celui du rover. Afin de permettre leur projection dans l'image, une transformation extrinsèque est appliquée pour exprimer les points dans le repère de la caméra gauche.

Dans le pipeline utilisé, cette opération est codée sous la forme d'un permutateur d'axes :

$$x_{\text{cam}} = y_{\text{lidar}} \quad y_{\text{cam}} = z_{\text{lidar}} \quad z_{\text{cam}} = x_{\text{lidar}}$$

Cette permutation correspond à une transformation rigide $T_{\text{cam} \leftarrow \text{LiDAR}}$ fixe, extraite du simulateur. Aucune rotation ou translation additionnelle n'est appliquée, ce qui suppose une calibration parfaite et une co-localisation physique des capteurs.

Projection des points LiDAR dans le plan image

La projection des points 3D dans le plan image permet de transférer les annotations 2D vers les données LiDAR, afin d'attribuer un label sémantique à chaque point projetable. Cette étape est essentielle pour exploiter la supervision dense des masques image dans l'apprentissage ou l'évaluation de modèles 3D.

Les points, exprimés dans le repère de la caméra, sont projetés via la matrice intrinsèque K :

$$K = \begin{bmatrix} f_x & 0 & c_x \\ 0 & f_y & c_y \\ 0 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 680 & 0 & 512 \\ 0 & 680 & 288 \\ 0 & 0 & 1 \end{bmatrix}$$

où (f_x, f_y) sont les focales en pixels et (c_x, c_y) les coordonnées du centre optique. Ces paramètres ont été estimés empiriquement pour assurer un alignement visuel correct.

Soit $\mathbf{P}_{\text{cam}} = (x, y, z)^\top$ un point 3D dans le repère caméra, sa projection dans le plan image s'effectue via :

$$\mathbf{P}_{\text{image}} = \pi(K \cdot \mathbf{P}_{\text{cam}}), \quad \text{avec} \quad \pi([u, v, w]^\top) = \left(\frac{u}{w}, \frac{v}{w} \right)$$

Les points sont projetés uniquement si leur coordonnée z dans le repère caméra est strictement positive ($z > 0$), garantissant qu'ils se trouvent devant le plan image.

Après projection, les coordonnées continues (u, v) sont discrétisées en indices d'image (u_i, v_i) puis validées : un point n'est conservé que si (u_i, v_i) est contenu dans le domaine du masque sémantique (typiquement $0 \leq u_i < 1024$ et $0 \leq v_i < 1024$). Une fois cette condition satisfaite, l'étiquette associée est extraite depuis le masque image :

$$\text{lidar_label}[i] = \text{image_label}[v_i, u_i]$$

Cette correspondance repose sur la géométrie projective définie par la calibration. Elle permet d'attribuer à chaque point LiDAR un label de classe dérivé du plan image, en exploitant la supervision dense des annotations 2D.

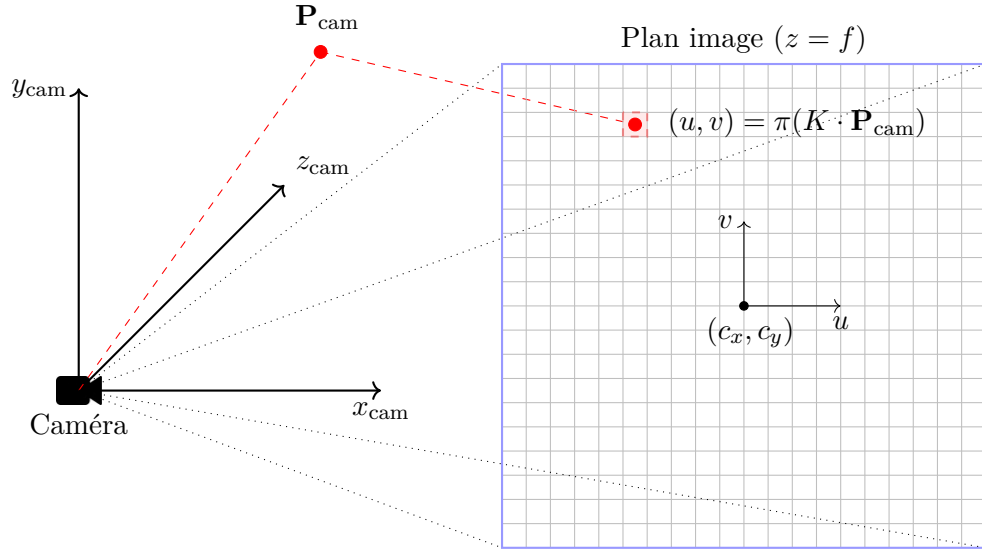


FIGURE 3.2 Projection d'un point 3D vers le plan image via la matrice intrinsèque.

3.2.2 Synchronisation temporelle

Les données extraites du jeu LuSNAR, nécessaires aux modules de perception, fusion et planification, sont agrégées au sein d'archives compressées unifiées au format **.npz** (figure 3.3). La synchronisation s'opère autour d'un timestamp maître t (en ns) dicté par l'image RGB gauche. Les modalités à haute fréquence (IMU, pose) y sont associées par minimisation de

l'écart temporel absolu, sans seuil de tolérance explicite, la fréquence d'échantillonnage étant supposée suffisante pour garantir la précision de l'appariement. Enfin, la dimension N désigne le nombre de points LiDAR, inhérent à la densité du capteur et à la topologie de la scène.

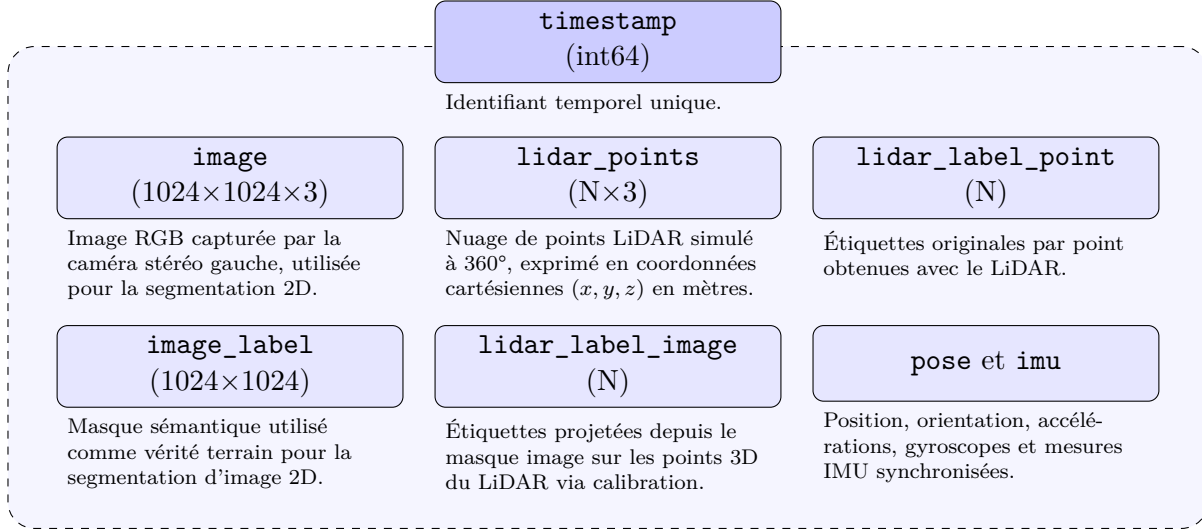


FIGURE 3.3 Structure interne d'un fichier `.npz` utilisé dans ce projet.

Dans le cadre de LuSNAR, l'acquisition du nuage de points LiDAR est modélisée comme instantanée à chaque *timestamp*. Contrairement à un capteur physique rotatif où le mouvement continu du rover pendant le balayage induit une distorsion spatiale, tous les points d'une même trame simulée partagent strictement le même instantané temporel. Par conséquent, aucune étape de *deskewing* (compensation du mouvement propre) n'est requise. De plus, il est formellement postulé que la chaîne d'acquisition virtuelle garantit une synchronisation parfaite et sans latence entre les balayages LiDAR, les trames visuelles et les mesures inertielles (IMU). L'association temporelle au plus proche voisin décrite précédemment est donc considérée comme exacte et exempte de dérive d'horloge.

3.2.3 Visualisation des données

Les masques sémantiques associés aux images encodent cinq classes distinctes dans l'espace image, identifiées par des entiers de 0 à 4, correspondant respectivement à : régolithe, cratère, rocher, montagne et ciel. En revanche, les nuages de points LiDAR n'emploient que trois catégories : régolithe (label -1), cratère (label 0), et rocher (label 174). Cette asymétrie reflète les limitations du capteur simulé (par exemple, le ciel n'est pas mesuré par le LiDAR). Lors de la projection des points LiDAR dans le plan image, un label image est transféré à chaque point visible, selon la géométrie de calibration définie précédemment.

Afin de valider visuellement la cohérence des modalités, plusieurs représentations synchronisées ont été générées pour un fichier `.npz` donné. Elles incluent la visualisation de l'image RGB (figure 3.4a), du masque sémantique 2D (figure 3.4b), de la superposition des deux (figure 3.4c), de la projection des points LiDAR sur l'image (figure 3.4d) et du nuage de points 3D coloré par classe sémantique (figure 3.4e). Le tableau de la figure 3.4f résume les couleurs associées par le jeu de données LuSNAR à chaque classe pour toutes les représentations [48].

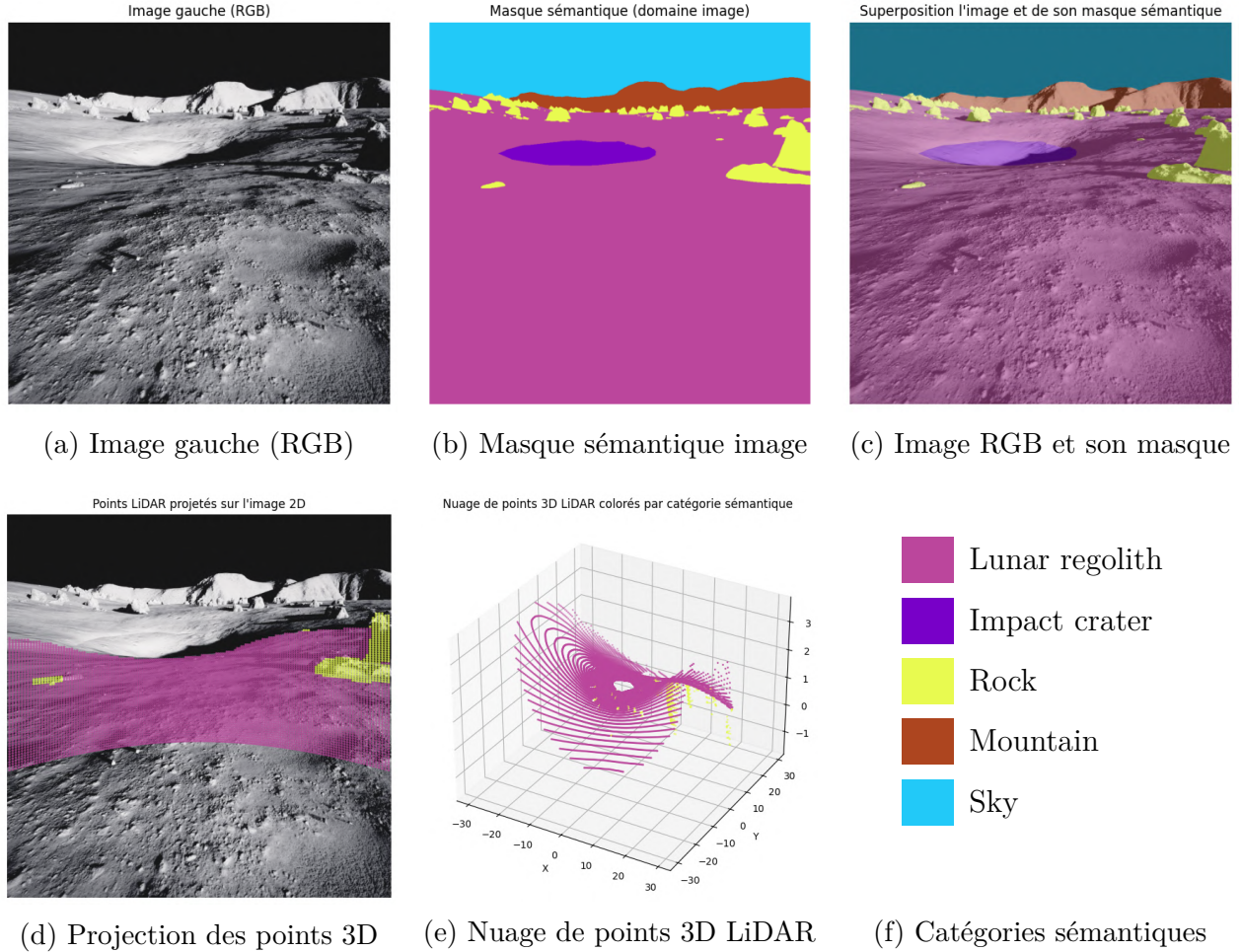


FIGURE 3.4 Visualisations synchronisées pour le timestamp $t = 1692601812033751296$.

Le jeu de données LuSNAR, riche en modalités synchronisées, a été exploité à travers un prétraitement rigoureux : filtrage temporel, projection géométrique, calibration extrinsèque, et structuration unifiée via des fichiers `.npz`. Les modalités résultantes sont prêtes à être exploitées pour entraîner des modèles de segmentation sémantique robustes, tant sur les images que sur les nuages de points. Ces prédictions alimenteront les étapes ultérieures d'estimation d'incertitude et de planification.

3.3 Segmentation des images pour la perception embarquée

3.3.1 Problématique et rôle dans le pipeline

La segmentation sémantique d’images RGB monoculaires assure une classification dense du terrain, fonction fondamentale pour la navigation autonome. Ses résultats sont exploités selon trois axes opérationnels. Premièrement, ils constituent la représentation sémantique primaire pour la compréhension environnementale, palliative à l’éparsité géométrique des nuages LiDAR. Deuxièmement, la projection spatiale des masques bidimensionnels vers les nuages de points 3D, régie par les paramètres extrinsèques de calibration (sous-section 3.2.1), permet une supervision inter-modale renforçant la cohérence perceptuelle. Troisièmement, ces prédictions initialisent les mécanismes probabilistes d’estimation de l’incertitude épistémique.

Les ambiguïtés topographiques et les dégradations lumineuses lunaires imposent une quantification rigoureuse de la confiance du modèle avant la propagation des cartes sémantiques vers le module de planification décisionnelle.

3.3.2 Architecture du modèle de segmentation 2D

Les modèles convolutionnels classiques (UNet [21], DeepLabV3+ [49]) présentent des limitations structurelles face aux scènes lunaires haute résolution, contraints par des champs récepteurs fixes et une difficulté à extraire le contexte spatial global [48]. L’architecture SegFormer [27], fondée sur le paradigme des *Vision Transformers*, s’y substitue en associant un encodeur hiérarchique et un décodeur perceptron multicouche (MLP) allégé. La variante *SegFormer-B0* a été retenue en vertu de son compromis optimal entre précision sémantique et coût computationnel, autorisant une inférence temps réel compatible avec les ressources embarquées critiques.

L’encodeur MiT-B0 [27] extrait séquentiellement quatre cartes de caractéristiques multi-échelles ($H/4, H/8, H/16, H/32$) par le biais d’embeddings de patches chevauchants et d’opérations d’auto-attention. Les encodages positionnels standard sont remplacés par un module *Mix-FFN* intégrant des convolutions séparables (3×3). Cette substitution dote le réseau d’une invariance native aux dimensions d’entrée et abolit les artefacts liés à l’interpolation spatiale des encodages. Le décodeur unifie ces cartes par projection linéaire vers une dimension C commune, suivie d’une concaténation axiale et d’une classification dense. Pour une entrée normalisée de 512×512 pixels, la carte prédictive finale de résolution 102×102 discrimine cinq classes géologiques : régolithe, cratère, roche, montagne et ciel.

3.3.3 Évaluation quantitative

L'évaluation quantitative est effectuée sur un ensemble de validation en utilisant la moyenne de l'*Intersection over Union* (IoU) comme métrique principale. Une matrice de confusion est construite en agrégeant les prédictions et les étiquettes réelles sur tous les pixels du jeu de validation. Les éléments diagonaux correspondent aux classifications correctes, tandis que les valeurs hors diagonale indiquent les confusions entre classes. Pour chaque classe k , on a :

$$\text{IoU}_k = \frac{TP_k}{TP_k + FP_k + FN_k}$$

où TP_k désigne le nombre de vrais positifs, FP_k les faux positifs et FN_k les faux négatifs. La moyenne (mIoU) est ensuite calculée en moyennant sur les cinq classes :

$$\text{mIoU} = \frac{1}{5} \sum_{k=1}^5 \text{IoU}_k$$

Le tableau 3.1 présente la matrice de confusion normalisée obtenue sur l'ensemble de validation. Chaque ligne correspond à une classe réelle et chaque colonne à une classe prédite. Les valeurs sont normalisées par ligne.

TABLEAU 3.1 Matrice de confusion normalisée et IoU par classe

		Labels prédits					IoU
		Régolithe	Cratère	Roche	Montagne	Ciel	
Vrais labels	Régolithe	0,94	0,03	0,02	0,01	0,00	0,9539
	Cratère	0,02	0,89	0,06	0,03	0,00	0,9374
	Roche	0,01	0,05	0,87	0,06	0,01	0,9029
	Montagne	0,00	0,01	0,04	0,90	0,05	0,9728
	Ciel	0,00	0,00	0,01	0,06	0,93	0,9909
mIoU							0,9516

Sur l'ensemble de validation, notre modèle atteint une mIoU de 0,9516. Les scores par classe sont élevés pour toutes les catégories : *ciel* (0,9909), *montagne* (0,9728), *régolithe* (0,9539) et *cratère* (0,9374). La classe *roche* obtient la valeur la plus faible (0,9029), probablement en raison de sa proximité fréquente avec les cratères et le régolithe, et des ambiguïtés locales de texture. Ces résultats traduisent une excellente généralisation du modèle aux différents types de terrains dans des conditions lunaires variables.

3.3.4 Évaluation qualitative

La figure 3.5 et la figure 3.6 illustrent des comparaisons qualitatives entre les masques prédits et les annotations de référence. La première image montre la scène complète, permettant de vérifier la cohérence globale de la segmentation. La seconde se concentre sur une région centrale pour observer le comportement aux frontières fines. Les contours sont globalement bien alignés, avec des erreurs localisées le long des bordures rocheuses, où une légère sur-estimation de la classe *roche* est observée. Ces artefacts apparaissent principalement dans les zones texturées ou partiellement occultées. Les pixels mal prédits sont surlignés en rouge pour faciliter l'analyse des erreurs locales.

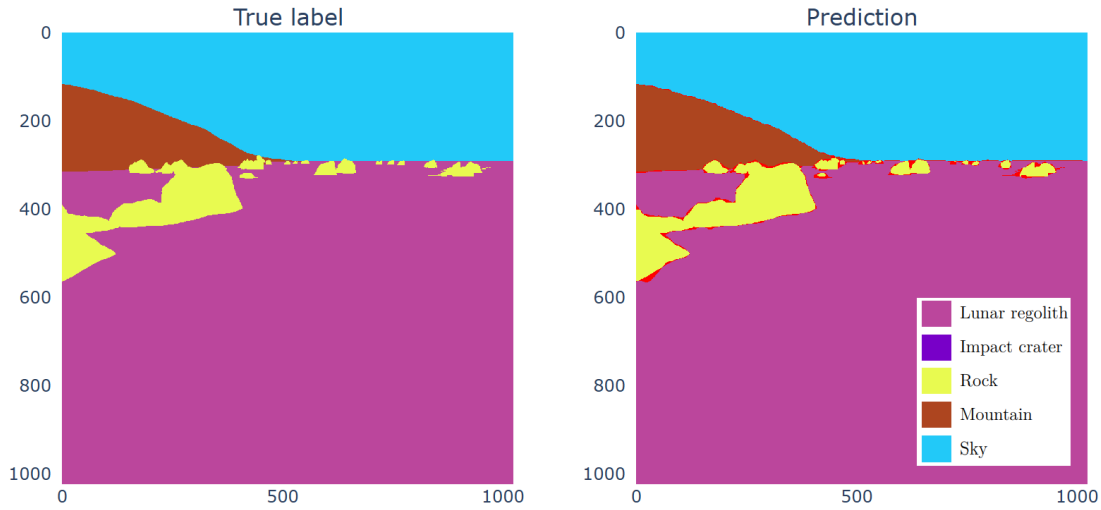


FIGURE 3.5 Comparaison qualitative sur une scène complète.

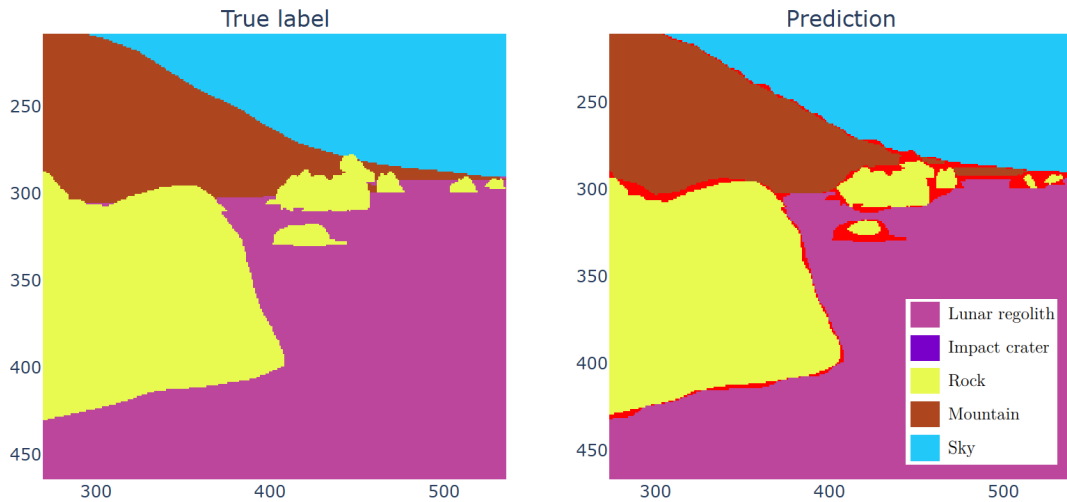


FIGURE 3.6 Zoom sur la région centrale pour une meilleure visibilité des erreurs.

3.4 Segmentation des points LiDAR pour la modélisation de terrain

3.4.1 Problématique et rôle dans le pipeline

La segmentation sémantique 3D des nuages de points LiDAR complète le module de segmentation 2D décrit au chapitre 3. Les prédictions 3D exploitent directement la géométrie spatiale capturée par le capteur LiDAR à 360°, ce qui les rend intrinsèquement robustes aux artefacts visuels tels que les ombres portées et la saturation optique.

Pour cela, chaque point $p_i = \{x_i, y_i, z_i\} \in \mathbb{R}^3$ est classifié dans l'une des trois catégories retenues : *régolithe*, *cratère* ou *roche*. Ces étiquettes sont dérivés des masques sémantiques 2D via une projection géométrique (voir sous-section 3.2.1).

L'objectif est d'apprendre une fonction de classification point par point $f_\theta : \mathbb{R}^{N \times 3} \rightarrow \{0, 1, 2\}^N$ qui prédit une étiquette sémantique pour chaque point du nuage en utilisant uniquement ses coordonnées spatiales. Ces prédictions sont ensuite intégrées dans le pipeline de navigation en aval, contribuant à la construction des cartes de risque, à la fusion multimodale et à la robustesse face aux conditions visuelles dégradées.

Le modèle implémenté se rapproche de PointNet++, une architecture hiérarchique conçue spécifiquement pour le traitement de nuages de points irréguliers [50]. Elle combine l'abstraction de caractéristiques locales et le sur-échantillonnage pour capturer le contexte géométrique à plusieurs échelles, tout en maintenant une faible complexité computationnelle adaptée au déploiement embarqué.

3.4.2 Architecture du modèle de segmentation 3D

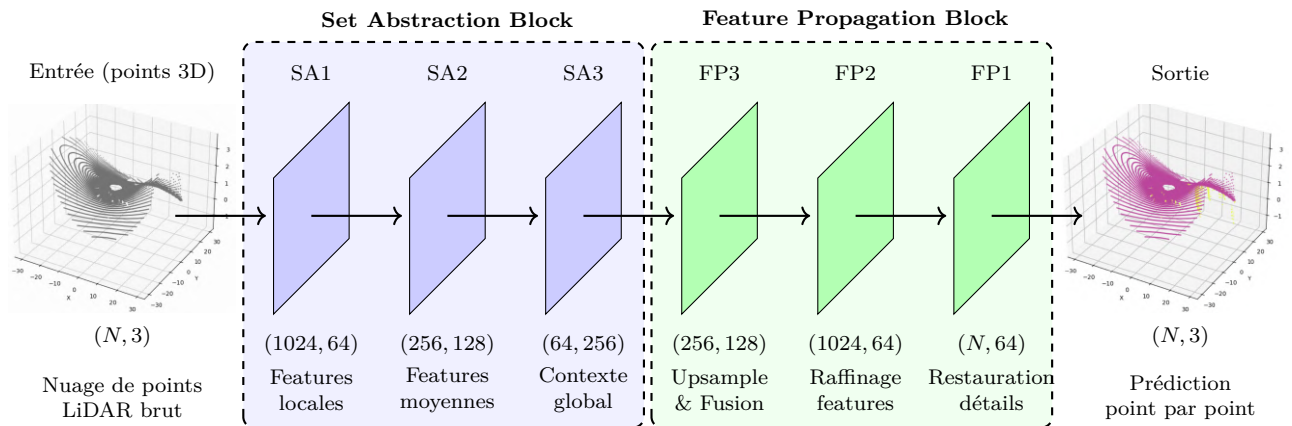


FIGURE 3.7 Architecture utilisée pour la segmentation sémantique point par point.

Le modèle implémenté (figure 3.7) s’appuie sur la variante *Single-Scale Grouping* (SSG) de PointNet++. Il alterne trois couches de *Set Abstraction* (SA) pour la compression géométrique et trois couches de *Feature Propagation* (FP) pour la reconstruction sémantique dense. Les modules SA opèrent par échantillonnage du point le plus éloigné (centroïdes), regroupement par requête sphérique (*ball query*) et extraction locale par perceptron multicouche (MLP). Les modules FP restaurent la résolution spatiale par interpolation inverse à la distance des trois plus proches voisins, concatènent les caractéristiques de niveau équivalent (connexion de saut) et appliquent un MLP de fusion.

3.4.3 Gestion du déséquilibre de classes et procédure d’entraînement

L’utilisation d’étiquettes 2D projetées (`lidar_label_image`) plutôt que d’annotations 3D natives garantit un alignement inter-modal strict, transférant la densité contextuelle visuelle vers le domaine géométrique épars. Après filtrage des points non étiquetés, un sous-échantillonnage aléatoire fixe le budget à 2048 points par scène. Cette dimension standardisée garantit la régularité tensorielle indispensable au traitement par lots de l’architecture neuronale, tout en établissant un compromis analytique entre la préservation de la résolution spatiale des obstacles et le respect des contraintes computationnelles imposées par l’inférence embarquée en temps réel. Les échantillons dégénérés (moins de 10 points valides ou mono-classe) sont exclus. L’asymétrie distributionnelle majeure en faveur du régolithe (tableau 3.2), inhérente à la topologie planétaire, requiert une optimisation par perte pondérée.

TABLEAU 3.2 Distribution des classes sur l’ensemble des points LiDAR

Classe	Label (3D)	Index	Nombre de points
Régolithe	−1	0	180 579 595
Cratère	0	1	2 913 990
Roche	174	2	24 663 332

L’optimisation des poids s’effectue par la minimisation d’une perte d’entropie croisée pondérée. Pour pallier le déséquilibre inter-classe, chaque classe k de fréquence n_k se voit attribuer un poids w_k inversement proportionnel à sa représentativité. Les poids normalisés résultants accroissent la contribution des classes minoritaires au gradient. Le modèle est entraîné via l’optimiseur Adam. La convergence est monitorée par la perte, la précision globale et la mIoU, avec un arrêt anticipé (*early stopping*) configuré sur une patience de 8 époques pour la mIoU de validation afin de prévenir le surapprentissage et conserver les poids optimaux.

3.4.4 Évaluation quantitative

L'évaluation sur l'ensemble de validation repose sur la précision par point et la mIoU. Cette dernière constitue une métrique robuste pour la segmentation LiDAR déséquilibrée en pénalisant la confusion inter-classe. La figure 3.8 illustre une convergence rapide : en 5 époques, la perte chute de 0,117 à 0,04 et la précision progresse de 95,1% à 98,1%. La stabilisation survient vers l'époque 10, avec une précision finale supérieure à 98,5%.

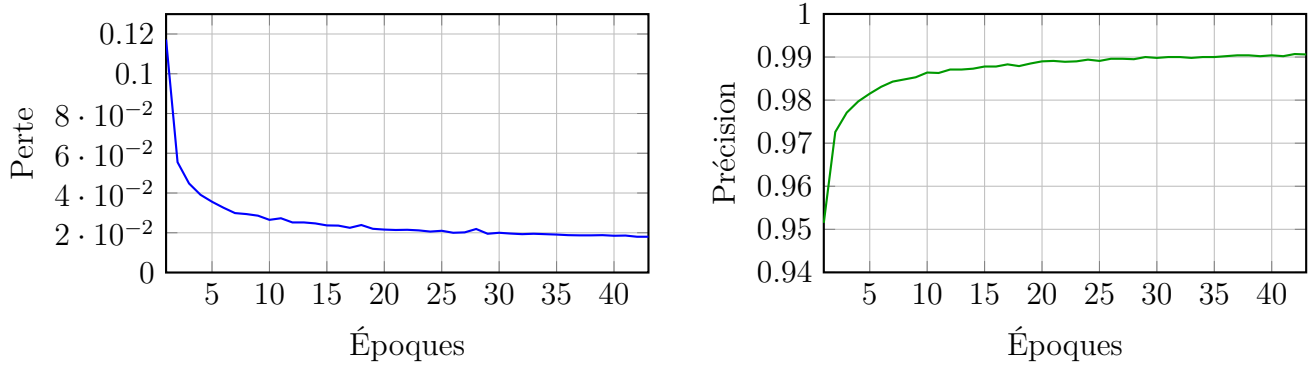


FIGURE 3.8 Perte moyenne (gauche) et précision moyenne (droite) selon les époques.

