

Faculté des sciences

DEEP HEDGING EN FINANCE ET ASSURANCE

pour la couverture en marché incomplet des produits complexes de type variable annuities ou Universal Life

Auteur: WEND KOUNI OLISEH DIENDERE

Promoteur : JEROME BARBARIN

Lecteur: DONATIEN HAINAUT

Année académique : 2025-2026

Filière :

MSc. Sciences Actuarielles, LSBA

Membres du Jury :

Pr. Jérôme Barbarin

Pr. Donatien Hainaut

Dédicace

À mes parents, SABINE et BRUNO, pour leur amour inconditionnel, leur affection et les valeurs dans lesquelles j'ai grandi et appris à me construire. Vous n'avez ménagé aucun effort tout au long de mon parcours académique jusqu'à ce jour. Puisse le Dieu Tout-Puissant vous accorder une longue vie afin que vous puissiez récolter les fruits de votre dévouement.

À mes frères et soeurs, particulièrement PRISCA, pour votre soutien constant, vos encouragements et vos prières qui m'ont toujours accompagné dans les moments de doute comme dans les moments de réussite.

À toutes les familles TENKODOGO et DIENDERE, pour votre soutien indéfectible, sous toutes ses formes.

Enfin, je dédie ce travail à ma très chère amie Yvette YAMEOGO, pour son soutien émotionnel, sa présence et ses encouragements tout au long de ce parcours.

Remerciements

Je tiens à exprimer ma profonde et sincère gratitude au Professeur JÉRÔME BARBARIN pour son encadrement, sa pédagogie, sa disponibilité et sa patience, qui ont constitué un repère essentiel dans la réalisation de ce travail de recherche. Ses orientations, ses conseils avisés ainsi que son expertise ont grandement contribué à l'aboutissement de ce mémoire.

J'adresse également mes remerciements au Professeur DONATIEN HAINAUT, lecteur de ce mémoire et coordonnateur de la filière, pour la qualité de ses enseignements, qui ont largement enrichi cette formation et nourri ce travail.

Je remercie sincèrement les Professeurs PIERRE DEVOLDER de l'UCLouvain et MARRI FOUAD de l'INSEA-Rabat, grâce à qui ce double cursus a pu voir le jour à travers le partenariat établi entre les deux institutions.

Enfin, j'adresse mes remerciements à l'ensemble du corps professoral ayant contribué, de près ou de loin, à ma formation académique et à l'accomplissement de ce parcours.

Table des matières

Table des matières	2
Introduction	4
1 Rappels de finance et de hedging	7
1.1 Cadre probabiliste et martingales	7
1.2 Mesures de risque convexes	8
1.3 Modèle de Black-Scholes et delta-hedging	9
1.4 Modèle de Heston et calcul stochastique	9
1.5 Méthodes classiques de couverture	11
1.5.1 Delta-hedging, gamma-hedging et vega-hedging	11
1.5.2 Superhedging	11
1.5.3 Couverture quadratique moyenne	11
1.5.4 Insuffisances structurelles	11
1.6 Produits d'assurance complexes	12
1.6.1 Universal Life	12
1.6.2 Variable Annuities et garanties	12
1.6.3 Path-dépendance et longue maturité	14
2 Architecture des réseaux de neurones	15
2.1 Origine et fonctionnement	15
2.2 Les FFNN	20
2.3 Les LSTM	21
3 Deep Hedging	26
3.1 Les fondamentaux	26
3.1.1 Approche <i>model-free</i> et <i>greek-free</i>	26
3.1.2 Optimisation directe du P&L	27
3.2 Formulation mathématique	28
3.2.1 Marché en temps discret avec frictions	28
3.3 Algorithmique	29
3.3.1 Algorithme deep hedging	29

4	Modélisation du marché	31
4.1	Hypothèses de la modélisation	31
4.2	Modélisation du marché	32
4.2.1	Pricing des instruments de couverture	32
4.2.2	Générateurs de scénarios	33
4.2.3	Discrétisation et simulation	34
4.2.4	Instruments de couverture et simulations sous deux mesures . .	37
4.3	Delta-hedging sous Heston	37
5	Application à un produit GMAB	39
5.1	Spécification et passif total du portefeuille de Variable Annuities	39
5.2	Optimisation des architectures	40
5.3	Résultats	41
5.4	Analyse mathématique	50
5.4.1	Formalisation du problème de couverture optimale	50
5.4.2	Sous-optimalité du delta-hedging en marché incomplet	51
5.4.3	Deep hedging FFNN et le delta-hedging	52
5.5	Incapacité structurelle du FFNN face à la path-dépendance	54
5.6	Reconstruction de la path-dépendance par LSTM	55
5.7	Les perte entropique et CVaR	57
	Bibliographie	61
A	Convexité des mesures de risque	63
A.1	Preuve : CVaR convexe	63
A.2	Preuve : risque entropique convexe	63
B	Portefeuille auto-financé en temps discret	63
C	Preuve du théorème 5.4.1	64
D	Preuve du proposition 5.4.1	67
E	Preuve du Théorème 5.5.1	69
F	Preuve du Théorème 5.6.1	72
G	Preuve du Corrolaire 5.6.2.1	73
H	Preuve du Théorème 5.7.1	74

Introduction

Durant ma formation en actuariat à l'INSEA-Rabat puis à l'UCLouvain, un paradoxe persistant dans le secteur de l'assurance-vie a retenu mon attention : alors que les produits modernes tels que *Universal Life* ou *Variable Annuities* et autres contrats unitaires à garanties, gagnent en sophistication, les outils de couverture utilisés en pratique demeurent souvent ancrés dans des cadres théoriques idéalisés, nés d'un monde sans frais, sans sauts, et sans comportement humain.

Ces produits, il est vrai, sont des hybrides fascinants : ils combinent la rigueur actuarielle à la complexité financière, intégrant des options cachées, des mécanismes de ratchet, des risques de mortalité et des réactions imprévisibles des assurés face à la performance de leur compte. Dans un tel environnement, les méthodes classiques telles que le delta-hedging, les grecques et la réplication parfaite se heurtent à leurs limites : elles supposent des marchés complets, des dynamiques calibrées à l'excès, et une vision locale du risque qui peine à capturer l'horizon long caractéristique de ces contrats.

Face à ce constat, ce mémoire ne prétend pas apporter une solution définitive, mais plutôt explorer une voie prometteuse : celle du *deep hedging*. En s'appuyant sur des architectures neuronales — les réseaux feedforward ou FFNN, les LSTM, voire générateurs adversariaux — cette approche remplace la quête de réplication parfaite par une gestion pragmatique du risque résiduel, directement ancrée dans les objectifs économiques de l'assureur : minimiser les pertes extrêmes, accepter les gains, et survivre à l'incertitude.

L'objectif est donc double : d'une part, implémenter et tester des stratégies de couverture pour un contrat *Variable annuité* ; d'autre part, comparer leurs performances à celles des méthodes classiques, en reconnaissant à la fois les gains en robustesse et les nouveaux défis qu'elles soulèvent, tels que l'instabilité numérique et la dépendance aux hyperparamètres.

Revue de littérature

La couverture des risques financiers s'est depuis lors reposée sur un idéal mathématique : celui d'un marché parfait, sans friction, où tout risque peut être éliminé grâce à une stratégie de réplication exacte. Ainsi naissent, dans les années 1970, les modèles de Black-Scholes et de Merton, offrant une formule pour valoriser et couvrir les options européennes. Le *delta-hedging*, fondé sur les « grecs », devient alors la pierre angulaire de la gestion des risques en salle de marché.

Cependant, rapidement, les praticiens réalisent que ce monde idéal sur quoi se basent les calculs est très loin de la réalité. Les frais de transaction, l'illiquidité, les sauts de marché ou encore les contraintes opérationnelles rendent la réplication parfaite des passifs quasi-impossible. C'est ainsi que dans les années 1990–2000, des travaux comme ceux de Davis, Panas & Zariphopoulou ou Whalley & Wilmott tentent d'adapter ces modèles en intégrant des frictions, mais les solutions deviennent rapidement inatteignables dès que la complexité du produit augmente.

Dans le secteur de l'assurance, où les produits tels que les *Variable Annuities* ou les *Universal Life* combinent plusieurs garanties financières longues, path-dépendantes et non linéaires avec un risque de mortalité, cette difficulté est confirmée par les travaux de Devolder (2016) et Perez (2022), qui soutiennent que les méthodes classiques capturent mal la complexité de ces structures. De ce fait, les assureurs se tournent alors vers des approches numériques comme le *Least Squares Monte Carlo*, popularisé par Krah (2020). Cependant celles-ci restent limitées par leur linéarité implicite et leur dépendance à des bases de fonctions prédéfinies.

C'est ainsi qu'émerge, en 2019, une rupture conceptuelle : l'article fondateur de Buehler, Gonon, Teichmann & Wood, intitulé « *Deep Hedging* ». Plutôt que de chercher à améliorer les modèles existants, les auteurs proposent de les abandonner et imaginent une autre approche. Le concept est radical : ne plus calculer de grecs, ni supposer de dynamique de marché, mais laisser un réseau de neurones apprendre directement la meilleure stratégie de couverture à partir de simulations. Le but n'est plus la réplication, mais la minimisation d'une mesure de risque convexe — comme le *CVaR* ou le risque entropique — dans un cadre réaliste, incluant frais, contraintes et instruments multiples.

Cette approche connaît rapidement un succès. En témoignent les travaux de Carboneau (2021) qui l'applique avec succès à une *Variable Annuity* avec garantie *ratchet*, démontrant que le *deep hedging* réduit les pertes extrêmes de plus de 80 % par rapport aux méthodes classiques. Plus récemment, Limmer & Horvath (2024) poussent la logique plus loin en introduisant les *Robust Hedging GANs*, une architecture adversariale qui rend la stratégie résistante à l'incertitude du modèle — un enjeu crucial en période de volatilité accrue. Aujourd'hui, le *deep hedging* n'est plus une simple cu-

riosity académique : c'est une réponse opérationnelle aux défis posés par les produits d'assurance modernes. Il s'agit d'une évolution profonde de la finance quantitative : du modèle vers les données, de la formule vers l'apprentissage, de la théorie vers la robustesse pratique.

Contrairement à ces approches, notre mémoire s'inscrit pleinement dans une logique visant non seulement à comprendre ces avancées, mais à les combiner, les comparer et les appliquer à des cas concrets, afin de proposer aux assureurs des outils de couverture à la fois performants, flexibles et résilients.

Chapitre 1

Rappels de finance et de hedging

Ce chapitre introduit les principaux outils théoriques mobilisés dans la suite du mémoire. L'objectif est de rappeler les concepts essentiels permettant de comprendre les stratégies de couverture étudiées par la suite, en particulier dans des marchés incomplets.

1.1 Cadre probabiliste et martingales

Définition 1.1.1. Soit $(\Omega, \mathcal{F}, \mathbb{P})$ un espace de probabilité. Une *filtration* $\mathbb{F} = (\mathcal{F}_t)_{t \in [0, T]}$ est une famille croissante de sous- σ -algèbres de $\mathcal{F}$: $\mathcal{F}_s \subseteq \mathcal{F}_t$ pour tout $0 \leq s \leq t \leq T$. Le quadruplet $(\Omega, \mathcal{F}, \mathbb{F}, \mathbb{P})$ est appelé *espace de probabilité filtré*.

Définition 1.1.2. Une mesure $\mathbb{Q}$ sur $(\Omega, \mathcal{F})$ est une *mesure martingale équivalente* (mesure risque-neutre) pour un processus de prix S si $\mathbb{Q} \sim \mathbb{P}$ et si S est une $\mathbb{Q}$ -martingale locale.

Définition 1.1.3. Une stratégie $\delta = (\delta_t)_{t \in [0, T]}$ est *auto-financée* si la valeur du portefeuille $V_t(\delta) := V_0 + \int_0^t \delta_u dS_u$ ne dépend que de δ et S , sans injection externe de cash après $t = 0$. La valeur actualisée $\tilde{V}_t := e^{-rt} V_t(\delta)$ est une $\mathbb{Q}$ -martingale sous la mesure risque-neutre si δ est admissible.

Définition 1.1.4. Une stratégie δ est une *opportunité d'arbitrage* si : $V_0 = 0$, $\mathbb{P}(V_T(\delta) \geq 0) = 1$ et $\mathbb{P}(V_T(\delta) > 0) > 0$.

Théorème 1.1.1. Soit $(S_t)_{t \in [0, T]}$ un processus de prix adapté à $\mathbb{F}$. Les assertions suivantes sont équivalentes :

- (1) Il n'existe pas d'opportunité d'arbitrage (définition 1.1.4),
- (2) Il existe $\mathbb{Q} \sim \mathbb{P}$ (définition 1.1.2) telle que (S_t) soit une $\mathbb{Q}$ -martingale locale.

Théorème 1.1.2. Le marché est complet (tout actif contingent borne est répliquable par une stratégie auto-financée, définition 1.1.3) si et seulement si il existe une unique mesure martingale équivalente $\mathbb{Q}$.

Définition 1.1.5. Dans un marché complet (théorème 1.1.2), le *prix arbitrage-free* a

t d'un payoff $\Phi(S_T) \in L^2$ est :

$$V_t = \mathbb{E}^{\mathbb{Q}}[e^{-r(T-t)}\Phi(S_T) \mid \mathcal{F}_t].$$

La stratégie de réplication est donnée par la *grecque delta* : $\Delta(t, S_t) := \partial V_t / \partial S_t$.

Remarque. Les théorèmes 1.1.1–1.1.2 s'étendent aux marchés *illiquides* en remplaçant l'intégrale stochastique par une intégrale stochastique non linéaire (Kunita, Bank–Baum).

1.2 Mesures de risque convexes

Les approches classiques reposant sur l'hypothèse de la réplication parfaite, qui devient irréaliste dans des marchés incomplets ou en présence de frictions. Dans ce contexte, la gestion du risque résiduel devient centrale. Notre travail, se reposant sur le marché réel ne garantissant pas la réplication parfaite, l'optimisation nécessite donc la prise en compte de mesure de risque convexe dont les plus pertinentes sont le CVaR ou le risque entropique.

Définition 1.2.1. Une application $\rho : L^1(\Omega, \mathcal{F}, \mathbb{P}) \rightarrow \mathbb{R}$ est une *mesure de risque convexe* si elle est convexe, monotone ($X \leq Y \Rightarrow \rho(X) \geq \rho(Y)$) et invariante par translation ($\rho(X + m) = \rho(X) - m$ pour tout $m \in \mathbb{R}$).

Définition 1.2.2. Le CVaR a niveau $\alpha \in (0, 1)$ est :

$$\text{CVaR}_\alpha(X) := \mathbb{E}^{\mathbb{P}}[X \mid X \geq \text{VaR}_\alpha(X)], \quad \text{VaR}_\alpha(X) = \inf\{x : \mathbb{P}(X \leq x) \geq \alpha\}.$$

Le CVaR appartient à la classe plus large des *Optimized Certainty Equivalents* (OCE), définis par :

$$\rho(X) = \inf_{w \in \mathbb{R}} \left\{ w + \mathbb{E}^{\mathbb{P}}[\ell(X - w)] \right\},$$

avec $\ell(x) = \frac{1}{1-\alpha} \max(x, 0)$ pour le CVaR. Ce qui donne la représentation duale suivante :

$$\text{CVaR}_\alpha(X) = \inf_{w \in \mathbb{R}} \left\{ w + \frac{1}{1-\alpha} \mathbb{E}[(X - w)_+] \right\}.$$

Le CVaR pénalise les **pertes extrêmes** au-delà du quantile α , ce qui le rend directement aligné avec les exigences prudentielles de Solvabilité II.

Définition 1.2.3. Le *risque entropique* avec aversion au risque $\lambda > 0$ est :

$$\rho_\lambda^{\text{ent}}(X) := \frac{1}{\lambda} \log \mathbb{E}^{\mathbb{P}}[e^{\lambda X}].$$

Il correspond au prix d'indifférence sous utilité exponentielle $U(x) = -e^{-\lambda x}$, est *time-consistent* et satisfait la convexité uniforme.

Proposition 1.2.1. Les mesures CVaR_α (définition 1.2.2) et $\rho_\lambda^{\text{ent}}$ (définition 1.2.3) sont des mesures de risque convexes au sens de la définition 1.2.1. Les preuves sont en annexe A.

Définition 1.2.4. On note $\mathcal{H}$ l'ensemble des stratégies $\delta = (\delta_{t_n})_{n=0}^{N-1}$ prédictibles et de carré intégrable. La sous-classe des *stratégies markoviennes* est :

$$\mathcal{H}^{\text{Markov}} := \left\{ \delta \in \mathcal{H} : \delta_{t_n} = f(t_n, S_{t_n}, V_{t_n}), f \text{ mesurable} \right\}.$$

Définition 1.2.5. Pour une mesure de risque convexe ρ et un passif $H_T \in L^1$, le *problème de couverture optimale* est :

$$\delta^* = \arg \min_{\delta \in \mathcal{H}} \rho(H_T - V_T^\delta). \quad (1.1)$$

1.3 Modèle de Black-Scholes et delta-hedging

Le modèle de Black-Scholes représente un cadre historique de référence pour la valorisation et la couverture des produits dérivés. Malgré ses hypothèses fortes, il fournit les fondements du delta-hedging et des méthodes basées sur les grecques. En effet, il comporte un actif sans risque $dB_t = rB_t dt$ et un actif risque $dS_t = \mu S_t dt + \sigma S_t dW_t^\mathbb{P}$. Par le théorème 1.1.2, ce marché est complet. Sous l'unique $\mathbb{Q}$ (définition 1.1.2), on a $dS_t = rS_t dt + \sigma S_t dW_t^\mathbb{Q}$. Le prix et la stratégie de réplication sont donnés par la définition 1.1.5.

Définition 1.3.1. Soit $P(t, S_t, V_t)$ le prix du passif sous $\mathbb{Q}$ (définition 1.1.2). La stratégie de *delta-hedging* est :

$$\delta_{t_n}^\Delta := \frac{\partial P}{\partial S}(t_n, S_{t_n}, V_{t_n}) \in \mathbb{R}.$$

Cette quantité ne minimise *aucune* mesure de risque globale (définition 1.2.1) : c'est une stratégie de réplication locale sous $\mathbb{Q}$, suboptimale en marché incomplet sous $\mathbb{P}$.

Définition 1.3.2. Pour $V(t, S_t)$ le prix d'une option, les *grecques* sont :

$$\Delta := \frac{\partial V}{\partial S}, \quad \Gamma := \frac{\partial^2 V}{\partial S^2}, \quad \mathcal{V} := \frac{\partial V}{\partial \sigma}.$$

1.4 Modèle de Heston et calcul stochastique

L'hypothèse de volatilité constante du modèle de Black-Scholes reste toutefois loin de la réalité des données du marché financier. Elles présentent des phénomènes de clustering de volatilité, de smiles implicites et de queues épaisses. Pour réaliser une modélisation qui se veut plus proche de la réalité, nous prendrons en compte une volatilité stochastique afin de mieux capturer ces dynamiques observées d'où le modèle de Heston

s'avère le mieux adapté pour la suite de ce mémoire.