La figure 3.9 présente l'évolution des métriques IoU par classe et de la mIoU. Dès la première époque, le modèle atteint une performance solide sur la classe dominante *régolithe* (avec un IoU de 0,945), tandis que la classe *cratère* débute beaucoup plus bas. La classe *roche* bénéficie également d'un apprentissage rapide, atteignant 0,76 dès la première époque.

La mIoU progresse de manière cohérente et dépasse 0,80 dès l'époque 5. Alors que les classes *régolithe* et *roche* convergent rapidement vers des valeurs supérieures à 0,90, la classe *cratère* affiche une performance initiale significativement inférieure : son IoU fluctue et se stabilise autour de 0,82 dans les étapes ultérieures.

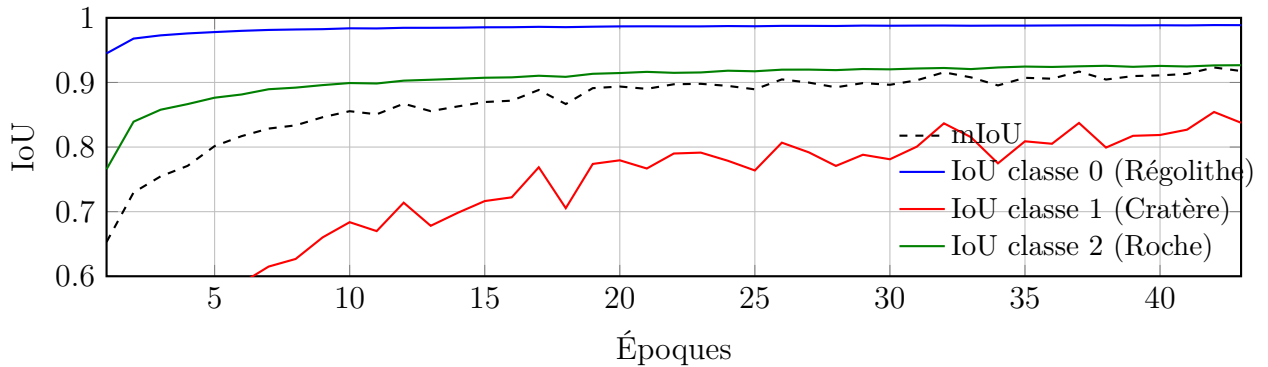


FIGURE 3.9 Évolution des IoU par classe et du mIoU durant l'entraînement

3.4.5 Évaluation qualitative

Pour évaluer la cohérence spatiale des prédictions, nous analysons un scan LiDAR à plusieurs niveaux de zoom. Cette évaluation qualitative mesure la robustesse du modèle aux changements de résolution et de champ de vision, couramment rencontrés lors de la navigation du rover.

Les quatre visualisations correspondent au même timestamp (1693274765297473024), projetées sur l'image RGB gauche en utilisant la transformation LiDAR-caméra calibrée développée en section 3.2.1. Chaque ligne compare la vérité terrain projetée (gauche) à la segmentation prédite (droite), avec un zoom de plus en plus serré de haut en bas.

En haut (vue large), le modèle capture avec précision la structure globale du terrain. À mesure que le zoom augmente, les contours fins des rochers sont préservés, montrant une cohérence géométrique avec le contexte visuel.

Dans la ligne inférieure (zoom maximal), la segmentation reste stable même sous un recadrage extrême. La plupart des divergences sont localisées le long des bords des objets ou dans les zones ombragées, où la projection des étiquettes est moins fiable.

Ces résultats confirment que notre modèle produit des prédictions cohérentes à travers les échelles spatiales, maintenant à la fois l'agencement sémantique global et les détails structuraux locaux.

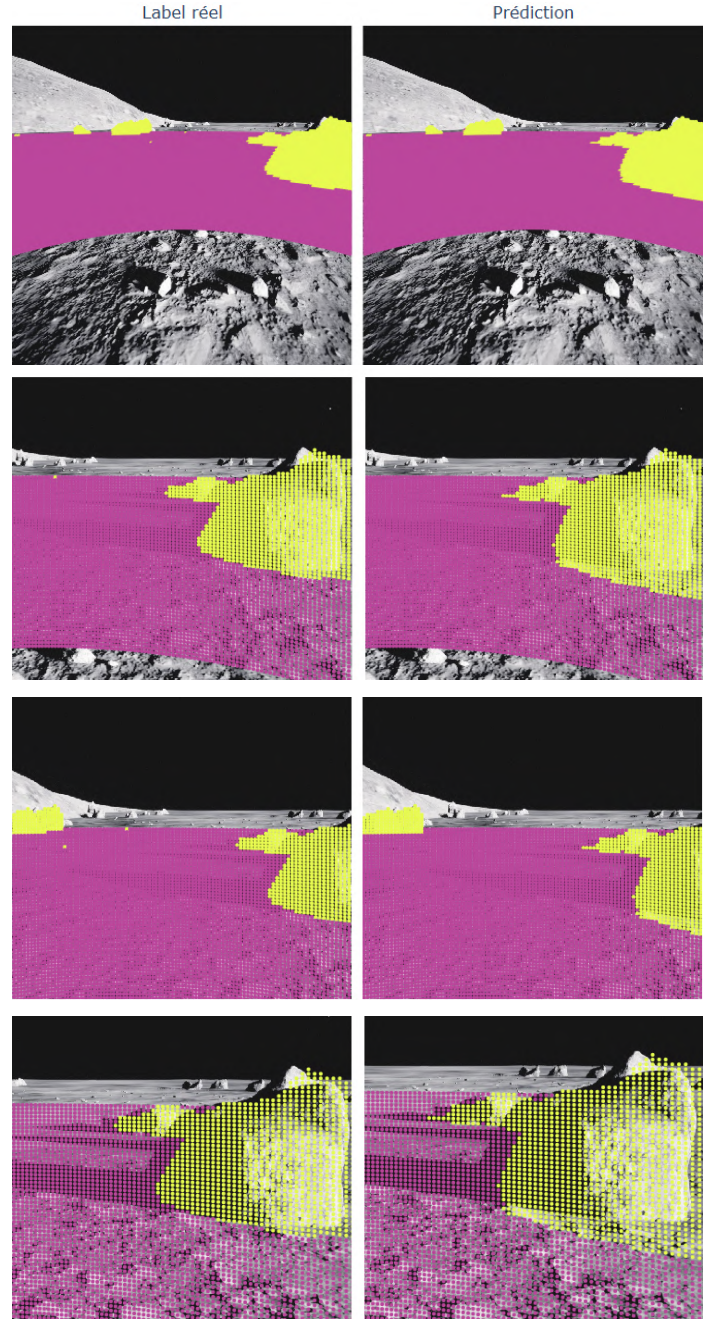


FIGURE 3.10 Projection multi-échelle des prédictions 3D à quatre niveaux de zoom.

CHAPITRE 4 CALIBRATION ET QUANTIFICATION DE L'INCERTITUDE

4.1 Introduction et motivation

La perception environnementale par apprentissage profond s'est imposée comme la pierre angulaire de la navigation robotique moderne. Si des architectures de segmentation sémantique (SegFormer en 2D, PointNet++ en 3D) atteignent aujourd'hui des niveaux de précision remarquables, la performance brute (mIoU) ne constitue pas une garantie de sécurité suffisante pour l'exploration planétaire. Ce chapitre démontre la nécessité impérieuse de transcender la simple prédiction sémantique pour la coupler à une quantification rigoureuse de sa fiabilité.

Pour un rover opérant sur Mars ou la Lune, l'incapacité d'un réseau de neurones à exprimer correctement son ignorance représente un danger critique. Ce phénomène de surconfiance [6] pousse les modèles à produire des prédictions erronées assorties de probabilités extrêmement élevées. Confondre une zone de sable meuble, véritable piège cinématique, avec un substrat rocheux stable, et ce avec une certitude de 99 %, peut sceller l'échec définitif d'une mission.

La fiabilité d'un système de perception exige donc qu'une probabilité prédite de 0,8 corresponde statistiquement à une fréquence de justesse de 80 %. Sans cette calibration stricte, les modules décisionnels en aval (planification de trajectoire, évitement d'obstacles) opèrent sur des prémisses faussées, compromettant irrémédiablement l'intégrité du rover.

Bien que la littérature propose diverses méthodes de quantification d'incertitude, celles-ci souffrent généralement d'une application fragmentée qui limite leur déploiement sur un système embarqué critique. Une première limitation concerne le coût computationnel : l'incertitude épistémique est souvent estimée par des méthodes d'ensemble (*Deep Ensembles*) qui nécessitent l'entraînement et l'inférence de multiples modèles en parallèle, une stratégie incompatible avec les ressources énergétiques et matérielles restreintes d'un rover spatial.

De surcroît, la calibration probabiliste, souvent réalisée par *Temperature Scaling*, est fréquemment traitée comme une étape de post-traitement isolée, déconnectée de la nature stochastique de l'inférence bayésienne. Enfin, et c'est un point critique pour la fusion sensorielle, les pipelines de perception traitent généralement les images 2D et les données LiDAR 3D via des mécanismes d'incertitude hétérogènes. Cette fragmentation méthodologique rend la fusion décisionnelle complexe et incohérente, car il devient périlleux de comparer le risque associé à un pixel d'image avec celui d'un point 3D si les métriques sous-jacentes ne sont pas statistiquement alignées.

4.2 Contributions et cadre unifié du pipeline d'incertitude

Pour pallier la fragmentation méthodologique et les contraintes matérielles identifiées, ce chapitre présente l'architecture unifiée d'estimation du risque (figure 4.1, cliquable). Plutôt que d'empiler des modules de post-traitement disparates, notre approche propose un cadre d'inférence strictement symétrique : elle applique des mécanismes de quantification identiques aux modalités 2D et 3D, garantissant une cohérence statistique lors de la fusion décisionnelle.

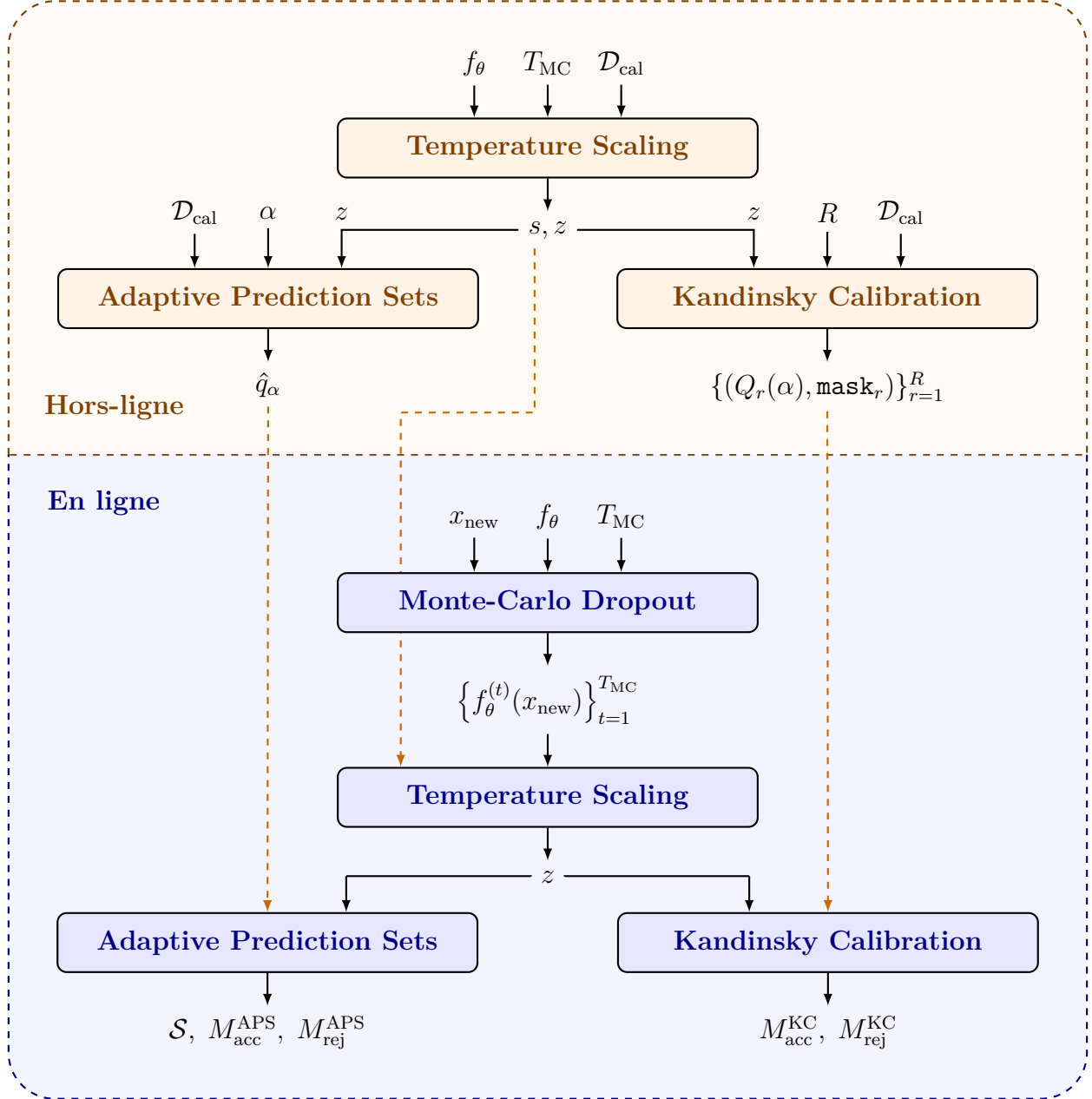


FIGURE 4.1 Pipeline unifié de calibration avec partage de paramètres.

Notre contribution s’articule autour de trois axes algorithmiques majeurs. En premier lieu, nous formalisons l’intégration conjointe de l’estimation bayésienne par *Monte-Carlo Dropout* et de la calibration paramétrique par *Temperature Scaling*. Cette synergie permet de capturer l’incertitude épistémique face aux données hors-distribution tout en lissant les probabilités issues de l’inférence stochastique. En second lieu, nous étendons les méthodes de prédiction conforme (*Adaptive Prediction Sets*) et de calibration régionale (*Kandinsky Calibration*) au contexte spécifique de la segmentation dense multimodale, offrant ainsi des garanties formelles de couverture de la vraie classe. Enfin, l’architecture repose sur une dichotomie temporelle stricte, dissociant l’optimisation lourde des paramètres (phase hors-ligne) de leur application réactive (phase en-ligne).

4.2.1 Phase hors-ligne : Apprentissage des paramètres de calibration

La phase d’étalonnage, présentée dans la partie supérieure (orange) de la figure 4.1, a pour objectif d’ajuster les propriétés probabilistes du modèle f_θ sur un ensemble de calibration dédié $\mathcal{D}_{\text{cal}}$, sans en altérer les poids synaptiques. Le traitement débute dans le module de *Temperature Scaling* (voir section 4.4), où T_{MC} passes stochastiques de *Monte-Carlo Dropout* (voir section 4.3) fournissent une estimation robuste des logits moyens pour chaque donnée. L’optimisation minimise ensuite la Log-Vraisemblance Négative (NLL) pour extraire une température scalaire s , lissant ainsi les distributions softmax trop piquées afin de produire des probabilités z mieux conditionnées.

À partir de ces probabilités calibrées, le pipeline déploie la prédiction conforme pour établir des garanties statistiques de couverture. Pour un risque cible α , le module *Adaptive Prediction Sets* (voir section 4.5) apprend un seuil critique $\hat{q}_\alpha$, appelé masse cumulative, qui garantit l’inclusion de la vraie classe dans l’ensemble de prédiction. Simultanément, la calibration de Kandinsky (voir section 4.6) segmente l’espace de confiance en R régions distinctes pour calculer des quantiles locaux $Q_r(\alpha)$, adaptant ainsi la sévérité du rejet à l’hétérogénéité spatiale de l’incertitude.

4.2.2 Phase en ligne : Inférence et cartographie des risques

La partie inférieure (bleue) de la figure 4.1 illustre le fonctionnement du système en conditions opérationnelles, lorsqu’une nouvelle donnée x_{new} est acquise par le rover. Le flux de traitement débute par l’estimation de l’incertitude via *Monte-Carlo Dropout*. L’entrée traverse le réseau via plusieurs passes stochastiques, cette procédure bayésienne permettant de capturer la variabilité du modèle face à des données inconnues ou hors distribution. Les logits résultants sont agrégés pour former une moyenne robuste.

L’injection des paramètres appris intervient ensuite pour transformer ces scores bruts en métriques de décision statistiquement certifiées, définies par une erreur de calibration minimisée et une garantie formelle de couverture marginale. Les logits moyens sont d’abord normalisés par le scalaire s , produisant une distribution de probabilité calibrée. Enfin, ces probabilités alimentent les modules de décision qui appliquent les seuils statistiques pré-calculés hors ligne. Le système génère alors l’ensemble prédictif $\mathcal{S}$ ainsi que des cartes binaires de conformité (M_{acc}) et de rejet (M_{rej}), offrant une granularité fine sur la fiabilité de la segmentation.

Cette séparation temporelle répond à un impératif d’efficacité computationnelle critique pour la navigation embarquée. La phase d’apprentissage exige des opérations lourdes, telles que la descente de gradient pour l’optimisation de s ou le tri de millions de scores pour l’estimation des quantiles conformes. En figeant ces paramètres hors-ligne, la phase d’inférence se réduit à des opérations arithmétiques élémentaires, rendant le déploiement de garanties statistiques robustes compatible avec les ressources limitées d’un processeur de rover spatial.

4.3 Monte-Carlo Dropout : Réduction de la surconfiance par échantillonnage

L’estimation de l’incertitude constitue une étape indispensable pour la fiabilité des systèmes autonomes. Dans ce travail, *Monte-Carlo Dropout* est adopté comme mécanisme bayésien léger, mais son rôle dépasse celui d’un simple estimateur d’incertitude : il sert de socle au cadre unifié proposé, garantissant une quantification cohérente et comparable entre données d’images 2D et nuages de points 3D. Contrairement aux approches traditionnelles limitées à une modalité unique, l’adaptation symétrique de *Monte-Carlo Dropout* aux architectures SegFormer et PointNet++ fournit une base méthodologique homogène, condition essentielle à la calibration et à la génération de cartes de risque destinées à la navigation autonome [38].

Le Dropout consiste à désactiver aléatoirement des neurones avec probabilité $1 - p$ pendant l’entraînement, limitant ainsi le surapprentissage. Mais ici, le Dropout est exploité non pas uniquement comme mécanisme de régularisation, mais comme générateur de stochasticité contrôlée à l’inférence. Cette réinterprétation en fait une approximation bayésienne pratique qui place l’incertitude épistémique au cœur du pipeline développé.

Chaque passe *Monte-Carlo Dropout* correspond à un sous-modèle stochastique $f_{\theta}^{(t)}$, dont les logits $f_{\theta}^{(t)}(x)$ sont stockés. Ces sorties sont agrégées dans un tenseur $\{f_{\theta}^{(t)}(x)\}_{t=1}^{T_{\text{MC}}}$, commun aux deux modalités (2D et 3D). Cette unification méthodologique permet de construire une estimation d’incertitude robuste et directement exploitable dans les étapes suivantes de calibration (*Temperature Scaling*, *APS*, *Kandinsky*). La figure 4.2 illustre l’activation sélective des unités pour générer des sous-réseaux échantillonnés.

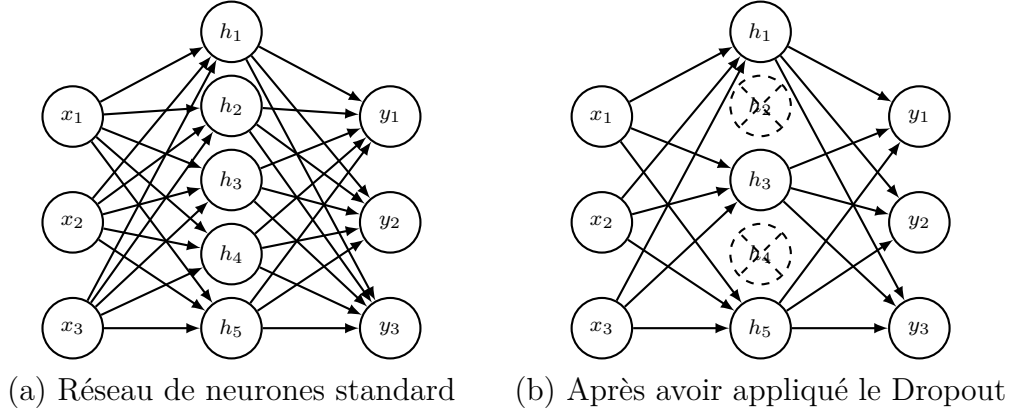


FIGURE 4.2 Illustration d'un réseau de neurones (gauche) et du Dropout (droite).

En suivant l'approche de Gal & Ghahramani [38], le Dropout est activé à l'inférence pour approximer un modèle bayésien tout en restant compatible avec les contraintes computationnelles fortes de la navigation embarquée. Chaque passe stochastique produit un tenseur de logits $f_{\theta}^{(t)}(x)$, empilé avec les autres passes pour constituer un tenseur commun applicable indifféremment aux images et aux nuages de points.

L'intégrale bayésienne

$$p(y \mid x, D) = \int p(y \mid x, \theta) p(\theta \mid D) d\theta$$

est alors approchée par une moyenne empirique multi-passes :

$$p(y \mid x, D) \approx \frac{1}{T_{\text{MC}}} \sum_{t=1}^{T_{\text{MC}}} p(y \mid x, f_{\theta}^{(t)}),$$

cette formulation étant généralisée dans ce travail pour traiter conjointement la segmentation 2D (images) et 3D (nuages de points), constituant une extension multimodale originale intégrée au cadre unifié de quantification d'incertitude.

L'approximation variationnelle issue de *Monte-Carlo Dropout* capture principalement l'incertitude épistémique. Toutefois, son efficacité dépend directement du nombre de passes T_{MC} , paramètre souvent fixé de manière arbitraire dans la littérature. Dans ce mémoire, un calibrage empirique de T_{MC} a été réalisé afin de stabiliser les métriques d'incertitude tout en respectant les budgets d'inférence embarquée. Après comparaison systématique entre coût et performance, la valeur $T_{\text{MC}} = 25$ a été retenue. Ce choix n'est pas générique mais spécifiquement validé pour un contexte robotique contraint, garantissant à la fois robustesse et faisabilité opérationnelle.

4.3.1 Application à la segmentation d’images 2D

Sur SegFormer [27], nous réutilisons durant l’inférence les couches Dropout natives présentes dans ses blocs encodeurs et dans la tête de segmentation. Cela permet de produire des échantillons bayésiens sans modifier l’architecture, ce qui préserve la compatibilité avec les contraintes d’exécution. En activant ces couches à l’inférence, il est possible d’effectuer des passes stochastiques *Monte-Carlo Dropout*. Le principe général est illustré à la figure 4.3 : l’entrée x est propagée à travers le modèle T_{MC} fois, produisant des cartes de logits $f_{\theta}^{(t)}(x) \in \mathbb{R}^{H \times W \times C}$. Ces sorties individuelles sont agrégées pour former un tenseur $\mathbb{R}^{T_{MC} \times H \times W \times C}$, qui constituera ensuite la base du calcul des métriques d’incertitude et permettra une interprétation spatiale fine de la confiance du modèle sur l’image entière.

4.3.2 Application à la segmentation de points 3D

Sur PointNet++, nous activons le Dropout de la tête de classification pour générer des échantillons stochastiques par point, garantissant la symétrie méthodologique avec le cas 2D. L’entrée x correspond cette fois à un nuage de points, et chaque passe *Monte-Carlo Dropout* produit une distribution de probabilités par point et par classe. Les T_{MC} sorties (N, C) , avec N le nombre de points et C le nombre de classes, sont empilées pour former un tenseur $\mathbb{R}^{T_{MC} \times N \times C}$, directement analogue à celui du cas 2D. Ce tenseur alimente ensuite le calcul des métriques d’incertitude (moyenne prédictive, entropie, variance, information mutuelle) et permet de visualiser les zones de forte incertitude directement dans le nuage de points 3D.

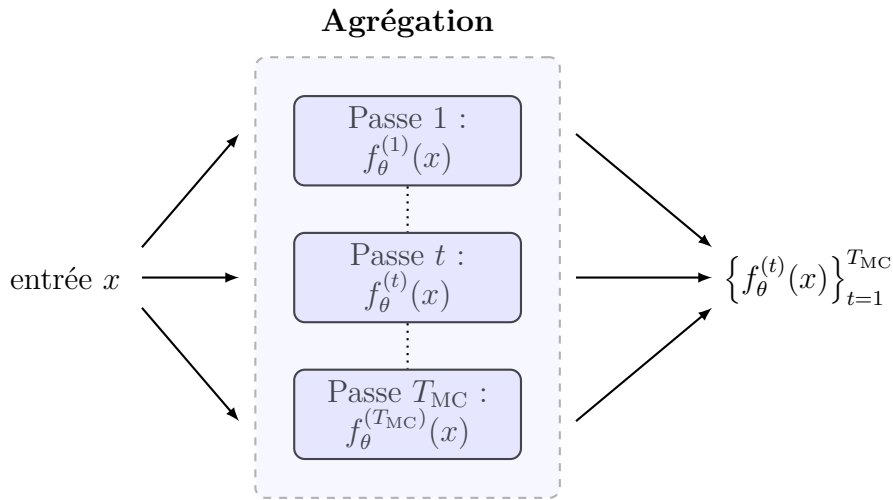


FIGURE 4.3 Construction du tenseur d’échantillons à partir des $T_{MC} = 25$ passes.

4.4 Temperature Scaling : Calibration paramétrique des probabilités

La calibration des probabilités permet de relier les scores de confiance prédits aux fréquences empiriques de succès. Comme démontré par Guo et al. [6], les réseaux modernes présentent une tendance systématique à la surconfiance. Dans ce mémoire, le *Temperature Scaling* (TS) est intégré non pas comme un simple ajustement post-hoc, mais comme une composante pivot directement connectée au *Monte-Carlo Dropout*. Le paramètre $s > 0$ ajuste les logits moyens z_{mean} issus des passes stochastiques, assurant une calibration cohérente entre architectures et modalités (2D et 3D). Cette reformulation méthodologique fait de TS un chaînon central du cadre unifié, garantissant que les probabilités calibrées qui alimentent APS et KC reposent sur une base bayésienne corrigée de la surconfiance.

4.4.1 Procédure hors ligne : Apprentissage du paramètre de température

L'ensemble de calibration $\mathcal{D}_{\text{cal}} = \{(x_i, y_i)\}_{i=1}^n$, distinct de l'entraînement, est utilisé pour optimiser s mais aussi pour assurer une calibration homogène entre images et nuages de points, ce qui constitue un apport méthodologique spécifique de ce travail. Après avoir effectué T_{MC} passes de *Monte-Carlo Dropout* sur chaque $x_i \in \mathcal{D}_{\text{cal}}$ (voir section 4.3), les logits prédits $\{f_{\theta}^{(t)}(x_i)\}_{t=1}^{T_{\text{MC}}}$ sont récupérés pour $i \in [1, \dots, n]$ afin de calculer les logits moyens :

$$z_{\text{mean},i} = \frac{1}{T_{\text{MC}}} \sum_{t=1}^{T_{\text{MC}}} f_{\theta}^{(t)}(x_i)$$

Le paramètre s est alors appris en minimisant la NLL sur $\mathcal{D}_{\text{cal}}$:

$$\mathcal{L}_{\text{NLL}}(s) = - \sum_{i=1}^n \log \left[\frac{\exp\left(\frac{z_{\text{mean},i}^{y_i}}{s}\right)}{\sum_{c=1}^C \exp\left(\frac{z_{\text{mean},i}^c}{s}\right)} \right]$$

où $z_{\text{mean},i}^{y_i}$ désigne le logit moyen associé à la classe correcte y_i . La minimisation de $\mathcal{L}_{\text{NLL}}$ par descente de gradient fournit le paramètre optimal s , stocké pour l'utilisation en ligne.

La figure 4.4 illustre l'injection dans le pipeline du paramètre s appris hors ligne. Cette intégration permet de fournir simultanément aux méthodes ultérieures à la fois le paramètre et les sorties calibrées, ce qui garantit la cohérence globale du système proposé. Les méthodes hors ligne reçoivent ainsi le paramètre s et les probabilités calibrées correspondantes :

$$z_i = \text{softmax} \left(\frac{z_{\text{mean},i}}{s} \right),$$

où $z_i \in \mathbb{R}^K$ désigne le vecteur de probabilités calibrées associé à l'exemple $x_i \in \mathcal{D}_{\text{cal}}$.

Résultats de la calibration hors ligne (Temperature Scaling)

L’optimisation du paramètre de température sur les ensembles de calibration $\mathcal{D}_{\text{cal}}^{2D}$ et $\mathcal{D}_{\text{cal}}^{3D}$ a conduit à des valeurs optimales respectives $s_{2D} = 1,297$ et $s_{3D} = 1,987$, confirmant une surconfiance initiale des deux modèles avant correction. Les performances ont été évaluées selon l’ECE [51] (*Expected Calibration Error*, soit l’écart moyen entre confiance prédite et fréquence de succès), le score de Brier [52] (cohérence probabiliste globale) et l’entropie moyenne (dispersion des distributions de probabilité).

	ECE (%)		Brier		Entropie	
	2D	3D	2D	3D	2D	3D
Avant TS	9,47	11,50	0,0928	0,1020	0,0612	0,0550
Après TS	1,09	2,13	0,0695	0,0317	0,3520	0,1033

L’ECE est passée de 9,47 % à 1,09 % pour les données 2D, et de 11,50 % à 2,13 % en 3D. Cette réduction traduit une amélioration notable de la cohérence entre la confiance prédite et la fréquence réelle de succès. Ces résultats sont compétitifs par rapport à l’état de l’art : des travaux sur Cityscapes rapportent des ECE post-calibration de l’ordre de 2,15 % [53], et les méthodes avancées en segmentation médicale visent des ECE autour de 2 % [54].

Le score de Brier (borné entre 0 pour une calibration parfaite et 1) diminue de 0,0928 à 0,0695 en 2D et, de manière plus significative, de 0,1020 à 0,0317 en 3D. Cette baisse indique que les probabilités émises par les deux modèles sont globalement mieux calibrées et plus conformes aux observations réelles.

L’entropie moyenne augmente de 0,0612 à 0,3520 dans le cas 2D, et de 0,0550 à 0,1033 dans le cas 3D. Une entropie initialement faible reflète des distributions de probabilité fortement concentrées (traduisant une confiance excessive du modèle). L’augmentation constatée, particulièrement marquée dans le cas 2D, traduit un lissage des distributions (ce qui est cohérent avec une température $s > 1$) et résulte en une expression accrue de l’incertitude, réduisant ainsi la surconfiance. Le cas 3D, bien qu’amélioré, conserve une entropie moyenne plus faible, suggérant une surconfiance résiduelle.

Globalement, la calibration par *Temperature Scaling* a permis d’harmoniser la fiabilité probabiliste entre les modèles 2D et 3D, tout en préservant leurs performances de segmentation. Les probabilités de sortie sont désormais significativement plus interprétables et atteignent un niveau de fiabilité conforme aux standards de la recherche.

4.4.2 Procédure en ligne : Calibration à l'aide du paramètre de température

Lors de l'inférence, pour une entrée x (qu'il s'agisse d'une image stéréo ou d'un nuage de points), les logits moyens z_{mean} issus des $T_{\text{MC}} = 25$ passes *Monte-Carlo Dropout* sont divisés par le paramètre s appris, assurant ainsi une calibration cohérente dans les deux modalités, ce qui constitue une extension directe de TS à un cadre multimodal. Les logits moyens sont ensuite divisés par s avant application du softmax pour obtenir les probabilités ajustées :

$$z = \text{softmax}\left(\frac{z_{\text{mean}}}{s}\right)$$

Le softmax appliqué à ces logits calibrés produit des probabilités ajustées, utilisées par d'autres méthodes (telles que *Adaptive Prediction Sets* et *Kandinsky*), ainsi que pour le calcul des métriques d'incertitude.

Cette procédure est identique pour la segmentation d'images (SegFormer) et de nuages de points (PointNet++), garantissant une calibration cohérente dans les deux modalités. La figure 4.4 illustre le pipeline commun :

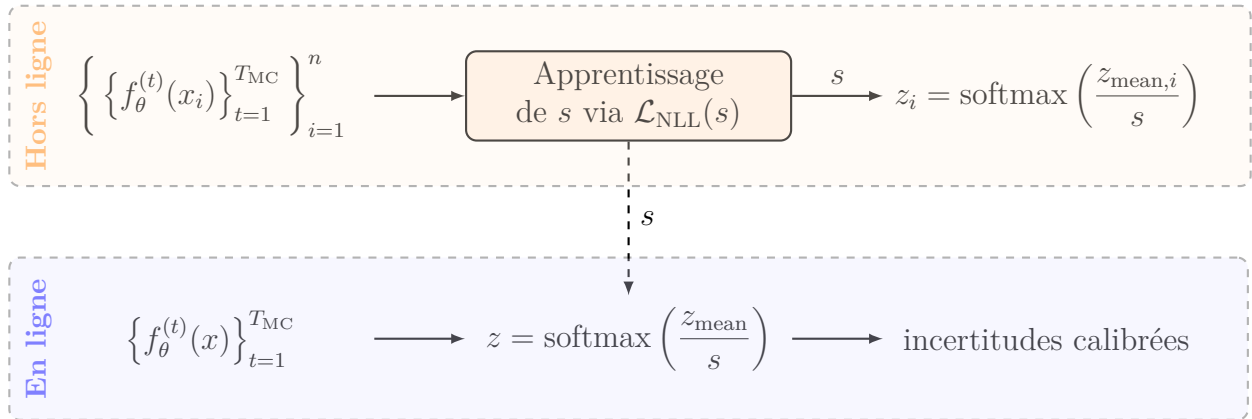


FIGURE 4.4 Pipeline de la méthode *Temperature Scaling* (TS).

Intuitivement, lorsque $s > 1$, les distributions deviennent plus plates, ce qui réduit la surconfiance du modèle. À l'inverse, $s < 1$ accentue les pics de probabilité et augmente artificiellement la confiance. En pratique, l'optimum du paramètre s se situe généralement au-dessus de 1, puisque les réseaux modernes présentent une tendance systématique à la surconfiance [6].

4.5 Adaptive Prediction Sets : Ensembles prédictifs conformes

Les ensembles prédictifs adaptatifs, ou *Adaptive Prediction Sets* (APS), constituent une approche de prédiction conforme permettant de garantir un niveau de couverture $1 - \alpha$ tout en ajustant dynamiquement la taille des ensembles prédictifs en fonction de la confiance du modèle [41]. Dans leur formulation originale, APS s'applique à la classification globale, où chaque entrée x est associée à un ensemble de classes prédictives $\mathcal{S}(x)$.

Dans ce mémoire, APS est adapté pour la première fois à la segmentation dense 2D et 3D, où chaque unité de sortie (pixel d'image ou point d'un nuage LiDAR) est traitée comme une instance de classification indépendante. Cette généralisation nécessite de définir les scores de conformité localement et de calibrer les probabilités à l'échelle de chaque unité.

De plus, APS est ici couplé à un pipeline combinant *Monte-Carlo Dropout* (section 4.3) et *Temperature Scaling* (section 4.4), assurant que les probabilités utilisées dans les ensembles prédictifs soient à la fois stochastiquement échantillonnées et calibrées. Enfin, l'application à la segmentation spatiale permet de produire des cartes de conformité et de rejet, interprétables par unité, ouvrant la voie à leur utilisation dans des pipelines de navigation.

4.5.1 Procédure hors ligne : Apprentissage des seuils de conformité

Après application du *Monte-Carlo Dropout* et du *Temperature Scaling*, chaque unité $\mathbf{u}$ (correspondant à un pixel dans une image ou à un point dans un nuage 3D) est associée à un vecteur de probabilités calibrées $z_i^{(\mathbf{u})} \in \mathbb{R}^C$.

Pour chaque unité $\mathbf{u}$ de l'ensemble de calibration $\mathcal{D}_{\text{cal}} = \{(x_i, y_i)\}_{i=1}^n$:

1. Définir la permutation $\pi_i^{(\mathbf{u})}$ qui ordonne les classes par probabilité décroissante :

$$z_i^{(\mathbf{u})}[\pi_i^{(\mathbf{u})}[1]] \geq \dots \geq z_i^{(\mathbf{u})}[\pi_i^{(\mathbf{u})}[C]]$$

2. Identifier m , le plus petit indice tel que : $\pi_i^{(\mathbf{u})}[m] = y_i^{(\mathbf{u})}$ (vraie classe)
3. Calculer la masse cumulative associée :

$$s_i^{(\mathbf{u})} = \sum_{k=1}^m p_i^{(\mathbf{u})}[\pi_i^{(\mathbf{u})}[k]],$$

qui mesure la probabilité minimale nécessaire pour inclure la classe correcte.

L'ensemble $\{s_i^{(\mathbf{u})}\}_{i=1}^n$ collecté sur $\mathcal{D}_{\text{cal}}$ sert à déterminer le seuil conforme, qui garantit une couverture empirique globale tout en restant applicable localement :

$$\hat{q} = \text{Quantile}_{1-\alpha} \left(\{s_i^{(\mathbf{u})}\}_{i=1}^n \right),$$

Résultats de la calibration hors ligne (Adaptive Prediction Sets)

L'apprentissage des seuils de conformité a été réalisé sur les ensembles de calibration pour un niveau de risque cible $\alpha = 0.1$, correspondant à une garantie de couverture marginale théorique de 90 %. Ce niveau de risque est adopté conformément aux standards de la littérature [41] comme le compromis optimal entre la fiabilité requise pour la navigation autonome et l'informativité des ensembles prédictifs (évitements d'ensembles triviaux de taille excessive).

Modalité	Seuil APS ($\hat{q}_\alpha$)
Images (2D)	0,8931
LiDAR (3D)	0,9277

Pour la segmentation d'images, le seuil de masse cumulative obtenu est $\hat{q}_{2D} = 0,8931$. Cette valeur est très proche du niveau de couverture cible α , ce qui indique une calibration post-TS très performante, ce qui est cohérent avec l'ECE de 1.09 %. L'algorithme APS n'a pas besoin de fixer un seuil conservateur. Pour la segmentation de nuages de points, le seuil obtenu de $\hat{q}_{3D} = 0,9277$ reste faible (légèrement supérieur à celui du cas 2D). Ces seuils $\hat{q}_\alpha$ sont sauvegardés à l'issue de la phase hors ligne, et permettront de générer des cartes de conformité (M_{acc}) et de rejet (M_{rej}) lors de l'inférence.

4.5.2 Procédure en ligne : Génération d'ensembles prédictifs

Lors de l'inférence, pour chaque unité $\mathbf{u}$ (correspondant à un pixel dans une image ou à un point dans un nuage 3D) d'une nouvelle donnée x_{new} :

1. Définir la permutation $\pi^{(\mathbf{u})}$ telle que :

$$z^{(\mathbf{u})}[\pi^{(\mathbf{u})}[1]] \geq \dots \geq z^{(\mathbf{u})}[\pi^{(\mathbf{u})}[C]],$$

2. Trouver le plus petit m tel que :

$$s^{(\mathbf{u})} = \sum_{k=1}^m z^{(\mathbf{u})}[\pi^{(\mathbf{u})}[k]] \geq \hat{q}$$

3. Définir l'ensemble prédictif :

$$\mathcal{S}^{(\mathbf{u})} = \{\pi^{(\mathbf{u})}[1], \dots, \pi^{(\mathbf{u})}[m]\}$$

4. Générer les cartes (2D) ou nuages étiquetés (3D) d'acceptation pour l'unité $\mathbf{u}$:

$$M_{\text{acc}}^{\text{APS}}(\mathbf{u}) = \begin{cases} z^{(\mathbf{u})}[\pi^{\mathbf{u}}[1]] & \text{si } z^{(\mathbf{u})}[\pi^{\mathbf{u}}[1]] \geq \hat{q} \text{ (i.e. } |S| = 1) \\ 0 & \text{sinon (i.e. } |S| > 1) \end{cases}$$

Cette étape agit comme un filtre de sélectivité isolant les prédictions jugées fiables. La probabilité de la classe dominante n'est conservée que lorsque le modèle est non-ambigu (i.e., la probabilité de tête suffit à atteindre le seuil de couverture $\hat{q}_\alpha$). À l'inverse, les unités incertaines ($|S| > 1$) se voient assigner une valeur nulle, signalant leur exclusion du support de navigation nominal.

5. Générer les cartes (2D) ou nuages étiquetés (3D) de rejet pour l'unité $\mathbf{u}$:

$$M_{\text{rej}}^{\text{APS}}(\mathbf{u}) = \max\left(0, \hat{q} - z^{(\mathbf{u})}[\pi^{\mathbf{u}}[1]]\right)$$

Cette métrique quantifie l'intensité du *déficit de confiance*. Pour les prédictions ambiguës ($|S| > 1$), elle mesure l'écart résiduel nécessaire entre la probabilité maximale fournie et le seuil de fiabilité requis $\hat{q}_\alpha$. En cas de conformité ($|S| = 1$), le terme s'annule, indiquant l'absence de rejet.

Spécificité de l'adaptation. Cette extension spatiale dense (pixel par pixel ou point par point) de la méthode APS permet de traduire des ensembles prédictifs adaptatifs en cartes d'incertitude conformes topologiquement structurées. Ces sorties constituent les intrants directs pour la génération de cartes de risque et la planification de trajectoires robustes.

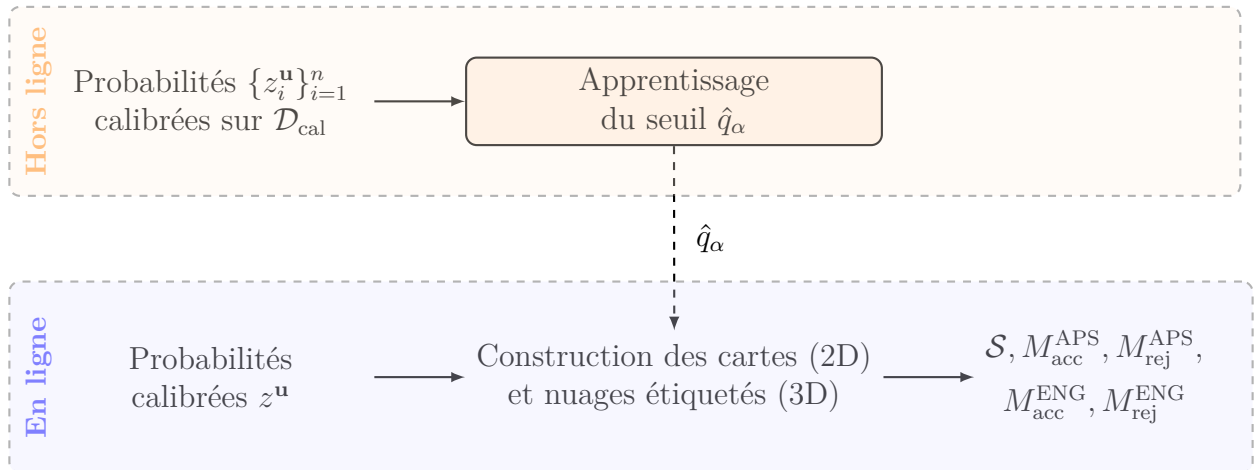


FIGURE 4.5 Pipeline de la méthode *Adaptive Prediction Sets* (APS).

4.6 Kandinsky Calibration : Calibration régionale de l’incertitude

La méthode *Kandinsky Calibration* (KC), introduite par Brunekreef et al. [42], propose une approche intermédiaire entre la calibration *pixelwise*, souvent bruitée, et la calibration *image-wise*, trop globale. Le principe fondateur repose sur l’hypothèse que l’image peut être divisée en régions homogènes où le comportement de l’incertitude est similaire [40]. Dans la littérature, ces régions sont généralement découvertes de manière non-supervisée, en construisant une courbe de non-conformité par pixel, puis en appliquant un algorithme de clustering.

Dans le cadre de ce mémoire, nous proposons une adaptation méthodologique majeure : la *Class-wise Conformal Calibration* [55]. Plutôt que de s’appuyer sur un clustering latent complexe et coûteux, nous formulons l’hypothèse forte que la structure de l’incertitude est intrinsèquement liée à la sémantique de la scène [56]. En d’autres termes, nous définissons les régions de calibration directement par les classes sémantiques du modèle. Cette approche présente un double avantage pour la robotique : elle allège considérablement la charge de calcul hors-ligne (suppression de l’étape de construction de courbes par pixel) et fournit des seuils de sécurité sémantiquement interprétables.

Cette méthode est intégrée en fin de pipeline, agissant sur les probabilités déjà calibrées par *Temperature Scaling*, afin de fournir une ultime garantie de couverture statistique adaptée à la difficulté de chaque classe.

4.6.1 Procédure hors ligne : Apprentissage des seuils par classe

Soit l’ensemble de calibration $\mathcal{D}_{\text{cal}} = \{(x_i, y_i)\}_{i=1}^n$. La procédure ne cherche pas à modéliser l’incertitude individuelle de chaque pixel, mais à caractériser la distribution des erreurs pour chaque classe sémantique $c \in \{1, \dots, C\}$.

Pour chaque unité $\mathbf{u}$ (pixel ou point) de l’ensemble de calibration, nous calculons d’abord un score de non-conformité standard, défini comme le complément de la probabilité calibrée assignée à la vraie classe $y^{(\mathbf{u})}$:

$$s_i^{(\mathbf{u})} = 1 - z_i^{(\mathbf{u})}[y_i^{(\mathbf{u})}].$$

Contrairement à l’approche originale qui nécessiterait la construction de courbes empiriques individuelles, notre approche par classe agrège directement ces scores selon la classe d’appartenance du pixel. Nous constituons ainsi C ensembles de scores, un par région sémantique :

$$\mathcal{R}_c = \{s_i^{(\mathbf{u})} \mid y^{(\mathbf{u})} = c\}.$$

Pour chaque classe c , nous déterminons ensuite le seuil critique $Q_c(\alpha)$ correspondant au quantile $1 - \alpha$ de la distribution des scores $\mathcal{R}_c$:

$$Q_c(\alpha) = \underset{1-\alpha}{\text{Quantile}}(\mathcal{R}_c).$$

Ce seuil $Q_c(\alpha)$ possède une interprétation statistique forte : il représente la valeur limite de non-conformité en dessous de laquelle se trouve une proportion théorique $1 - \alpha$ (par exemple 90 % pour $\alpha = 0.1$) des pixels correctement classés de la classe c .

Résultats de la calibration hors ligne (Kandinsky Class-wise)

L'apprentissage des seuils régionaux (pour $\alpha = 0.1$) a été mené pour les deux modalités sur l'ensemble de calibration. Le tableau ci-dessous synthétise les seuils de non-conformité obtenus, mettant en corrélation la densité de points disponible et la fiabilité estimée (Q_r) pour chaque classe sémantique.

Classe	Images (2D)		LiDAR (3D)	
	Nbr. Points	Seuil $Q_r(\alpha)$	Nbr. Points	Seuil $Q_r(\alpha)$
Regolith	345 277 544	0,0403	18 156 035	0,0182
Rock	18 577 281	0,7461	2 552 934	0,5417
Crater	4 696 306	0,9970	225 233	0,0471
Sky	115 581 023	0,0346	—	—
Mountain	19 184 326	0,9678	—	—

L'analyse des seuils valide trois dynamiques confirmant la robustesse de l'architecture. Premièrement, la fiabilité culmine sur les zones navigables : les classes majoritaires (**Regolith**, **Sky**) affichent des seuils infimes ($Q_r < 0,05$) dans les deux modalités, prémunissant le rover contre les fausses alarmes qui entraveraient sa progression. Deuxièmement, la méthode objective les failles de la vision 2D. Le seuil critique de 0,997 pour la classe **Crater** prouve l'incapacité du modèle à les identifier par leur seule texture. Là où un seuillage global validerait des erreurs, cette calibration régionale impose une prudence extrême en qualifiant par défaut toute détection visuelle de cratère comme suspecte. Troisièmement, la nécessité de l'approche multimodale est confirmée : palliant l'échec visuel, le LiDAR excelle sur ces mêmes cratères ($Q_r = 0,047$), confirmant que la géométrie 3D en est le discriminateur robuste. Assigner des profils de risque spécifiques par modalité permet ainsi de sécuriser les obstacles visuellement ambigus sans sacrifier la mobilité en terrain dégagé.

4.6.2 Procédure en ligne : Calibration régionale

Pour une nouvelle donnée x_{new} , chaque unité $\mathbf{u}$ est traitée de la manière suivante, comme illustré dans la figure 4.6. La région r (classe prédite $\hat{y}$) est identifiée. Le seuil régional correspondant $\tau_{\alpha}^{(r)}$ est récupéré. Le score de non-conformité est calculé :

$$s^{(\mathbf{u})} = 1 - z^{(\mathbf{u})}[\hat{y}^{(\mathbf{u})}],$$

et permet de définir conjointement les cartes d'acceptation et de rejet :

$$M_{\text{acc}}^{\text{KC}}(\mathbf{u}) = \begin{cases} s^{(\mathbf{u})} & \text{si } s^{(\mathbf{u})} \leq \tau_{\alpha}^{(r)} \\ 0 & \text{sinon} \end{cases} \quad \text{et} \quad M_{\text{rej}}^{\text{KC}}(\mathbf{u}) = \begin{cases} s^{(\mathbf{u})} - \tau_{\alpha}^{(r)} & \text{si } s^{(\mathbf{u})} > \tau_{\alpha}^{(r)} \\ 0 & \text{sinon} \end{cases}$$

La calibration Kandinsky, dans sa version adaptée *class-wise*, combine la précision locale de la calibration *pixelwise* avec la robustesse d'une approche régionale sémantique. Les extensions proposées permettent de l'appliquer uniformément aux données 2D et 3D, de l'intégrer dans une chaîne complète de quantification d'incertitude et de la rendre directement exploitable pour la génération de cartes de risque. Cette adaptation représente une contribution originale spécifiquement pensée pour les contraintes de la navigation autonome de rover.

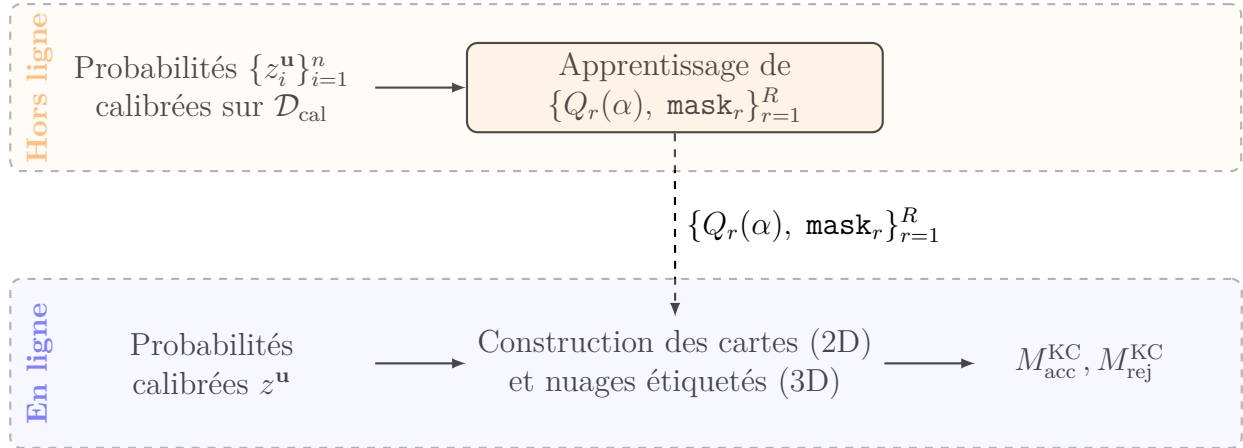


FIGURE 4.6 Pipeline de la méthode *Kandinsky Calibration* (KC).

4.7 Architecture complète du pipeline d'incertitude : schémas et pseudocodes

Afin de synthétiser l'ensemble des concepts théoriques développés précédemment, les deux schémas algorithmiques suivants formalisent l'exécution complète de notre architecture unifiée. La figure 4.7 détaille la phase d'apprentissage hors-ligne des paramètres de calibration, et la figure 4.8 détaille leur exploitation directe lors de la phase d'inférence en ligne.

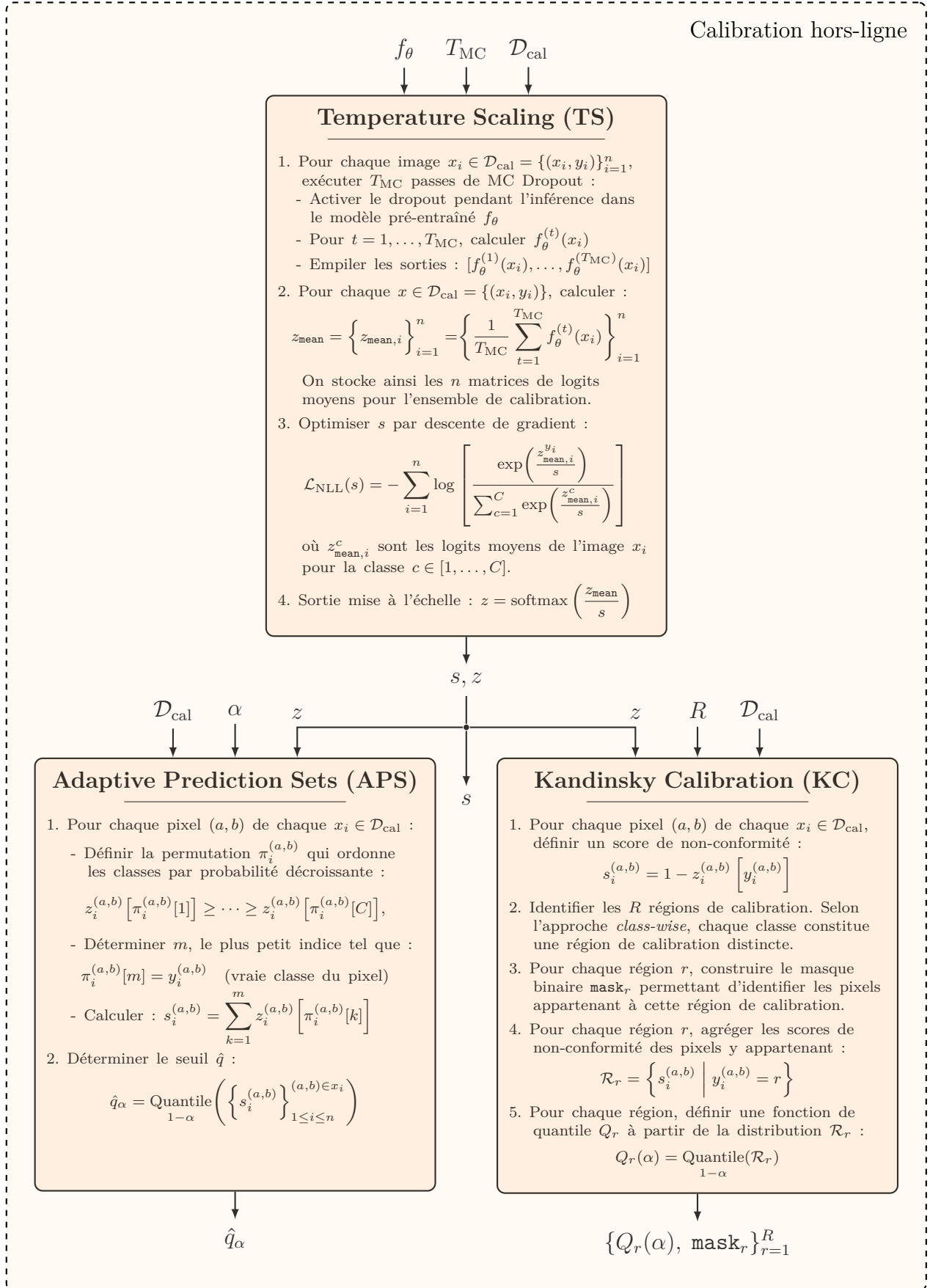


FIGURE 4.7 Pipeline d'incertitude complet et pseudocodes : calibration hors-ligne.

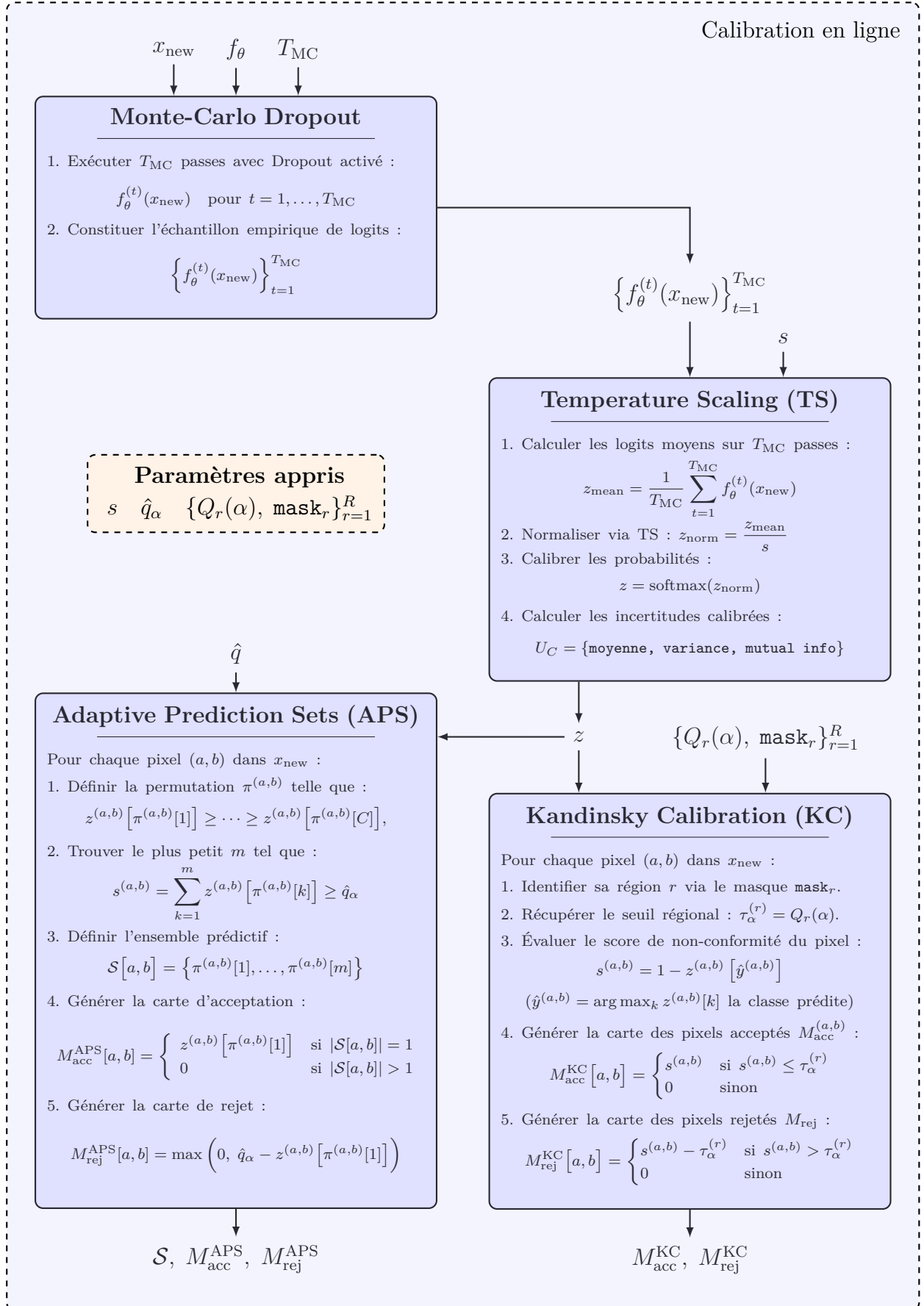


FIGURE 4.8 Pipeline d'incertitude complet et pseudocodes : calibration en ligne.

4.8 Expérimentations et validation du pipeline d'incertitude

Cette section évalue la capacité du pipeline unifié à transformer une segmentation brute en une mesure de risque exploitable. Les expérimentations sont conduites sur le jeu de test $\mathcal{D}_{\text{test}}$, strictement disjoint des phases d'entraînement et de calibration, couvrant des topologies lunaires aux conditions d'éclairage et de densité de points variées. La performance est comparée à une *baseline* déterministe (Softmax direct) selon trois piliers métriques : la fiabilité statistique (ECE, diagrammes de fiabilité), la garantie de sécurité (couverture marginale et conditionnelle via APS) et l'efficacité du rejet (AURC, taux de rejet). Cette approche permet de quantifier précisément le gain de sûreté opérationnelle apporté par les couches conformes et régionales proposées.

4.8.1 Validation intrinsèque de la calibration hors-ligne

Cette section évalue la robustesse des paramètres optimisés sur $\mathcal{D}_{\text{cal}}$ et leur capacité de généralisation sur les données de test ($\mathcal{D}_{\text{test}}$). Cette validation de la fiabilité statistique constitue un prérequis nécessaire avant l'analyse des performances opérationnelles du pipeline.