Définition 1.4.1. La dynamique de l'actif S_t et de sa variance V_t sous $\mathbb{P}$:

$$dS_t = \mu S_t dt + \sqrt{V_t} S_t dW_t^{(1)}, \quad (1.2)$$

$$dV_t = \kappa(\theta - V_t) dt + \sigma \sqrt{V_t} dW_t^{(2)}, \quad (1.3)$$

avec $\text{corr}(dW_t^{(1)}, dW_t^{(2)}) = \rho dt$. Via la décomposition de Cholesky : $dW_t^{(1)} = dB_t^{(1)}$ et $dW_t^{(2)} = \rho dB_t^{(1)} + \sqrt{1 - \rho^2} dB_t^{(2)}$ avec $B^{(1)}, B^{(2)}$ browniens indépendants.

Lemme 1.4.1 (Lemme d'Ito multidimensionnel). *Pour $f(t, S_t, V_t, r_t)$ suffisamment régulière :*

$$\begin{aligned} df &= \frac{\partial f}{\partial t} dt + \frac{\partial f}{\partial S} dS_t + \frac{\partial f}{\partial V} dV_t + \frac{\partial f}{\partial r} dr_t \\ &+ \frac{1}{2} \left(\frac{\partial^2 f}{\partial S^2} (dS_t)^2 + \frac{\partial^2 f}{\partial V^2} (dV_t)^2 + \frac{\partial^2 f}{\partial r^2} (dr_t)^2 \right) \\ &+ \frac{\partial^2 f}{\partial S \partial V} dS_t dV_t + \frac{\partial^2 f}{\partial S \partial r} dS_t dr_t + \frac{\partial^2 f}{\partial V \partial r} dV_t dr_t. \end{aligned}$$

Appliqué à $\ln S_t$ dans le modèle (1.2) : $d \ln S_t = (\mu - V_t/2) dt + \sqrt{V_t} dB_t^{(1)}$, d'où :

$$S_t = S_0 \exp \left(\int_0^t (\mu - \frac{V_s}{2}) ds + \int_0^t \sqrt{V_s} dB_s^{(1)} \right).$$

Définition 1.4.2. Pour les prix de risque λ_S et $\lambda \sqrt{V_t}$, la densité de Radon-Nikodym est :

$$\frac{d\mathbb{Q}}{d\mathbb{P}} \Big|_{\mathcal{F}_T} = \exp \left(- \int_0^T \lambda_S dB_s^{(1)} - \int_0^T \lambda \sqrt{V_s} dB_s^{(2)} - \frac{1}{2} \int_0^T (\lambda_S^2 + \lambda^2 V_s) ds \right).$$

Sous $\mathbb{Q}$, la dynamique de Heston conserve sa structure avec :

$$dS_t = r S_t dt + \sqrt{V_t} S_t dB_t^{1,\mathbb{Q}}, \quad (1.4)$$

$$dV_t = \tilde{\kappa}(\tilde{\theta} - V_t) dt + \sigma \sqrt{V_t} (\rho dB_t^{1,\mathbb{Q}} + \sqrt{1 - \rho^2} dB_t^{2,\mathbb{Q}}), \quad (1.5)$$

avec $\tilde{\kappa} = \kappa + \sigma \lambda \sqrt{1 - \rho^2}$ et $\tilde{\theta} = (\kappa \theta - \sigma \rho \lambda_S) / \tilde{\kappa}$.

1.5 Méthodes classiques de couverture

1.5.1 Delta-hedging, gamma-hedging et vega-hedging

En marché de Black-Scholes (section 1.3), le delta-hedging (définition 1.3.1) fournit une réplication parfaite (définition 1.1.5). En présence de volatilité stochastique (modèle 1.4.1) ou de sauts, le risque de convexité (*gamma*) et le risque de volatilité (*vega*) (définition 1.3.2) ne sont pas couverts. Le **gamma-vega hedging** utilise des options vanilles pour satisfaire :

$$\sum_{i=0}^D \delta_t^{(i)} \Delta_t^{(i)} \approx \Delta_t^{\text{cible}}, \quad \sum_{i=0}^D \delta_t^{(i)} \Gamma_t^{(i)} \approx \Gamma_t^{\text{cible}}, \quad \sum_{i=0}^D \delta_t^{(i)} \mathcal{V}_t^{(i)} \approx \mathcal{V}_t^{\text{cible}}.$$

Cette approche reste dépendante du modèle et ne minimise aucun critère global.

1.5.2 Superhedging

En marché incomplet (théorème 1.1.2 non satisfait), le **superhedging** cherche la domination presque sûr du passif :

$$\pi^{\text{super}}(H) := \inf \{x : \exists \delta \in \mathcal{H}, x + (\delta \cdot S)_T \geq H \text{ } \mathbb{P}\text{-p.s.}\}.$$

En présence de frictions due à des sauts, une volatilité stochastique, ou des frictions, le coût computationnel de superhedging devient souvent prohibitif. Cette explosion du coût provient du fait que le superhedging exige une couverture **dans tous les états du monde**, y compris les scénarios extrêmes de probabilité négligeable. Il s'agit donc d'une approche extrêmement prudente, incompatible avec les objectifs économiques des assureurs, qui cherchent à **minimiser le risque résiduel de manière efficiente**, et non à l'éliminer à tout prix [1].

1.5.3 Couverture quadratique moyenne

La **couverture quadratique moyenne** minimise $\mathbb{E}^{\mathbb{P}}[(H - V_0 - (\delta \cdot S)_T)^2]$. Elle pénalise *symétriquement* gains et pertes, contrairement aux objectifs asymétriques des assureurs. Elle est également myope sur les longues maturités.

1.5.4 Insuffisances structurelles

Les méthodes classiques présentent quatre limites fondamentales : (1) **dépendance au modèle** : en marché incomplet, les grecques (définition 1.3.2) dépendent du choix arbitraire de $\mathbb{Q}$; (2) **myopie temporelle** : pas d'exploitation de la diversification sur longue maturité; (3) **inadaptation à la path-dépendance** : les grecques ne capturent pas l'historique; (4) **pénalisation symétrique** : incompatible avec les objectifs de l'assureur. Ces limites justifient l'optimisation directe du problème (1.1).

1.6 Produits d'assurance complexes

Les produits d'assurance modernes tels que les Variable Annuities ou les contrats Universal Life combinent des caractéristiques financières et actuarielles complexes. Leur valorisation et leur couverture nécessitent de prendre en compte simultanément le risque de marché, la mortalité, les comportements des assurés et les garanties embarquées.

1.6.1 Universal Life

Le **Universal Life (UL)** combine protection et épargne. La valeur du compte évolue selon : $V_{t+1} = (V_t + P_{t+1})(1 + r_{t+1}^g + r_{t+1}^{pb}) - C_{t+1}$, où :

V_t est la valeur du compte à la date t ,

C_t est la prime de risque : coût de l'assurance décès et frais,

r_t^g représente le taux garanti révisé à de périodes successifs par exemple chaque 8 ans en Belgique,

r_t^{pb} est la participation bénéficiaire.

Le contrat inclut une **garantie de décès** CD_t , payable au bénéficiaire en cas de décès de l'assuré. La prestation décès versée en cas de décès de l'assuré sur l'intervalle $[t, t+1]$ est donnée par :

$$CD_{t+1} = 1.3 V_t.$$

Le coût total prélevé sur la valeur du compte au cours de la période s'écrit :

$$C_{t+1} = F_{t+1} + \pi_{t+1}, \quad \text{avec} \quad \pi_{t+1} = 0.3 V_t q_{x+t}.$$

Quatre risques interdépendants [8] le caractérisent : marché, mortalité, comportemental et taux.

1.6.2 Variable Annuities et garanties

Guaranteed Minimum Death Benefit (GMDB)

Définition 1.1. Un assuré âgé de x années investit un montant S_0 en actions.

Si le décès de l'assuré survient avant l'échéance T_n , le contrat garantit le versement d'un montant minimal :

$$\mathbf{1}_{\{\tau_x < T_n\}} \left(S_0 e^{r^g \tau_x} + \left(S_{\tau_x} - S_0 e^{r^g \tau_x} \right)_+ \right)$$

où τ_x est le temps aléatoire de décès et r^g est le taux garanti minimal. On suppose que τ_x est indépendant des actifs financiers.

La valeur de ce contrat à l'instant t est donnée par l'espérance actualisée sous la mesure risque-neutre $\mathbb{Q}$:

$$\begin{aligned}
V_t &= \mathbb{E}^{\mathbb{Q}} \left[\mathbf{1}_{\{\tau_x < T_n\}} e^{-\int_t^{\tau_x} r_s ds} \left(S_0 e^{r^g \tau_x} + (S_{\tau_x} - S_0 e^{r^g \tau_x})_+ \right) \middle| \mathcal{F}_t \right] \\
&= \sum_{u=t}^{T_n} {}_{u-t}q_{x+t} \mathbb{E}^{\mathbb{Q}} \left[e^{-\int_t^u r_s ds} \left(S_0 e^{r^g u} + (S_u - S_0 e^{r^g u})_+ \right) \middle| \mathcal{F}_t \right],
\end{aligned}$$

où ${}_{u-t}q_{x+t}$ est la probabilité de décès entre t et u .

Guaranteed Minimum Accumulation Benefit (GMAB)

Définition 1.2. Un assuré âgé de x années investit un montant S_0 en actions.

Si la durée du contrat est T_n et que l'assuré reçoit à l'échéance un payoff :

$$\mathbf{1}_{\{\tau_x \geq T_n\}} \left(S_0 e^{r^g T_n} + (S_{T_n} - S_0 e^{r^g T_n})_+ \right)$$

où τ_x est le temps aléatoire de décès et r^g est un rendement garanti minimal. τ_x est indépendant des actifs financiers et

$$\mathbb{E}^{\mathbb{Q}} \left(\mathbf{1}_{\{\tau_x \geq T_n\}} \middle| \mathcal{F}_t \right) = {}_{T_n-t}p_{x+t}.$$

La valeur de ce contrat est égale à la valeur actualisée espérée des flux monétaires :

$$\begin{aligned}
V_t &= \mathbb{E}^{\mathbb{Q}} \left(\mathbf{1}_{\{\tau_x \geq T_n\}} e^{-\int_t^{T_n} r_s ds} \left(S_0 e^{r^g T_n} + (S_{T_n} - S_0 e^{r^g T_n})_+ \right) \middle| \mathcal{F}_t \right) \\
&= {}_{T_n-t}p_{x+t} \left[S_0 e^{r^g T_n} P(t, T_n) + \mathbb{E}^{\mathbb{Q}} \left(e^{-\int_t^{T_n} r_s ds} (S_{T_n} - S_0 e^{r^g T_n})_+ \middle| \mathcal{F}_t \right) \right]
\end{aligned}$$

Guaranteed Minimum Withdrawal Benefit (GMWB)

Définition 1.3. Un assuré âgé de x années investit un montant S_0 dans un fonds sous-jacent.

Le GMWB garantit à l'assuré la possibilité de retirer un montant minimal annuel W pendant la durée du contrat, indépendamment de la performance du fonds. Si le solde du compte tombe à zéro, les rétractions garantis continuent jusqu'à épuisement de la garantie.

Le payoff à la date t_n peut s'écrire comme suit :

$$\mathbf{1}_{\{\tau_x \geq t_n\}} \min \left(W, S_{t_n} \right),$$

où S_{t_n} est la valeur du fonds sous-jacent à la date t_n et τ_x le temps de décès.

La valeur de ce contrat à l'instant t est la valeur actualisée espérée des flux de retraits garantis :

$$\begin{aligned}
V_t &= \mathbb{E}^{\mathbb{Q}} \left[\sum_{n:t_n \geq t} e^{-\int_t^{t_n} r_s ds} \mathbf{1}_{\{\tau_x \geq t_n\}} \min(W, S_{t_n}) \middle| \mathcal{F}_t \right] \\
&= \sum_{n:t_n \geq t} {}_{t_n-t}p_{x+t} \mathbb{E}^{\mathbb{Q}} \left[e^{-\int_t^{t_n} r_s ds} \min(W, S_{t_n}) \middle| \mathcal{F}_t \right],
\end{aligned}$$

où ${}_{t_n-t}p_{x+t}$ est la probabilité de survie entre t et t_n .

Définition 1.4. Guaranteed Lifetime Withdrawal Benefit (GLWB)

Le GLWB étend le GMWB à vie. Le payoff en t_n :

$$L^{\text{GLWB}} = \sum_{t=t^*+1}^{\tau_x} \mathbf{1}_{\{\tau_x \geq t_n\}} \min(W, S_{t_n}),$$

où τ_x est maintenant la durée de vie aléatoire de l'assuré. Ce produit combine risque financier via $S(t)$ et risque de longévité via τ_x , ce qui le rend particulièrement complexe à couvrir .

1.6.3 Path-dépendance et longue maturité

Le GMDB avec ratchet a un payoff $L^{\text{GMDB}} = \mathbf{1}_{\{\tau_x < T_n\}} \max_{n \leq n_\tau} S_{t_n}$, qui dépend de l'historique complet. L'état du système à t nécessite $Z_t = \max_{s \leq t} S_s$, non capture par les grecques (définition 1.3.2). La longue maturité amplifie les erreurs de couverture dans les approches myopes et rend les garanties non répliquables (théorème 1.1.2 non satisfait).

Dans la suite de notre mémoire, nous considererons uniquement le produit GMAB avec ratchet annuel qui sera notre principal produit à couvrir.

Chapitre 2

Architecture des réseaux de neurones

2.1 Origine et fonctionnement

Les réseaux de neurones artificiels (ANN) trouvent leur origine dans les recherches des biologistes qui, dès les années 1940, cherchaient à comprendre le fonctionnement du cerveau humain. En observant la manière dont les *neurones biologiques* communiquent entre eux à travers des signaux électriques, certains scientifiques ont eu l'intuition qu'il serait possible de *modéliser cette activité sous forme mathématique*. Cette vision, initiée notamment par **McCulloch et Pitts (1943)**, visait à représenter le raisonnement humain à l'aide de connexions simples entre unités artificielles. Ainsi, les premiers modèles de neurones artificiels sont nés d'une tentative d'*imiter les mécanismes cognitifs du cerveau* pour en reproduire, de manière simplifiée, la capacité d'apprentissage et de décision.

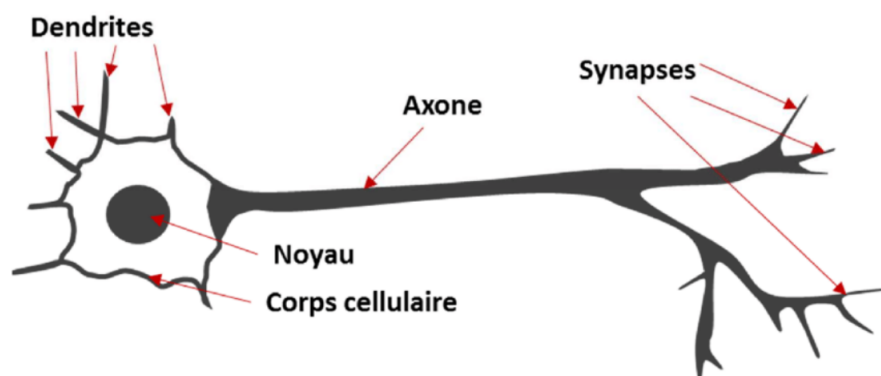


FIGURE 2.1 – Neurone humain

Les dendrites sont les portes d'entrée d'un neurone. Le neurone reçoit des signaux :

Excitateur : +1

Inhibiteur : -1

Quand la somme des signaux dépasse un seuil, le neurone s'active et produit un signal

électrique. Le signal circule le long de l'axone et se dirige vers les synapses, qui font passer le courant électrique aux autres neurones.

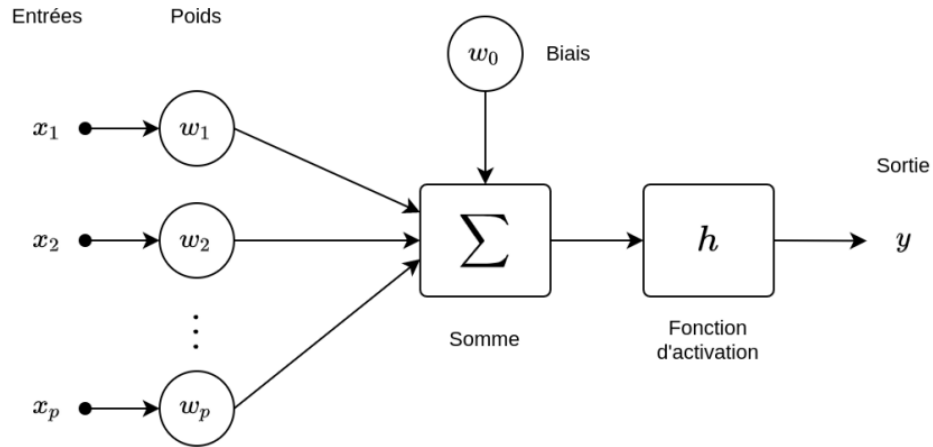


FIGURE 2.2 – Artificial Neural Network : perceptron

Pour connaître ce prédicteur, on passe par deux étapes :

Agrégation : $z = w_1x_1 + w_2x_2 + w_3x_3$, où h représente l'activité du neurone : +1 ou -1.

Activation : $y = 1$ si $h \geq 0$ et $y = 0$ sinon.

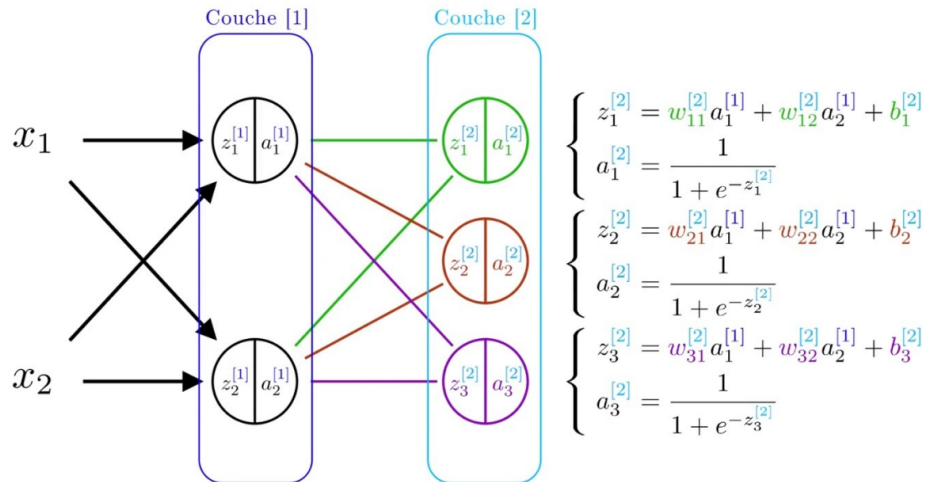


FIGURE 2.3 – Artificial Neural Network : perceptron 2 couches

Le deep Learning

Lorsque le nombre de couches cachées devient important, le réseau est qualifié de réseau de neurones profond (*Deep Neural Network*). Ce type d'architecture constitue la base de l'apprentissage profond ou *Deep Learning*, une approche qui permet au modèle

d'apprendre des représentations de plus en plus abstraites et complexes à mesure que les informations traversent les différentes couches.