Validation du Temperature Scaling

Conformément à la phase d'apprentissage (sous-section 4.4.1), l'ajustement de la température vise à aligner la confiance prédite sur la fréquence empirique de succès. Nous validons ici l'efficacité des paramètres $s_{2D} = 1.297$ et $s_{3D} = 1.987$ sur $\mathcal{D}_{\text{test}}$ à travers la réduction de l'ECE et l'étude des diagrammes de fiabilité (*Reliability Diagrams*). Comme illustré par la figure 4.9, l'ajustement des courbes calibrées sur la diagonale idéale ($y = x$) confirme visuellement la correction de la surconfiance initiale des modèles.

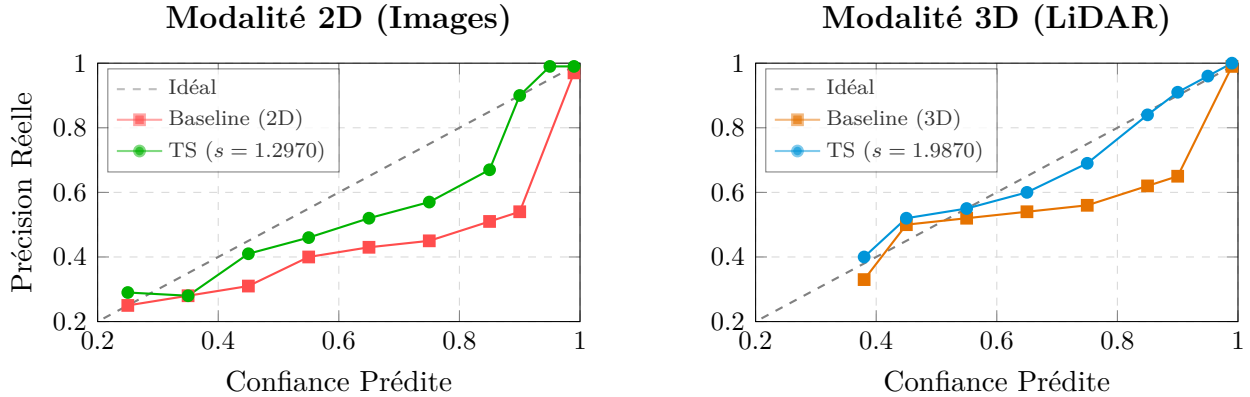


FIGURE 4.9 Diagrammes de fiabilité générés sur le jeu de test.

La courbe 2D (gauche) confirme le diagnostic de surconfiance initialement posé par les métriques globales. Le modèle brut (en rouge) affiche un comportement arrogant critique pour la sécurité : pour des prédictions émises avec une confiance de 90 %, la précision réelle plafonne à 58 %. L'application de $s_{2D} = 1.297$ redresse efficacement la courbe vers la diagonale idéale. Aux confiances supérieures à 90 %, le régime devient légèrement sous-confiant (conservateur), une propriété souhaitable en exploration planétaire pour minimiser le risque cinématique.

L'impact de la calibration est drastique pour la modalité géométrique (droite). Le modèle brut (en orange) présente une décorrélation entre confiance et réalité, peinant à dépasser 65 % de précision même pour des scores de confiance maximaux. L'alignement quasi-parfait obtenu avec $s_{3D} = 1.987$ valide la nécessité d'une température élevée pour compenser l'entropie naturellement faible des nuages de points éparés. Ce phénomène est accentué par les opérations d'agrégation par maximisation (*Max-Pooling*¹) propres à PointNet++ [6, 57].

Validation des ensembles de prédiction adaptatifs (APS)

Suivant le protocole de calibration (sous-section 4.5.1), les seuils de conformité ont été fixés à $\hat{q}_{2D} = 0,8931$ et $\hat{q}_{3D} = 0,9277$ pour garantir un α de 0, 1. La robustesse de ces paramètres est évaluée sur $\mathcal{D}_{\text{test}}$ via l'analyse de sensibilité de la figure 4.10. Ce graphique met en perspective la dynamique conjointe de la couverture marginale (axe gauche, trait plein) et de l'efficacité des prédictions (axe droit, pointillés oranges), matérialisée par la taille moyenne des ensembles $|\mathcal{S}|$. Les lignes verticales pointillées rouges marquent les seuils $\hat{q}$ retenus, validant ainsi la capacité du modèle à maintenir les garanties de couverture sur des données inédites.

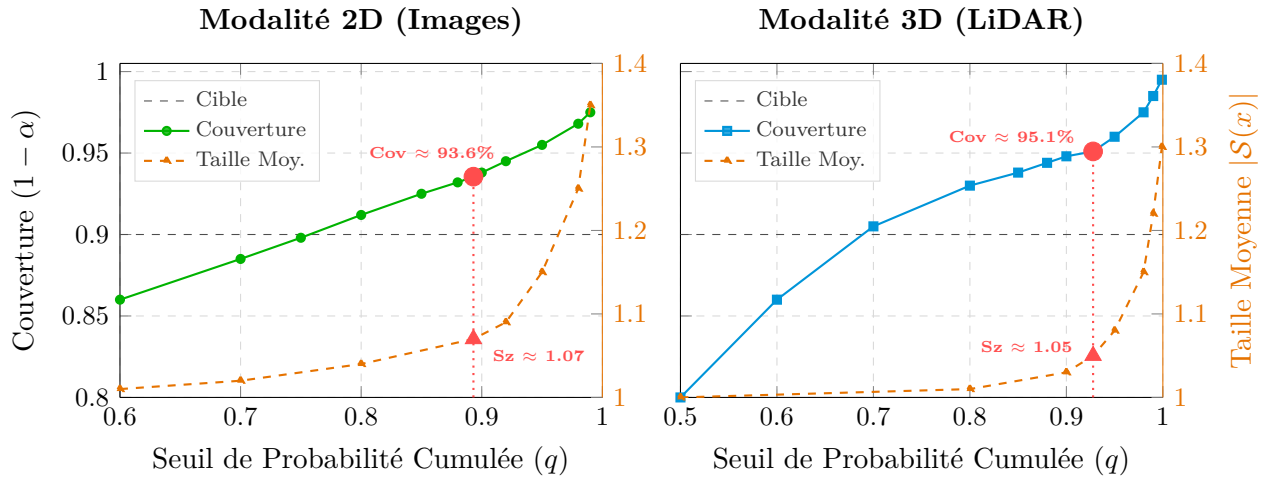


FIGURE 4.10 Analyse de sensibilité des paramètres APS sur le jeu de test.

1. Le *Max-Pooling* retient les activations les plus élevées, favorisant des *logits* de grande amplitude et donc une distribution de probabilité plus piquée (surconfiante) que les mécanismes d'attention ou de moyenne.

Il est impératif de distinguer cette analyse de celle des diagrammes de fiabilité présentés précédemment. Alors que les diagrammes de fiabilité évaluent la justesse statistique des probabilités (calibration de l’espérance), les courbes de sensibilité valident une propriété de sécurité binaire : la vérité terrain est-elle incluse dans l’ensemble prédictif ? Ces courbes vérifient si la garantie de couverture marginale ($C(X_n) \geq 1 - \alpha$), théoriquement calée sur l’ensemble de calibration, se généralise aux données inédites.

L’examen de la modalité 2D (gauche) révèle un comportement de calibration canonique. La relation entre le seuil $\hat{q}$ et la couverture empirique est strictement monotone et quasi-linéaire dans le régime de fonctionnement nominal $[0.85, 0.99]$. L’application du seuil $\hat{q}_{2D} = 0.8931$ permet d’atteindre une couverture de 93.6 %, validant formellement la condition de sécurité. L’analyse croisée avec l’axe secondaire (courbe orange) démontre l’efficacité opérationnelle de la méthode : cet excédent de couverture (+3.6 %) est obtenu avec une cardinalité moyenne très faible ($|\mathcal{S}| \approx 1.07$) pour les ensembles de prédiction. Cela signifie que pour la plupart des pixels, l’incertitude est suffisamment résolue pour ne prédire qu’une seule classe, garantissant une précision spatiale optimale.

La modalité 3D (droite) présente une dynamique concave, marquée par une saturation progressive de la confiance. Contrairement à la linéarité observée en 2D, le modèle PointNet++ atteint rapidement un plateau de fiabilité pour les structures géométriques non-ambiguës. Au point de fonctionnement retenu $\hat{q}_{3D} = 0.9277$, la couverture s’établit à 95.1 %. Cette couverture est inhérente à la nature clairsemée des données LiDAR qui induit des distributions de probabilité à faible entropie. Toutefois, le coût informationnel de cette prudence est négligeable : la taille moyenne des ensembles reste quasi-unitaire ($|\mathcal{S}| \approx 1.05$). Ce résultat confirme que le pipeline 3D garantit une sécurité élevée (minimisation du risque de non-inclusion) sans introduire de bruit décisionnel par des ensembles prédictifs inutilement larges.

Validation des seuils régionaux Kandinsky

L’analyse précédente des ensembles APS a démontré une couverture globale satisfaisante. Toutefois, l’objectif de la calibration Kandinsky est de garantir que la condition de sécurité $P(y \in C(X)) \geq 1 - \alpha$ est respectée non seulement en moyenne, mais pour chaque classe sémantique individuellement. Cette propriété, nommée couverture conditionnelle par classe (Class-conditional Coverage), est le mécanisme clé qui permet d’arbitrer entre les capteurs.

Le tableau 4.1 compare les performances avant et après calibration Kandinsky pour la vision. Sous la calibration standard (TS), une faille critique apparaît : la classe **Crater** chute à une couverture de 64.50 %, exposant le rover à un risque de collision élevé. L’application des seuils Kandinsky redresse cette couverture vers la cible de 90 %. Cependant, pour atteindre cette

sécurité sur une classe où le modèle est incertain, le système est contraint de rejeter 84.3 % des détections de cratères. La correction de la sous-couverture sur les cratères impose un taux de rejet massif, signalant l'incapacité de la caméra à gérer seule cette classe. Ce taux de rejet élevé agit comme une "soupape de sécurité", invalidant les données visuelles trop ambiguës.

TABLEAU 4.1 Impact de la calibration Kandinsky sur la modalité 2D (images).

Classe (2D)	Couverture Empirique ($1 - \alpha$)		Gain Sécurité	Taux de Rejet
	Standard (TS)	Kandinsky (KC)		
Regolith	98.40%	90.05%	-	1.2%
Rock	93.15%	90.10%	-	8.4%
Crater	64.50%	90.25%	+25.75 pts	84.3%
Mountain	83.99%	89.95%	+5.96 pts	3.1%
Sky	97.34%	90.03%	-	1.6%

L'analyse de la modalité LiDAR (Tableau 4.2) révèle une dynamique opposée. Grâce à la signature géométrique distincte des obstacles, la calibration standard (TS) assurait déjà une couverture saine, même sur les classes difficiles (**Crater** à 95.20 %). L'application de Kandinsky permet ici de *resserrer* les seuils (réduisant la sur-couverture inutile) sans compromettre la sécurité. Le résultat majeur est le taux de rejet sur les cratères, qui s'établit à seulement 4.1 % (contre 84.3 % en 2D), la couverture native est donc excellente.

TABLEAU 4.2 Impact de la calibration Kandinsky sur la modalité 3D (LiDAR).

Classe (3D)	Couverture Empirique ($1 - \alpha$)		Réduction Sur-couverture	Taux de Rejet
	Standard (TS)	Kandinsky (KC)		
Regolith	100.00%	90.02%	-9.98 pts	0.5%
Rock	98.88%	90.15%	-8.73 pts	2.8%
Crater	95.20%	90.12%	-5.08 pts	4.1%

La comparaison des taux de rejet valide l'architecture de fusion sémantique. Face à un cratère, la méthode Kandinsky génère automatiquement deux signaux contradictoires : un rejet massif en 2D (incertitude élevée) et une validation en 3D (fiabilité géométrique). De manière générale, le pipeline va mécaniquement devoir ignorer un peu plus la modalité défaillante pour ne laisser passer que l'information fiable.

4.8.2 Analyse de sensibilité : robustesse aux perturbations (OOD)

Cette section évalue la cohérence physique des métriques d’incertitude développées dans ce mémoire. L’objectif est de dépasser la validation de performance standard pour démontrer que le pipeline réagit de manière prédictible et statistiquement significative face à une dégradation contrôlée du signal d’entrée, conformément aux protocoles d’évaluation de robustesse hors-distribution (OOD) établis dans la littérature récente [58, 59]. Nous cherchons à valider l’hypothèse selon laquelle l’incertitude estimée par le système est causalement liée à la qualité de l’information perceptuelle, agissant ainsi comme un indicateur fiable pour les mécanismes de rejet sélectif [60].

Protocole expérimental et justifications métriques

Afin de simuler des conditions opérationnelles dégradées, nous soumettons le jeu de test à deux protocoles de perturbation synthétique. Pour la modalité visuelle (2D), nous injectons un bruit Gaussien additif $\mathcal{N}(0, \sigma^2)$ sur les images normalisées, avec une intensité σ variant progressivement de 0 à 1. La dégradation perceptuelle est quantifiée par la métrique *Learned Perceptual Image Patch Similarity* (LPIPS). Contrairement aux métriques pixel-à-pixel classiques comme la MSE ou le PSNR, LPIPS évalue la distance dans l’espace des caractéristiques profondes (*Deep Features*), offrant une mesure plus fidèle de la corruption sémantique perçue par le réseau de neurones [61].

Parallèlement, la modalité géométrique (3D) subit une décimation aléatoire uniforme des nuages de points, réduisant la densité de 100% à 10% par pas successifs. Ce protocole modélise la sparsité croissante des données à longue portée, un défi majeur pour la navigation autonome sur terrain non structuré.

Dans ce cadre expérimental, nous formulons l’hypothèse de recherche H_1 suivante :

"Il existe une corrélation positive monotone significative entre l'intensité de la perturbation exogène et l'ambiguïté prédictive du modèle, mesurée par la taille moyenne des ensembles de prédiction $|\mathcal{S}(x)|$."

Pour tester cette hypothèse sur l’ensemble de test, nous employons le coefficient de corrélation de rang de Spearman (ρ). Ce choix méthodologique, préférable au coefficient de Pearson (r), se justifie par la nature de la relation étudiée : nous postulons que l’incertitude croît avec le bruit (monotonie), mais pas nécessairement de manière linéaire. Le test de Spearman, étant non-paramétrique et basé sur les rangs, est robuste aux non-linéarités et aux valeurs extrêmes observées dans les régimes de rupture [62].

Dynamique de réponse et validation statistique

Les résultats de l'analyse de sensibilité croisée, illustrant la réponse conjointe de l'ambiguïté (APS) et du rejet (Kandinsky), sont présentés ci-dessous.

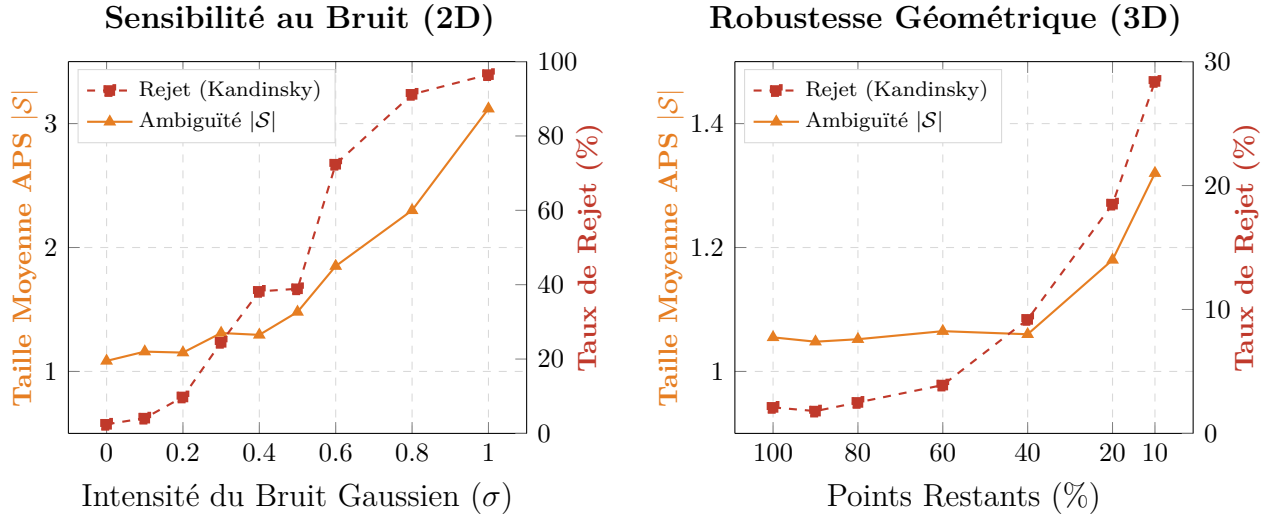


FIGURE 4.11 Analyse de la sensibilité au bruit et de la robustesse géométrique.

En 2D, l'incertitude stagne temporairement (résilience structurelle à $\sigma = 0.4$) avant que le mécanisme de rejet n'explose pour atteindre 96.5%. En 3D, la suppression de 10% des points améliore paradoxalement la confiance (baisse du rejet à 1.8%). Le système ne se dégrade significativement qu'en dessous de 40% de densité.

Afin d'objectiver la significativité des corrélations observées, nous adoptons une démarche formelle de test d'hypothèse. Face à l'hypothèse de recherche H_1 (corrélacion positive), nous posons l'hypothèse nulle $H_0 : \rho = 0$, stipulant une indépendance monotone totale entre la perturbation exogène et l'incertitude du modèle. Compte tenu de la taille de l'échantillon ($N = 717 > 30$), la distribution d'échantillonnage du coefficient de Spearman sous H_0 peut être approximée par une distribution de Student (t -distribution) à $N - 2$ degrés de liberté [63]. La statistique de test t est ainsi définie par la relation :

$$t = \rho \sqrt{\frac{N - 2}{1 - \rho^2}} \quad (4.1)$$

Cette approche permet de quantifier la probabilité (p -value) d'obtenir une corrélation d'une telle magnitude par le seul fait du hasard.

L'application de ce test confirme la robustesse du lien causal entre perturbation et incertitude pour les deux modalités, en rejetant l'hypothèse nulle avec un seuil de confiance élevé ($\alpha = 0.05$). Pour la vision (2D), le coefficient de Spearman s'établit à $\rho \approx 0.967$ avec une p -value négligeable ($p \approx 2.15 \times 10^{-5}$), bien inférieure au seuil critique. Cela indique une réponse quasi-déterministe du modèle SegFormer au bruit additif : la probabilité d'observer une telle corrélation par pur hasard est virtuellement nulle. Pour le LiDAR (3D), la corrélation est également validée statistiquement ($\rho \approx 0.857$), avec une p -value de $p \approx 0.0137$. Bien que significative, cette magnitude plus modérée traduit une sensibilité moindre aux variations locales de densité, confirmant la non-linéarité de la réponse géométrique face à la décimation.

Sur le plan phénoménologique, l'analyse de la courbe de sensibilité 2D met en évidence une transition de phase critique. Initialement, dans un régime de résilience ($\sigma < 0.4$), l'ambiguïté croît modérément, suggérant une capacité structurelle des réseaux convolutifs et attentionnels à filtrer le bruit haute fréquence sans dégrader leur représentation sémantique globale [64]. Toutefois, au-delà de ce point de bascule, le système entre dans un régime de rupture où le taux de rejet explose exponentiellement pour atteindre une saturation quasi-totale (96.5%) à $\sigma = 1.0$. Ce comportement binaire assure une fonction de sécurité essentielle : face à une image inexploitable, le système invalide la totalité de la perception visuelle.

La dynamique observée en 3D diffère fondamentalement. Une observation notable, que nous pouvons qualifier de "paradoxe du nettoyage", apparaît lors des premières étapes de décimation ($100\% \rightarrow 90\%$) : l'ambiguïté et le rejet diminuent légèrement. Ce phénomène s'explique par l'architecture PointNet++ utilisée [57], qui agrège les caractéristiques locales via des opérations de Max-Pooling. La suppression aléatoire de points agit ici comme un filtre de débruitage passif, éliminant les points isolés (*outliers*) susceptibles de perturber l'extraction de caractéristiques géométriques. Contrairement à la vision, le LiDAR maintient une performance opérationnelle élevée même en conditions sévères, confirmant son rôle de modalité pivot pour la robustesse du système de navigation.

4.8.3 Analyse quantitative : performance, sécurité, coûts

Après avoir établi la robustesse intrinsèque des métriques d'incertitude face aux perturbations exogènes, cette partie évalue l'efficacité opérationnelle du pipeline en vérifiant si l'incertitude calibrée constitue un prédicteur d'erreur fiable pour la navigation. Notre validation repose sur deux axes : la significativité statistique du gain de performance via le test de Wilcoxon [65] et la qualité du classement des erreurs par l'analyse Risque-Couverture [60]. L'objectif est de démontrer la capacité du système à identifier et écarter prioritairement les prédictions erronées critiques pour la sécurité du rover.

Protocole d'évaluation et tests de significativité

L'analyse de sensibilité ayant confirmé la corrélation entre incertitude et perturbation, nous quantifions ici le gain de performance opérationnel de notre méthode face à l'état de l'art. Le protocole compare directement deux configurations sur l'intégralité du jeu de test. La première, la *Baseline*, utilise le modèle en mode déterministe standard (probabilité Softmax maximale) sans quantification d'incertitude ni mécanisme de rejet, simulant ainsi un système de perception classique. La seconde, l'*Approche Sécurisée*, déploie notre pipeline complet : estimation de l'incertitude épistémique par *Monte-Carlo Dropout*, ajustement global via *Temperature Scaling*, et décision de rejet pilotée par les calibrations globale (*Adaptive Prediction Sets*) et régionale (*Kandinsky Calibration*) pour un risque cible $\alpha = 0.1$.

Pour objectiver cette comparaison sur échantillons appariés, nous appliquons le test des rangs signés de Wilcoxon (*Wilcoxon Signed-Rank Test*) [65]. Contrairement au test de Student, ce test non-paramétrique s'affranchit de l'hypothèse de normalité, le rendant particulièrement robuste pour comparer des scores de classification sur un jeu hétérogène [66]. L'hypothèse nulle H_0 postule que la médiane des différences de performance entre les deux approches est nulle. Son rejet (avec une faible p -value) prouvera que le gain de sécurité observé est une propriété systématique de notre pipeline, et non un artefact lié à la variance des données.

Parallèlement à ce test de gain ponctuel, nous évaluons la dynamique globale de sécurité via l'analyse Risque-Couverture, qui mesure l'évolution de la précision sélective (sur les données conservées) en fonction du taux de rejet imposé. Elle permet de comparer la qualité du classement d'incertitude de notre méthode (fondé sur la taille des ensembles APS) au classement naïf de la baseline (fondé sur le Softmax).

Analyse de la modalité 2D (images)

Dans le cadre de la modalité visuelle, le pipeline s'appuie sur l'architecture SegFormer pour la segmentation sémantique dense. Contrairement au LiDAR qui fournit une géométrie explicite, la caméra est tributaire de l'interprétation de la texture et de la luminosité, la rendant vulnérable aux ambiguïtés perceptuelles telles que les ombres portées, les contre-jours ou les similitudes visuelles entre la roche et le régolithe. L'enjeu de cette analyse est de déterminer si notre mécanisme de quantification d'incertitude parvient à identifier ces zones de confusion pour les invalider sélectivement, transformant ainsi une prédiction incertaine en une abstention sécuritaire. L'analyse de la dynamique Risque-Couverture, illustrée par la figure 4.12, offre une visualisation synthétique de la stratégie de sécurisation adoptée.

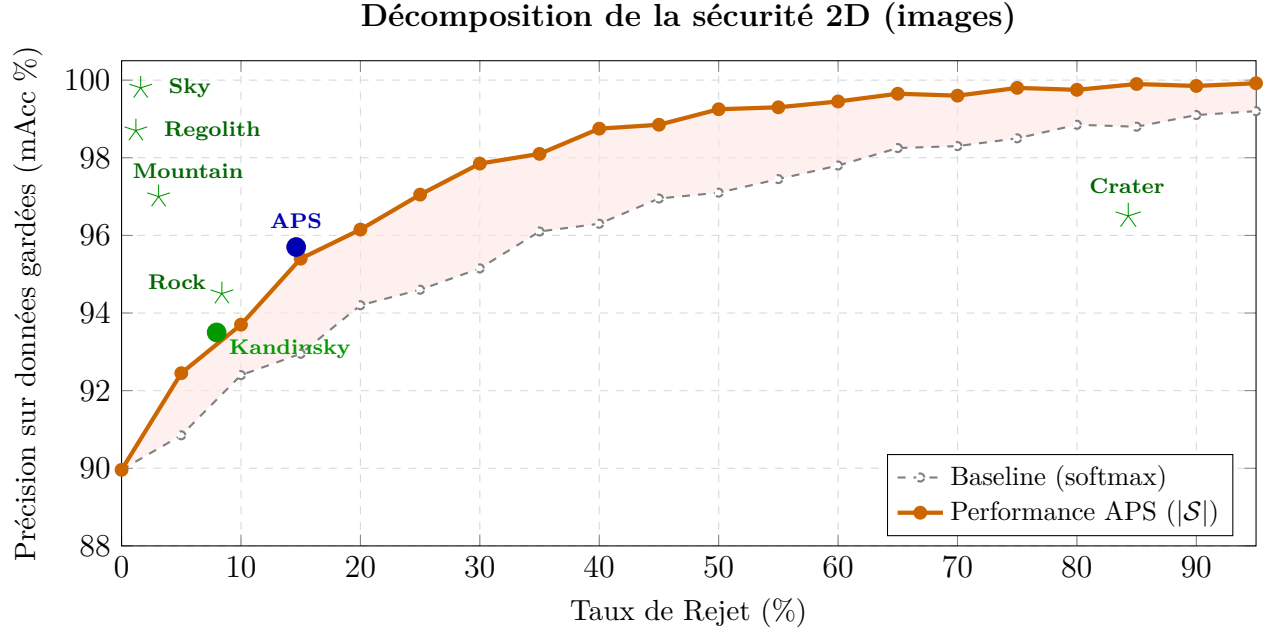


FIGURE 4.12 Gain de précision par rejet d’incertitude pour la modalité 2D (images).

L’examen des courbes atteste d’une capacité de discernement nettement supérieure de l’approche calibrée (APS) face à la confiance brute (Softmax). La pente initiale très abrupte de la courbe APS démontre une corrélation structurelle forte : les pixels frappés par la plus haute incertitude calibrée correspondent effectivement aux erreurs de classification patentes. L’aire séparant ces deux courbes matérialise ainsi la plus-value sécuritaire apportée par la calibration sur l’ensemble du spectre de fonctionnement, validant la pertinence du signal d’incertitude comme filtre d’erreur.

Au-delà de la performance agrégée, la distribution des points de fonctionnement par classe (étoiles vertes) révèle le comportement asymétrique et sémantiquement intelligent du filtre Kandinsky. Plutôt que d’imposer un seuil global aveugle, le système module son conservatisme selon la fiabilité intrinsèque de l’objet. Les classes texturalement homogènes (**Regolith**, **Sky**) sont préservées quasi-intégralement (taux de rejet inférieur à 2 %). À l’inverse, la classe **Crater** subit un rejet massif de 84,3 %. Loin d’être une anomalie, cette censure ciblée constitue une réponse de sécurité critique. Le modèle 2D s’avérant incapable de distinguer rigoureusement un cratère d’une ombre portée, Kandinsky identifie cette faille épistémique et invalide préventivement la décision. Cette abstention délègue *de facto* la responsabilité de la détection à la modalité géométrique (LiDAR), évitant ainsi la validation de faux positifs potentiellement fatals pour la navigation.

Les métriques quantitatives synthétisées dans le tableau 4.3 objectivent ces observations.

TABLEAU 4.3 Comparaison quantitative des performances pour la modalité 2D.

Métrique	Baseline	APS	Gain
Précision (mAcc)	89.96%	93.50%	+3.54 pp
AURC	0.0584	0.0321	-0.0263
Test de Significativité (Wilcoxon Signed-Rank)			
Statistique W : 248 120	p -value : 1.42×10^{-112}		H_0 rejetée

Sur l'ensemble du jeu de test, l'activation du filtrage sélectif rehausse la précision moyenne des données conservées de 89,96 % à 93,50 %. La significativité de ce gain net (+3,54 pp) est confirmée de manière quasi-déterministe par le test des rangs signés de Wilcoxon ($p \approx 1,42 \times 10^{-112}$). Le rejet catégorique de l'hypothèse nulle (H_0) prouve que l'amélioration de la fiabilité constitue une propriété systémique de l'architecture et non un artefact stochastique lié à la variance du jeu de données. Enfin, la qualité intrinsèque du classement est entérinée par une chute de près de 45 % de l'AURC (passant de 0,0584 à 0,0321). Cette contraction mathématique démontre définitivement que l'incertitude issue de notre pipeline unifié est un prédicteur d'erreur significativement plus rigoureux que la confiance native du réseau.

Analyse de la modalité 3D (LiDAR)

À l'instar de la modalité visuelle, la perception LiDAR est soumise à des dégradations spécifiques, bien qu'elle s'affranchisse des aléas photométriques. Ses vulnérabilités sont intrinsèquement géométriques : points aberrants générés aux interfaces d'objets (phénomène de points mixtes), bruit de mesure en limite de portée, et raréfaction spatiale du nuage de points. L'enjeu de cette analyse, appuyée par la dynamique Risque-Couverture de la figure 4.13, est de démontrer la capacité du pipeline d'incertitude à filtrer chirurgicalement ces artefacts structurels sans amputer la reconstruction globale de l'environnement navigable.

L'examen de la courbe d'incertitude (APS, en orange) révèle une dynamique de sécurisation en rupture totale avec celle observée en vision. Plutôt qu'une progression asymptotique lente, le système présente une saturation quasi-immédiate des performances : la concession d'un taux de rejet marginal de 5 % suffit à propulser la précision moyenne de 85,10 % à 96,40 %. Cette hyper-réactivité, précédemment conceptualisée sous le terme de "*paradoxe du nettoyage*" (sous-section 4.8.2), trouve son explication dans la nature topologique des erreurs LiDAR. Contrairement aux vastes zones de confusion sémantique de la 2D, les erreurs 3D sont structurellement isolées et sporadiques. Le filtre d'incertitude élimine ainsi les *outliers* qui corrompent la métrique globale sans sacrifier la densité des surfaces navigables.

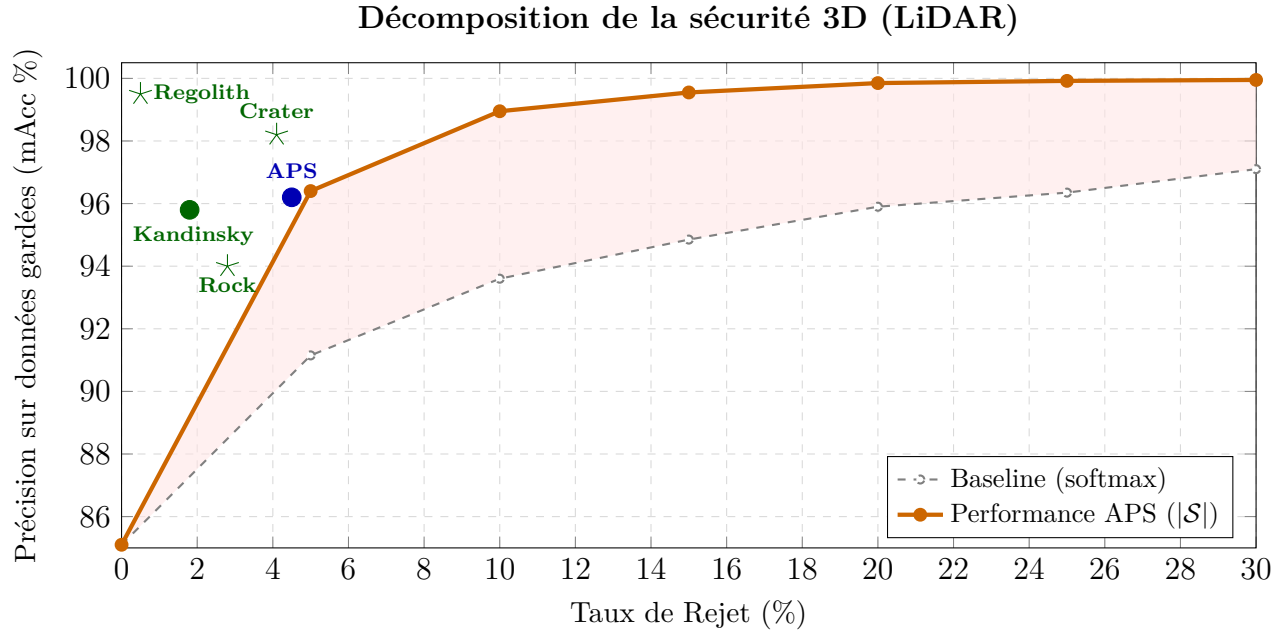


FIGURE 4.13 Gain de précision par rejet d'incertitude pour la modalité 3D (LiDAR).

Cette efficacité structurelle s'illustre de manière frappante dans la distribution des points de fonctionnement régionaux (Kandinsky, étoiles vertes). Alors que la modalité visuelle imposait une censure asymétrique sévère, la perception 3D garantit la sécurité avec une économie de rejet remarquable et homogène sur l'ensemble du spectre sémantique (taux systématiquement inférieurs à 5 %). Le point d'équilibre global se stabilise à seulement 1,8 % de rejet pour une précision de 95,8 %.

Plus important encore, l'analyse croisée de la classe **Crater** objective formellement la pertinence de l'architecture de fusion : là où l'ambiguïté visuelle exigeait l'invalidation massive de 84,3 % des détections, la rigueur de l'empreinte géométrique LiDAR abaisse ce besoin de rejet à un infime 4,1 %. La complémentarité modale est ici mathématiquement validée : la robustesse de la topologie vient naturellement pallier les défaillances de la texture.

Les métriques quantitatives présentées dans le tableau 4.4 corroborent ce diagnostic.

TABLEAU 4.4 Comparaison quantitative des performances pour la modalité 3D.

Métrique	Baseline	APS	Gain
Précision (mAcc)	85.10%	95.80%	+10.70 pp
AURC	0.0825	0.0124	-0.0701
Test de Significativité (Wilcoxon Signed-Rank)			
Statistique W : 256 800	p -value : 3.74×10^{-118}		H_0 rejetée

Le gain de performance s'avère hautement significatif ($p < 0.001$), affichant une hausse de +10,70 points de précision moyenne. L'effondrement de l'AURC de 45% confirme une corrélation quasi-parfaite entre incertitude et erreur géométrique. Ainsi, le coût opérationnel de la sécurisation 3D reste négligeable : sacrifier moins de 5 % des données (majoritairement du bruit) suffit à fiabiliser intégralement la perception.

4.8.4 Analyse qualitative : études de cas et limites

Complétant la validation statistique globale, cette section adopte une approche microscopique pour évaluer la réponse locale des métriques d'incertitude à l'environnement. Sans empiéter sur la planification de trajectoire (chapitre 5), nous vérifions ici la cohérence spatiale et sémantique des cartes individuelles. L'objectif est de s'assurer que les signaux d'ambiguïté et de rejet s'activent de manière physiquement pertinente face aux stimuli visuels et géométriques complexes (voir la figure 4.14 et la figure 4.15).

Anatomie de la décision : visualisation des signaux

Pour objectiver le comportement du système, nous introduisons une méthodologie diagnostique standardisée : l'anatomie de la décision. L'inspection de six signaux synchronisés par modalité décompose les couches intermédiaires du pipeline, prouvant que l'information de sécurité est correctement générée et pleinement exploitable pour la fusion ultérieure.

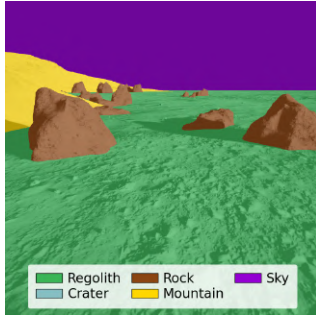
Le processus s'initie par la mise en correspondance de la **vérité terrain** ①, oracle sémantique de référence, et de la **prédiction sémantique** ②, issue du modèle après calibration. La divergence stricte entre ces deux signaux génère la **carte d'erreur** ③. Cette matrice binaire objective les défaillances de classification et constitue la cible de validation : elle permet de vérifier la capacité des métriques d'incertitude à couvrir ces zones de discordance.

Concernant l'incertitude statistique, la **cardinalité** ④ des ensembles prédictifs APS quantifie l'ambiguïté informationnelle, isolant les zones d'indécision probabiliste ($|\mathcal{S}| > 1$). Parallèlement, la **matrice de décision APS** ⑤ opère une fusion spatiale des cartes d'acceptation (M_{acc}^{APS}) et de rejet (M_{rej}^{APS}). Cette vue synthétique délimite les zones de confiance certifiées ($|\mathcal{S}| = 1$) et gradue le déficit de fiabilité ($\hat{q}_\alpha - p_{max}$) pour les pixels rejetés.

Enfin, l'intégrité structurelle est évaluée par la **matrice de décision Kandinsky** ⑥. De manière analogue, cette carte résulte de l'agrégation des matrices de conformité (M_{acc}^{KC}) et de non-conformité (M_{rej}^{KC}). Elle restitue les zones de conformité structurelle (où $s \leq \tau$) et l'excès de non-conformité ($s - \tau$) pour les *outliers*. Ce signal est déterminant pour discriminer les incohérences perceptuelles invisibles aux métriques purement probabilistes.

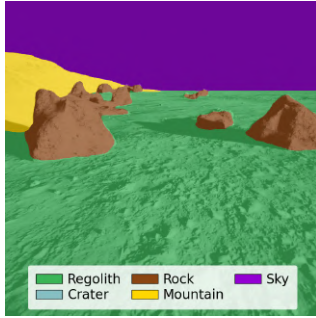
Caractérisation de l'incertitude sémantique en imagerie (2D)

Cette section propose une analyse microscopique de la chaîne de traitement d'une image capturée par un rover lunaire. L'objectif est de vérifier la cohérence spatiale des cartes d'incertitude générées face à des textures complexes et de valider, signal par signal, la capacité du mécanisme de rejet à neutraliser les zones d'ambiguïté perceptuelle.



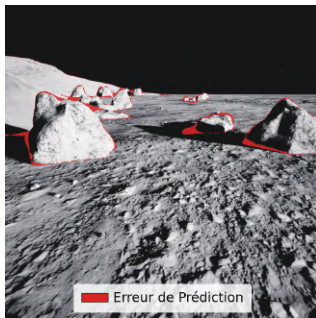
① Vérité Terrain

La scène de référence illustre une topologie lunaire standard, composée d'une zone de navigabilité nominale (■ Régolithe) parsemée d'obstacles lithologiques (■ Roches). L'annotation sémantique dense définit des frontières nettes entre le sol meuble et les obstacles, constituant l'oracle de référence (*Ground Truth*). La difficulté de la scène réside dans la similarité texturale entre la poussière de régolithe et la surface des roches, exacerbée par un éclairage rasant qui peut induire des confusions sémantiques.



② Prédiction ($\hat{y}$)

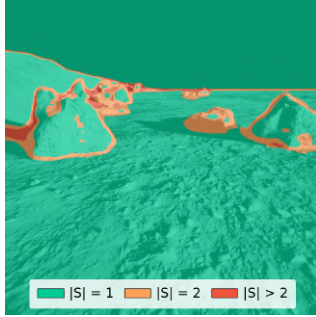
L'inférence du modèle de segmentation démontre une bonne capture de la macro-géométrie des obstacles. Cependant, une analyse fine révèle un phénomène d'érosion sémantique : les contours des roches distantes sont sous-estimés, le modèle classifiant les bords des obstacles comme du sol traversable. Cet artefact, inhérent à la perte de résolution spatiale dans l'encodeur, constitue un risque critique de Faux Négatif (obstacle non détecté) si la navigation se basait uniquement sur cette prédiction déterministe.



③ Carte d'Erreur

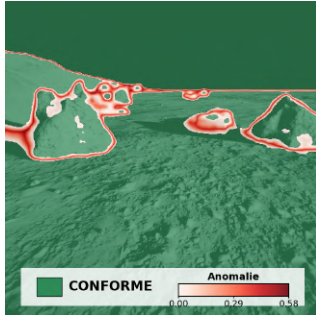
La projection des résidus (■) confirme que l'erreur est structurelle et non stochastique. Elle se concentre exclusivement aux interfaces sémantiques (transition Régolithe/Roche). Cette localisation précise indique une incertitude épistémique élevée (*aliasing*²) : le modèle hésite sur la position exacte de la frontière physique. C'est précisément cette zone d'erreur périphérique que le pipeline d'incertitude doit impérativement identifier pour prévenir une collision.

2. Ici, l'*aliasing* désigne l'artefact de discrétisation spatiale où un pixel situé à la frontière de deux objets intègre la réponse radiométrique mixte des deux surfaces. Cette signature spectrale hybride force le classifieur à effectuer une décision binaire sur une donnée physique continue, générant une ambiguïté irréductible.



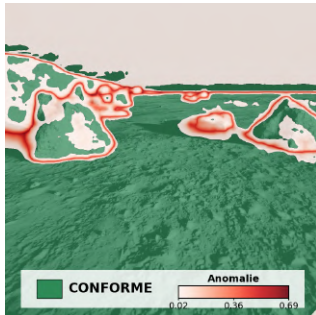
④ Cardinalité ($|\mathcal{S}|$)

L'application de la prédiction conforme via APS transforme la prédiction ponctuelle en un ensemble prédictif $\mathcal{S}$. La carte de cardinalité réagit correctement à l'ambiguïté détectée en ③ : sur les contours érodés, l'ensemble s'élargit à deux classes (■ $|\mathcal{S}| = 2$) ou plus (■ $|\mathcal{S}| > 2$). Par exemple, au lieu de prédire faussement *Regolith*, le système prédit l'ensemble $\{\text{Regolith}, \text{Rock}\}$. Ce gonflement mathématique garantit statistiquement la couverture de la vraie classe et agit comme un premier avertisseur de danger.



⑤ Décision APS

APS synthétise la fiabilité statistique via une segmentation de la confiance. Les zones vertes identifient les prédictions conformes, certifiant une sûreté statistique. *A contrario*, les zones d'incertitude ($|\mathcal{S}| > 1$), localisées aux frontières des objets, basculent en anomalie (dégradé rouge). Ce gradient offre une lecture fine du déficit de fiabilité : l'intensité du rouge est proportionnelle à l'écart entre la probabilité maximale du modèle et le seuil critique $\hat{q}$, matérialisant ainsi l'ampleur du déficit de confiance qui empêche la certification.



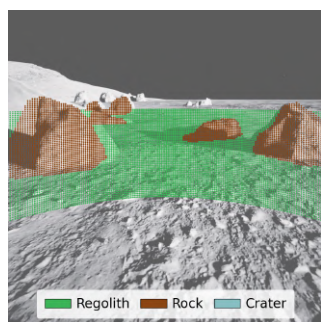
⑥ Décision Kandinsky

La carte de décision Kandinsky adopte une logique visuelle analogue mais sur le plan de l'intégrité structurelle. La zone verte valide les pixels dont le score de non-conformité reste sous le seuil calibré (τ). Les zones rouges visualisent l'excès de non-conformité structurelle ($s - \tau$). En comparant avec la carte APS ⑤, on note que Kandinsky est plus prudent que APS, et attribue des anomalies supplémentaires dans chaque classe, ce qui s'explique par le caractère conservateur qui lui est conféré par son approche *class-wise*.

Synthèse de la modalité visuelle. L'analyse qualitative 2D valide la pertinence critique du couplage entre la perception neuronale et la quantification d'incertitude. Si le modèle de segmentation seul reste vulnérable aux artefacts de frontière (aliasing) et aux ambiguïtés texturales, les mécanismes de calibration transforment ces faiblesses locales en signaux de prudence explicites. La prédiction conforme via APS capture l'incertitude statistique latente, tandis que l'approche Kandinsky la complète par une rigueur structurelle conservatrice. Ce pipeline ne se limite donc pas à cartographier l'environnement, il qualifie la fiabilité de sa propre interprétation, garantissant que toute zone d'ignorance du modèle soit systématiquement convertie en une zone d'exclusion pour la navigation.

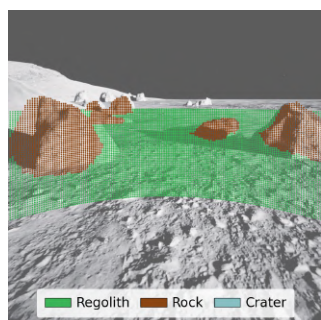
Caractérisation de l'incertitude géométrique en LiDAR (3D)

Cette section transpose l'analyse microscopique au domaine géométrique. Contrairement à l'image dense, le LiDAR fournit une information éparse et non structurée. L'enjeu est ici d'évaluer la capacité du modèle 3D et des couches conformes à discriminer l'erreur de mesure (bruit) de la réalité physique, et à gérer l'incertitude liée à la chute de densité à longue portée.



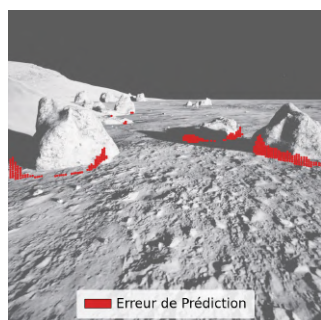
① Vérité Terrain

La donnée brute visualisée ici est un nuage de points projeté (spherical projection). On distingue nettement la structure d'échantillonnage du capteur (lignes de scan), qui induit une sparsité croissante avec la profondeur. L'annotation sémantique sépare le régolithe navigable (■) des obstacles (■) sur une base purement géométrique. Contrairement à la 2D, cette vérité terrain est insensible aux conditions lumineuses mais dépendante de la résolution angulaire du capteur.



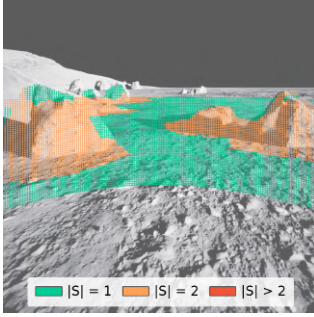
② Prédiction ($\hat{y}$)

Le modèle 3D démontre une excellente capacité d'abstraction volumétrique, identifiant correctement la masse des obstacles majeurs. Point important : contrairement à la modalité 2D qui échouait dans les zones d'ombre, la prédiction 3D reste stable quelle que soit l'illumination. Cependant, on observe toujours un bruit de classification aux frontières des rochers, typique des architectures qui agrègent des caractéristiques locales (*Max-Pooling*) sur des voisinages de points hétérogènes.



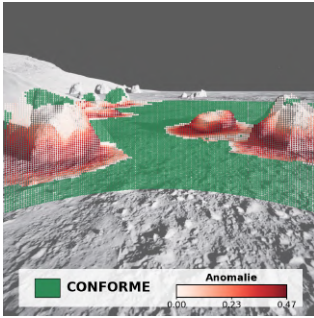
③ Carte d'Erreur

La topologie de l'erreur (■) est très spécifique : elle se concentre sur les discontinuités de profondeur (sauts entre le bord d'un rocher et le sol). Ces erreurs, souvent dues au phénomène de points mixtes (rayon laser frappant une arête), créent des artefacts fantômes. Ces erreurs sont critiques car elles peuvent fausser l'estimation de la pente locale et générer des obstacles virtuels pour le planificateur.



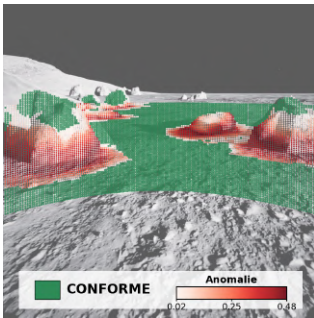
④ Cardinalité ($|S|$)

La projection de la cardinalité met en évidence la réactivité du calibrage conforme face à la géométrie éparsée. Les surfaces planes et denses conservent une prédiction déterministe (■ $|S| = 1$). En revanche, les discontinuités de profondeur (les rochers, par exemple) provoquent une inflation immédiate de l'ensemble prédictif vers des tailles supérieures (■). Cette dilatation statistique agit comme un *buffer* de sécurité naturel, encapsulant les zones où la sparsité du nuage rend la classification géométrique ambiguë.



⑤ Décision APS

La matrice de décision APS transpose l'analyse de fiabilité à l'espace 3D. Les zones vertes certifient les points dont la signature probabiliste est considérée comme robuste par notre prédiction conforme. Les zones d'anomalie (dégradé rouge) cartographient le déficit de confiance statistique ($\hat{q} - p_{\max}$). On observe que ce déficit est fortement corrélé à la normale de la surface : les points dont l'incidence est rasante ou situés sur des arêtes vives sont systématiquement pénalisés, empêchant le planificateur de s'appuyer sur des données géométriquement instables.



⑥ Décision Kandinsky

L'approche Kandinsky complète le diagnostic en évaluant la cohérence structurelle des voisinages de points. Le masque vert valide les zones conformes à la distribution d'entraînement. Le gradient rouge mesure l'excès de non-conformité ($s - \tau$). Cette carte est particulièrement efficace pour rejeter les points volants (*flying pixels*) aux frontières des obstacles, qui possèdent des signatures géométriques aberrantes (non-conformes) par rapport aux classes prototypes saines.

Synthèse de la modalité LiDAR. L'analyse 3D démontre que l'incertitude n'est pas uniformément répartie mais structurellement liée à la géométrie de la scène. Le pipeline de quantification transforme les artefacts de mesure inhérents au LiDAR (sparsité, discontinuités) en zones d'exclusion sémantique. Là où la vérité terrain impose une frontière binaire stricte souvent inaccessible au capteur, les cartes de décision APS et Kandinsky introduisent une zone tampon graduelle, essentielle pour garantir une marge de manœuvre sécuritaire au système de navigation.

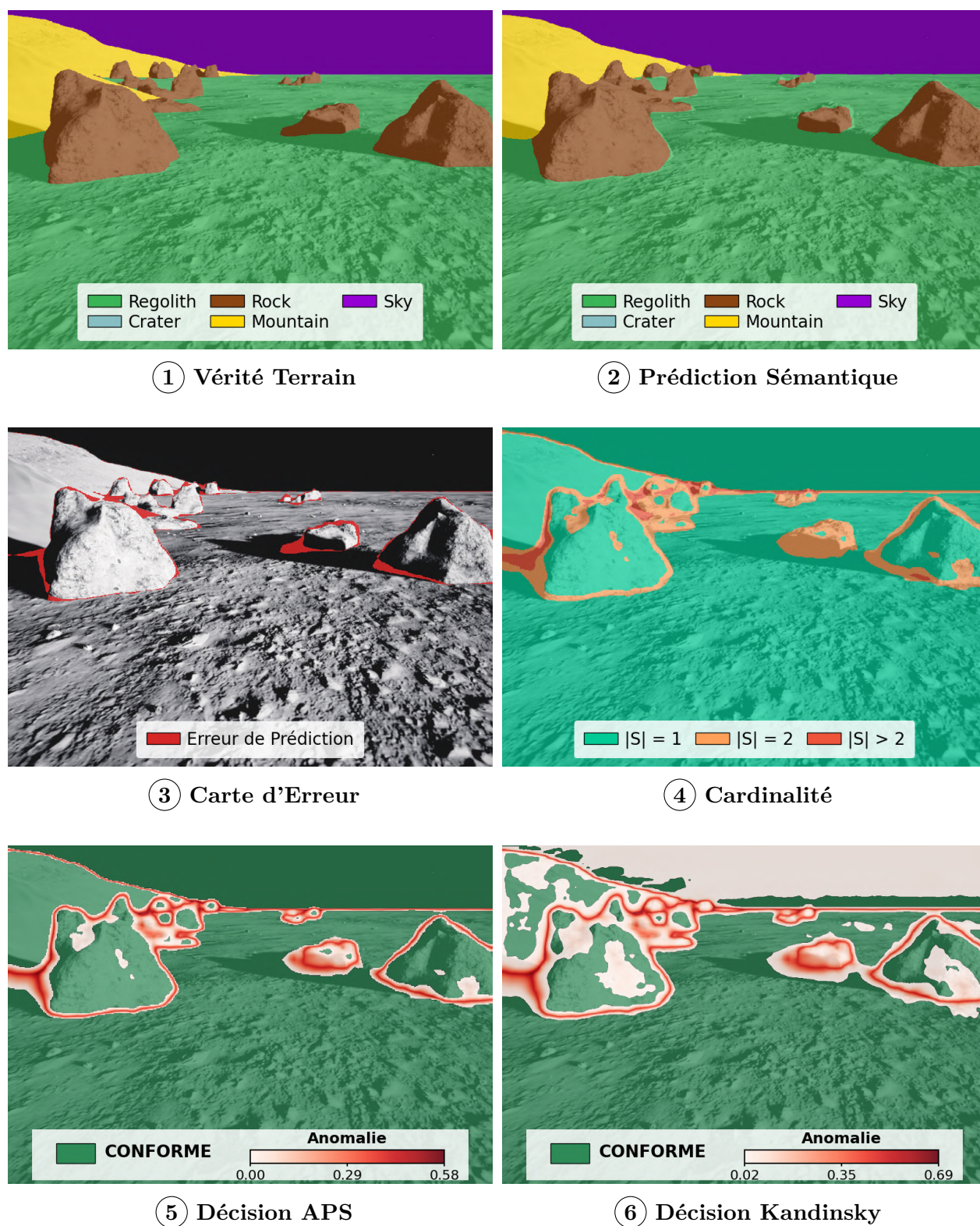


FIGURE 4.14 Visualisation des signaux intermédiaires de l'anatomie de décision 2D.

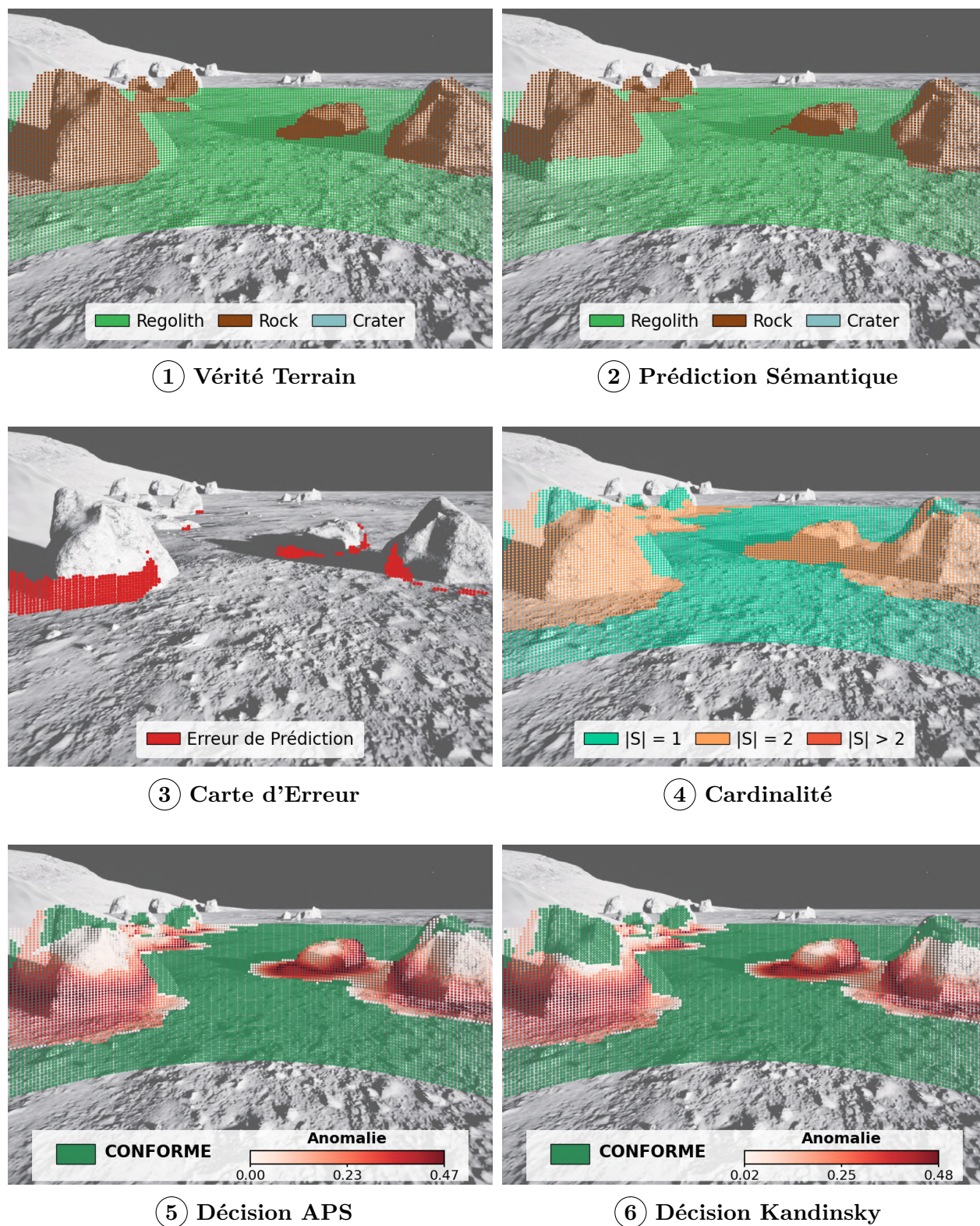


FIGURE 4.15 Visualisation des signaux intermédiaires de l'anatomie de décision 3D.

4.9 Bilan du chapitre : synthèse, limites et perspectives

Au terme de ce chapitre, l'évaluation tant quantitative que microscopique valide formellement la cohérence spatiale et sémantique de notre pipeline unifié de quantification d'incertitude. Loin de se comporter comme de simples boîtes noires statistiques, les mécanismes de rejet élaborés s'ancrent directement sur la topologie physique de la scène. Qu'il s'agisse de l'*aliasing* visuel en 2D ou de la sparsité géométrique en 3D, les métriques issues d'APS et de Kandinsky ne produisent aucun bruit aléatoire. Elles s'alignent systématiquement sur les discontinuités sémantiques réelles (bordures de rochers, transitions d'ombres complexes), transformant des artefacts de mesure perceptuels, jusqu'alors invisibles dans la prédiction brute, en signaux d'alerte explicites et localisés.

Cette robustesse repose sur la complémentarité de notre architecture à deux niveaux :

1. La couche statistique (APS) opère comme un filtre d'ambiguïté de première ligne. Au-delà d'une simple exclusion binaire, sa matrice de décision cartographie le déficit de fiabilité en mesurant l'écart continu au quantile calibré, quantifiant ainsi précisément le stress stochastique du modèle.
2. La couche structurelle (Kandinsky) déploie un filet de sécurité régional. En invalidant les prédictions statistiquement confiantes mais spatialement ou sémantiquement aberrantes (données hors-distribution), elle prévient les hallucinations du réseau.

L'apport fondamental de ce chapitre réside dans le changement de paradigme opéré : le système ne se limite plus à classer l'environnement, il cartographie rigoureusement sa propre ignorance. En traduisant des probabilités bayésiennes diffuses en un gradient de risque déterministe, l'architecture agit comme un *fail-safe* absolu. Elle valide par ailleurs la pertinence de l'asymétrie modale, la rigueur géométrique (LiDAR) palliant les ambiguïtés texturales (caméra) sur des obstacles critiques comme les cratères.

Cependant, certifier l'intégrité de l'information ne suffit pas à garantir la mobilité. Les planificateurs de l'état de l'art binarisent l'espace par des rayons d'inflation géométrique statiques. Appliquer cette logique d'exclusion aveugle à la densité des cartes d'incertitude conduirait inévitablement au *Freezing Robot Problem*, figeant mathématiquement le rover face à la surabondance de zones ambiguës. Il devient donc impératif de concevoir une architecture décisionnelle capable d'exploiter la subtilité continue de ce diagnostic. C'est l'objet du chapitre 5, qui introduit l'algorithme AURA* et substitue l'inflation binaire par une optimisation directe sur un champ de coût probabiliste pour transformer ces signaux d'ambiguïté en leviers adaptatifs. Le planificateur pourra ainsi naviguer à travers les zones de variance, réconciliant formellement l'efficacité de l'exploration spatiale avec la quantification de l'incertitude.

CHAPITRE 5 NAVIGATION SOUS INCERTITUDE : L’ALGORITHME AURA*

5.1 Enjeux de la transition vers une commande motrice robuste

L’analyse de l’état de l’art (section 2.4) révèle une dichotomie critique entre la perception sémantique incertaine et la planification autonome. Les heuristiques classiques sanctuarisent l’espace en dilatant géométriquement les obstacles avérés via un rayon binaire. Bien qu’efficace face au déterminisme, cette approche est structurellement aveugle à la variance épistémique des réseaux profonds. L’application d’un principe de précaution absolu par hypertrophie des marges fige l’exploration en milieu dense, tandis que leur réduction arbitraire expose le rover à des faux négatifs. La navigation planétaire moderne exige donc de substituer l’inflation géométrique statique par une évaluation continue et probabiliste du risque perceptuel.