Fonction d'activation

La fonction d'activation, notée h ou σ , est utilisée pour déterminer l'état du neurone, c'est-à-dire la sortie qu'il produit. Elle prend en entrée la somme pondérée des stimuli et effectue une transformation avant de fournir la sortie finale.

Fonction sigmoïde

La fonction sigmoïde transforme la somme pondérée en une valeur comprise entre 0 et 1 :

$$h(z) = \frac{1}{1 + e^{-z}}$$

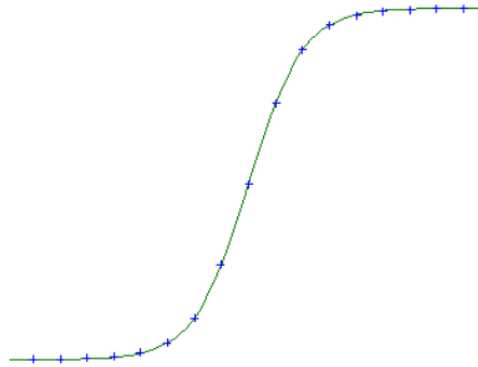


FIGURE 2.4 – Fonction sigmoïde

Fonction tangente hyperbolique

La fonction $\tanh$ est similaire à la fonction sigmoïde, mais elle produit des valeurs entre -1 et 1 :

$$h(z) = \frac{e^z - e^{-z}}{e^z + e^{-z}}$$

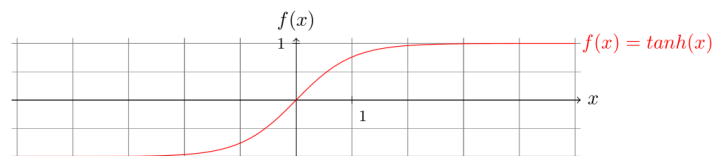


FIGURE 2.5 – Fonction tanh

La fonction $\tanh$ est couramment utilisée dans les réseaux de neurones où les valeurs positives et négatives sont importantes.

Fonction Rectified Linear Unit (ReLU)

La fonction ReLU transforme toutes les valeurs négatives en zéro et laisse les valeurs positives inchangées :

$$h(z) = \max(0, z)$$

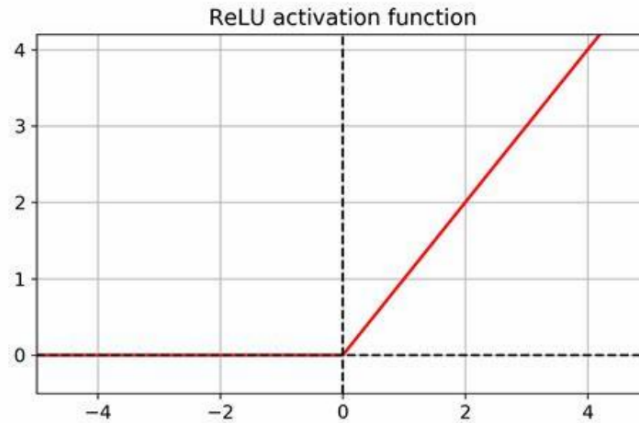


FIGURE 2.6 – Fonction ReLU

La fonction ReLU est largement utilisée dans les réseaux de neurones convolutifs (CNN) en raison de sa simplicité et de ses bons résultats empiriques.

Fonction Leaky ReLU

La fonction Leaky ReLU est une variation de la fonction ReLU qui permet une petite pente pour les valeurs négatives :

$$h(z) = \begin{cases} z, & \text{si } z > 0 \\ |\alpha| \cdot z, & \text{sinon} \end{cases}$$

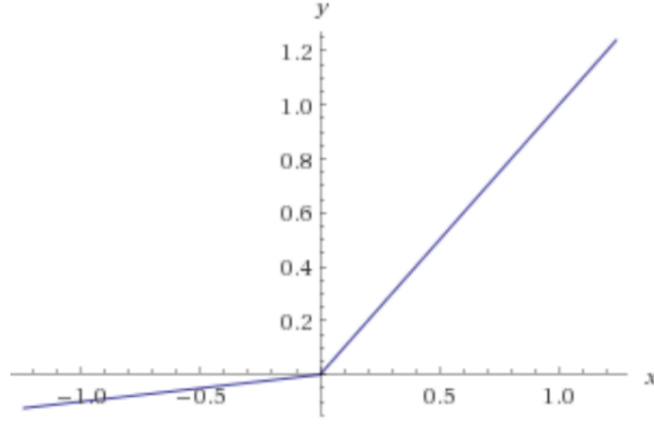


FIGURE 2.7 – Leaky ReLU

Théorème 2.1.1 (Théorème de Cybenko). *Soit $\phi(\cdot)$ une fonction bornée et continue. Soit I_m l'ensemble $[0, 1]^m$. Étant donné $\varepsilon > 0$ et toute fonction $f \in C(I_m)$, il existe $N \in \mathbb{N}$, $v_i, \omega_{0,i} \in \mathbb{R}$ et des vecteurs $\omega_i \in \mathbb{R}^m$, où $i = 1, \dots, N$, tels que*

$$F(x) = \sum_{i=1}^N v_i \phi(\omega_i^T x + \omega_{0,i})$$

approche f avec une erreur inférieure à ε pour tout $x \in I_m$, c'est-à-dire :

$$|F(x) - f(x)| < \varepsilon \quad \text{pour tout } x \in I_m.$$

Théorème 2.1.2 (Théorème d'approximation universelle, Hornik, Stinchcombe & White, 1991). *Soit $h : \mathbb{R} \rightarrow \mathbb{R}$ une fonction d'activation non polynomiale et mesurable bornée (ReLU, sigmoïde, tanh). Pour tout compact $K \subset \mathbb{R}^d$, toute fonction continue $g \in \mathcal{C}(K, \mathbb{R})$, et tout $\varepsilon > 0$, il existe un réseau feedforward f_θ à une couche cachée tel que :*

$$\sup_{x \in K} |f_\theta(x) - g(x)| < \varepsilon.$$

En d'autres termes : *Toute fonction continue peut être approximée par un réseau de neurones à une couche.*

Nous détaillerons dans la suite, les différentes architectures que nous appliquerons au hedging en marché incomplet.

2.2 Les FFNN

Le *feedforward neural network* (FFNN) constitue l'architecture de base du *deep hedging*, introduite par Buehler et al. [1]. Il s'agit d'un réseau neuronal artificiel composé de couches successives, où l'information circule **uniquement dans le sens direct** sans boucle de rétroaction. Cette structure en fait un outil idéal pour approximer des fonctions déterministes, comme une stratégie de couverture dépendant de l'état du marché à un instant donné.

Définition 2.1 (Feedforward Neural Network). Un *FFNN* est une fonction composée

$$F : \mathbb{R}^{N_0} \longrightarrow \mathbb{R}^{N_L}$$

définie par

$$F(x) = Z_L \circ h \circ Z_{L-1} \circ \cdots \circ h \circ Z_1(x),$$

où :

$L \geq 2$ est le nombre total de couches du réseau,

pour $\ell = 1, \dots, L$, la transformation affine Z_ℓ est donnée par

$$Z_\ell(x) = W_\ell x + b_\ell,$$

avec $W_\ell \in \mathbb{R}^{N_\ell \times N_{\ell-1}}$ et $b_\ell \in \mathbb{R}^{N_\ell}$,

$h : \mathbb{R} \rightarrow \mathbb{R}$ est une fonction d'activation non linéaire, appliquée composante par composante.

Les couches 1 à $L - 1$ sont appelées *couches cachées*, et leurs dimensions $N_1, \dots, N_{L-1}$ déterminent la capacité expressive du réseau.

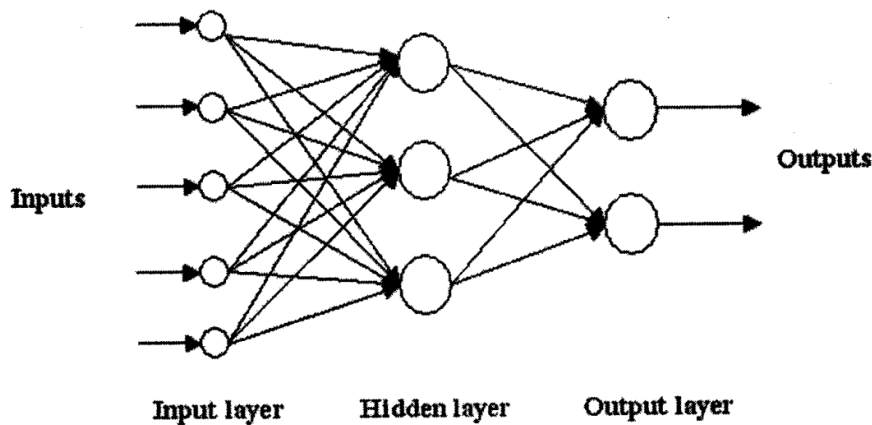


FIGURE 2.8 – Multi-layer feedforward neural network [20]

L'entraînement de ce modèle se fera par la méthode standard de *backpropagation*,

qui consiste à comparer, à chaque sortie d'un neurone, la sortie et la valeur réelle en minimisant une fonction de coût.

2.3 Les LSTM

Les réseaux de neurones récurrents (RNN)

Les réseaux de neurones récurrents (RNN) sont des réseaux de neurones conçus pour traiter des données séquentielles en conservant une « mémoire » de l'information précédente dans la séquence. Ils sont particulièrement utiles pour les tâches impliquant des séries temporelles ou des séquences de texte. Cependant, ils ont du mal à gérer les dépendances à long terme, un problème atténué dans des variantes plus sophistiquées comme les LSTM.

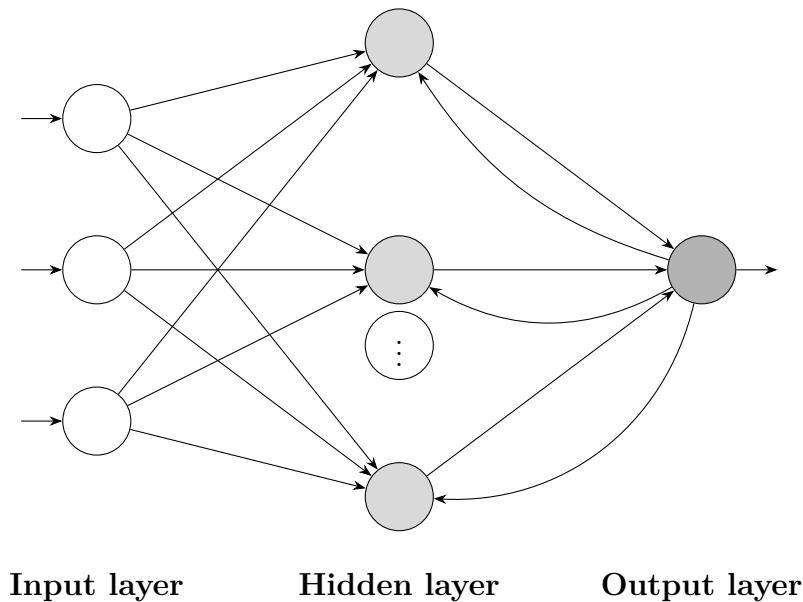


FIGURE 2.9 – Architecture dun réseau de neurones récurrent (RNN)

Long Short-Term Memory (LSTM)

Le réseau de neurones LSTM est une variante du réseau de neurones récurrent (RNN) spécifiquement conçue pour capturer les relations à long terme dans des séries chronologiques. Contrairement aux RNN traditionnels, qui ne peuvent prendre en compte que des dépendances à court terme, les LSTM sont capables de conserver et d'utiliser des informations sur de plus longues périodes grâce à une structure spéciale.

Les RNN sont confrontés au problème de la disparition du gradient, ce qui limite leur capacité à apprendre à partir de séquences de données longues. Les gradients contiennent les informations nécessaires pour mettre à jour les paramètres du RNN lors de l'apprentissage. Lorsque les gradients deviennent très petits, les mises à jour

des paramètres deviennent négligeables, ce qui signifie que l'apprentissage réel ne se produit pas de manière significative.

Notations et abréviations

- σ : fonction sigmoïde,
- $\tanh$: tangente hyperbolique,
- $\odot$: produit de Hadamard (élément par élément),
- $X_t \in \mathbb{R}^D$: vecteur d'entrée à l'instant t ,
- $H_t \in \mathbb{R}^d$: état caché à l'instant t ,
- $C_t \in \mathbb{R}^d$: état de cellule à l'instant t ,
- $W = (W_c, W_x)$: matrices de poids,
- $\theta = (W, b)$: paramètres entraînaables du réseau,
- $\delta_{t_n} \in \mathbb{R}^{D+1}$: vecteur de positions en actifs risqués à la date t_n ,
- E_k : erreur de prédiction à l'instant k ,
- η : taux d'apprentissage.

Problème du gradient vanishing [19]

À titre de simplification, nous allons travailler dans la suite avec un simple RNN à couche cachée avec une seule séquence de sortie.

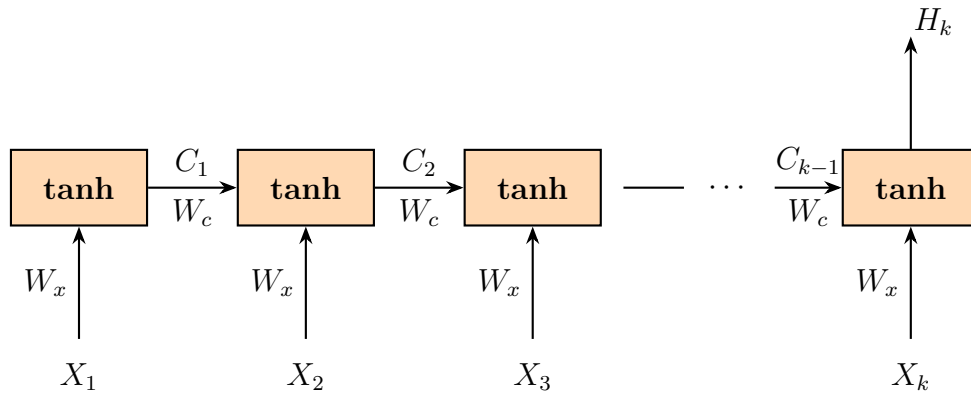


FIGURE 2.10 – Schéma d'un RNN simple à couche cachée.

Le réseau de neurones récurrents (RNN) que nous considérons est structuré pour traiter une séquence d'entrée de vecteurs, que nous notons $(X_1, X_2, \dots, X_k)$. À chaque pas de temps t , nous avons un vecteur d'entrée X_t .

En plus de l'entrée actuelle, le RNN maintient un vecteur d'état C_{t-1} . L'entrée complète au réseau à chaque pas de temps est donc le vecteur concaténé (C_{t-1}, X_t) .

Le RNN utilise deux matrices de poids, W_c et W_x , combinées en une seule matrice $W = (W_c, W_x)$. La dynamique du vecteur d'état s'écrit :

$$C_t = \tanh(W_c C_{t-1} + W_x X_t)$$

Pour mettre à jour les paramètres du réseau sur k pas de temps, nous calculons le gradient :

$$\frac{\partial E_k}{\partial W} = \frac{\partial E_k}{\partial H_k} \frac{\partial H_k}{\partial C_k} \left(\prod_{t=2}^k \frac{\partial C_t}{\partial C_{t-1}} \right) \frac{\partial C_1}{\partial W}$$

Alors :

$$\frac{\partial C_t}{\partial C_{t-1}} = \tanh'(W_c C_{t-1} + W_x X_t) \cdot W_c$$

La dernière expression tend à disparaître lorsque k est grand, car la dérivée de la fonction $\tanh$ est inférieure ou égale à 1 : $\frac{\partial E_k}{\partial W} \rightarrow 0$. Cela montre qu'aucun apprentissage significatif ne sera fait.

Pour résoudre ce problème, nous utilisons une variante de RNN appelée Long Short-Term Memory (LSTM). Un LSTM maintient un « état de cellule » C_t à chaque pas de temps, permettant au réseau de conserver les informations pertinentes sur une longue séquence de pas de temps.

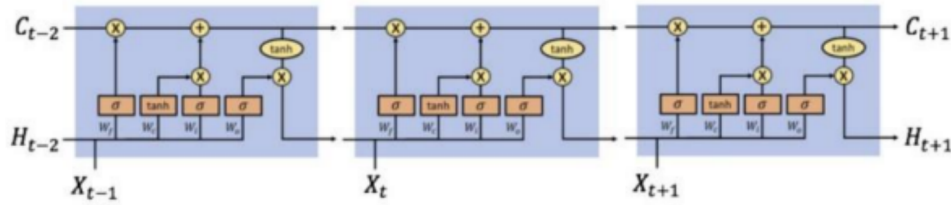


FIGURE 2.11 – Cellules LSTM à différents instants

Les réseaux LSTM utilisent trois « portes » spéciales pour contrôler et mettre à jour l'état de la cellule à chaque pas de temps : la porte d'oubli, la porte d'entrée et la porte de sortie.