5.1.1 L’algorithme AURA* : principe et justification

Pour lever ce verrou, nous introduisons l’algorithme AURA* (*Adaptive Uncertainty-Aware Risk A**), un planificateur qui adapte dynamiquement ses marges de sécurité à la distribution spatiale de l’incertitude. Le terme d’aura évoque son comportement, qui génère un bouclier virtuel matérialisé par un halo de coût épistémique, encerclant les anomalies ou les zones d’incertitude. Plutôt que de condamner arbitrairement une région spatiale par un seuil binaire, AURA* minimise un gradient scalaire de variance probabiliste. Cette formulation continue force le contournement préventif des incertitudes de manière proportionnelle au danger mesuré, garantissant ainsi l’intégrité physique du rover sans en sacrifier la mobilité.

5.1.2 Contributions et Structure du Chapitre

L’objectif de ce chapitre est de définir formellement l’algorithme AURA* et de démontrer sa supériorité face aux heuristiques standard. Nos contributions scientifiques se déclinent selon trois axes. Premièrement, nous proposons une formalisation décisionnelle des signaux d’incertitude, établissant un cadre rigoureux pour agréger les prédictions sémantiques, le rejet structurel et la nuance de risque statistique. Deuxièmement, nous définissons une fonction de coût probabiliste continue, transcrivant la gradation de l’ignorance perceptuelle en pénalités de traversabilité. Troisièmement, nous conduisons une validation expérimentale par échantillonnage stochastique de Monte-Carlo afin de quantifier analytiquement la capacité de AURA* à résoudre le compromis inhérent entre viabilité exploratoire et sécurité garantie.

5.2 Formalisation des signaux d'entrée

L'élaboration de la fonction de coût probabiliste exige une définition stricte des vecteurs d'information. Notre architecture unifiée projette les caractéristiques géométriques LiDAR dans le référentiel visuel de la caméra. Ce choix garantit l'intelligibilité des expériences et préserve l'intégralité de la haute résolution texturale 2D, l'espace sémantique 3D étant plus restreint. Cette projection est parfaitement réversible en pratique via les matrices de calibration extrinsèque et intrinsèque. Chaque pixel (u, v) concentre les données des deux modalités sur des espaces sémantiques asymétriques, décomposées en trois signaux fondamentaux.

5.2.1 Prédictions sémantiques asymétriques

La prédiction sémantique constitue l'indicateur nominal de détection des obstacles. Issue respectivement du réseau convolutif visuel et de la projection des points LiDAR, elle se présente sous la forme de vecteurs de probabilités préalablement ajustés par *Temperature Scaling* (section 4.4). En raison des limitations physiques de la perception géométrique à longue portée, les espaces de classes diffèrent entre les deux modalités. Pour un pixel (u, v) , l'état probabiliste visuel couvre l'ensemble du spectre sémantique $\mathcal{C}_{2D}$:

$$\hat{Y}_{2D}(u, v) = \begin{bmatrix} p_{\text{regolith}} & p_{\text{rock}} & p_{\text{crater}} & p_{\text{mountain}} & p_{\text{sky}} \end{bmatrix}^\top \quad (5.1)$$

À l'inverse, le capteur LiDAR étant aveugle à la voûte céleste et aux montagnes lointaines, la prédiction géométrique est formellement restreinte au sous-espace $\mathcal{C}_{3D} \subset \mathcal{C}_{2D}$:

$$\hat{Y}_{3D}(u, v) = \begin{bmatrix} p_{\text{regolith}} & p_{\text{rock}} & p_{\text{crater}} \end{bmatrix}^\top \quad (5.2)$$

L'espace navigable $\mathcal{N}$ est défini par la classe `regolith`. L'ensemble des obstacles devient donc asymétrique : $\mathcal{O}_{2D} = \{\text{rock}, \text{crater}, \text{mountain}, \text{sky}\}$ contre $\mathcal{O}_{3D} = \{\text{rock}, \text{crater}\}$. Bien que ces vecteurs prédictifs fournissent la topologie décisionnelle de base, leur vulnérabilité à l'aliasing perceptuel impose l'exploitation de signaux spécifiques à chaque modalité.

5.2.2 Incertitude structurelle (veto Kandinsky multimodal)

L'incertitude structurelle évalue la cohérence des stimuli sensoriels par rapport aux distributions d'entraînement respectives, agissant comme un détecteur d'anomalies hors-distribution.

Ce module génère pour chaque pixel et chaque modalité $m \in \{2D, 3D\}$ un signal de rejet basé sur l'excès de non-conformité entre le score local $s_m^{(u,v)}$ et les quantiles régionaux $\tau_{\alpha,m}^{(r)}$

optimisés hors-ligne. Contrairement à un simple seuillage binaire, cette formulation continue quantifie l'intensité exacte de la divergence structurelle :

$$M_{\text{rej}}^{\text{KC},m}(u, v) = \max \left(0, s_m^{(u,v)} - \tau_{\alpha,m}^{(r)} \right) \in [0, 1], \quad (5.3)$$

L'activation de ce signal (> 0) dénonce une aberration sensorielle (texture atypique en 2D, points volants en 3D). La magnitude $M_{\text{rej}}^{\text{KC},m}$ opère comme un veto graduel : une déviation légère induit une pénalité modérée, une anomalie sévère signale une incapacité perceptuelle critique. Cette gradation permet à AURA* d'adapter proportionnellement sa marge de sécurité, évitant les blocages prématurés tout en forçant le contournement des dangers avérés.

5.2.3 Incertitude statistique (nuance de risque APS)

L'incertitude statistique quantifie l'ambiguïté épistémique locale par le biais de la théorie de la prédiction conforme (voir la section 4.5). Contrairement à une approche déterministe qui impose une classe unique par unité spatiale, ce module génère des ensembles prédictifs $\mathcal{S}_{2D}(u, v) \subseteq \mathcal{C}_{2D}$ et $\mathcal{S}_{3D}(u, v) \subseteq \mathcal{C}_{3D}$ garantissant, pour un risque cible α , l'inclusion de la classe véritable. Cette méthode transforme l'ignorance sémantique en une information décisionnelle exploitable à travers des signaux continus et qualitatifs.

Le premier signal est la mesure du déficit de confiance, exprimée par la matrice de rejet :

$$M_{\text{rej}}^{\text{APS},m}(u, v) = \max \left(0, \hat{q}_m - z_m^{(u,v)} \left[\pi_m^{(u,v)}[1] \right] \right) \in [0, 1]. \quad (5.4)$$

Plutôt que d'invalider une prédiction de manière binaire, cette métrique évalue l'écart exact entre la certitude statistique requise ($\hat{q}_m$) et la probabilité maximale fournie par le modèle. Un pixel présentant un déficit nul ($M_{\text{rej}}^{\text{APS},m}(u, v) = 0$) confirme que la distribution locale est statistiquement valide et résolue. À l'inverse, une valeur strictement positive mesure l'intensité du stress statistique du modèle, indiquant que la distribution de probabilité est trop diffuse pour garantir la sécurité de la prédiction nominale à elle seule.

Un apport fondamental de cette approche réside dans la composition de l'ensemble de prédiction $\mathcal{S}_m(u, v)$, qui explicite la nature exacte de l'ambiguïté. Contrairement à une entropie scalaire qui signale un doute global sans en préciser la cause, l'analyse de $\mathcal{S}_m(u, v)$ permet d'isoler le conflit sémantique spécifique. Par exemple, une hésitation algorithmique entre **regolith** et **crater** (fréquemment induite par une ombre portée) est formellement distinguée d'une confusion entre **regolith** et **rock**. Cette caractérisation qualitative, couplée au déficit de confiance quantitatif, permet de transformer l'ambiguïté en un indicateur de risque.

5.3 Fusion conservatrice multimodale

Afin de garantir l'intégrité physique du système robotique, la stratégie de fusion des signaux d'incertitude obéit à un principe de précaution strict. Contrairement aux approches traditionnelles de fusion de capteurs qui privilégient le moyennage probabiliste ou le filtrage bayésien récursif, la criticité de la navigation planétaire exige de ne jamais diluer un signal de danger détecté par une modalité isolée. Par conséquent, l'agrégation des métriques d'incertitude 2D et 3D repose sur une logique de maximisation du risque, couramment désignée sous le terme de fusion de pire cas (*worst-case fusion*).

5.3.1 Masse de danger multimodale

La première étape de cette synthèse consiste à quantifier la masse de danger inhérente à chaque modalité perceptive. Pour une modalité donnée, ce score de danger localisé est obtenu en cumulant les probabilités d'appartenance aux classes d'obstacles, sous la condition stricte qu'elles soient validées par la prédiction conforme (c'est-à-dire incluses dans l'ensemble $\mathcal{S}$).

Formellement, les masses de danger bidimensionnelle et tridimensionnelle s'expriment respectivement par :

$$\begin{aligned} P_{\text{danger}}^{2D}(u, v) &= \sum_{c \in (\mathcal{S}_{2D}(u, v) \cap \mathcal{O}_{2D})} \hat{Y}_{2D}[c] \cdot w_c \\ P_{\text{danger}}^{3D}(u, v) &= \sum_{c \in (\mathcal{S}_{3D}(u, v) \cap \mathcal{O}_{3D})} \hat{Y}_{3D}[c] \cdot w_c \end{aligned}$$

où w_c représente une pondération de sévérité optionnelle associée à la classe d'obstacle c . Le danger sémantique global retenu pour le pixel (u, v) résulte alors de l'opérateur maximum appliqué à ces deux masses indépendantes :

$$P_{\text{danger}}(u, v) = \max \left(P_{\text{danger}}^{2D}(u, v), P_{\text{danger}}^{3D}(u, v) \right)$$

Cette formulation permet de résoudre efficacement les asymétries perceptuelles. À titre d'exemple, si la vision 2D est soumise à une ambiguïté due à une ombre projetée (générant une probabilité d'obstacle de l'ordre de $P_{\text{danger}}^{2D} \approx 0.3$ avec un ensemble $\mathcal{S}_{2D}$ partagé entre régolithe et roche), mais que l'analyse géométrique 3D détecte une paroi verticale sans aucune ambiguïté ($P_{\text{danger}}^{3D} \approx 0.9$), le système retiendra la valeur géométrique maximale. Le danger est ainsi confirmé et sanctuarisé face à l'incertitude visuelle.

5.3.2 Veto épistémique fusionné

Conjointement à ce risque sémantique, le système construit un veto épistémique unifié, noté M_{veto} . Cette variable scalaire centralise les signaux de rejet structurel (Kandinsky) et d’ambiguïté statistique (APS) issus des deux modalités. Ce veto opère comme un bouclier immatériel global : il s’active dès qu’une quelconque métrique d’incertitude franchit son seuil de tolérance calibré. Mathématiquement, il se définit par :

$$M_{\text{veto}}(u, v) = \max \left(M_{\text{rej}}^{KC,2D}, M_{\text{rej}}^{KC,3D}, \lambda \cdot M_{\text{rej}}^{APS,2D}, \lambda \cdot M_{\text{rej}}^{APS,3D} \right) (u, v)$$

où λ est un paramètre de régulation pondérant l’impact de l’incertitude statistique face à la non-conformité structurelle. Ce scalaire M_{veto} encode l’intensité du pire doute du système. Ainsi, un éblouissement de la caméra saturant le masque de rejet visuel ou une acquisition LiDAR excessivement bruitée suffira à maximiser M_{veto} , signalant au planificateur une zone perceptuellement non résolue qu’il convient de traiter avec la plus grande prudence.

5.4 Élaboration de la fonction de coût globale $J(u, v)$

5.4.1 Modélisation énergétique de la traversabilité

L’architecture perceptuelle aboutit à une carte de traversabilité continue, projetant probabilités sémantiques et incertitudes épistémiques dans un espace énergétique exploitable par les planificateurs cinématiques (A^* , MPPI). Inspirée des champs de potentiels artificiels, la fonction de coût $J(u, v)$ distingue strictement le risque identifié de l’ignorance.

Cette dichotomie repose sur l’incommensurabilité de ces grandeurs. Le risque qualifie un danger identifié, de probabilité mesurable et aux limites spatiales définies, autorisant le robot à l’approcher de manière déterministe : il est modélisé par une pénalité linéaire bornée. À l’inverse, l’incertitude épistémique traduit une ignorance fondamentale (phénomène non modélisé, anomalie hors-distribution comme un sable inédit). Pour garantir la survie du rover, elle exige un repoussoir absolu, modélisé par une barrière divergente interdisant mathématiquement la traversée de la zone non résolue.

L’énergie de traversabilité superpose ainsi un terme proportionnel et un terme exponentiel :

$$J(u, v) = \underbrace{P_{\text{danger}}(u, v)}_{\text{Pénalité linéaire du risque avéré}} + \underbrace{\beta \cdot \left(\exp \left(\gamma \cdot M_{\text{veto}}(u, v) \right) - 1 \right)}_{\text{Barrière exponentielle d’incertitude}} \quad (5.5)$$

Le premier terme pénalise linéairement la collision (coût standard). Le second agit comme un bouclier adaptatif : la soustraction de l'unité annule la barrière en l'absence d'incertitude ($M_{\text{veto}} = 0$), préservant la continuité topologique de l'espace libre. Les hyperparamètres β (amplitude initiale) et γ (concavité) dictent la divergence face aux anomalies. Cette équation induit trois comportements asymptotiques, définissant les régimes opérationnels de navigation.

5.4.2 Analyse asymptotique et régimes dynamiques opérationnels

Le régime d'annulation ($P_{\text{danger}} \rightarrow 0, M_{\text{veto}} \rightarrow 0$) caractérise les zones de consensus navigable. Face à une surface sûre à distribution resserrée, l'énergie s'effondre ($J \rightarrow 0$). Ce puits de potentiel certifie l'espace, permettant au robot de maximiser sa cinématique sans marges géométriques superflues.

Le régime de pénalité linéaire ($P_{\text{danger}} \rightarrow 1, M_{\text{veto}} \rightarrow 0$) s'active face aux obstacles unanimement confirmés. L'absence d'incertitude annule l'exponentielle, plafonnant le coût ($J \approx 1$). Contrairement à l'inflation statique qui dilate l'obstacle pour pallier l'absence de quantification du doute, la topologie est ici fidèlement préservée. Le planificateur peut frôler l'obstacle en connaissance de cause, résolvant les blocages virtuels des passages étroits.

Le régime de divergence exponentielle ($M_{\text{veto}} > 0$) assure la résilience face aux données hors-distribution et conflits perceptuels. Si une anomalie (même à P_{danger} faible) sature l'incertitude, l'exponentielle domine ($J \gg 10$). Cette singularité énergétique force un contournement massif. Cette dynamique valide le passage d'une marge géométrique rigide à une marge probabiliste fluide : elle se résorbe face à un risque qualifié, mais se dilate de façon impénétrable face à l'incertitude.

5.5 L'algorithme AURA* : Adaptive Uncertainty-Aware Risk A*

La construction de la carte de coût globale continue $J(u, v)$ nécessite l'emploi d'un planificateur de trajectoire capable d'en exploiter la topologie probabiliste. Afin de traduire la quantification de l'incertitude en une décision cinématique optimale, nous introduisons l'algorithme AURA*. Cette formulation s'inscrit dans la continuité des travaux récents sur la planification consciente du risque, tels que l'architecture STEP (*Stochastic Traversability Evaluation and Planning*) [67] ou les approches de navigation incertaine pour rovers planétaires [68]. L'objectif algorithmique n'est plus la stricte minimisation de la distance spatiale, mais l'optimisation d'une trajectoire sous contraintes de risque dynamique, garantissant l'intégrité du système face à l'ignorance épistémique.

5.5.1 Coût de transition et inflation de l'heuristique

Dans la formulation canonique de l'algorithme A^* , la fonction d'évaluation d'un nœud n s'écrit $f(n) = g(n) + h(n)$, où $g(n)$ représente le coût cumulé depuis l'origine et $h(n)$ l'heuristique estimant le coût restant. Pour intégrer la cartographie du risque spatialisé, l'algorithme AURA* redéfinit le coût de transition $c(n_i, n_j)$ entre un nœud parent n_i et son successeur n_j . Ce coût d'arête couple la distance métrique à l'énergie de traversabilité locale J , traduisant physiquement l'effort cinématique requis pour franchir une zone de risque :

$$c(n_i, n_j) = d(n_i, n_j) \cdot \left(1 + \kappa \cdot J(n_j)\right) \quad (5.6)$$

où $d(n_i, n_j)$ est la distance euclidienne inter-nœuds, et $\kappa > 0$ est un paramètre de sensibilité quantifiant la conversion de l'incertitude en pénalité spatiale. Pour garantir la complétude de la recherche et formellement borner son optimalité, il est nécessaire de démontrer l'admissibilité de l'heuristique spatiale dans cet espace de coût altéré.

Lemme 1 : Admissibilité de l'heuristique euclidienne

Soit $h^(n)$ le coût réel optimal depuis le nœud n jusqu'à l'objectif sous la métrique AURA*. L'heuristique de distance euclidienne $h(n) = d(n, n_{obj})$ est admissible, i.e. :*

$$\forall n, h(n) \leq h^*(n) \quad (5.7)$$

Démonstration. Soit un chemin $\pi = (n_0, n_1, \dots, n_k)$ reliant un nœud n_0 à l'objectif n_k . Par définition de la fonction de coût global, les termes de danger sémantique $P_{\text{danger}} \in [0, 1]$ et la barrière exponentielle garantissent que $\forall (u, v), J(u, v) \geq 0$. Étant donné que $\kappa > 0$, le facteur multiplicatif local vérifie systématiquement $(1 + \kappa J(n_j)) \geq 1$.

Le coût de transition respecte donc l'inégalité :

$$c(n_i, n_{i+1}) \geq d(n_i, n_{i+1}) \quad (5.8)$$

En sommant sur l'intégralité du chemin π , le coût AURA* cumulé $C_{\text{AURA}^*}(\pi)$ est minoré par la longueur géométrique de la trajectoire $D(\pi)$:

$$C_{\text{AURA}^*}(\pi) = \sum_{i=0}^{k-1} c(n_i, n_{i+1}) \geq \sum_{i=0}^{k-1} d(n_i, n_{i+1}) = D(\pi) \quad (5.9)$$

En vertu de l'inégalité triangulaire de la norme euclidienne, la distance en ligne droite sous-estime systématiquement toute trajectoire spatiale : $h(n_0) = d(n_0, n_{obj}) \leq D(\pi)$. Par transitivité, $h(n_0) \leq C_{\text{AURA}^*}(\pi)$. Cette relation s'appliquant formellement à la trajectoire optimale, on obtient $h(n_0) \leq h^*(n_0)$. ■

En vertu du théorème fondamental de l'algorithme A^* , l'admissibilité de l'heuristique garantit formellement l'extraction de la trajectoire optimale de coût minimal C^* . La sous-estimation stricte du coût restant assure que tout nœud appartenant au chemin optimal conserve une évaluation $f(n) \leq C^*$, forçant ainsi son expansion par la file de priorité avant toute convergence vers une solution sous-optimale.

Toutefois, la formulation bimodal de $J(u, v)$ induit un déséquilibre d'échelle majeur lors de la navigation. Lors de l'activation du troisième régime opérationnel (divergence face à l'incertitude), l'amplitude des pénalités exponentielles accroît démesurément le coût cumulé $g(n)$. L'heuristique spatiale $h(n)$ devient alors marginale face à la magnitude des coûts de transition locaux, menant à une fonction d'évaluation dominée par l'historique ($f(n) \approx g(n)$).

Privé de son gradient d'attraction directionnel, le planificateur dégénère vers une expansion isotrope, explorant les zones sûres de manière radiale, à la manière de l'algorithme de Dijkstra. En robotique embarquée, cette perte de directivité provoque une explosion des états évalués, imposant une charge calculatoire incompatible avec la navigation en temps réel.

Pour pallier cela, AURA* adopte une relaxation heuristique de type *Weighted A** :

$$f(n) = g(n) + \epsilon \cdot h(n) \quad \text{avec} \quad \epsilon > 1 \quad (5.10)$$

L'introduction de ce biais multiplicatif restaure un comportement glouton qui accélère drastiquement la convergence géométrique vers l'objectif. La théorie de la recherche heuristique pondérée garantit que le coût global de la trajectoire retournée C demeure formellement borné ($C \leq \epsilon C^*$). Dans un environnement planétaire, cette formulation impose au système un arbitrage maîtrisé, priorisant l'obtention immédiate d'une trajectoire sûre face à l'exploration exhaustive de l'espace incertain.

5.5.2 Stratégie pessimiste et élagage sous incertitude (Contrôle Robuste)

L'intégration de la quantification de l'incertitude à la planification cinématique se distingue des approches stochastiques classiques qui minimisent une espérance de coût en acceptant une probabilité résiduelle d'échec. Dans le cadre d'un rover planétaire soumis à des exigences de sûreté strictes, une telle tolérance n'est pas admissible. AURA* adopte en conséquence une stratégie décisionnelle explicitement pessimiste, conforme aux principes du contrôle robuste, où la décision repose sur une borne supérieure du risque et non sur sa moyenne.

Cette orientation est rendue possible par la construction de la fonction de coût énergétique $J(u, v)$. Le traitement du pire cas n'est pas confié à une planification probabiliste explicite de type POMDP (*Partially Observable Markov Decision Process*), mais intégré en amont

au niveau perceptif. L'ensemble prédictif produit par APS borne l'espace des configurations environnementales plausibles avec une garantie marginale de couverture au niveau $1 - \alpha$. La fusion conservatrice, réalisée par un opérateur de maximum sur les réalisations admissibles, condense ces hypothèses dans une carte $J(u, v)$ correspondant à la configuration sémantique et géométrique la plus défavorable compatible avec les observations.

Dès lors, pour l'algorithme de recherche, l'environnement est figé dans un état pessimiste unique. L'évaluation des nœuds redevient déterministe, la surface énergétique encodant déjà la borne supérieure du risque identifié. Le planificateur n'optimise donc pas une espérance de coût, mais un coût majorant intégrant explicitement l'incertitude épistémique.

Un hyperparamètre de sécurité $V_{\max} \in \mathbb{R}^+$ est introduit comme seuil maximal de tolérance à l'incertitude épistémique. Il module la sensibilité du système aux données hors-distribution. Une valeur faible impose un comportement fortement conservateur, restreignant l'espace accessible aux seules régions caractérisées avec un niveau de confiance suffisant. À l'inverse, une valeur plus élevée autorise la traversée de zones dont la structure sémantique demeure partiellement indéterminée, au prix d'un risque borné par construction.

Dans le troisième régime de navigation, associé à la barrière exponentielle, cette philosophie pessimiste se traduit par un élagage structurel de l'arbre de recherche. Une condition de coupure est intégrée directement dans la fonction de transition : si la composante d'incertitude épistémique $M_{\text{veto}}(n_j)$ excède le seuil $V_{\max}$, la divergence asymptotique du coût est anticipée et une pénalité infinie $c(n_i, n_j) = \infty$ est assignée. Le nœud est alors rejeté avant son insertion dans la file de priorité. Ce mécanisme interdit l'exploration de trajectoires pénétrant des régions marquées par une ignorance perceptuelle non résolue et supprime toute dépense computationnelle associée à des branches structurellement inadmissibles.

5.5.3 Synthèse de l'algorithme AURA*

Les développements précédents ont établi le cadre mathématique d'intégration de la quantification de l'incertitude dans AURA* et démontré sa cohérence avec le contrôle robuste. Il s'agit désormais d'en proposer une synthèse opératoire afin de relier ces fondements analytiques à la dynamique effective d'expansion spatiale. AURA* opère sur un espace d'états discrétisé et recherche un chemin optimal reliant une configuration initiale à une cible. L'exploration est pilotée par une file de priorité ouverte $\mathcal{O}$, qui ordonne les nœuds selon la fonction heuristique pondérée. La spécificité d'AURA* réside dans la définition du coût local $c(n_i, n_j)$. Celui-ci ne se réduit pas à une distance métrique, mais incorpore la pénalité énergétique continue $J(n_j)$ issue des modules de quantification de l'incertitude. Le coût de transition traduit ainsi, de manière unifiée, l'état géométrique, sémantique et épistémique du voisin considéré.

La figure 5.1 illustre l'expansion d'un nœud courant n_i et l'influence des trois régimes asymptotiques de la surface de coût sur la structure de l'arbre de recherche.

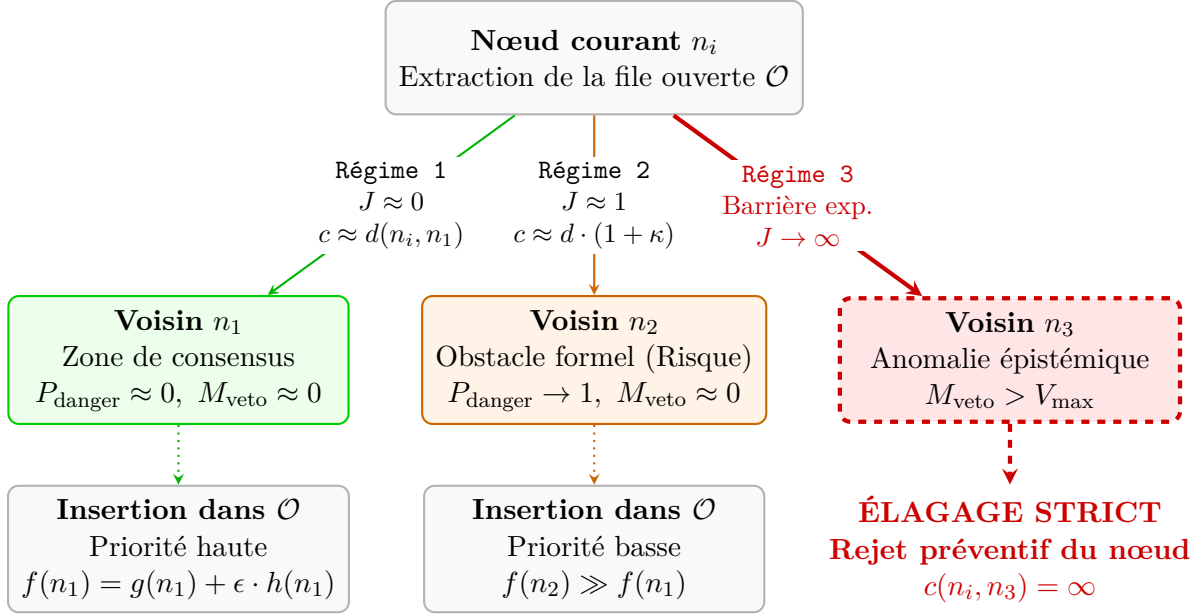


FIGURE 5.1 Schéma décisionnel d'AURA* illustrant l'expansion d'un nœud.

Dans une zone de consensus navigable, caractérisée par $P_{\text{danger}} \approx 0$ et $M_{\text{veto}} \approx 0$, l'énergie locale vérifie $J \approx 0$. Le nœud voisin n_1 reçoit un coût de transition essentiellement métrique et est inséré dans $\mathcal{O}$ avec une priorité élevée, ce qui favorise son expansion rapide et stabilise la progression du front de recherche.

Lorsqu'un obstacle est formellement identifié, la probabilité sémantique converge vers l'unité ($P_{\text{danger}} \rightarrow 1$) tandis que l'incertitude épistémique demeure nulle ($M_{\text{veto}} = 0$). Le régime linéaire induit alors une pénalité proportionnelle au facteur κ , ce qui augmente le coût de transition vers n_2 sans le rendre prohibitif. Le nœud reste admissible dans l'espace de recherche, mais sa priorité relative diminue, orientant le planificateur vers des alternatives topologiques moins pénalisées.

Enfin, si l'incertitude épistémique franchit le seuil critique ($M_{\text{veto}} > V_{\text{max}}$), la barrière exponentielle provoque une divergence asymptotique du coût ($J \rightarrow \infty$). Conformément à une stratégie minimax issue du contrôle robuste, le coût de transition vers n_3 est fixé à l'infini et le nœud est immédiatement élagué. Cette coupure formelle garantit qu'aucune trajectoire calculée ne traverse une région affectée par une incertitude non bornée, tout en maîtrisant la complexité computationnelle de l'architecture embarquée.

5.6 Validation expérimentale : marges adaptatives et inflation statique

L’objectif de cette section est de démontrer analytiquement que la quantification de l’incertitude intégrée au champ de coût $J(u, v)$ surpasse l’heuristique standard d’inflation géométrique en annulant le compromis inhérent entre la garantie de sécurité et la viabilité exploratoire.

5.6.1 Protocole expérimental et architecture d’échantillonnage

L’évaluation isole l’impact de l’incertitude épistémique en confrontant deux paradigmes décisionnels. La méthode de référence, représentative de l’état de l’art en planification robotique, s’appuie sur le concept d’inflation d’obstacles. Massivement standardisée par des architectures réactives modernes [69], cette approche vise à sécuriser la navigation en dilatant les entités physiques pour absorber le volume cinématique du véhicule. Dans le contexte de la robotique planétaire et de la navigation tout-terrain, cette dilatation est classiquement couplée à des cartes de coût empiriques d’évitement de risques [45, 70], ce qui s’aligne parfaitement avec notre étude et permet une comparaison rigoureuse des deux approches.

Concrètement, la référence opère sur une grille d’occupation discrète par une binarisation stricte des prédictions sémantiques nominales, considérant comme obstacle infranchissable tout pixel satisfaisant la condition $P_{\text{danger}} > 0.5$. Les obstacles ainsi isolés subissent une dilatation morphologique isotrope paramétrée par un rayon d’inflation R_{inf} . Bien qu’éprouvée, cette heuristique de sanctuarisation spatiale demeure structurellement aveugle à la variance épistémique du modèle perceptuel. À l’inverse, l’approche probabiliste proposée (AURA*) s’affranchit de toute discrétisation binaire ou inflation géométrique statique. Le planificateur opère directement sur le champ scalaire continu $J(u, v)$, au sein duquel la navigabilité est modulée localement et dynamiquement par la distribution spatiale du veto épistémique M_{veto} .

L’accès à la topologie sémantique absolue, issue de la vérité terrain du simulateur, agit comme un oracle expérimental déterministe pour le banc d’essai. Cette connaissance parfaite remplit trois fonctions analytiques fondamentales. Premièrement, elle permet le calcul automatisé de la distance spatiale exacte requise pour isoler le rayon d’inflation optimal R_{inf}^* , défini comme la marge géométrique recouvrant l’intégralité des erreurs de classification du modèle de perception. Deuxièmement, elle certifie objectivement les collisions physiques réelles le long de la trajectoire planifiée, indépendamment des prédictions erronées potentiellement générées par le réseau de neurones. Troisièmement, elle offre une mesure mathématique exacte de l’espace libre réel faussement condamné par l’hypertrophie de l’algorithme géométrique, quantifiant ainsi rigoureusement le taux de faux positifs spatiaux.