La **porte d'oubli** détermine la proportion de l'état de la cellule précédente qui sera conservée :

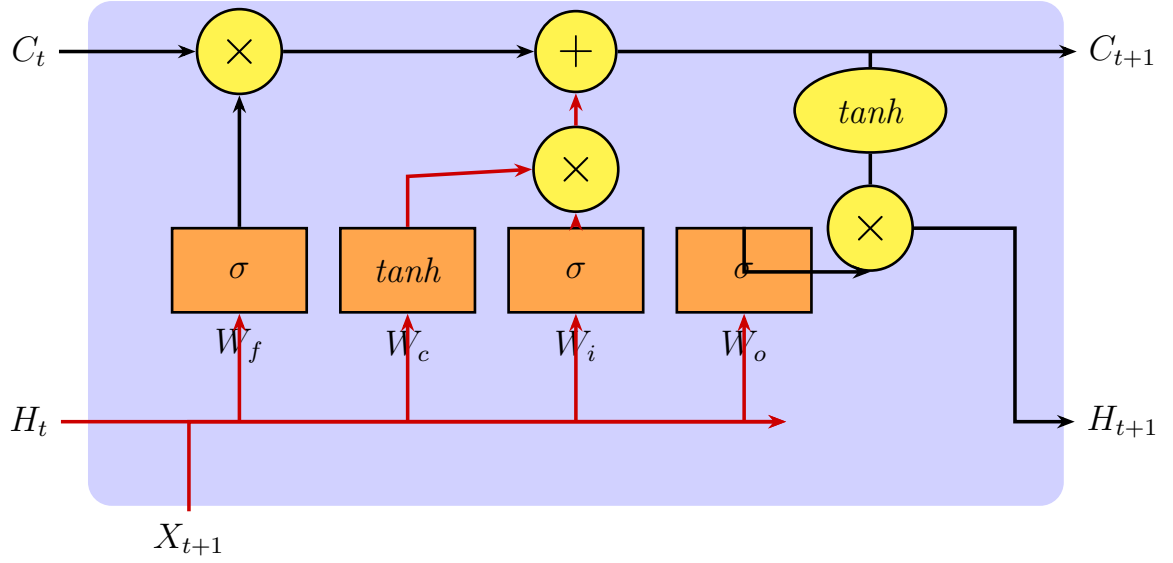
$$f_{t+1} = \sigma(W_f \cdot (H_t, X_{t+1}))$$

La **porte d'entrée** contrôle les nouvelles informations ajoutées à l'état de la cellule :

$$\tilde{C}_{t+1} = \tanh(W_c \cdot (H_t, X_{t+1})), \quad i_{t+1} = \sigma(W_i \cdot (H_t, X_{t+1}))$$

L'état de la cellule à l'étape $t + 1$ est alors mis à jour par :

$$C_{t+1} = f_{t+1} \odot C_t + i_{t+1} \odot \tilde{C}_{t+1}$$

FIGURE 2.12 – État de la cellule à l'étape de temps $t + 1$

La **porte de sortie** contrôle la quantité d'information transmise au réseau lors de l'étape suivante :

$$o_t = \sigma(W_o \odot (H_t, X_{t+1})), \quad H_{t+1} = o_t \odot \tanh(C_{t+1})$$

Backpropagation dans le temps dans les réseaux LSTM [19, 21]

Dans un LSTM, le vecteur d'état C_t est défini par :

$$C_t = f_t \cdot C_{t-1} + i_t \cdot \tilde{C}_t,$$

où :

$f_t = \sigma(W_f \cdot [H_{t-1}, X_t])$ est la porte d'oubli ;

$i_t = \sigma(W_i \cdot [H_{t-1}, X_t])$ est la porte d'entrée ;

$\tilde{C}_t = \tanh(W_c \cdot [H_{t-1}, X_t])$ est l'état candidat de la cellule.

Nous approximations la dérivée de C_t par rapport à C_{t-1} en utilisant la porte d'oubli f_t :

$$\frac{\partial C_t}{\partial C_{t-1}} \approx f_t$$

Ceci démontre que le comportement du gradient est aligné avec celui de la porte d'oubli. Ainsi, si la porte d'oubli choisit de retenir une information, le gradient correspondant ne tendra pas vers zéro.

Dans notre cas pratique, la stratégie de couverture est paramétrée par :

$$\delta_{t_n} = f_\theta(X_{t_n}),$$

où le vecteur d'entrée X_{t_n} correspond aux variables de marché observées à la date t_n .

La récursivité du LSTM s'écrit :

$$\begin{aligned}
 i_{t_n} &= \sigma(W_i[h_{t_{n-1}}, X_{t_n}] + b_i), \\
 f_{t_n} &= \sigma(W_f[h_{t_{n-1}}, X_{t_n}] + b_f), \\
 \tilde{c}_{t_n} &= \tanh(W_c[h_{t_{n-1}}, X_{t_n}] + b_c), \\
 c_{t_n} &= f_{t_n} \odot c_{t_{n-1}} + i_{t_n} \odot \tilde{c}_{t_n}, \\
 o_{t_n} &= \sigma(W_o[h_{t_{n-1}}, X_{t_n}] + b_o), \\
 h_{t_n} &= o_{t_n} \odot \tanh(c_{t_n}), \\
 \delta_{t_n} &= W_y h_{t_n} + b_y.
 \end{aligned}$$

Chapitre 3

Deep Hedging

3.1 Les fondamentaux

3.1.1 Approche *model-free* et *greek-free*

Le deep hedging, introduit par Buehler [1], se distingue radicalement des méthodes classiques par son caractère à la fois **model-free** et **greek-free**. Ces deux propriétés en font une approche particulièrement adaptée aux marchés incomplets et aux produits d'assurance complexes.

L'approche est dite **model-free** car elle ne repose sur aucune hypothèse structurelle concernant la dynamique du marché au-delà du générateur de scénarios utilisé pour l'entraînement. Contrairement aux méthodes traditionnelles qui exigent la spécification d'une mesure martingale équivalente (ex. Black-Scholes, Heston), le deep hedging ne nécessite pas de modèle de pricing. L'optimisation se fait directement sur des trajectoires simulées sous la mesure physique $\mathbb{P}$, ce qui permet d'intégrrer des *stylized facts* tels que les sauts, la volatilité stochastique ou l'impact de liquidité.

L'approche est également **greek-free**, car elle n'utilise aucune dérivée analytique du prix par rapport aux facteurs de risque. Dans les marchés incomplets, les grecques dépendent arbitrairement du choix de la mesure martingale [3], ce qui rend les stratégies basées sur elles suboptimales. Le deep hedging contourne ce problème en représentant la stratégie de couverture $\delta = (\delta_{t_n})_{n=0}^{N-1}$ comme la sortie d'un réseau de neurones, dont les poids sont optimisés par descente de gradient sur une fonction de perte appliquée au P&L terminal :

$$\min_{\theta} \mathbb{E}^{\mathbb{P}} \left[\ell \left(H - V_0 - (\delta^{\theta} \cdot S)_T \right) \right],$$

où θ désigne l'ensemble des paramètres entraînaables du réseau.

En résumé, le caractère model-free et greek-free du deep hedging lui confère une grande flexibilité : il s'adapte à n'importe quelle dynamique simulable, à n'importe quel payoff path-dépendant, et à n'importe quelle mesure de risque convexe.

3.1.2 Optimisation directe du P&L

Le deep hedging repose sur une optimisation globale et directe du profit et perte (P&L) terminal. L'objectif n'est plus de répliquer un payoff, mais de minimiser une mesure de risque appliquée à l'erreur de couverture finale.

Soit H le risque à couvrir et V_T^δ la valeur terminale du portefeuille de couverture sous la stratégie δ . Le *P&L terminal* en tant qu'assureur est défini par :

$$\text{P\&L} = H - V_T^\delta.$$

L'assureur cherche à minimiser une **fonction de perte convexe** ℓ appliquée à ce P&L :

$$\min_{\delta \in \mathcal{H}} \mathbb{E}^\mathbb{P} \left[\ell \left(H - V_T^\delta \right) \right],$$

où $\mathbb{P}$ est la mesure physique utilisée pour générer des scénarios réalistes [1, 3].

Cette formulation présente plusieurs avantages :

- **Globale** : les décisions de trading sur $[0, T]$ sont optimisées conjointement, permettant d'exploiter la diversification temporelle du risque sur longue maturité [3].
- **Alignée sur les objectifs économiques** : en choisissant $\ell(x) = x^2 \mathbf{1}_{\{x > 0\}}$ (SMSE), on pénalise uniquement les pertes, conduisant à des stratégies « bullish » captant la prime de risque actions [3].
- **Indépendante de toute mesure martingale** : l'optimisation se fait sous $\mathbb{P}$, évitant l'arbitraire du choix d'une mesure risque-neutre en marché incomplet [3].

En résumé, l'optimisation directe du P&L terminal constitue le principal du deep hedging : elle remplace la logique de réplication parfaite par une logique de gestion optimale du risque résiduel, conforme aux contraintes et objectifs réels des assureurs.

3.2 Formulation mathématique

3.2.1 Marché en temps discret avec frictions

Le cadre du deep hedging est formulé dans un *marché en temps discret* avec **frictions réalistes**. On considère un horizon fini $T > 0$ divisé en N périodes de rebalancement : $0 = t_0 < t_1 < \dots < t_N = T$.

Le marché comporte $D + 1$ actifs dont un actif sans risque $B = (B_{t_n})_{n=0,\dots,N}$ avec $B_{t_n} = e^{rt_n}$, où r est le taux sans risque constant ; et D actifs risqués $\bar{S}^{(b)} = (S^{(1,b)}, \dots, S^{(D,b)})$, qui peuvent être actions, des options ou des actifs fixed income tels que les obligations et les swaps . Contrairement aux modèles idéalisés, le deep hedging intègre explicitement des **frictions de marché** :

- **Coûts de transaction proportionnels** : le coût d'achat/vente d'une quantité $\Delta\delta_{t_n}$ d'actif j est $\lambda^{(j)}|\Delta\delta_{t_n}|S_{t_n}^{(j,b)}$, où $\lambda^{(j)} \geq 0$ est le taux de friction [1].
- **Impact de liquidité** : Picha [9] modélise dans une extension récente le grand investisseur dont les ordres influencent les prix .
- **Contraintes de liquidité ou de risque** : des limites peuvent être imposées sur les positions δ_{t_n} (VaR, levier maximal) [1].

Toutefois, notre travail se basera sur **l'intégration des coûts de transactions** et **une régularisation ou pénalisation** des stratégies de trading.

Définition 3.2.1. Sur la grille $0 = t_0 < \dots < t_N = T$, la valeur du portefeuille de couverture à la date t_n , notée $V_{t_n}^\delta$, évolue selon une dynamique **auto-financée avec frictions** :

$$V_{t_{n+1}}^\delta = e^{r\Delta t}V_{t_n}^\delta + \delta_{t_n}^\top (\mathbf{S}_{t_{n+1}} - e^{r\Delta t}\mathbf{S}_{t_n}) - C_{t_{n+1}}(\delta),$$

ou $C_{t_{n+1}}(\delta) = \sum_j \lambda^{(j)}|\delta_{t_n}^{(j)} - \delta_{t_{n-1}}^{(j)}|S_{t_n}^{(j)}$ sont les couts de transaction proportionnels cumulés sur la période. La dérivation depuis le cas continu est en annexe B.

3.3 Algorithmique

3.3.1 Algorithme deep hedging

Algorithme 1 Deep Hedging pour la couverture d'un passif GMAB path-dépendant

Entrée :

$X_{i,t_0} = (S_0^{(i)}, V_0^{(i)}, Op_0^{(i)}, r_0^{(i)}, t_0)$: état initial de la trajectoire i ,

N_{batch} : taille du mini-batch,

T : nombre de pas de temps,

$\mathcal{F}_\theta$: réseau neuronal (LSTM ou FFNN¹),

V_0 : valeur initiale du portefeuille,

θ_j : poids du réseau après le mini-batch $j - 1$.

Sortie :

$\delta_{i,t} \in \mathbb{R}^{n_{\text{assets}}}$: stratégie de couverture,

$V_{i,T}$: valeur terminale du portefeuille,

$\text{PnL}_i = H_i - V_{i,T}$: profit et perte final,

θ_{j+1} .

1 : **for** $i = 1, \dots, N_{\text{batch}}$ **do**

2 : $V_{i,0} \leftarrow V_0$

3 : $\delta_{i,-1} \leftarrow \mathbf{0}$

4 : **for** $t = 0, \dots, T - 1$ **do**

5 : Construire le vecteur d'entrée :

$$X_{i,t} = [\log S_t^{(i)}, \log Op_{t,1}^{(i)}, \dots, \log Op_{t,5}^{(i)}, V_t^{(i)}, r_t^{(i)}, t/T]$$

6 : Ajouter la richesse courante :

$$X_{i,t}^{(V)} = [S_t^{(i)}, V_t^{(i)}, V_{i,t}/V_0]$$

7 : Prédire les positions :

$$\delta_{i,t} \leftarrow \mathcal{F}_{\theta_j}(X_{i,t}^{(V)})$$

8 : Calculer le coût de transaction :

$$\text{cost}_{i,t} \leftarrow \lambda_{\text{cost}} \cdot \|\delta_{i,t} - \delta_{i,t-1}\|_1 \cdot A_{i,t}$$

9 : Calculer le gain sur $[t, t + 1]$:

$$\text{gain}_{i,t} \leftarrow \delta_{i,t} \cdot (A_{i,t+1} - e^{r_t^{(i)} \Delta t} A_{i,t})$$

10 : Mettre à jour le portefeuille :

$$V_{i,t+1} \leftarrow e^{r_t^{(i)} \Delta t} V_{i,t} + \text{gain}_{i,t} - \text{cost}_{i,t}$$

11 : **end for**

12 : Calculer le passif terminal :

$$H_i \leftarrow N \cdot \max_{0 \leq s \leq T} S_s^{(i)}$$

13 : Calculer le PnL :

$$\text{PnL}_i \leftarrow H_i - V_{i,T}$$

14 : **end for**

15 : Calculer la perte :

$$\mathcal{L} = \frac{1}{\lambda} \log \mathbb{E} \left[e^{\lambda \cdot \max(\text{PnL}, 0)} \right] \quad \text{ou} \quad \text{CVaR}_\alpha(\text{PnL})$$

16 : Mettre à jour les poids θ_j via Adam :

$$\theta_{j+1} \leftarrow \theta_j - \eta_t \nabla_{\theta_j} \mathcal{L}$$

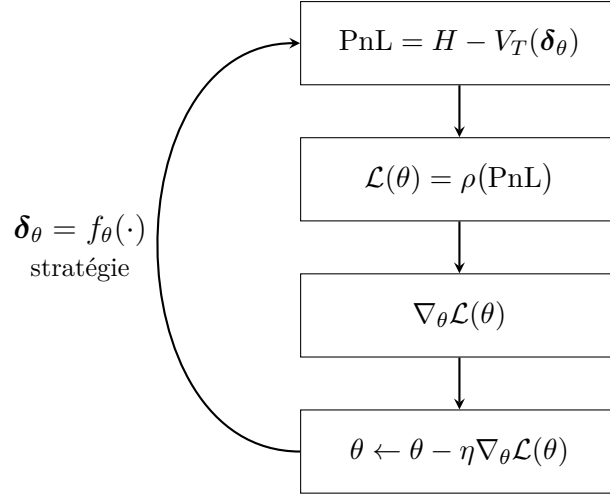


FIGURE 3.1 – Pipeline d'entraînement du deep hedging

Échantillonnage Monte Carlo et descente de gradient stochastique

Pour le mini-batch j avec poids $\theta = \theta_j$, soit $B_j := \{PnL_{i,j}\}_{i=1}^{N_{\text{batch}}}$ constitué d'erreurs de couverture simulées :

$$PnL_{i,j} := H_{T,i} - V_{T,i}^{\delta_{\theta_j}}.$$

L'estimateur empirique de la fonction objectif s'écrit :

$$\hat{J}(\theta_j) := \frac{1}{N_{\text{batch}}} \sum_{i=1}^{N_{\text{batch}}} \rho(PnL_{i,j}).$$

À chaque itération j , l'algorithme procède comme suit :

- (1) Simuler un mini-batch B_j de N_{batch} chemins de marché,
- (2) Calculer les erreurs de couverture $\{PnL_{i,j}\}_{i=1}^{N_{\text{batch}}}$,
- (3) Évaluer $\hat{J}(B_j)$ et son gradient $\nabla_{\theta} \hat{J}(B_j)$,
- (4) Calculer η par Adam,
- (5) Mettre à jour les paramètres : $\theta_{j+1} \leftarrow \theta_j - \eta \nabla_{\theta_j} \hat{J}(\theta_j)$.

1. Les architectures LSTM et FFNN sont implémentées à l'aide de la bibliothèque **TensorFlow** sous Python.

Chapitre 4

Modélisation du marché

Nous présentons dans ce chapitre une implémentation du deep hedging pour la couverture du produit d'assurance dans un marché incomplet discret (voir section 3.2.1), une stratégie paramétrée par *différentes techniques de deep learning à complexité graduelle*, et deux mesures de risque convexes distinctes (voir section 1.2). Ensuite, nous comparerons ces résultats à des benchmarks classiques.

4.1 Hypothèses de la modélisation

Pour la suite du travail, nous allons considérer les hypothèses suivantes.

Hypothèses démographiques

- **Absence de mortalité** durant la période du contrat.
- **Absence de rachat anticipé.**

Hypothèses contractuelles

- **Prime unique** : chaque assuré verse une prime unique P_0 à la souscription ($t = 0$).
- **Produit GMAB** avec un mécanisme de ratchet annuel.
- **Frais et charges** : les frais de gestion, frais de garantie et charges opérationnelles sont supposés inclus implicitement dans la dynamique du compte, ou négligés.

Hypothèses de marché

- **Marché incomplet.**
- **Absence d'arbitrage** : le marché est supposé sans opportunité d'arbitrage et admet au moins une mesure risque-neutre $\mathbb{Q}$.
- **Trajectoire du marché** : on suppose que le marché est simulé sous le modèle d'Heston.

Hypothèses de couverture

- **Rebalancement discret** : les stratégies de couverture sont ajustées à des dates discrètes $t_0 < t_1 < \dots < t_N = T$.

- **Actifs de couverture** : l'assureur peut investir dans l'indice sous-jacent et des options.
- **Absence de contraintes réglementaires** : les contraintes prudentielles (Solvabilité II, limites d'investissement) ne sont pas prises en compte.

Hypothèses numériques

- **Simulation Monte Carlo** : les espérances conditionnelles et les distributions de pertes sont estimées par simulation de Monte Carlo.
- **Convergence numérique** : le nombre de trajectoires simulées et la discrétisation temporelle sont choisis de manière à garantir la stabilité des résultats.

4.2 Modélisation du marché

Nous reprenons ici la modélisation stochastique du modèle gouvernant notre marché, le modèle d'Heston 1.4 qui, du fait de sa volatilité stochastique, est mieux adapté à la couverture de produits d'assurance-vie à long terme. Ce cadre permet de simuler de manière cohérente et réaliste les trajectoires de marché.

4.2.1 Pricing des instruments de couverture

Pricing des options

Principe fondamental

Notons $f_{t,T}(x)$ la densité de probabilité du rendement logarithmique X_T . La fonction caractéristique $\Upsilon_{t,T}(i\omega)$ de $f_{t,T}(x)$ coïncide avec la transformée de Fourier de cette densité :

$$\Upsilon_{t,T}(i\omega) = \mathbb{E}^{\mathbb{Q}}(e^{i\omega X_T} \mid \mathcal{F}_t) = \int_0^\infty e^{i\omega x} f_{t,T}(x) dx.$$

La densité se retrouve par transformation inverse :

$$f_{t,T}(x) = \frac{1}{\pi} \Re \left(\int_0^{+\infty} \Upsilon_{t,T}(i\omega) e^{-i\omega x} d\omega \right).$$

La fonction caractéristique de X_s sous $\mathbb{Q}$ s'écrit :

$$\mathbb{E}^{\mathbb{Q}}(e^{\omega X_s} \mid \mathcal{F}_t) = \left(\frac{S_s}{S_0} \right)^\omega \exp(A(\omega, t, s) + B(\omega, t, s) V_t),$$

Avec :

$$d = \sqrt{(\rho\sigma\omega - \kappa)^2 + \sigma^2(\omega - \omega^2)}, \quad g = \frac{\kappa - \rho\sigma\omega + d}{\kappa - \rho\sigma\omega - d},$$

$$A(\omega, t, s) = r\omega(s-t) + \frac{\kappa\theta}{\sigma^2} \left[(\kappa - \rho\sigma\omega + d)(s-t) - 2 \ln \left(\frac{1 - g e^{d(s-t)}}{1 - g} \right) \right],$$

$$B(\omega, t, s) = \frac{\kappa - \rho\sigma\omega + d}{\sigma^2} \frac{1 - e^{d(s-t)}}{1 - g e^{d(s-t)}}.$$

Approximation de la densité par FFT

Soit M le nombre de points de la transformée discrète de Fourier. On définit :

$$\Delta_\omega = \frac{2\pi}{M\Delta_x}, \quad \omega_j = (j-1)\Delta_\omega, \quad x_k = -\frac{M}{2}\Delta_x + (k-1)\Delta_x.$$

La densité $f_{t,T}(x_k)$ est approchée par :

$$f(x_k, \delta_t) \approx \frac{2}{M\Delta_x} \Re \left(\sum_{m=1}^M \varrho_m \Upsilon_{t,T}(i\omega_m) (-1)^{m-1} e^{-i\frac{2\pi}{M}(m-1)(k-1)} \right),$$

avec $\varrho_m = \frac{1}{2} \mathbf{1}_{\{m=1\}} + \mathbf{1}_{\{m \neq 1\}}$.

Remarque : Pour accélérer les calculs lors de la simulation Monte Carlo, les prix d'options sont pré-calculés *une seule fois* sur une grille tridimensionnelle (S, v, τ) via la FFT, puis obtenus par interpolation linéaire tridimensionnelle pendant la simulation.

Calcul des prix d'options par FFT

Les prix de calls et puts européens sont calculés par sommation discrète :

$$C(t, k, \delta_t) = e^{-r_t \delta_t} \sum_{k=1}^M \left(S_0 e^{x_k} - K \right) \mathbb{I}_{\{S_0 e^{x_k} \geq K\}} f_{t,T}(x_k) \Delta_x, \quad (4.1)$$

$$P(t, k, \delta_t) = e^{-r_t \delta_t} \sum_{k=1}^M \left(K - S_0 e^{x_k} \right) \mathbb{I}_{\{S_0 e^{x_k} \leq K\}} f_{t,T}(x_k) \Delta_x. \quad (4.2)$$

4.2.2 Générateurs de scénarios

Le deep hedging repose sur la simulation de trajectoires de marché sous la mesure physique $\mathbb{P}$. Ces trajectoires sont produites par un **générateur de scénarios**.

Un générateur de scénarios est une application mesurable

$$G_\xi : \mathbb{R}^{N \times d} \longrightarrow \mathbb{R}^{(N+1) \times d}, \quad (\Delta B_n)_{n=0, \dots, N} \longmapsto (S_n)_{n=0, \dots, N},$$

où $\xi = (\kappa, \theta, \sigma, \rho, v_0, \mu)$.

Générateur classique d'Heston Sous $\mathbb{P}$, la dynamique simulée par schéma d'Euler est :

$$S_{n+1} = S_n \exp \left[\left(\mu - \frac{1}{2} v_n \right) \Delta t + \sqrt{v_n} \sqrt{\Delta t} Z_n^{(1)} \right], \quad (4.3)$$

$$v_{n+1} = \max \left(0, v_n + \kappa(\theta - v_n) \Delta t + \sigma \sqrt{v_n} \sqrt{\Delta t} Z_n^{(2)} \right). \quad (4.4)$$

Générateur neuronal : Neural SDE Pour pallier cette limitation, on peut remplacer le modèle paramétrique par une *équation différentielle stochastique pilotée par un réseau de neurones* :

$$dS_t = \mu_\phi(t, S_t) dt + \sigma_\psi(t, S_t) dW_t,$$

où μ_ϕ et σ_ψ sont des approximateurs universels paramétrés par des réseaux de neurones.

Remarque. On utilise les trajectoires générées sous $\mathbb{Q}$ du prix de l'indice S_t pour pricer les options. Celles de S_t sous $\mathbb{P}$ serviront pour le hedging dans le marché incomplet réel.

4.2.3 Discrétisation et simulation

Nous considérons un marché discret avec un pas de temps mensuel $\Delta t = 1/12$.

Génération des chocs gaussiens Pour chaque trajectoire $m = 1, \dots, M$ et pas de temps $k = 0, \dots, N - 1$:

$$Z_{m,k}^{(1)}, Z_{m,k}^{(2)} \sim \mathcal{N}(0, 1).$$

Les innovations corrélées sous $\mathbb{P}$ sont :

$$\varepsilon_{m,k}^{(1),\mathbb{P}} = Z_{m,k}^{(1)}, \quad \varepsilon_{m,k}^{(2),\mathbb{P}} = \rho Z_{m,k}^{(1)} + \sqrt{1 - \rho^2} Z_{m,k}^{(2)}.$$

1. Simulation sous $\mathbb{P}$

$$v_{k+1}^{(m),\mathbb{P}} = \max \left(0, v_k^{(m),\mathbb{P}} + \kappa(\theta - v_k^{(m),\mathbb{P}}) \Delta t + \sigma \sqrt{v_k^{(m),\mathbb{P}}} \sqrt{\Delta t} \varepsilon_{m,k}^{(2),\mathbb{P}} \right),$$

$$S_{k+1}^{(m),\mathbb{P}} = S_k^{(m),\mathbb{P}} \exp \left(\left(\mu - \frac{1}{2} v_k^{(m),\mathbb{P}} \right) \Delta t + \sqrt{v_k^{(m),\mathbb{P}}} \sqrt{\Delta t} \varepsilon_{m,k}^{(1),\mathbb{P}} \right).$$

2. Détermination des prix de risque Le changement de mesure s'effectue via le théorème de Girsanov. Les chocs sous $\mathbb{Q}$ sont obtenus par :

$$Z_{m,k}^{(1),\mathbb{Q}} = Z_{m,k}^{(1)} + \lambda_S \Delta t, \quad Z_{m,k}^{(2),\mathbb{Q}} = Z_{m,k}^{(2)} + \lambda \sqrt{v_k^{(m),\mathbb{Q}}} \sqrt{\Delta t}.$$

3. Simulation sous $\mathbb{Q}$ Avec les paramètres ajustés $\tilde{\kappa}$ et $\tilde{\theta}$:

$$v_{k+1}^{(m),\mathbb{Q}} = \max\left(0, v_k^{(m),\mathbb{Q}} + \tilde{\kappa}(\tilde{\theta} - v_k^{(m),\mathbb{Q}})\Delta t + \sigma\sqrt{v_k^{(m),\mathbb{Q}}}\sqrt{\Delta t}\varepsilon_{m,k}^{(2),\mathbb{Q}}\right),$$

$$S_{k+1}^{(m),\mathbb{Q}} = S_k^{(m),\mathbb{Q}} \exp\left((r - \frac{1}{2}v_k^{(m),\mathbb{Q}})\Delta t + \sqrt{v_k^{(m),\mathbb{Q}}}\sqrt{\Delta t}\varepsilon_{m,k}^{(1),\mathbb{Q}}\right).$$

Pour la suite, nous avons généré les trajectoires avec les initialisations suivantes :

κ	θ	σ	ρ	V_0	μ	S_0
1.5	0.04	0.3	-0.7	0.04	0.08	10

TABLE 4.1 – Initialisation des paramètres sous $\mathbb{P}$

$\tilde{\kappa}$	$\tilde{\theta}$	σ	ρ	V_0	r
2.0	0.03	0.3	-0.7	0.04	0.02

TABLE 4.2 – Initialisation des paramètres sous $\mathbb{Q}$

Il s'agit d'un nombre total de $M = 20\,000$ trajectoires sous $\mathbb{P}$ et leurs correspondantes sous $\mathbb{Q}$, qui serviront de base pour l'entraînement, la validation et le test.

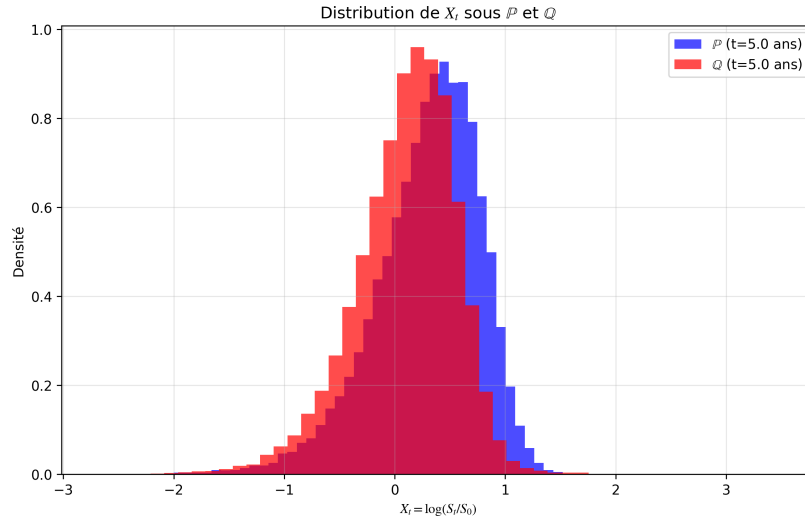


FIGURE 4.1 – Distribution du log-rendement

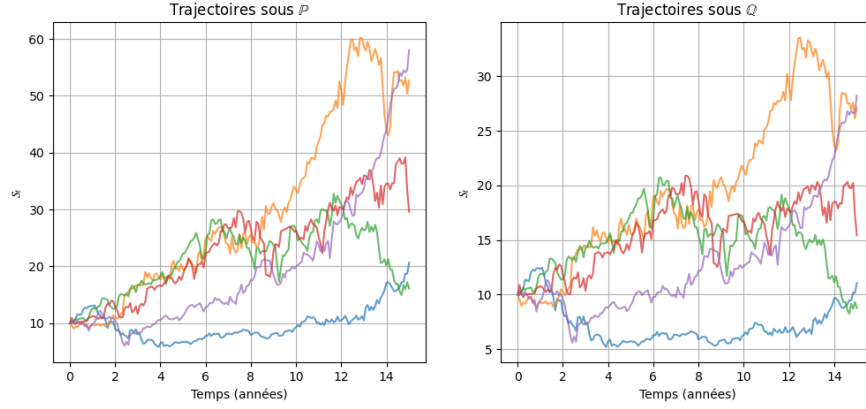


FIGURE 4.2 – Quelques trajectoires sous Heston

Asymétrie et dispersion des rendements Les figures 4.2 et 4.1 mettent en évidence une dynamique non triviale du processus de prix sous les deux mesures.

Sous la mesure physique $\mathbb{P}$, la distribution de $X_t = \log(S_t/S_0)$ présente une asymétrie positive marquée : une queue droite étendue jusqu'à $X_t \approx +1.2$ avec une queue gauche modérée. Cela reflète la probabilité non négligeable de fortes hausses de marché, caractéristique des modèles avec volatilité stochastique et prime de risque sur la variance $\lambda > 0$.

En revanche, sous la mesure risque-neutre $\mathbb{Q}$, la distribution de X_t est plus concentrée, avec un pic nettement plus étroit et un support plus centré. Ce phénomène s'explique par la modification de la dynamique de la variance lors du changement de mesure :

$$\tilde{\kappa} = \kappa + \sigma\lambda\sqrt{1 - \rho^2} > \kappa, \quad \tilde{\theta} = \frac{\kappa\theta}{\tilde{\kappa}} < \theta.$$

Conséquence pour le pricing et la couverture Cette réduction de la queue droite sous $\mathbb{Q}$ a une implication directe sur le pricing des options lookback : le payoff $(Z_T - S_T)^+$ dépend du maximum historique Z_T , qui est sensible aux hausses extrêmes. Or, sous $\mathbb{Q}$, ces hausses sont moins fréquentes et moins intenses que sous $\mathbb{P}$, rendant la formule fermée du delta-hedging inefficace sous $\mathbb{P}$.

4.2.4 Instruments de couverture et simulations sous deux mesures

Notre objectif est de couvrir le passif de l'assureur qui ressemble à une option lookback de maturité $T = 15$ ans. Notre marché simulé contient, en plus de l'actif sous-jacent S_t , deux ou six options de maturité 1 an, introduites à chaque début d'année t_a et tradées plusieurs fois dans l'année (ajustement mensuel de la position). Le premier lot d'options comprend des puts de strikes $K \in \{1.0 S_{t_a}, 1.1 S_{t_a}, 1.2 S_{t_a}\}$; dans la deuxième simulation, on ajoute un lot de trois calls de strikes $K \in \{0.8 S_{t_a}, 0.9 S_{t_a}, 1.0 S_{t_a}\}$.

La simulation des scénarios de marché pour la couverture de notre produit repose sur un **cadre hybride** distinguant deux mesures probabilistes :

- La *mesure historique* $\mathbb{P}$, utilisée pour simuler les trajectoires réelles telles que S_t , Z_t , H_T , afin de refléter la fréquence réelle des événements extrêmes.
- La *mesure risque-neutre* $\mathbb{Q}$ utilisée pour simuler les prix des **instruments de couverture** tels que les options dans notre cas, garantissant l'absence d'arbitrage.

Ainsi, le payoff $H_T = \max(Z_T, S_T)$ est simulé sous $\mathbb{P}$, tandis que la stratégie de couverture $\delta_t = f_\theta(X_t)$ est optimisée via la mesure de risque ρ appliquée au P&L $H_T - V_T^\delta$. Cette séparation est cruciale car elle permet d'évaluer le risque réel sous $\mathbb{P}$ tout en construisant une couverture dynamique cohérente avec les prix du marché via $\mathbb{Q}$.

4.3 Delta-hedging sous Heston

Sous la mesure risque-neutre $\mathbb{Q}$, le prix à l'instant t d'une option put européenne de strike K et maturité T est :

$$P(t, S_t, V_t) = e^{-r(T-t)} \mathbb{E}^{\mathbb{Q}}[(K - S_T)^+ | \mathcal{F}_t].$$

En notant $X = \ln(S_T/S_t)$, le prix s'écrit :

$$P(t, S_t) = e^{-r\delta_t} \int_{-\infty}^{+\infty} (K - S_t e^x)^+ f_{t,T}(x) dx.$$

Par dérivation sous le signe intégral, le delta du put est :

$$\Delta_P := \frac{\partial P}{\partial S_t} = -e^{-r\tau} \int_{-\infty}^{\ln(K/S_t)} e^x f_{t,T}(x) dx = -\frac{e^{-r\tau}}{S_t} \mathbb{E}^{\mathbb{Q}}[S_T \mathbf{1}_{\{S_T \leq K\}} | \mathcal{F}_t].$$

Delta-hedging d'une option path-dépendante

Dans le cadre du GMAB, le passif à l'échéance T est :

$$L_T = S_T + (Z_T - S_T)^+, \quad Z_T := \max_{0 \leq s < T} S_s.$$

À un instant $t < T$, le delta de la composante optionnelle comporte deux cas :

— **Cas 1** : $S_t < Z_t$. Alors Z_t est indépendant de S_t et :

$$\Delta_P(t) = -e^{-r(T-t)} \mathbb{E}^{\mathbb{Q}} \left[\frac{S_T}{S_t} \mathbf{1}_{\{S_T < Z_T\}} \mid \mathcal{F}_t \right].$$

— **Cas 2** : $S_t = Z_t$. Le delta comporte un terme supplémentaire dû à la sensibilité de Z_T à S_t :

$$\Delta_P(t) = -e^{-r(T-t)} \mathbb{E}^{\mathbb{Q}} \left[\frac{S_T}{S_t} \mathbf{1}_{\{S_T < Z_T\}} \mid \mathcal{F}_t \right] + e^{-r(T-t)} \mathbb{Q}(Z_T = S_t \mid \mathcal{F}_t).$$

Approximation numérique via FFT

L'approximation du delta path-dépendant par FFT procède comme suit :

- (1) Simuler l'historique $\{S_s\}_{s \leq t}$ sous $\mathbb{P}$,
- (2) Calculer $Z_t = \max_{s \leq t} S_s$,
- (3) À chaque pas t , traiter le payoff résiduel comme une option européenne avec strike fixe $K = Z_t$,
- (4) Appliquer la formule de delta du put européen :

$$\Delta_P(t) \approx -e^{-r(T-t)} \sum_{k: S_t e^{x_k} < Z_t} e^{x_k} f_{t,T}(x_k) \Delta x.$$

Cette approximation est valide tant que $S_t < Z_t$, ce qui est le cas dans la majorité des scénarios.

Chapitre 5

Application à un produit GMAB

5.1 Spécification et passif total du portefeuille de Variable Annuities

Soit un portefeuille homogène de N assurés âgés de $x = 40$ ans à la souscription $t = 0$, chacun versant une prime unique P suffisant pour acheter un indice dont $S_0 = 10$, pendant $T = 15$ ans. La valeur du compte est investie en unités de compte liées à un indice boursier. Le payoff est :

$$L_T = N \left(S_T + (Z_T - S_T)_+ \right), \quad \text{avec } Z_T = \max_{0 \leq s \leq T-1} S_s.$$

Pour la suite de la modélisation, on prend $N = 1$. Le passif de l'assureur est :

$$\mathcal{L}_t = \mathbb{E}^{\mathbb{Q}} \left[e^{-\int_t^T r_s ds} L_T \mid \mathcal{F}_t \right].$$

La valeur initiale du portefeuille

La valeur initiale du portefeuille de couverture, notée V_0 , représente le capital économique nécessaire à l'émetteur pour financer la stratégie de couverture :

$$V_0 = \mathbb{E}^{\mathbb{Q}} \left[e^{-rT} \cdot H_T \right],$$

où $H_T = \max(Z_T - S_T, 0)$ est le payoff de la garantie GMAB lookback.