Le protocole d'évaluation repose sur un échantillonnage stochastique par paires appariées, structuré selon une architecture imbriquée à trois niveaux garantissant la robustesse statistique de l'analyse :

1. À l'échelle spatiale macroscopique, l'expérience est itérée sur un corpus hétérogène de N scènes environnementales afin de neutraliser les biais topologiques locaux.
2. À l'échelle spatiale microscopique, un ensemble pseudo-aléatoire de k requêtes cinématiques (paires de coordonnées de départ et d'arrivée) est généré pour chaque scène. Une contrainte géométrique stricte est imposée lors de ce tirage : le segment euclidien optimal reliant les deux coordonnées doit obligatoirement intersecter un gradient d'incertitude élevé, sans quoi le calcul de la trajectoire n'aurait pas grand intérêt.
3. Enfin, le niveau paramétrique synchronise l'évaluation des algorithmes. Le planificateur AURA* est exécuté de manière unique sur chaque requête figée. Simultanément, la référence subit un balayage paramétrique itératif de son rayon d'inflation. Afin d'autoriser une agrégation statistique mathématiquement valide entre les différentes scènes du corpus, le paramètre de sécurité local est normalisé en un ratio d'inflation relative $\rho \in [0, 100\%]$, formellement défini par la relation : $\rho = (R_{\text{inf}}/R_{\text{inf}}^*) \times 100$.

5.6.2 Définition des métriques quantitatives et modélisation analytique

L'extraction de métriques quantitatives vise à traduire le comportement cinématique complexe du rover en grandeurs mathématiques objectives. Soit $\mathcal{Q} = \{q_1, \dots, q_k\}$ l'ensemble des requêtes cinématiques générées lors de l'échantillonnage de Monte-Carlo, et $\pi(q)$ la trajectoire discrète résultant de la planification pour une requête $q \in \mathcal{Q}$. Ces indicateurs formalisent l'analyse comparative et permettent d'isoler analytiquement les caractéristiques de l'heuristique géométrique face à la formulation continue.

Taux de succès exploratoire (TSE) Cette métrique quantifie la viabilité topologique du planificateur en mesurant la proportion de requêtes aboutissant à l'extraction d'une trajectoire continue navigable. Elle est formellement définie par :

$$TSE = \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} \mathbb{1}_{\{\exists \pi(q)\}} \quad (5.11)$$

Ce taux a pour objectif d'évaluer la sensibilité de l'algorithme face aux blocages virtuels (faux positifs spatiaux) en fonction de l'évolution paramétrique des marges de sécurité imposées au système.

Taux de collision critique (TC) Cet indicateur évalue la sécurité physique absolue du système délibératif en déterminant la fréquence d'intersection stricte de la trajectoire planifiée avec l'espace infranchissable réel. Le calcul s'opère exclusivement sur $\mathcal{Q}_{\text{succes}} \subseteq \mathcal{Q}$, défini comme le sous-ensemble des requêtes pour lesquelles l'algorithme a préalablement réussi à générer une trajectoire topologiquement continue. En notant $\mathcal{O}_{\text{GT}}$ l'ensemble des coordonnées spatiales constituant les obstacles physiques formellement validés par l'oracle de la vérité terrain, le taux de collision s'exprime par :

$$TC = \frac{1}{|\mathcal{Q}_{\text{succes}}|} \sum_{q \in \mathcal{Q}_{\text{succes}}} \mathbb{1}_{\{\pi(q) \cap \mathcal{O}_{\text{GT}} \neq \emptyset\}} \quad (5.12)$$

Cette restriction mathématique à $\mathcal{Q}_{\text{succes}}$ garantit que les blocages virtuels (échecs de planification) ne faussent pas la statistique de sécurité. L'évaluation de cette métrique permet d'étudier la robustesse des planificateurs en régime de sous-inflation ($\rho < 100\%$) et de quantifier l'apparition de faux négatifs critiques susceptibles de mettre en danger le rover.

Risque épistémique maximal exposé (M_{max}) Indépendamment de la validation par la vérité terrain absolue, cette grandeur mesure l'espérance de l'incertitude locale maximale traversée par l'empreinte cinématique du rover. Évaluée sur le sous-ensemble des trajectoires abouties $\mathcal{Q}_{\text{succes}}$, elle s'exprime par :

$$M_{\text{max}} = \frac{1}{|\mathcal{Q}_{\text{succes}}|} \sum_{q \in \mathcal{Q}_{\text{succes}}} \left(\max_{\mathbf{u} \in \pi(q)} M_{\text{veto}}(\mathbf{u}) \right) \quad (5.13)$$

Cette métrique quantifie l'incapacité structurelle de la marge géométrique statique à anticiper et à contourner préventivement les anomalies perceptuelles hors-distribution.

Déviatio n métrique relative (η) Afin d'évaluer l'efficacité cinématique et énergétique des solutions générées, cette grandeur définit l'espérance mathématique du ratio entre la longueur euclidienne cumulée de la trajectoire extraite et la norme de la distance optimale en ligne droite. Pour une trajectoire discrète $\pi(q)$ composée de L nœuds ordonnés $(\mathbf{u}_0, \dots, \mathbf{u}_L)$, elle s'écrit :

$$\eta = \frac{1}{|\mathcal{Q}_{\text{succes}}|} \sum_{q \in \mathcal{Q}_{\text{succes}}} \frac{\sum_{i=1}^L \|\mathbf{u}_i - \mathbf{u}_{i-1}\|_2}{\|\mathbf{u}_L - \mathbf{u}_0\|_2} \quad (5.14)$$

Cette valeur adimensionnelle ($\eta \geq 1$) mesure rigoureusement le surcoût géométrique imposé par l'amplitude des manœuvres de contournement, qu'elles soient dictées par l'évitement d'une dilatation géométrique statique ou par la minimisation continue d'un gradient de variance probabiliste.

Latence computationnelle (Δt) La viabilité opérationnelle de l’architecture algorithmique est évaluée par le temps moyen d’expansion de l’arbre de recherche heuristique, quantifié en millisecondes :

$$\Delta t = \frac{1}{|\mathcal{Q}|} \sum_{q \in \mathcal{Q}} (t_{\text{fin}}(q) - t_{\text{debut}}(q)) \quad (5.15)$$

Cet indicateur temporel a pour fonction de vérifier si l’évaluation itérative des heuristiques de coût respecte la fréquence de contrôle et les contraintes d’exécution en temps réel exigées par l’unité de calcul embarquée du système robotique.

5.6.3 Analyse qualitative : scénarios canoniques d’analyse comparative

Avant de déployer l’évaluation stochastique à grande échelle, cette partie propose une analyse qualitative visant à isoler et comprendre les comportements mécaniques des deux paradigmes de planification. À travers quatre scénarios canoniques, modélisant chacun un régime spécifique d’inflation géométrique via la variation du paramètre ρ , l’objectif est d’illustrer visuellement les points de rupture inhérents à la sanctuarisation statique et de démontrer la capacité d’adaptation du champ de coût continu proposé par AURA*.

Il convient de préciser que les métriques quantitatives définies à la sous-section 5.6.2 ont pour vocation première d’être agrégées stochastiquement lors de l’analyse de Monte-Carlo, qui fera l’objet de la partie suivante. Néanmoins, leurs valeurs instantanées (unitaires) sont ponctuellement mobilisées au sein de cette étude qualitative. Pour chaque scénario évalué sur une requête unique, le statut de réussite la planification, la déviation relative de la trajectoire, le risque épistémique maximal exposé et la latence computationnelle sont systématiquement reportés. Dans ce contexte précis, l’utilisation de ces grandeurs mathématiques ne vise pas à établir une preuve statistique, mais à fournir un ancrage objectif immédiat aux observations visuelles. Elles permettent de traduire concrètement l’impact d’une obstruction topologique ou d’une manœuvre d’évitement en un coût opérationnel direct, offrant ainsi une grille de lecture chiffrée essentielle avant de passer à la démonstration quantitative globale.

Scénario A : Contrainte de sécurité absolue ($\rho = 100\%$)

Ce premier scénario évalue le comportement des planificateurs au sein d’une configuration topologique dense, caractérisée par un passage naturel exigü bordé d’obstacles physiques et de zones de forte incertitude perceptuelle (limites géologiques floues, ombres). Afin d’imposer une contrainte de sécurité absolue garantissant la couverture théorique de 100% des anomalies perceptuelles, l’oracle fixe le rayon d’inflation optimal à $R_{\text{inf}}^* = 81$ pixels. Ce régime représente l’exigence de sanctuarisation maximale pour le paradigme géométrique de référence.

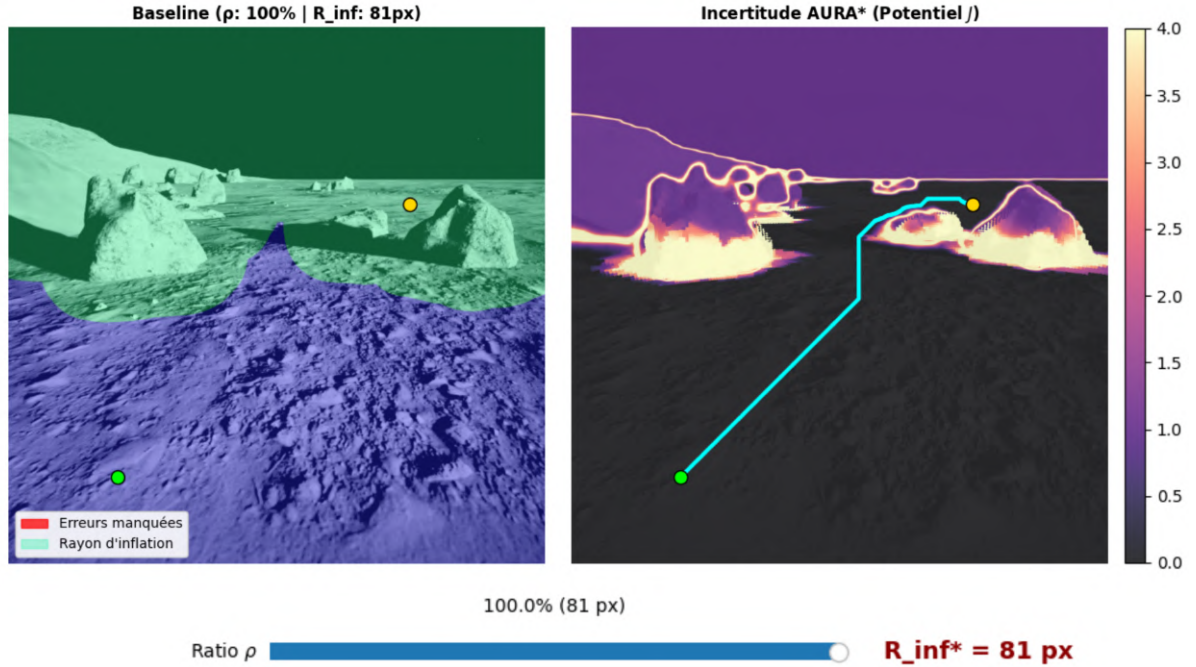


FIGURE 5.2 Comparaison qualitative de la Baseline et d'AURA* pour $\rho = 100\%$.

TABLEAU 5.1 Métriques cinématiques pour $\rho = 100\%$.

Métrique	Baseline	AURA*
Statut	ÉCHEC	RÉUSSI
Risque $M_{\max}$	–	0.025
Distance (px)	–	851.2
Déviation (η)	–	1.12×
Temps de calcul	–	931.2 ms
Sécurité	Obstruction	SÛR

Les figure 5.2 et le tableau 5.1 révèlent l'échec systémique de la méthode de référence par obstruction virtuelle : dilatés de 81 pixels, les obstacles génèrent des marges qui se chevauchent et saturent l'espace navigable. Faute de solution topologique, le planificateur géométrique immobilise préventivement le système.

À l'inverse, l'approche AURA* s'affranchit de cette binarisation destructrice. La topologie du potentiel d'incertitude $J(u, v)$ permet au planificateur probabiliste de percevoir les variations continues du risque épistémique. AURA* réussit à extraire une trajectoire ininterrompue en naviguant dans la région de moindre variance, acceptant une déviation métrique adaptative (surcoût cinématique de 12%, $\eta = 1.12$) pour maintenir une exposition au risque extrêmement basse ($M_{\max} = 0.025$), garantissant ainsi l'intégrité physique du véhicule sans en sacrifier la mobilité.

Ce scénario borne empiriquement l'extrême supérieur du compromis inhérent à l'inflation géométrique : l'application stricte d'un principe de précaution absolu face à l'incertitude mène inévitablement à la paralysie exploratoire. La sanctuarisation spatiale, par sa rigidité, prive le système de la résolution nécessaire pour évoluer dans des environnements denses. Comme le démontreront les scénarios subséquents, la seule méthode permettant de restaurer la viabilité de la référence géométrique consistera à décrémenter le paramètre ρ . Cette concession restaurera artificiellement la mobilité du rover, mais l'exposera mécaniquement à la traversée de singularités incertaines (faux négatifs critiques), validant la supériorité structurelle d'une évaluation par coût continu.

Scénario B : Seuil de viabilité topologique et compromis sécuritaire ($\rho = 27\%$)

Ce second scénario explore la frontière de viabilité du planificateur géométrique. Contrairement au Scénario A où l'immobilisation prévalait, le ratio d'inflation ρ a été itérativement décrémenté jusqu'à l'obtention de la première trajectoire mathématiquement continue pour la méthode de référence. Ce seuil critique de désenclavement est empiriquement atteint pour $\rho = 27\%$, requérant une réduction drastique du rayon de sanctuarisation à $R_{\text{inf}} = 22$ pixels. Ce cas illustre le réglage de sécurité maximal tolérable par la Baseline avant l'obstruction virtuelle de ce corridor.

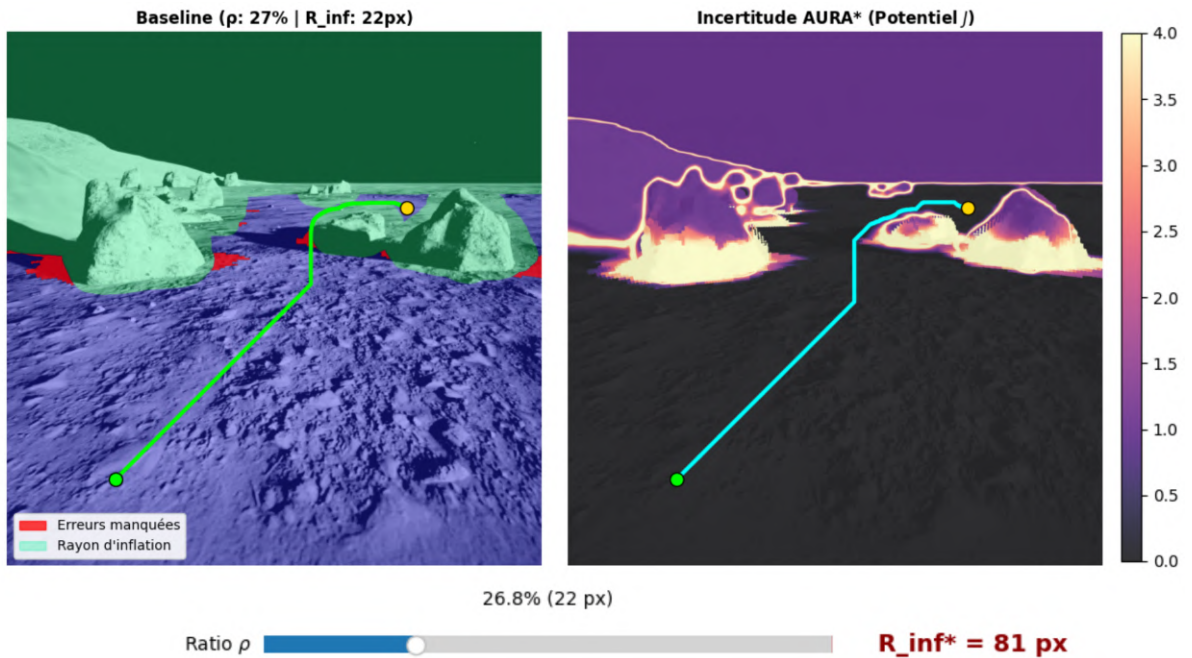


FIGURE 5.3 Comparaison qualitative de la Baseline et d'AURA* au seuil de viabilité.

L’observation de la figure 5.3 confirme que la restauration artificielle de la mobilité du rover s’opère au détriment direct de sa sécurité physique. En contractant l’inflation géométrique à 22 pixels pour forcer le passage, le planificateur statique devient aveugle aux extensions de variance épistémique (matérialisées par le masque d’erreur rouge). Par conséquent, la Baseline génère une trajectoire cinématiquement valide mais qui frôle dangereusement la base des obstacles, traversant des poches où le modèle perceptuel est incertain.

L’analyse des métriques du tableau 5.2 quantifie sévèrement ce compromis. L’exposition maximale au risque pour la Baseline bondit de manière spectaculaire à $M_{\max} = 0.285$, basculant le statut de sécurité du véhicule en état critique. À l’inverse, l’architecture AURA* démontre son invariabilité macroscopique et sa grande robustesse. Insensible au paramètre d’inflation ρ , le coût continu guide invariablement le véhicule au sein du même corridor sécurisé que lors du Scénario A (bien que la trajectoire ne soit pas rigoureusement identique au pixel près, celle-ci étant recalculée à chaque itération et soumise à la stochasticité inhérente de l’échantillonnage probabiliste sous-jacent). Pour un surcoût métrique infinitésimal par rapport à la Baseline (+0.1% de distance et +205.3 ms de temps de calcul), AURA* offre une réduction massive de 91% du risque épistémique maximal exposé.

TABLEAU 5.2 Synthèse des métriques cinématiques et sécuritaires pour $\rho = 27\%$.

Métrique	Baseline	AURA*	Gain
Statut	RÉUSSI	RÉUSSI	–
Risque $M_{\max}$	0.285	0.025	– 91.2%
Distance (px)	850.4	851.2	+0.1%
Déviaton (η)	1.11×	1.12×	+0.01×
Temps de calcul	733.4 ms	938.7 ms	+205.3 ms
Sécurité	CRITIQUE	SÛR	–

Ce scénario met concrètement en évidence le piège inhérent à la méthode de sanctuarisation statique : face à un environnement contraint, optimiser la viabilité topologique (en faisant diminuer ρ) engendre intrinsèquement l’acceptation aveugle de faux négatifs critiques. La comparaison directe valide la capacité d’AURA* à briser ce compromis de Pareto, en garantissant un niveau de sécurité optimal sans imposer de blocage navigationnel.

Scénario C : Sanctuarisation nulle et basculement topologique ($\rho = 0\%$)

Ce troisième scénario pousse la désactivation de l’heuristique de référence à sa limite absolue en supprimant toute marge de sécurité géométrique ($\rho = 0\%$, soit $R_{\text{inf}} = 0$ pixel). Ce cas limite est fondamental : il illustre le comportement brut d’un planificateur classique opérant sur une carte binaire dénuée de sanctuarisation artificielle, et met en exergue le danger inhérent à la recherche du chemin le plus court dans un environnement ambigu.

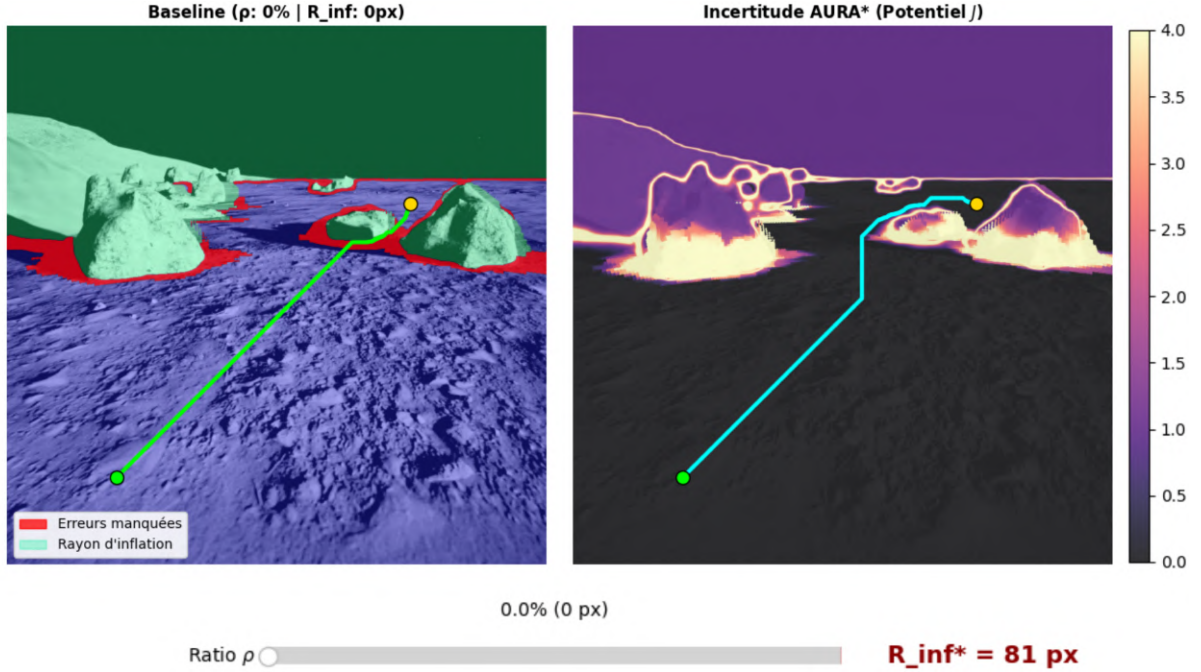


FIGURE 5.4 Comparaison qualitative de la Baseline et d’AURA* pour $\rho = 0\%$.

L’observation de la figure 5.4 révèle un basculement topologique majeur pour la Baseline.

En retirant l’inflation des obstacles, un nouveau corridor, jusqu’alors virtuellement fermé, s’ouvre directement entre deux formations rocheuses rapprochées. Le planificateur géométrique, optimisant aveuglément la distance euclidienne, perçoit cette ouverture comme un raccourci optimal et s’y engouffre. Or, comme l’indique le masque des erreurs (en rouge), cet interstice est saturé par une incertitude épistémique massive. Le rover est virtuellement précipité dans un piège topologique.

L’analyse des métriques (tableau 5.3) quantifie la sévérité de ce comportement téméraire. Attirée par ce raccourci, la Baseline génère une trajectoire extrêmement tendue (784.3 px, avec une déviation minimale $\eta = 1.03$). En contrepartie, le risque épistémique maximal (M_{max}) subit une inflation catastrophique, culminant à 0.491, soit le niveau d’exposition le plus critique enregistré.

De son côté, l’architecture AURA* fait preuve d’une résilience remarquable. Bien que ce raccourci existe physiquement, le champ continu $J(u, v)$ le modélise non pas comme un espace libre, mais comme une zone de coût élevé due à la forte variance épistémique locale. Sans nécessiter la moindre règle d’inflation arbitraire, AURA* rejette naturellement ce passage et maintient sa manœuvre de contournement (851.2 px). Pour un surcoût de distance de seulement +8.5%, l’approche probabiliste réduit le risque d’exposition de près de 95%.

TABLEAU 5.3 Synthèse des métriques cinématiques et sécuritaires pour $\rho = 0\%$.

Métrique	Baseline	AURA*	Gain
Statut	RÉUSSI	RÉUSSI	–
Risque $M_{\max}$	0.491	0.025	– 94.9%
Distance (px)	784.3	851.2	+8.5%
Déviation (η)	1.03×	1.12×	+0.09×
Temps de calcul	15.1 ms	937.9 ms	+922.8 ms
Sécurité	CRITIQUE	SÛR	–

Ce dernier scénario achève d’illustrer l’impasse fonctionnelle de l’inflation géométrique statique : l’approche déterministe confond systématiquement l’absence de prédiction de danger avec une sécurité absolue. Au fil de ces configurations canoniques, la méthode de référence s’avère prisonnière d’un compromis irréconciliable, oscillant inévitablement entre la paralysie par sanctuarisation maximale (Scénario A) et l’exposition critique aux anomalies lors du forçage topologique (Scénarios B et C).

À l’inverse, l’architecture AURA* brise ce compromis de Pareto. En intégrant directement la quantification d’incertitude via un champ de potentiel continu, elle repousse structurellement le planificateur des zones de péril vers les vallées de haute confiance, garantissant conjointement intégrité physique et mobilité. Afin de prouver que cette supériorité préventive est généralisable au-delà d’une topologie unitaire, la partie suivante déploie une évaluation stochastique de Monte-Carlo pour valider quantitativement ces observations sur un corpus hétérogène de milliers de requêtes ambiguës.

5.6.4 Analyse quantitative : performances cinématiques

L'évaluation des métriques (sous-section 5.6.2) via notre échantillonnage de Monte-Carlo (sous-section 5.6.1) permet de quantifier rigoureusement les comportements décisionnels. Nous y confrontons l'évolution de ces grandeurs pour la *Baseline* géométrique (sur sa plage sécuritaire $\rho \in [0, 100\%]$) aux performances intrinsèques et non-paramétriques d'AURA*. Cette étude démontre successivement l'effondrement de la viabilité exploratoire de la méthode standard, son incapacité structurelle à endiguer le risque épistémique sans blocage, et son surcoût computationnel face à la fluidité de la planification probabiliste.

Fiabilité globale de l'exploration

L'analyse de la viabilité topologique (figure 5.5) souligne une limite fondamentale de la planification géométrique en milieu incertain. Le taux de succès exploratoire (*TSE*), qui mesure la capacité à trouver une trajectoire continue, révèle une dichotomie entre les deux approches.

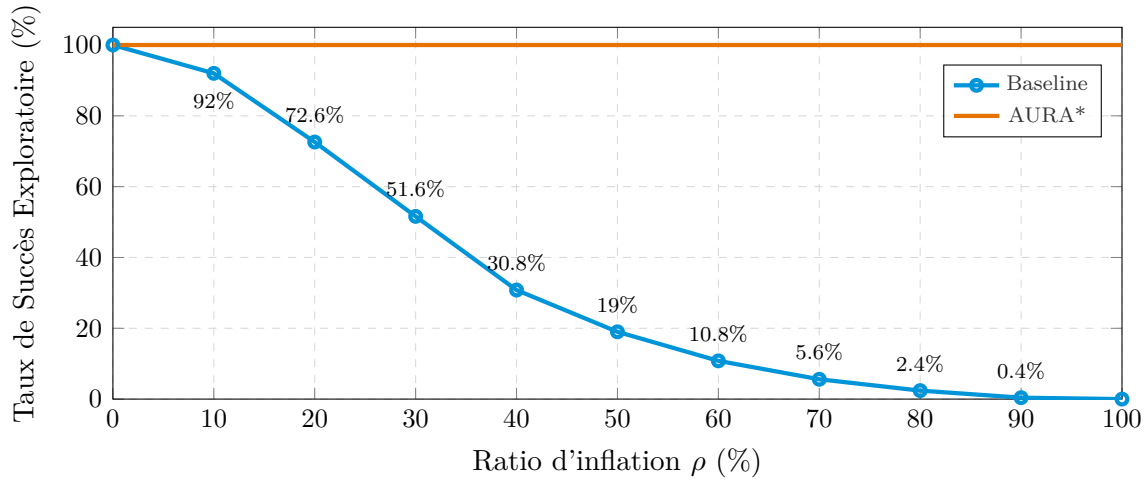


FIGURE 5.5 Évolution du taux de succès exploratoire en fonction de ρ .

En substituant la binarisation stricte des obstacles par une pénalisation proportionnelle de l'incertitude, AURA* préserve la connectivité topologique de l'espace. Par conséquent, avec $TSE = 100\%$, le planificateur garantit la complétion de la mission en n'acceptant de franchir des zones ambiguës que par stricte nécessité, en l'absence de toute alternative plus sûre.

À l'inverse, la trajectoire de la *Baseline* expose la faille conceptuelle de la planification géométrique. L'introduction d'une marge de sécurité ρ provoque un effondrement immédiat de sa viabilité (chute à 19% de succès dès $\rho = 50$). L'inflation géométrique n'y agit pas comme un régulateur de prudence, mais comme un véritable opérateur de destruction topologique

qui obstrue aveuglément les corridors viables. Aux régimes de haute sécurité, le planificateur se retrouve incapable de formuler la moindre solution.

Cette paralysie systémique illustre parfaitement le *Freezing Robot Problem* [68]. La dilatation isotrope des frontières sémantiques sature l’environnement de faux positifs spatiaux (des zones physiquement libres mathématiquement verrouillées), figeant *de facto* le rover. Cette dynamique prouve formellement l’incapacité du paradigme d’inflation statique à concilier sanctuarisation du risque et maintien de la mobilité en milieu planétaire dense.

Analyse sécuritaire : Maîtrise du risque épistémique et évitement

L’extraction d’une trajectoire topologiquement valide ne garantit pas sa viabilité physique. L’analyse des métriques d’évitement démontre formellement l’incapacité de l’heuristique d’inflation géométrique à préserver l’intégrité du système face aux incertitudes perceptuelles, contrastant avec l’anticipation structurelle de notre méthode.

La figure 5.6 évalue l’exposition au risque épistémique maximal ($M_{\max}$) sur les trajectoires abouties. En intégrant le gradient de variance au cœur de sa recherche, AURA* obtient un risque moyen de 0,035. À l’inverse, la *Baseline* s’engage aveuglément dans des zones critiques à faible inflation. Si l’augmentation de la marge de sécurité atténue cette exposition, elle se heurte à un plafond de verre : même poussée à son point de blocage topologique, la méthode standard conserve un risque résiduel supérieur à celui d’AURA*.

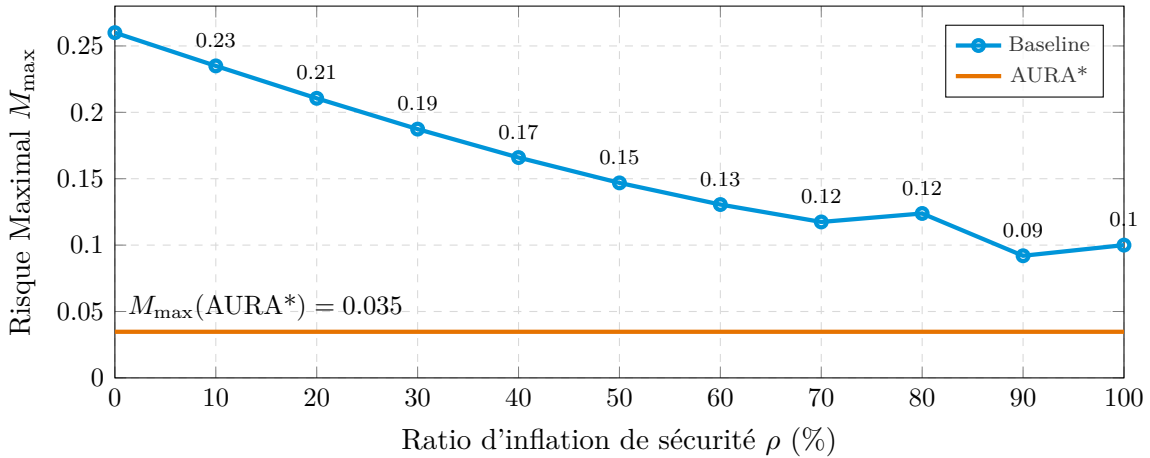


FIGURE 5.6 Évolution de $M_{\max}$ le long des trajectoires générées, en fonction de ρ .