Cette définition assure trois propriétés essentielles :

- (1) **Auto-financement** : le portefeuille débute avec un capital suffisant pour couvrir, en espérance, l'engagement futur.
- (2) **Équité comptable** : le prix facturé au client reflète le coût réel de la couverture.
- (3) **Comparabilité** : les stratégies de deep hedging et de delta-hedging sont évaluées à partir du même niveau de capital initial.

Dans notre implémentation, V_0 est estimé par :

$$V_0 \approx \frac{1}{M} \sum_{i=1}^M e^{-rT} H_T^{(i)},$$

avec $M = 17\,000/3\,000$ trajectoires respectivement pour l'entraînement, la validation. Ensuite, nous générons $M = 5\,000$ pour chacun des 100 seeds différents du générateur pour tester notre modèle déjà entraîné.

5.2 Optimisation des architectures

Notre payoff $H = f(Z_T, S_T) = (Z_T - S_T)_+$ est une fonction continue. De ce fait, d'après les théorèmes 2.1.1 et 2.1.2, une seule couche peut approximer notre payoff. Néanmoins, pour optimiser la profondeur de notre architecture dans le but de choisir la mieux convenable, nous avons construit plusieurs architectures sur la base du **nombre de couches cachées**, **nombre de neurones par couches** et du **taux d'apprentissage**.

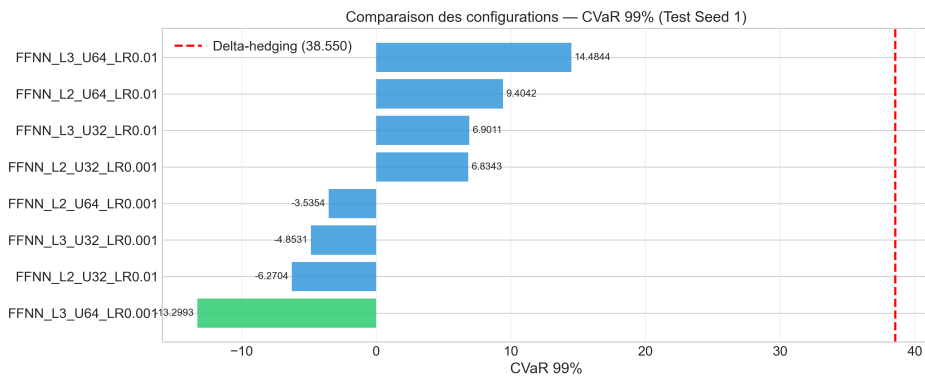


FIGURE 5.1 – Performance des différentes architectures de FFNN

Architecture	1	2	3	4	5	6	7	8
Mean	-48.32	-43.62	-34.22	-38.83	-25.86	-38.22	-39.39	-32.47
CVaR95	-18.90	-14.31	-10.37	-10.66	-1.31	0.09	1.76	8.63
CVaR99	-13.29	-6.27	-4.85	-3.53	6.83	6.90	9.40	14.48
VaR99	-16.77	-10.95	-9.11	-8.80	1.33	3.67	5.94	11.13

TABLE 5.1 – Performance des modèles FFNN

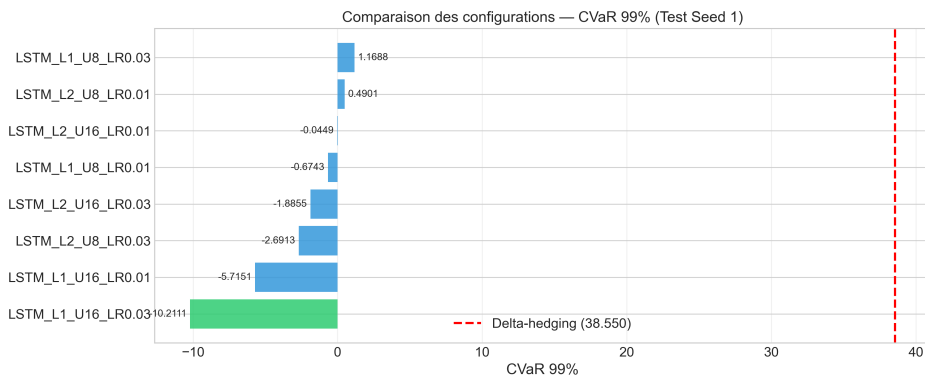


FIGURE 5.2 – Performance des différentes architectures de LSTM

Architecture	1	2	3	4	5	6	7	8
Mean	-39.04	-32.72	-44.73	-40.49	-27.34	-38.92	-31.47	-30.67
CVaR95	-16.15	-11.33	-13.69	-9.72	-6.75	-7.79	-6.18	-7.03
CVaR99	-10.21	-5.71	-2.69	-1.88	-0.67	-0.04	0.49	1.16
VaR99	-14.24	-9.57	-10.79	-6.83	-4.91	-4.71	-3.28	-4.74

TABLE 5.2 – Performance des modèles LSTM

Sur la base des figures et tableaux ci-dessus et sur la base du critère CVaR99, il s'avère que l'architecture 1 correspondant à *FFNN_L3_U64_LR0.001* soit 3 couches cachées dans cet ordre 64, 32, 16 et un taux d'apprentissage de 0.001 dans le cas FFNN et dans le cas LSTM, il correspond à *LSTM_L1_U16_LR0.03* soit une couche cachée de 16 neurones et un taux de 0.03. Ces architectures choisies issus de cette optimisation constitueront le socle de notre étude à posteriori.

5.3 Résultats

Tous les résultats ci-dessous sont les moyennes issus d'une simulation monte carlo sur 100 seeds différents sur le générateur pour garantir la robustesse des résultats. De ce fait, chaque seed est exécuté sur 150 epochs avec taille de batch 2048, pour garantir un bon apprentissage du réseau et rappelons que ces résultats sont d'un échantillon différent de celui d'entraînement. Chaque entraînement a une complexité de 45 minutes malgré le nombre raisonnablement bas de neurones pour chaque modèle, le nombre de paramètre explose facilement.

Dans toute la suite, le profit et perte terminal de l'assureur est défini par $\text{PnL} = H_T - V_T^\delta$ où H_T est le passif à couvrir et V_T^δ la valeur terminale du portefeuille de couverture. Un **PnL positif** correspond donc à une **perte pour l'assureur**, donc le portefeuille ne suffit pas à couvrir le passif, tandis qu'un PnL négatif représente un surplus. De ce fait, les mesures de risque sont appliquées au PnL selon la définition 1.2.2 : un **CVaR $_\alpha$ négatif** signifie que, dans les $1 - \alpha$ pires scénarios, l'assureur réalise en moyenne un **gain**, ce qui est favorable. À l'inverse, un CVaR $_\alpha$ positif indique une perte moyenne dans les queues extrêmes. Le **taux de succès** ou *success rate* est défini comme la proportion de scénarios où $\text{PnL} < 0$, c'est-à-dire où le portefeuille excède le passif.

PnL	V0	Mean	Std	CVaR95	CVaR99	VaR95	VaR99	SucRate
Delta Hedging	4.23	0.51	7.62	23.79	38.53	15.20	28.73	60.43%

TABLE 5.3 – Résultat du delta-hedging

Performance du FFNN

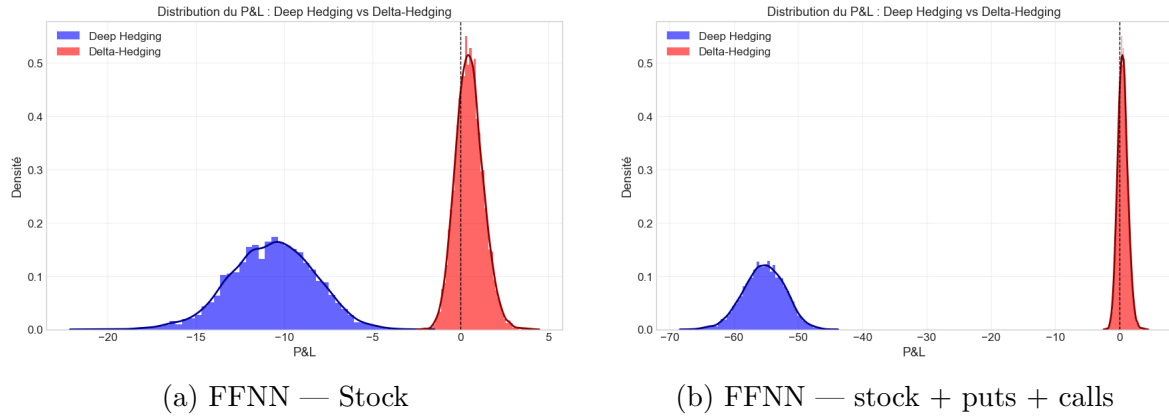


FIGURE 5.3 – Distribution du PnL FFNN avec perte entropique

PnL	V0	Mean	Std	CVaR95	CVaR99	VaR95	VaR99	SucRate
Stock	4.23	-10.68	23.14	22.46	33.62	15.87	26.04	68.03%
+3 puts	4.23	-43.69	26.58	-10.69	-6.36	-13.22	-9.31	99.95%
+3 calls	4.23	-55.32	31.13	-14.83	-11.02	-17.44	-13.58	99.98%

TABLE 5.4 – Performance du modèle FFNN avec perte entropique

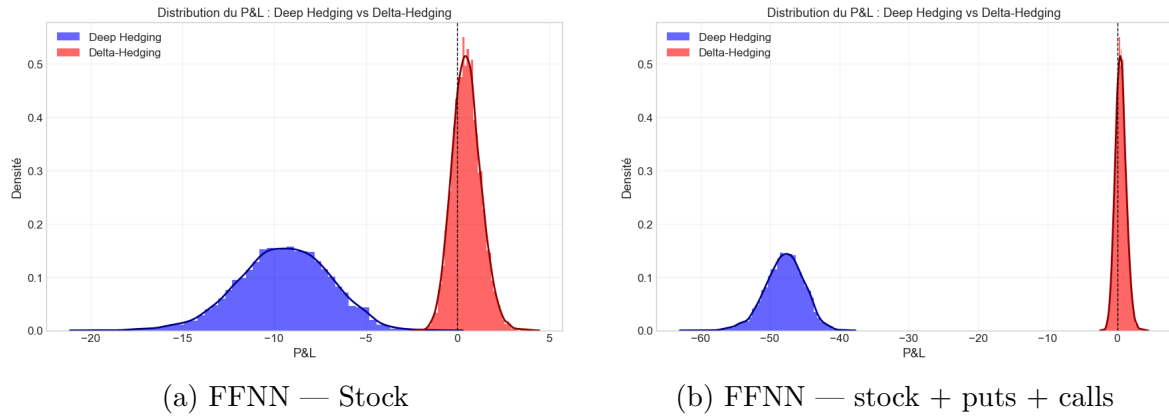


FIGURE 5.4 – Distribution du PnL FFNN avec perte CVaR

PnL	V0	Mean	Std	CVaR95	CVaR99	VaR95	VaR99	SucRate
Stock	4.23	-9.45	24.27	30.16	44.11	20.88	35.91	66.5%
+3 puts	4.23	-37.15	23.83	-12.04	-7.44	-14.60	-10.85	99.94%
+3 calls	4.23	-47.91	27.22	-18.48	-11.58	-21.88	-17.26	99.97%

TABLE 5.5 – Performance du modèle FFNN avec perte CVaR 99%

Les résultats présentés dans les tableaux 5.4 et 5.5 suggèrent que le modèle FFNN, sans instruments de couverture additionnels, a légèrement surpassé le delta-hedging

classique sur certaines métriques et dépendamment de la fonction de perte. En effet, le FFNN atteint un taux de succès de 66.5-68 % contre 60.43 % pour le delta-hedging, tout en réduisant significativement le CVaR99 % de +12.7% pour la perte entropique et une dégradation de -14.5% pour la fonction de perte CVaR. Cette amélioration s'explique par la capacité du réseau à apprendre une stratégie non linéaire qui anticipe partiellement la dynamique du ratchet et l'aspect cohérent de la perte entropique contrairement à la CVaR.

Toutefois, force est de constater que l'introduction de calls puis des puts améliore la stabilité du portefeuille, indiquant une meilleure gestion des pertes extrêmes. En effet, le taux de succès de couverture a bondit à 99.98%, le CVaR99 % s'améliore de près de +130.1% et devenant négatif, montrant que même dans les 1% de scénario extrême, l'assureur arrive à se faire des gains. Cependant, pendant que toutes les autres métriques affichent une amélioration progressive en fonction de la richesse des instruments de couverture, l'écart type du PnL montre une tendance croissance.

Ce comportement paradoxal montre que le FFNN, ne possédant pas de mémoire, structurellement procède par une optimisation étape par étape, et ses stratégies sont locales au détriment de la stratégie optimale globale assurant une meilleure stabilité de toutes les métriques. En somme, bien que le FFNN domine le delta-hedging en configuration minimale, il révèle une *incapacité structurelle à intégrer des instruments path-dépendants* lorsqu'il ne dispose pas d'un état suffisant. Cette limitation souligne l'importance critique de la représentation de l'état dans les architectures de deep hedging.

Performance du LSTM

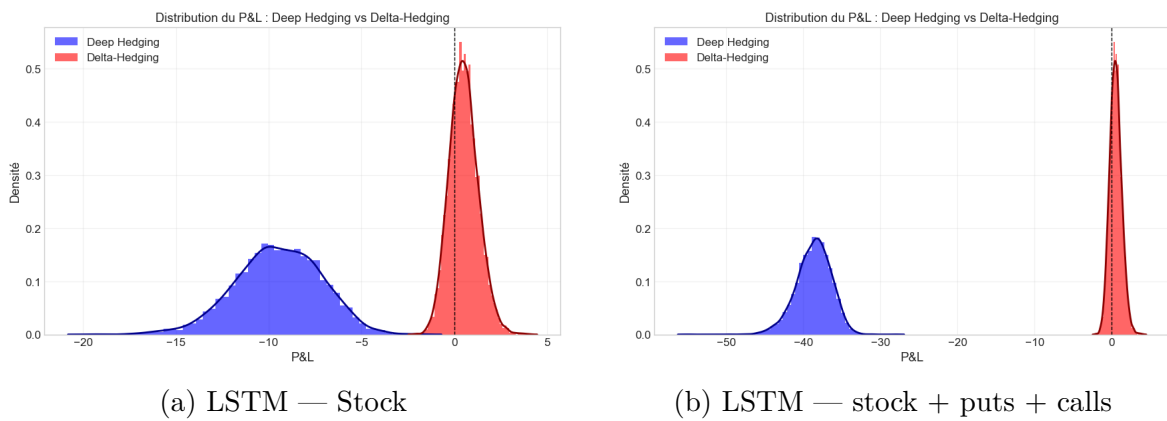


FIGURE 5.5 – Distribution du PnL LSTM avec CVaR 99%

PnL	V0	Mean	Std	CVaR95	CVaR99	VaR95	VaR99	SucRate
Stock	4.23	−9.49	22.69	25.19	36.66	17.57	29.83	68.02%
+3 puts	4.23	−26.15	19.35	−5.78	−0.50	−9.06	−3.77	99.68%
+3 calls	4.23	−38.65	22.26	−16.08	−9.72	−19.70	−14.24	99.96%

TABLE 5.6 – Performance du modèle LSTM avec perte CVaR 99%

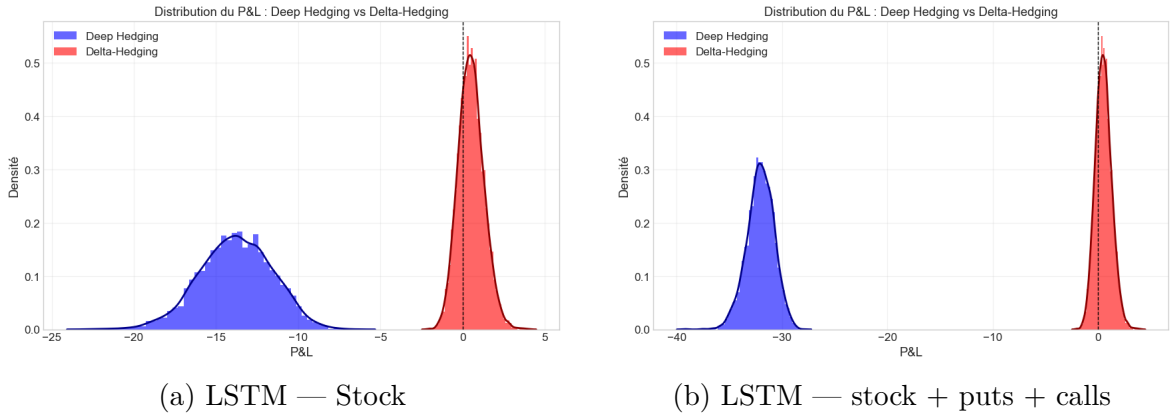


FIGURE 5.6 – Distribution du PnL LSTM avec perte entropique

PnL	V0	Mean	Std	CVaR95	CVaR99	VaR95	VaR99	SucRate
Stock	4.23	−13.79	21.66	15.60	24.93	9.86	19.04	75.25%
+3 puts	4.23	−25.83	12.41	−9.09	−5.56	−11.28	−7.96	99.96%
+3 calls	4.23	−32.05	12.67	−13.63	−9.663	−16.32	−12.08	99.99%

TABLE 5.7 – Performance du modèle LSTM avec perte entropique

Le modèle LSTM, contrairement au FFNN, démontre une capacité robuste à tirer parti des instruments de couverture additionnels. Avec le stock seulement, le LSTM réduit CvaR99 de 4.8-35.3% et avec l'ajout progressif des puts et calls, la réduction va jusqu'à 125.2%. La volatilité du PnL par rapport au FFNN, est plus stable voir en baisse, montrant un meilleur compromis stabilité des gains. Cette stabilité intrinsèque provient de sa mémoire interne, qui lui permet d'inférer approximativement l'information historique nécessaire pour fournir une stabilité globale dans le temps. En plus, l'ajout de trois puts européens amplifie cet avantage : le taux de succès bondit à 99.96 % sous perte entropique.

Ces résultats démontrent que le LSTM résout le problème fondamental rencontré par le FFNN : la mémoire récurrente compense l'information manquante en reconstruisant dynamiquement la path-dépendance.

Performance de la perte entropique

Les résultats expérimentaux obtenus avec un paramètre d'aversion au risque $\lambda = 0.5$ démontrent clairement que la fonction de perte entropique induit des stratégies de couverture plus robustes et mieux équilibrées que celles générées par l'optimisation de la fonction de perte $CVaR_{99\%}$. Cette supériorité s'explique par la nature fondamentale des deux critères : alors que le $CVaR$ se concentre exclusivement sur la moyenne conditionnelle des pertes extrêmes, la perte entropique intègre une *aversion au risque constante*, permettant une gestion globale de la distribution du PnL et sa structure via le développement de Taylor révèle qu'il contient à la fois les moments de tous les ordres, permettant de les optimiser globalement.

Sous perte entropique, le modèle LSTM atteint un taux de succès de 99.99 % avec l'écart-type du PnL en baisse avec l'ajout progressif de produits de couverture (tableau 5.7), tandis que sous $CVaR$ à 99 %, le taux de succès est nettement inférieur et l'écart type du PnL se stabilise néanmoins à défaut (tableau 5.6). La perte entropique, grâce à son paramètre λ , permet un *ajustement continu de l'aversion au risque* : avec $\lambda = 0.5$, le trade-off entre rendement moyen et dispersion est optimal. Cet aspect est particulièrement précieux dans un contexte d'assurance, où la stabilité du capital économique est aussi cruciale que la couverture du risque de ruine.

Robustesse et sensibilité aux Seeds

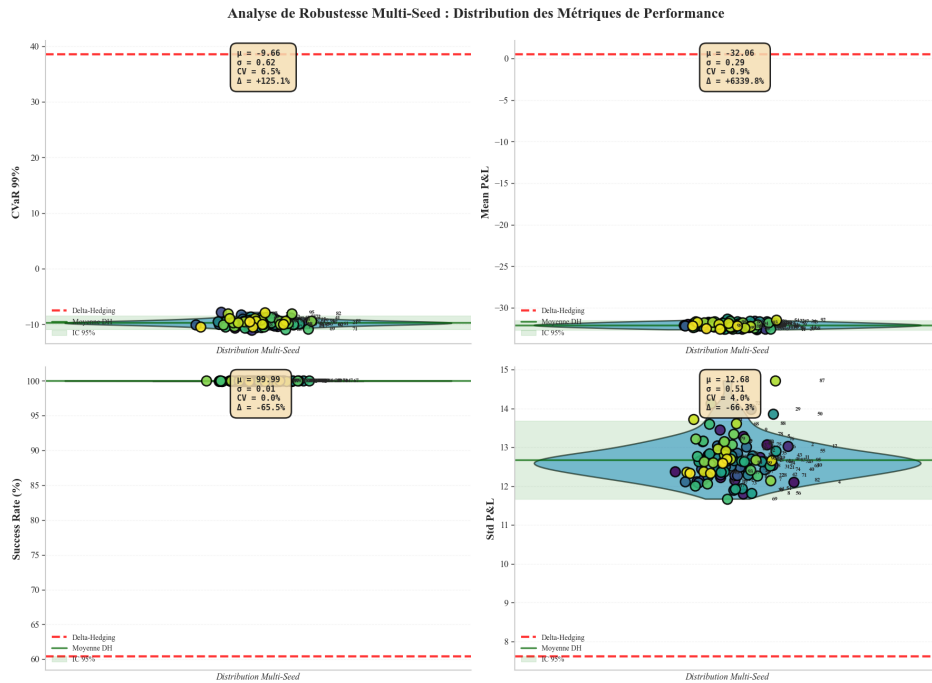


FIGURE 5.7 – Cas LSTM-Entropique avec Stock+puts+Calls

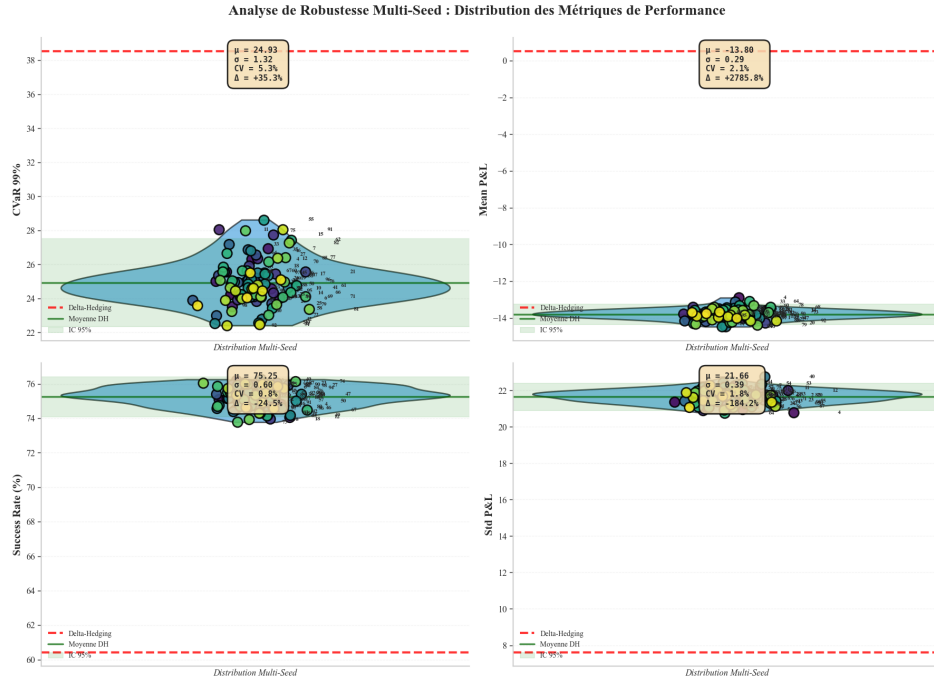


FIGURE 5.8 – Cas LSTM-Entropique avec Stock

Ecart-Type	Mean	CVaR95	CVaR99	VaR95	VaR99	SucRate
Delta Hedging	0.05	0.77	2.05	0.42	1.26	0.52%
FFNN& CVaR	0.43	0.64	2.61	0.28	0.47	0.02%
FFNN& Entr	0.44	0.37	1.44	0.24	0.28	0.02%
LSTM&CVaR	0.33	0.57	2.49	0.21	0.43	0.03%
LSTM&Entr	0.29	0.24	0.62	0.19	0.31	0.01%

TABLE 5.8 – Ecart-type des métriques : Cas stock+puts+calls

Les résultats présentés dans le tableau 5.8 et les figures 5.7 et 5.8 confirment la robustesse statistique de l'approche de deep hedging mise en oeuvre. En effet, l'analyse de la variabilité des performances à travers 100 initialisations aléatoires distinctes révèle des écarts-types remarquablement faibles, en particulier pour le modèle LSTM couplé à la perte entropique. Par exemple, le taux de succès présente un écart-type inférieur à 0.1 %, tandis que le $CVaR_{99\%}$, métrique critique pour un assureur, affiche une dispersion quasi-négligeable. Cette faible sensibilité aux conditions initiales du générateur de scénario démontre que la stratégie apprise n'est pas le fruit d'un heureux hasard d'initialisation, mais bien le reflet d'une convergence stable vers une politique quasi-optimale. En comparaison, les architectures plus simples telles que FFNN exhibent une variabilité nettement plus élevée, soulignant leur instabilité intrinsèque dans ce contexte de marché incomplet et path-dépendant. Pour ce qui est des figures, elles montrent clairement que la performance du delta hedging est exclue de l'intervalle de confiance celle du deep hedging, suggérant une différence significative et une meilleure gestion des queues

de distribution. Ces observations permettent d'affirmer avec un degré de confiance que le pipeline proposé, combinant un générateur de scénarios réaliste, une architecture récurrente et une fonction de perte cohérente, fournit une solution de couverture à la fois performante, fiable et reproductible. Cette robustesse est une condition sine qua non pour toute application opérationnelle en gestion actif-passif, où la prévisibilité et la stabilité des résultats sont aussi cruciales que leur performance moyenne.

Sensibilité aux paramètres des fonctions de pertes

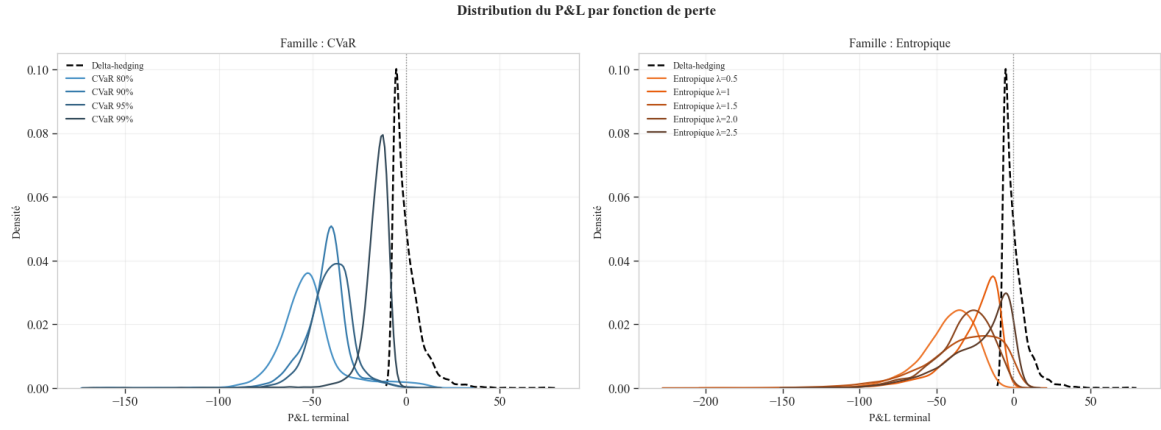


FIGURE 5.9 – Sensibilité au hyperparamètre des fonctions de perte

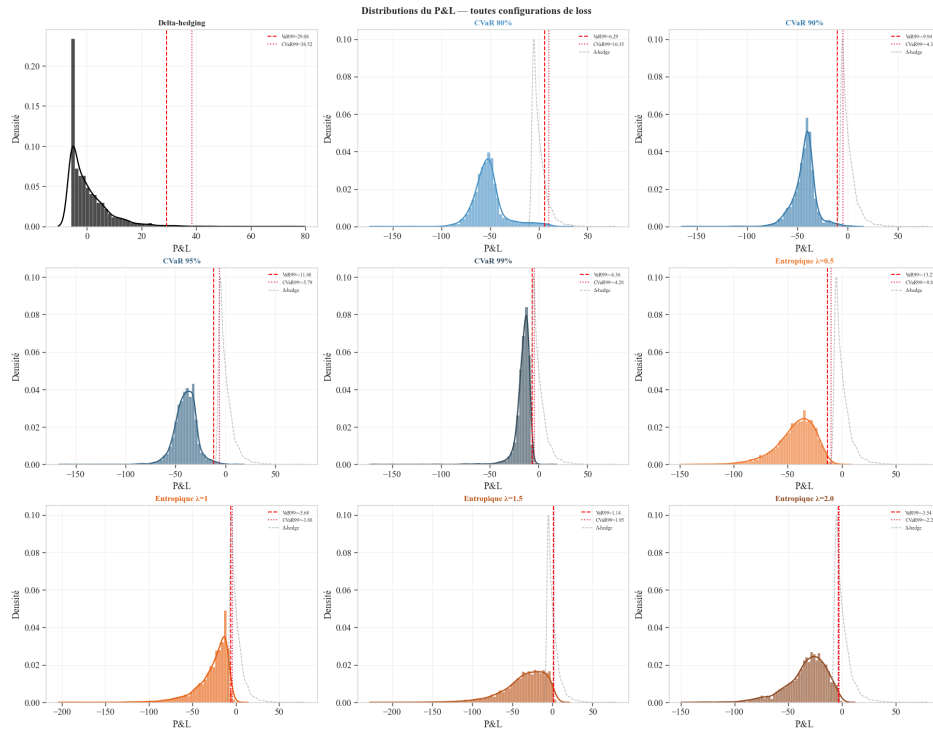


FIGURE 5.10 – Sensibilité au hyperparamètre des fonctions de perte

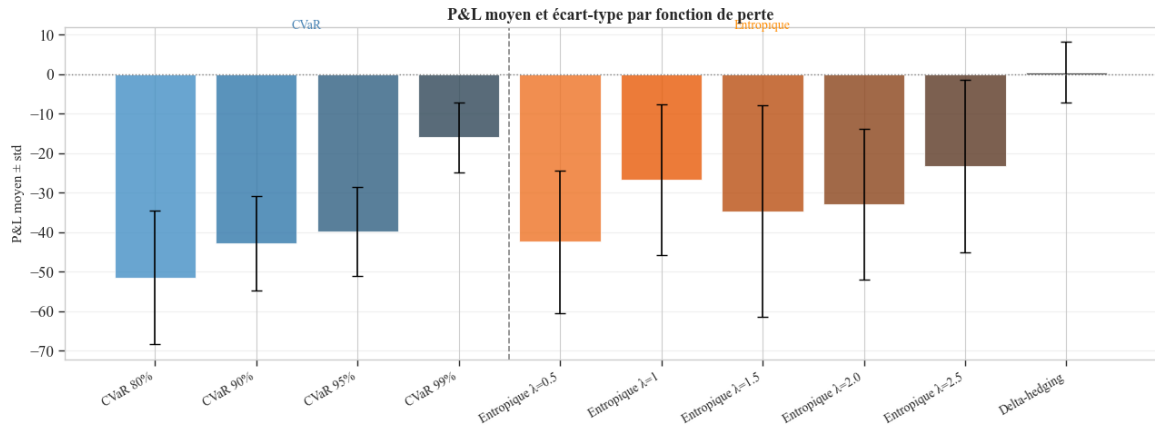


FIGURE 5.11 – PnL moyen par hyperparamètre des fonctions de perte

Afin d'étudier la sensibilité des fonctions de perte et procéder à une éventuelle comparaison, nous établissons deux grands groupes de familles de critères d'optimisation :

- la *perte CVaR*, paramétrée par le niveau de confiance $\alpha \in \{0.80, 0.90, 0.95, 0.99\}$,
- la *perte entropique*, paramétrée par le coefficient de risque $\lambda \in \{0.5, 1.0, 1.5, 2.0, 2.5\}$,

contre la référence classique *Delta-Hedging*. Les résultats sont obtenus sur un ensemble de test de $N = 5000$ scénarios hors échantillons.

La figure 5.9 montre clairement que contrairement à une intuition commune, la fonction de perte CVaR à 80% donne une distribution PnL plus décalée à gauche que toutes les autres valeurs de α . Cela s'explique par le fait que une telle fonction de perte se concentre sur les 20% de scénario donnant une vue plus ou moins globale par rapport à $\alpha = 99\%$ qui n'a que 1% des scénarios. Cet aspect explique la figure 5.11 où le gain moyen diminue plus α croît car l'optimisation se concentre de plus en plus sur une partie réduite de la distribution, entraînant une myopie et réduisant les gains. De même, pour la perte entropique, plus la valeur de λ augmente, on observe une queue de distribution de plus en plus lourde liée à la pénalisation plus forte des pertes. La 5.11 confirme que plus l'aversion au risque λ augmente, plus les profits générés baissent légèrement assurant ainsi un bon compromis entre aversion au risque et gains. Toutefois, force est de constater que les profits moyens dans le cas de la perte CVaR subit une baisse plus brutale que la perte entropique, suggérant une meilleure stabilité des gains entropiques liés à ses stratégies globalistes.

En Conclusion, le choix de la fonction de perte n'est pas un simple détail technique car elle incarne en elle l'appétit au risque de l'assureur. Le Choix de la perte CVaR est approprié dans le cadre du risque management et de la solvabilité 2 lié à sa structure qui se concentre à la réduction des risques extrêmes tandis que l'usage de la perte entropique peut être à but purement spéculative pour des résultats équilibrés entre rendement et aversion au risque. Elle est pertinente en ce sens qu'elle est équivalente à la fonction d'utilité exponentielle.

Entraînement et validation

L'algorithme de deep hedging repose sur une séparation stricte entre trois ensembles de scénarios simulés sous $\mathbb{P}$:

- un *jeu d'entraînement* utilisé pour mettre à jour les poids du réseau via la descente de gradient,
- un *jeu de validation* employé pour surveiller la généralisation,
- un *jeu de test* réservé exclusivement à l'évaluation finale et composé de 100 ensembles hors échantillons.

Les diagnostics d'entraînement confirment que l'entraînement est *stable*, *convergent* et *généralisable*, ce qui valide la fiabilité des stratégies apprises.

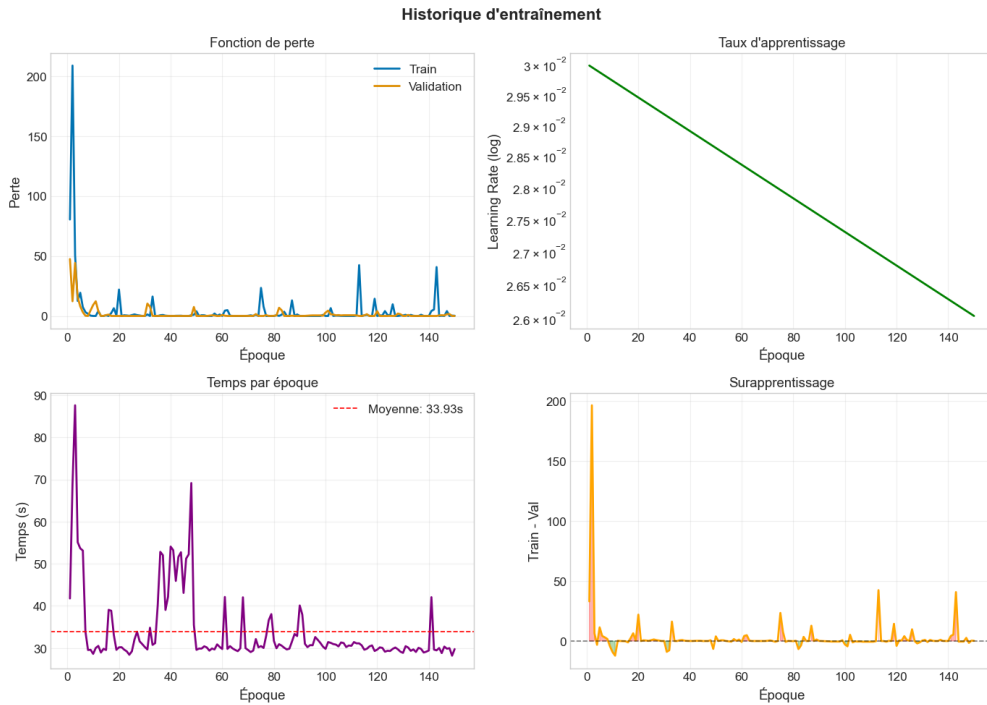


FIGURE 5.12 – Performance de généralisation : cas perte entropique

Le résultat ci dessus illustre la capacité de généralisation de notre modèle avec la perte entropique qui s'avère très satisfaisant en général, ici LSTM comme dans les autres configurations. Cependant, il faut noter que il n'en est pas le cas pour la perte CVaR qui montre un surapprentissage chronique comme le montre la figure ci dessous. Cela montre que même si le modèle réussit dans notre hedging en donnant de performance remarquable, il peut souffre d'un problème statistique.