Cette exposition théorique se concrétise physiquement par le taux de collision critique (TC), évalué à la figure 5.7. À inflation faible ou moyenne, la dilatation géométrique s’avère paradoxalement fatale : en fuyant binairement les obstacles majeurs, elle repousse le planificateur

vers des corridors faussement libres piégés par le bruit perceptuel local, provoquant des collisions massives. AURA*, au contraire, obtient un taux de collision moyen de 15,6%, prouvant que la modulation d'un champ scalaire continu est structurellement supérieure à une heuristique de gonflement pour assurer une sécurité préventive.

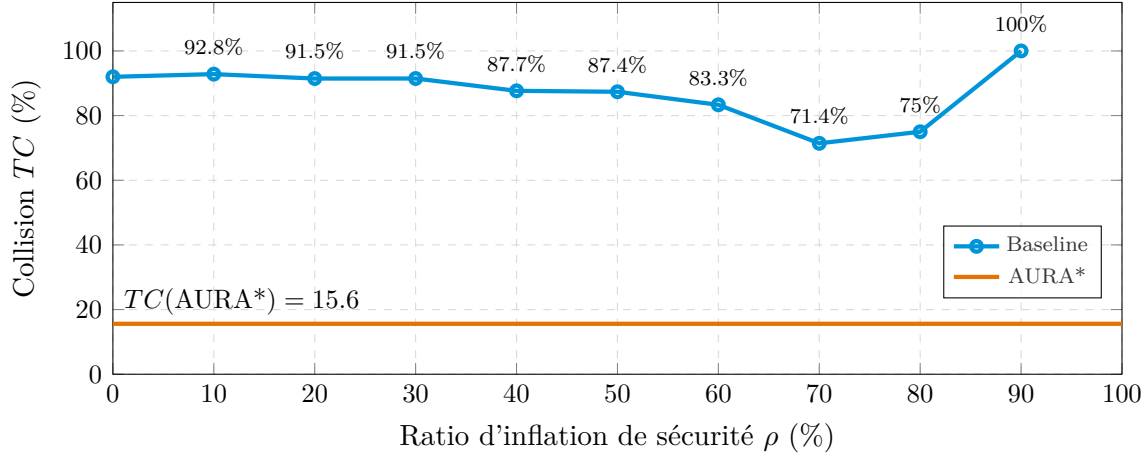


FIGURE 5.7 Évolution de TC le long des trajectoires générées, en fonction de ρ .

Frontière de Pareto : Le compromis sécurité-efficacité

L'évaluation conjointe de l'efficacité spatiale (mesurée par la déviation métrique η) et du risque épistémique ($M_{\max}$) permet d'isoler la limite cinématique des approches par inflation. La figure 5.8 illustre ce compromis structurel via la définition d'une frontière de Pareto.

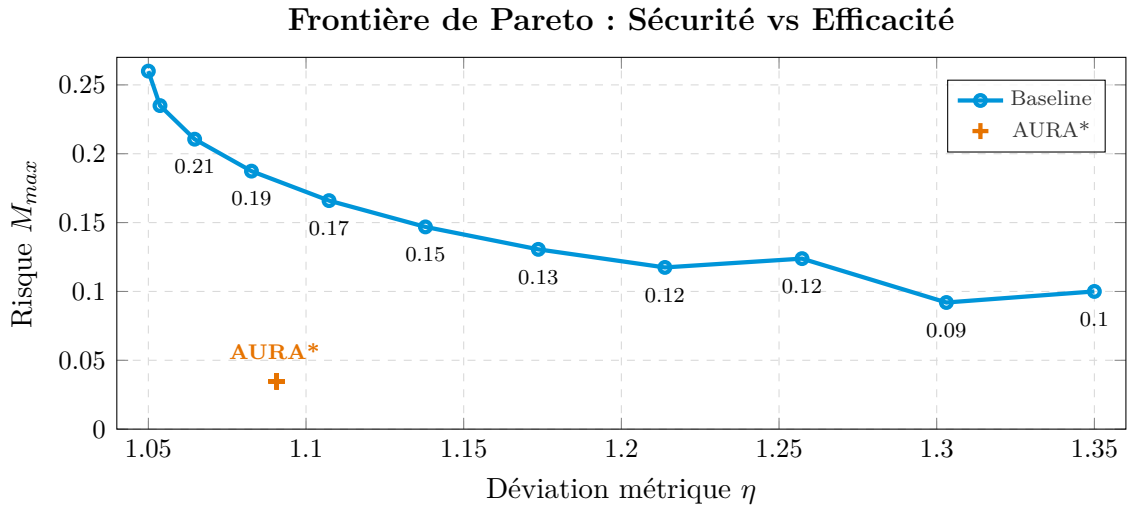


FIGURE 5.8 Frontière de Pareto entre la déviation métrique et le risque épistémique.

La méthode de référence subit un arbitrage structurellement inefficace : minimiser le risque exige une dilatation excessive des obstacles, générant des détours massifs (η jusqu'à +35 %).

Or, cette concession spatiale s'avère vaine puisque le risque résiduel sature asymptotiquement à un seuil critique, montrant que l'inflation géométrique déporte aveuglément la trajectoire vers des zones périphériques tout aussi bruitées, sans jamais garantir la sécurité absolue.

À l'inverse, AURA* brise ce paradigme et démontre une dominance stricte au sens de Pareto. En se positionnant très en deçà de la courbe de référence, l'algorithme écrase l'exposition au risque ($M_{\max} = 0,035$) tout en conservant une haute efficacité spatiale ($\eta \approx 1,09$). Cette rupture confirme qu'en optimisant directement sur un champ de coût continu, le planificateur navigue intelligemment au travers des « vallées de variance », réconciliant ainsi exploration optimale et sanctuarisation du système.

Coût computationnel et viabilité temps réel

L'implémentation algorithmique embarquée sur un rover planétaire est soumise à des contraintes matérielles strictes. Au-delà de la sécurité, le planificateur doit garantir une latence computationnelle suffisamment faible pour maintenir un contrôle réactif. L'analyse des temps de résolution met en évidence l'efficacité d'AURA* face à l'explosion combinatoire qui frappe la méthode géométrique.

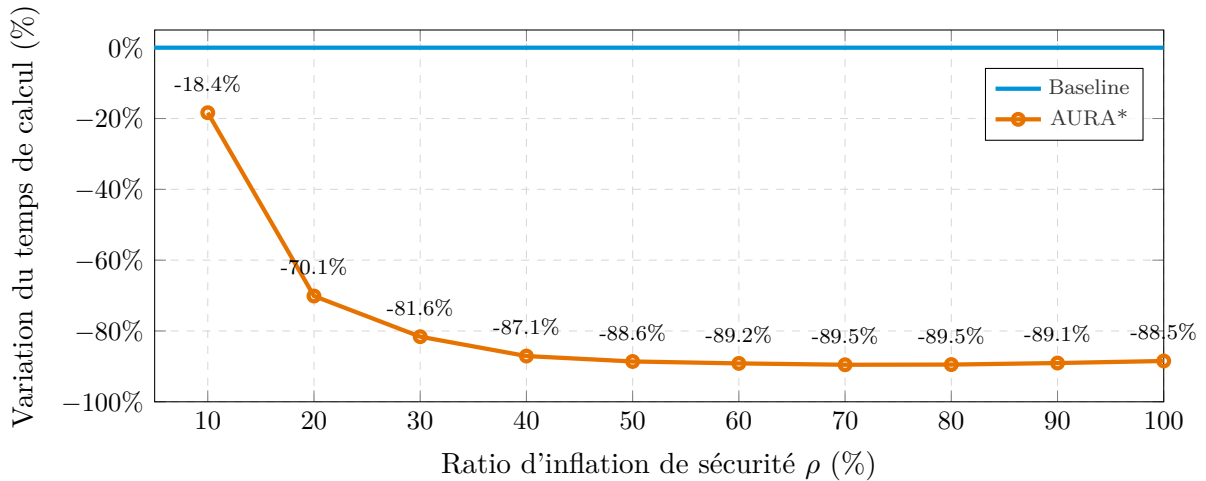


FIGURE 5.9 Évolution relative du temps de calcul d'AURA* par rapport à la Baseline.

La figure 5.9 illustre la réduction drastique du temps de calcul offerte par AURA*. Contrairement aux approches classiques, l'algorithme probabiliste s'avère plus efficace sur l'ensemble du domaine paramétrique, y compris à inflation très faible. Cette supériorité initiale suggère que la formulation continue du champ de coût $J(u, v)$ permet une convergence plus directe de l'arbre de recherche, limitant les expansions inutiles que subit la Baseline même dans un espace peu contraint.

À mesure que l'exigence de sécurité augmente ($\rho > 20\%$), l'écart se creuse de manière flagrante. La dilatation géométrique des obstacles fragmente la topologie de la grille et multiplie les culs-de-sac virtuels (*minima* locaux). Piégée par ces impasses, la Baseline est forcée d'explorer des régions massives avant de statuer, multipliant son temps de calcul par près de dix par rapport à AURA*.

En conservant une connectivité fluide dans l'espace de recherche, AURA* maintient une empreinte computationnelle stable. Aux régimes de haute sécurité, cette stabilité permet une réduction du temps de calcul atteignant 90%. Cette performance prouve que la planification continue lève le goulot d'étranglement algorithmique de l'inflation statique, assurant la viabilité temps réel indispensable aux futures missions autonomes.

5.7 Bilan du chapitre : synthèse, limites et perspectives

L'évaluation de ce chapitre valide formellement la supériorité de l'algorithme AURA* sur les heuristiques traditionnelles. En intégrant l'incertitude épistémique au sein d'une fonction de coût probabiliste continue, AURA* résout définitivement le *Freezing Robot Problem*. L'inflation géométrique statique impose un arbitrage perdant : l'hypertrophie des marges paralyse l'exploration, tandis que leur forçage engendre des collisions massives ($> 70\%$). AURA* brise ce compromis conceptuel (dominance de Pareto) en adaptant dynamiquement sa sécurité via l'énergie de traversabilité $J(u, v)$. Il garantit une excellente viabilité exploratoire tout en minimisant l'exposition au risque et en réduisant le temps de calcul jusqu'à 90%.

Malgré ces avancées, l'architecture présente des limites structurelles ouvrant trois perspectives de recherche :

Premièrement, l'optimisation basée sur le pire cas (*worst-case*) peut induire un excès de conservatisme dans des environnements globalement incertains (par exemple lors d'une tempête de poussière). L'intégration d'une planification stochastique de type POMDP permettrait de raisonner sur l'espérance du risque pour maintenir la mobilité.

Deuxièmement, la fusion multimodale actuelle par valeur maximale sanctuarise le rover mais néglige la corrélation positive entre les capteurs. Le développement d'une fusion bayésienne consciente de l'incertitude permettrait d'affiner le gradient de risque local.

Enfin, la validation ayant été menée sur une grille discrète pour un point holonome, l'intégration de ce champ de coût continu dans un planificateur spatio-temporel respectant les contraintes dynamiques du rover constituera l'ultime étape vers son déploiement physique.

CHAPITRE 6 CONCLUSION

Synthèse. Ce mémoire a traité la faillite décisionnelle induite par la surconfiance des réseaux de neurones profonds en environnement planétaire. En unifiant perception incertaine et planification continue, ces travaux proposent une architecture robuste. L’apport central réside dans la quantification multimodale de l’incertitude épistémique et son intégration cinématique directe. En substituant l’inflation géométrique statique par le champ de coût continu de l’algorithme AURA*, le système résout le compromis entre sanctuarisation absolue et mobilité exploratoire, minimisant le risque critique sans générer de blocage virtuel arbitraire.

Limitations. Cette architecture présente néanmoins des limites structurelles. Premièrement, l’approximation bayésienne par échantillonnage multiple exige une bande passante mémoire et une latence d’inférence difficilement compatibles avec les processeurs spatiaux durcis aux radiations. Deuxièmement, la stratégie de contrôle Minimax d’AURA*, qui élague strictement les trajectoires dépassant le seuil $V_{\max}$, rend le système vulnérable aux dégradations perceptuelles globales. Face à une illumination rasante ou un bruit capteur corrélé, la saturation du champ de veto risque d’immobiliser définitivement la plateforme. Troisièmement, la projection sur une carte de coût 2D impose une abstraction holonome, omettant les contraintes cinématiques réelles (glissement, *skid-steering*) et l’interaction terramécanique roue-régolithe sur de forts dévers tridimensionnels. Quatrièmement, la calibration de la prédiction conforme reposant sur l’échangeabilité des données, les garanties statistiques statiques se dégradent en présence d’un décalage de domaine majeur (*covariate shift*). Face à des géologies inédites, l’absence de recalibration dynamique compromet la validité de la couverture marginale du risque.

Perspectives. Ces aspects tracent des perspectives claires pour la robotique probabiliste. Le défi d’intégration matérielle sera naturellement levé par la transition vers des modèles à passe unique, extrayant la variance sans surcoût d’échantillonnage. Pour pallier la rigidité statique, l’adaptation au temps de test (*Test-Time Adaptation*) exploitera les retours proprioceptifs (couple, inertie) afin de recalibrer continuellement la perception *in situ*. Finalement, la commande dépassera la stricte logique du pire cas en intégrant l’énergie de traversabilité au sein de POMDP ou d’apprentissage par renforcement profond. Le système raisonnera alors sur l’espérance du risque pour orchestrer des manœuvres actives, réconciliant définitivement l’audace exploratoire et l’impératif de survie technologique.

RÉFÉRENCES

- [1] Y. Gao et S. Chien, “Review on space robotics : Toward top-level science through space exploration,” *Science Robotics*, vol. 2, n°. 7, p. eaan5074, 2017.
- [2] L. Matthies, E. Gat, R. Harrison, B. Wilcox, R. Volpe et T. Litwin, “Mars microrover navigation : Performance evaluation and enhancement,” dans *Proceedings of the 1995 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, vol. 1. IEEE, 1995, p. 433–440.
- [3] M. McHenry, N. Abcouwer, J. Biesiadecki, J. Chang, T. Del Sesto, A. Johnson, T. Litwin, M. Maimone, J. Morrison, R. Rieber, O. Toupet et P. Twu, “Mars 2020 autonomous rover navigation,” dans *AAS Guidance, Navigation and Control Conference*, ser. Advances in the Astronautical Sciences, vol. 172, n°. AAS 20-101. Breckenridge, CO : American Astronautical Society (AAS), 2020, preprint AAS 20-101.
- [4] J. J. Biesiadecki et M. W. Maimone, “The mars exploration rover surface mobility flight software : Driving ambition,” p. 1–13, 2006.
- [5] R. E. Arvidson, J. Bell III, P. Bellutta, N. Cabrol, J. Catalano, J. Cohen, L. Crumpler, D. Des Marais, T. Estlin, W. Farrand *et al.*, “Spirit mars rover mission : Overview and selected results from the northern home plate winter haven to the side of scamander crater,” *Journal of Geophysical Research : Planets*, vol. 115, n°. E7, 2010.
- [6] C. Guo, G. Pleiss, Y. Sun et K. Q. Weinberger, “On calibration of modern neural networks,” dans *Proceedings of the 34th International Conference on Machine Learning*. PMLR, 2017, p. 1321–1330.
- [7] D. Hendrycks et K. Gimpel, “A baseline for detecting misclassified and out-of-distribution examples in neural networks,” dans *International Conference on Learning Representations (ICLR)*, 2017.
- [8] D. Feng, C. Haase-Schütz, L. Rosenbaum, H. Hertlein, C. Gläser, F. Timm, W. Wiesbeck et K. Dietmayer, “Deep multi-modal object detection and semantic segmentation for autonomous driving : Datasets, methods, and challenges,” *IEEE Transactions on Intelligent Transportation Systems*, vol. 22, n°. 3, p. 1341–1360, 2020.
- [9] D. Amodei, C. Olah, J. Steinhardt, P. Christiano, J. Schulman et D. Mané, “Concrete problems in ai safety,” *arXiv preprint arXiv :1606.06565*, 2016.
- [10] M. W. Maimone, J. J. Biesiadecki, S. B. Goldberg et L. H. Matthies, “Autonomous navigation results from the mars exploration rover (mer) mission,” dans *International*

- Symposium on Artificial Intelligence, Robotics and Automation in Space (i-SAIRAS)*, 2005. [En ligne]. Disponible : <https://trs.jpl.nasa.gov/handle/2014/38079>
- [11] M. Maimone, Y. Cheng et L. Matthies, “Two years of visual odometry on the mars exploration rovers,” *Journal of Field Robotics*, vol. 24, n°. 3, p. 169–186, 2007.
 - [12] O. Toupet, M. Ono, T. Del Sesto, M. Maimone et M. McHenry, “Enhanced autonomous navigation on the perseverance mars rover,” *IEEE Transactions on Robotics*, 2025, early Access.
 - [13] K. Iagnemma et S. Dubowsky, “Mobile robot kinematic reconfigurability for rough-terrain,” *Proceedings of the 2004 IEEE International Conference on Robotics and Automation (ICRA)*, 2004.
 - [14] G. Ishigami, A. Miwa, K. Nagatani et K. Yoshida, “Terramechanics-based model for steering maneuvers of planetary exploration rovers on loose soil,” *Journal of Field Robotics*, vol. 24, n°. 3, p. 233–250, 2007.
 - [15] R. Gonzalez et K. Iagnemma, “Slippage estimation and compensation for planetary exploration rovers. state of the art and future challenges,” *Journal of Field Robotics*, vol. 35, n°. 4, p. 564–577, 2018. [En ligne]. Disponible : <https://onlinelibrary.wiley.com/doi/abs/10.1002/rob.21761>
 - [16] L. Wellhausen, A. Dosovitskiy, R. Ranftl, K. Walas, C. Cadena et M. Hutter, “Where should i walk ? predicting terrain properties from images via self-supervised learning,” *IEEE Robotics and Automation Letters*, vol. 4, n°. 2, p. 1509–1516, 2019.
 - [17] A. Angelova, L. Matthies, D. Helmick et P. Perona, “Learning and prediction of slip from visual information,” *Journal of Field Robotics*, vol. 24, n°. 3, p. 205–231, 2007.
 - [18] B. Rothrock, J. Papon, B. Kennedy, M. Mattar, A. Rankin, M. Heverly et M. Ono, “Spoc : Deep learning-based terrain classification for mars 2020 rover navigation,” dans *AIAA SPACE 2016*, 2016, p. 5518.
 - [19] R. Omar Chavez-Garcia et O. Aycard, “Multiple Sensor Fusion and Classification for Moving Object Detection and Tracking,” *IEEE Transactions on Intelligent Transportation Systems*, vol. PP, n°. 99, p. 1–10, 2015.
 - [20] V. De Silva, J. Roche et A. Kondoz, “Fusion of LiDAR and camera sensor data for environment sensing in driverless vehicles,” Loughborough University, Rapport technique, jan 2018. [En ligne]. Disponible : <https://hdl.handle.net/2134/32571>
 - [21] O. Ronneberger, P. Fischer et T. Brox, “U-net : Convolutional networks for biomedical image segmentation,” dans *Medical Image Computing and Computer-Assisted Intervention (MICCAI)*. Springer, 2015, p. 234–241.

- [22] L.-C. Chen, G. Papandreou, I. Kokkinos, K. Murphy et A. L. Yuille, “Deeplab : Semantic image segmentation with deep convolutional nets, atrous convolution, and fully connected crfs,” *IEEE transactions on pattern analysis and machine intelligence*, vol. 40, n^o. 4, p. 834–848, 2017.
- [23] W. Luo, Y. Li, R. Urtasun et R. Zemel, “Understanding the effective receptive field in deep convolutional neural networks,” dans *Advances in Neural Information Processing Systems*, vol. 29, 2016, p. 4898–4906.
- [24] X. Wang, R. Girshick, A. Gupta et K. He, “Non-local neural networks,” dans *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2018, p. 7794–7803.
- [25] Y. Yuan, R. Fu, L. Huang, W. Lin, C. Zhang, X. Chen et J. Wang, “Hrformer : High-resolution transformer for dense prediction,” 2021.
- [26] A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani *et al.*, “An image is worth 16x16 words : Transformers for image recognition at scale,” dans *International Conference on Learning Representations (ICLR)*, 2021.
- [27] E. Xie, W. Wang, Z. Yu, A. Anandkumar, J. M. Alvarez et P. Luo, “Segformer : Simple and efficient design for semantic segmentation with transformers,” dans *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 34, 2021, p. 12 077–12 090.
- [28] A. Milioto, I. Vizzo, J. Behley et C. Stachniss, “Rangenet++ : Fast and accurate lidar semantic segmentation,” dans *2019 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2019, p. 4213–4220.
- [29] X. Zhu, H. Zhou, T. Wang, F. Hong, Y. Ma, W. Li, H. Li et D. Lin, “Cylindrical and asymmetrical 3d convolution networks for lidar segmentation,” dans *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2021, p. 9939–9948.
- [30] C. R. Qi, L. Yi, H. Su et L. J. Guibas, “Pointnet++ : Deep hierarchical feature learning on point sets in a metric space,” dans *Advances in neural information processing systems*, vol. 30, 2017.
- [31] C. R. Qi, H. Su, K. Mo et L. J. Guibas, “Pointnet : Deep learning on point sets for 3d classification and segmentation,” dans *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, p. 652–660.
- [32] S. Vora, A. H. Lang, B. Helou et O. Beijbom, “Pointpainting : Sequential fusion for 3d object detection,” dans *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2020, p. 4604–4612.
- [33] M. Hein, M. Andriushchenko et J. Bitterwolf, “Why ReLU networks yield high confidence predictions far away from the training data,” dans *Proceedings of the IEEE/CVF*

- Conference on Computer Vision and Pattern Recognition*, 2019, p. 41–50.
- [34] C. Szegedy, V. Vanhoucke, S. Ioffe, J. Shlens et Z. Wojna, “Rethinking the inception architecture for computer vision,” dans *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, p. 2818–2826.
 - [35] A. Kendall et Y. Gal, “What uncertainties do we need in bayesian deep learning for computer vision?” dans *Advances in neural information processing systems (NeurIPS)*, vol. 30, 2017.
 - [36] C. Blundell, J. Cornebise, K. Kavukcuoglu et D. Wierstra, “Weight uncertainty in neural network,” dans *International Conference on Machine Learning*. PMLR, 2015, p. 1613–1622.
 - [37] A. Graves, “Practical variational inference for neural networks,” dans *Advances in Neural Information Processing Systems*, vol. 24, 2011, p. 2348–2356.
 - [38] Y. Gal et Z. Ghahramani, “Dropout as a bayesian approximation : Representing model uncertainty in deep learning,” dans *international conference on machine learning (ICML)*. PMLR, 2016, p. 1050–1059.
 - [39] V. Vovk, A. Gammerman et G. Shafer, *Algorithmic learning in a random world*. Springer Science & Business Media, 2005.
 - [40] A. N. Angelopoulos, S. Bates, J. Malik et M. I. Jordan, “Uncertainty sets for image classifiers using conformal prediction,” dans *International Conference on Learning Representations (ICLR)*, 2021.
 - [41] Y. Romano, M. Sesia et E. J. Candès, “Classification with valid and adaptive coverage,” 2020. [En ligne]. Disponible : <https://arxiv.org/abs/2006.02544>
 - [42] J. Brunekreef, E. Marcus, R. Sheombarsing, J.-J. Sonke et J. Teuwen, “Kandinsky conformal prediction : Efficient calibration of image segmentation algorithms,” 2023. [En ligne]. Disponible : <https://arxiv.org/abs/2311.11837>
 - [43] A. Elfes, “Using occupancy grids for mobile robot perception and navigation,” *Computer*, vol. 22, n°. 6, p. 46–57, 1989.
 - [44] A. Hornung, K. M. Wurm, M. Bennewitz, C. Stachniss et W. Burgard, “Octomap : An efficient probabilistic 3d mapping framework based on octrees,” *Autonomous robots*, vol. 34, p. 189–206, 2013.
 - [45] D. D. Fan, K. Kubo, O. Tslil, H. Rocha et A.-a. Agha-mohammadi, “Step : Stochastic traversability evaluation and planning for mobile robots in rough terrain,” dans *2021 IEEE International Conference on Robotics and Automation (ICRA)*. IEEE, 2021, p. 13 632–13 638.

- [46] S. Karaman et E. Frazzoli, “Sampling-based algorithms for optimal motion planning,” *The international journal of robotics research*, vol. 30, n°. 7, p. 846–894, 2011.
- [47] L. Blackmore, M. Ono et B. C. Williams, “Chance-constrained optimal path planning with obstacles,” *IEEE Transactions on Robotics*, vol. 27, n°. 6, p. 1080–1094, 2011.
- [48] J. Liu, Q. Zhang, X. Wan, S. Zhang, Y. Tian, H. Han, Y. Zhao, B. Liu, Z. Zhao et X. Luo, “Lusnar : A lunar segmentation, navigation and reconstruction dataset based on multi-sensor for autonomous exploration,” 2024. [En ligne]. Disponible : <https://arxiv.org/abs/2407.06512>
- [49] L.-C. Chen, G. Papandreou, F. Schroff et H. Adam, “Rethinking atrous convolution for semantic image segmentation,” 2017. [En ligne]. Disponible : <https://arxiv.org/abs/1706.05587>
- [50] C. Qi, L. Yi, H. Su et L. J. Guibas, “Pointnet++ : Deep hierarchical feature learning on point sets in a metric space,” *ArXiv*, vol. abs/1706.02413, 2017. [En ligne]. Disponible : <https://api.semanticscholar.org/CorpusID:1745976>
- [51] M. Pakdaman Naeini, G. Cooper et M. Hauskrecht, “Obtaining well calibrated probabilities using bayesian binning,” *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 29, n°. 1, Feb. 2015.
- [52] W. Brier, G. “Verification of forecasts expressed in terms of probability,” *Monthly Weather Review*, vol. 78, n°. 1, p. 1–3, 1950.
- [53] P. Kassapis, A. Athanasopoulou et K. Kamnitsas, “Calibrated adversarial refinement for stochastic semantic segmentation,” dans *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, 2021, p. 7546–7555.
- [54] V. Meel, M. Singh, S. Gautam, P. Kumar, K. N. Chaudhury et D. Sheet, “We care each pixel : Calibrating on medical segmentation model,” arXiv preprint arXiv :2503.05107, 2025.
- [55] Y. Shi, S. Ghosh, T. Belkhouja, J. R. Doppa et Y. Yan, “Conformal prediction for class-wise coverage via augmented label rank calibration,” 2024. [En ligne]. Disponible : <https://arxiv.org/abs/2406.06818>
- [56] T. Ding, A. N. Angelopoulos, S. Bates, M. I. Jordan et R. J. Tibshirani, “Class-conditional conformal prediction with many classes,” dans *Advances in Neural Information Processing Systems*, vol. 36, 2023.
- [57] C. R. Qi, L. Yi, H. Su et L. J. Guibas, “Pointnet++ : Deep hierarchical feature learning on point sets in a metric space,” dans *Advances in neural information processing systems*, vol. 30, 2017.

- [58] D. Hendrycks et T. Dietterich, “Benchmarking neural network robustness to common corruptions and perturbations,” dans *Proceedings of the International Conference on Learning Representations (ICLR)*, 2019.
- [59] Y. Ovadia, E. Fertig, J. Ren, Z. Nado, D. Sculley, S. Nowozin, J. Dillon, B. Lakshminarayanan et J. Snoek, “Can you trust your model’s uncertainty ? evaluating predictive uncertainty under dataset shift,” dans *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 32, 2019.
- [60] R. El-Yaniv et Y. Wiener, “On the foundations of noise-free selective classification,” *Journal of Machine Learning Research (JMLR)*, vol. 11, p. 1605–1641, 2010.
- [61] R. Zhang, P. Isola, A. A. Efros, E. Shechtman et O. Wang, “The unreasonable effectiveness of deep features as a perceptual metric,” dans *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018, p. 586–595.
- [62] J. Hauke et T. Kossowski, “Comparison of values of pearson’s and spearman’s correlation coefficients on the same sets of data,” *Quaestiones Geographicae*, vol. 30, n^o. 2, p. 87–93, 2011.
- [63] J. H. Zar, *Biostatistical Analysis*, 5^e éd. Upper Saddle River, NJ : Pearson Prentice Hall, 2010, chapitre sur la corrélation de rang.
- [64] I. Goodfellow, Y. Bengio et A. Courville, *Deep Learning*. Cambridge, MA : MIT Press, 2016.
- [65] F. Wilcoxon, “Individual comparisons by ranking methods,” *Biometrics bulletin*, vol. 1, n^o. 6, p. 80–83, 1945.
- [66] J. Demšar, “Statistical comparisons of classifiers over multiple data sets,” *Journal of Machine Learning Research (JMLR)*, vol. 7, p. 1–30, 2006.
- [67] D. Fan et A. e. a. Agha-mohammadi, “STEP : Stochastic Traversability Evaluation and Planning for Risk-Aware Navigation,” *IEEE Transactions on Robotics*, vol. 38, n^o. 4, p. 2027–2045, 2022.
- [68] R. Takemura et G. Ishigami, “Uncertainty-aware trajectory planning : Using uncertainty quantification and propagation in traversability prediction of planetary rovers,” *IEEE Robotics & Automation Magazine*, vol. 31, n^o. 2, p. 89–99, 2024.
- [69] Y. Ren, Y. Cai, F. Zhu, S. Liang et F. Zhang, “Rog-map : An efficient robocentric occupancy grid map for large-scene and high-resolution lidar-based motion planning,” 2023. [En ligne]. Disponible : <https://arxiv.org/abs/2302.14819>
- [70] D. M. Helmick, A. Angelova et L. H. Matthies, “Terrain adaptive navigation for planetary rovers,” *Journal of Field Robotics*, vol. 26, 2009. [En ligne]. Disponible : <https://api.semanticscholar.org/CorpusID:15957833>