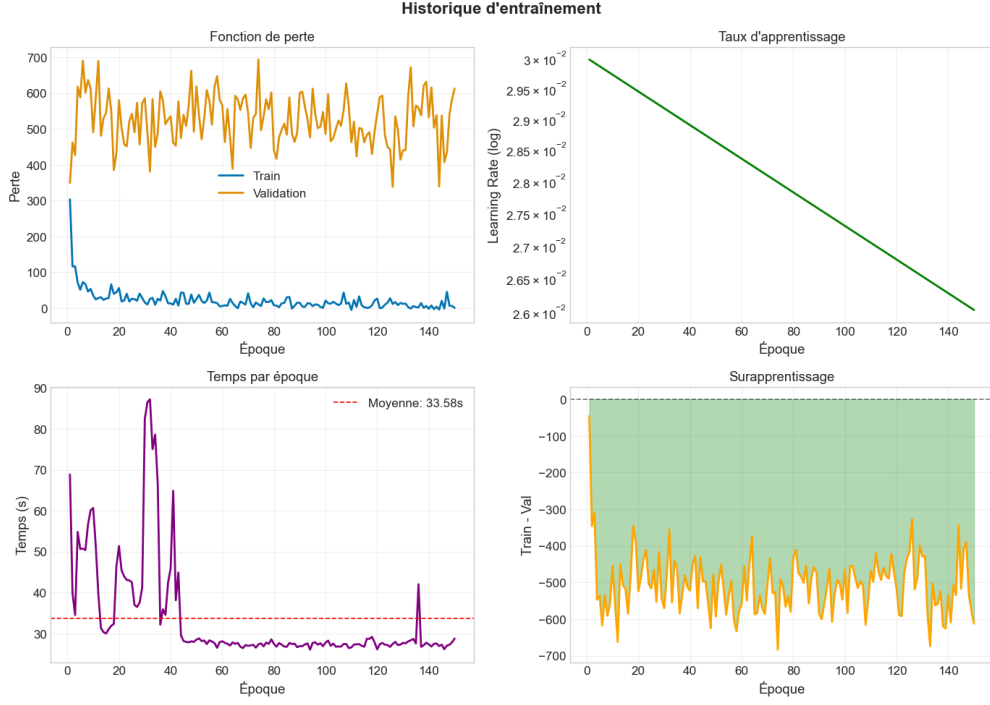


FIGURE 5.13 – Performance de généralisation cas : Perte CVaR

5.4 Analyse mathématique

Dans cette partie, l'on compare trois familles de stratégies :

- (1) le **delta-hedging** classique, fondé sur le calcul des grecques du portefeuille par différentiation du prix par rapport au sous-jacent sous la mesure risque-neutre $\mathbb{Q}$;
- (2) le **deep hedging FFNN**, stratégie markovienne paramétrée par un réseau feed-forward ;
- (3) le **deep hedging LSTM**, stratégie path-dépendante paramétrée par un réseau récurrent.

5.4.1 Formalisation du problème de couverture optimale

D'après définition 1.2.4 on note $\mathcal{H}$ l'ensemble des stratégies $\delta = (\delta_{t_n})_{n=0}^{N-1}$ telles que :

- (1) chaque $\delta_{t_n} \in \mathbb{R}^{D+1}$ est $\mathcal{F}_{t_n}$ -mesurable,
- (2) $\mathbb{E}^{\mathbb{P}} \left[\sum_{n=0}^{N-1} \|\delta_{t_n}\|^2 \right] < \infty$.

La valeur du portefeuille auto-financé avec frictions est :

$$V_{t_{n+1}}^{\delta} = e^{r\Delta t} V_{t_n}^{\delta} + \delta_{t_n}^{\top} (\mathbf{S}_{t_{n+1}} - e^{r\Delta t} \mathbf{S}_{t_n}) - C_{t_{n+1}}(\delta),$$

où $C_{t_{n+1}}(\delta) = \lambda_c \|\delta_{t_n} - \delta_{t_{n-1}}\|_1 \cdot \mathbf{S}_{t_n}$ désigne les coûts de transaction.

5.4.2 Sous-optimalité du delta-hedging en marché incomplet

Le delta-hedging classique consiste, à chaque date t_n , à calculer la sensibilité du prix du passif par rapport au sous-jacent sous une mesure martingale $\mathbb{Q}$ fixée, puis à détenir cette quantité dans le portefeuille.

D'après la définition 1.1.5 et 1.3.1, soit $P(t, S_t, V_t)$ le prix du passif résiduel $(Z_T - S_T)_+$ sous la mesure risque-neutre $\mathbb{Q}$ du modèle de Heston, calculé par FFT (voir section 4.2.1). La stratégie de delta-hedging est définie par la grecque du portefeuille :

$$\delta_{t_n}^\Delta := \frac{\partial P}{\partial S}(t_n, S_{t_n}, V_{t_n}) \in \mathbb{R},$$

Cette quantité est obtenue par différentiation numérique de la surface de prix FFT, et non par minimisation d'une fonctionnelle de risque.

Théorème 5.4.1 (Décomposition de l'erreur du delta-hedging). *Soit $H_T = (Z_T - S_T)_+$ le passif GMAB lookback, et soit $\delta^\Delta = (\delta_{t_n}^\Delta)_{n=0}^{N-1}$ la stratégie de delta-hedging de la définition 1.3.1. Sous le modèle de Heston avec $\rho < 0$ et en l'absence de coûts de transaction, l'erreur de couverture terminale admet la décomposition :*

$$\tilde{H}_T - \tilde{V}_T^{\delta^\Delta} = \sum_{n=0}^{N-1} \tilde{\varepsilon}_{t_n}^{(1)} + \sum_{n=0}^{N-1} \tilde{\varepsilon}_{t_n}^{(2)} + \sum_{n=0}^{N-1} \tilde{\varepsilon}_{t_n}^{(3)},$$

où :

- $\tilde{\varepsilon}_{t_n}^{(1)}$ représente l'erreur de modèle et de discrétisation,
- $\tilde{\varepsilon}_{t_n}^{(2)}$ représente l'erreur due à la dépendance au chemin du payoff,
- $\tilde{\varepsilon}_{t_n}^{(3)}$ représente le risque de volatilité non couvert.

Et elles sont toutes des processus F_T -mesurables

$$\varepsilon_{t_n}^{(1)} = R_n = O((\Delta t)^{3/2}), \quad (5.1)$$

$$\varepsilon_{t_n}^{(2)} = \frac{\partial P}{\partial S}(t_n, S_{t_n}, V_{t_n}) - \delta_{t_n}^\Delta, \quad (5.2)$$

$$\tilde{\varepsilon}_{t_n}^{(3)} = \int_{t_n}^{t_{n+1}} \sigma \sqrt{V_s} \frac{\partial \tilde{P}}{\partial V}(s, \tilde{S}_s, V_s) dW_s^{(2), \mathbb{P}} + \int_{t_n}^{t_{n+1}} \frac{\partial \tilde{P}}{\partial V} [\kappa(\theta - V_s) - \tilde{\kappa}(\tilde{\theta} - V_s)] ds. \quad (5.3)$$

Démonstration : On procède par décomposition télescopique sur $[0, T]$, voir annexe C.

Complexité des erreurs :

Erreur de modèle $\varepsilon^{(1)}$: elle provient de l'écart entre la dérivée discrète et le terme différentiel continu d'Itô, et de la discrétisation du schéma d'Euler pour V_t . Cet écart est d'ordre $O((\Delta t)^{3/2})$ par pas.

Erreur de path-dépendance $\varepsilon^{(2)}$: le delta FFT est calculé en traitant Z_t comme un strike fixe. Or la véritable dérivée du passif conditionnel $\mathbb{E}^{\mathbb{Q}}[(Z_T - S_T)_+ | \mathcal{F}_{t_n}]$ par rapport à S_{t_n} contient un terme supplémentaire :

$$\frac{\partial P}{\partial S}(t_n, S_{t_n}, V_{t_n}) - \delta_{t_n}^{\Delta} = \mathbb{Q}(Z_T = S_{t_n} | \mathcal{F}_{t_n}) \cdot \frac{\partial Z_{t_n}}{\partial S_{t_n}},$$

le second terme étant non nul uniquement quand $S_{t_n} = Z_{t_n}$, lorsqu'il ya nouveau maximum. L'approximation FFT ne prenant pas en compte ce terme produit un biais chaque fois qu'un nouveau maximum est atteint, soit en moyenne $\mathbb{E}[\#\{n : S_{t_n} = Z_{t_n}\}]$ fois sur l'horizon. Ce nombre est de l'ordre de $\log N$.

Erreur de mesure $\varepsilon^{(3)}$: le terme $\partial P / \partial V$ correspond au vega du passif. Ce terme est non couvert car le delta-hedging ne détient que le sous-jacent. Sous $\mathbb{P}$, la volatilité V_t évolue selon la dynamique de Heston avec $\sigma_V > 0$, de sorte que l'intégrale stochastique (5.3) est non nulle et représente le risque de volatilité résiduel.

Corollaire 5.4.1.1. *Pour notre passif GMAB lookback avec ratchet mensuel sur $N = 180$ périodes, le terme d'erreur de path-dépendance satisfait :*

$$\mathbb{E}^{\mathbb{P}} \left[\sum_{n=0}^{N-1} \varepsilon_{t_n}^{(2)} \right] \asymp \log N \cdot \mathbb{E}^{\mathbb{P}} [\mathbb{Q}(Z_T = S_{t_\tau} | \mathcal{F}_{t_\tau})] > 0,$$

où $\tau = \inf\{n : S_{t_n} = Z_{t_n}\}$ est le premier instant de nouveau maximum. Ce biais positif explique le success rate de seulement 61% du delta-hedging, contre 73% pour le FFNN.

5.4.3 Deep hedging FFNN et le delta-hedging

Théorème 5.4.2 (Représentation de Fenchel–Moreau). *Soit $\rho : L^1(\Omega, \mathcal{F}, \mathbb{P}) \rightarrow \mathbb{R} \cup \{+\infty\}$ une fonctionnelle convexe, propre et semi-continue inférieurement pour la topologie faible $\sigma(L^1, L^\infty)$. Alors ρ admet la représentation duale :*

$$\rho(X) = \sup_{Q \in \mathcal{Q}} \left(\mathbb{E}^Q[X] - \alpha(Q) \right), \quad \forall X \in L^1,$$

où $\mathcal{Q}$ est un ensemble de mesures de probabilité absolument continues par rapport à $\mathbb{P}$, et $\alpha : \mathcal{Q} \rightarrow \mathbb{R} \cup \{+\infty\}$ est la pénalité minimale définie par $\alpha(Q) = \sup_{X \in L^1} (\mathbb{E}^Q[X] - \rho(X))$.

En particulier, ρ est entièrement déterminée par son sous-différentiel $\partial \rho$, et le supergradient $\rho'(X) \in \partial \rho(X)$ vérifie :

$$\rho(Y) \geq \rho(X) + \mathbb{E}^{\mathbb{P}}[\rho'(X)(Y - X)], \quad \forall X, Y \in L^1.$$

Le CVaR $_{\alpha}$ et le risque entropique $\rho_{\lambda}^{\text{ent}}$ satisfont tous deux les hypothèses du théorème 5.4.2 (voir annexe A).

Théorème 5.4.3 (Weierstrass généralisé, Rockafellar & Wets). *Soit X un espace topologique et $f : X \rightarrow \mathbb{R} \cup \{+\infty\}$ une fonctionnelle propre, convexe et semi-continue inférieurement. Si f est coercive, c'est-à-dire si*

$$\|x_n\| \rightarrow +\infty \implies f(x_n) \rightarrow +\infty,$$

alors f admet au moins un minimiseur : $\exists x^ \in X$ tel que $f(x^*) = \inf_{x \in X} f(x) > -\infty$.*

Théorème 5.4.4. *Soit $\mathcal{H}^{Markov} \subset \mathcal{H}$ la sous-classe des stratégies markoviennes :*

$$\mathcal{H}^{Markov} := \left\{ \delta \in \mathcal{H} : \delta_{t_n} = f(t_n, S_{t_n}, V_{t_n}), f \text{ mesurable} \right\}.$$

Soit ρ une mesure de risque convexe satisfaisant le théorème 5.4.2 Alors :

(1) Il existe une stratégie optimale markovienne $\delta^{,Markov} \in \mathcal{H}^{Markov}$ telle que :*

$$\delta^{*,Markov} = \arg \min_{\delta \in \mathcal{H}^{Markov}} \rho(H_T - V_T^\delta).$$

(2) Pour tout $\varepsilon > 0$, il existe un ensemble de poids FFNN f_θ tel que :

$$\rho(H_T - V_T^{f_\theta}) - \rho(H_T - V_T^{\delta^{*,Markov}}) < \varepsilon.$$

(3) La stratégie de delta-hedging δ^Δ n'appartient pas à $\arg \min_{\mathcal{H}^{Markov}} \rho$ en marché incomplet avec $\sigma_V > 0$.

Démonstration :

Existence : $\min_{\delta \in \mathcal{H}^{Markov}} \rho(H_T - V_T^\delta)$ est un problème de contrôle stochastique en horizon fini sur un espace d'état $(S, V) \in \mathbb{R}_+^2$. La valeur de V_T^δ est linéaire en $(\delta_{t_n})_{n=0}^{N-1}$, et ρ est convexe et semi-continue inférieurement par hypothèse. D'après le théorème 5.4.3 et la structure markovienne du modèle de Heston, il existe un minimiseur dans la fermeture de $\mathcal{H}^{Markov}$.

Approximation par FFNN : La stratégie optimale $\delta_{t_n}^{*,Markov} = g^*(t_n, S_{t_n}, V_{t_n})$ est une fonction continue sur le compact des états typiques de marché. D'après le théorème 2.1.2, on peut donc approcher g^* uniformément par un FFNN. Comme V_T^δ est linéaire en δ , et ρ est Lipschitz sur L^1 , la proximité uniforme des stratégies se transfère en proximité de la valeur du risque :

$$\left| \rho(H_T - V_T^{f_\theta}) - \rho(H_T - V_T^{\delta^{*,Markov}}) \right| \leq C_\rho \sum_{n=0}^{N-1} \mathbb{E} \mathbb{P} \left[\|f_\theta(X_{t_n}) - g^*(X_{t_n})\| \cdot \|A_{t_n}\| \right] \leq C_\rho N \varepsilon \mathbb{E}[\|A\|_\infty],$$

ce qui peut être rendu arbitrairement petit pour $\varepsilon \rightarrow 0$.

Non-optimalité du delta-hedging : Le delta-hedging est défini comme $\delta_{t_n}^\Delta = \partial P / \partial S(t_n, S_{t_n}, V_{t_n})$ sous $\mathbb{Q}$. Or l'optimum de ρ sous $\mathbb{P}$ satisfait les conditions du premier

ordre :

$$\mathbb{E}^{\mathbb{P}} \left[\rho' \left(H_T - V_T^{\delta^*} \right) \cdot \left(S_{t_{n+1}} - e^{r\Delta t} S_{t_n} \right) \right] = 0, \quad \forall n,$$

où ρ' est le supergradient de ρ . Cette condition fait intervenir la mesure $\mathbb{P}$ et le supergradient ρ' , tandis que δ^Δ est déterminé par la dérivée partielle du prix sous $\mathbb{Q}$, qui n'est pas le supergradient de ρ sous $\mathbb{P}$ en général dès lors que $\mathbb{P} \neq \mathbb{Q}$ et que $\sigma_V > 0$ (le risque de volatilité n'étant pas couvert).

Proposition 5.4.1. *On suppose que :*

- (1) *Les erreurs de couverture vérifient $H_T - V_T^\delta \in L^2(\Omega, \mathcal{F}, \mathbb{P})$ pour toute stratégie $\delta \in \mathcal{H}$, car les trajectoires (S_{t_n}, V_{t_n}) sont sur un compact K_ε .*
- (2) *La mesure de risque ρ est uniformément convexe sur L^2 avec constante $\kappa_\rho > 0$:*

— *Pour $\rho = \rho_\lambda^{\text{ent}}$: $\kappa_\rho = \lambda$, par le développement de Taylor*

$$\rho_\lambda^{\text{ent}}(X) = \mathbb{E}[X] + \frac{\lambda}{2} \text{Var}(X) + O(\lambda^2).$$

— *Pour $\rho = \text{CVaR}_\alpha$: $\kappa_\rho > 0$ sur L^2 , sous l'hypothèse que la loi de $H_T - V_T^\delta$ est absolument continue, i.e., sans atome au niveau du VaR_α .*

Alors le gain de l'optimisation satisfait :

$$\Delta\rho := \rho(H_T - V_T^{\delta^\Delta}) - \rho(H_T - V_T^{\delta^{*,\text{Markov}}}) \geq \frac{\kappa_\rho}{2} \mathbb{E}^{\mathbb{P}} \left[\|\delta^\Delta - \delta^{*,\text{Markov}}\|_{\ell^2([0,T])}^2 \right] > 0,$$

avec inégalité stricte dès que $\delta^\Delta \neq \delta^{*,\text{Markov}}$ $\mathbb{P}$ -p.s., ce qui est le cas dès que $\sigma_V > 0$ et $\rho \neq \text{MSE}$.

Démonstration : voir Annexe D

5.5 Incapacité structurelle du FFNN face à la path-dépendance

Définition 5.5.1 (Processus d'information suffisante). Pour le passif GMAB lookback $H_T = (Z_T - S_T)_+$, l'état *suffisant* à la date t_n est le triplet augmenté :

$$\Xi_{t_n} := (S_{t_n}, V_{t_n}, Z_{t_n}), \quad Z_{t_n} = \max_{0 \leq k \leq n} S_{t_k}.$$

La filtration naturelle engendrée par Ξ est $\mathcal{G}_{t_n} = \sigma(\Xi_{t_0}, \dots, \Xi_{t_n}) = \sigma(S_{t_0}, \dots, S_{t_n}, V_{t_n}) = \mathcal{F}_{t_n}$. Cependant, la stratégie FFNN est conditionnée à $\hat{\mathcal{G}}_{t_n} = \sigma(S_{t_n}, V_{t_n}) \subsetneq \mathcal{F}_{t_n}$.

Théorème 5.5.1. *Soit $\delta^{*,\Xi}$ la stratégie optimale sous information complète Ξ_{t_n} , et $\delta^{*,\text{FFNN}}$ la stratégie optimale markovienne sous information réduite (S_{t_n}, V_{t_n}) . Alors :*

- (1) **Gap d'information :**

$$\rho(H_T - V_T^{\delta^{*,\text{FFNN}}}) - \rho(H_T - V_T^{\delta^{*,\Xi}}) \geq \kappa_\rho \text{Var}^{\mathbb{P}} \left(\mathbb{E}[H_T | \mathcal{F}_{t_n}] - \mathbb{E}[H_T | \hat{\mathcal{G}}_{t_n}] \right).$$

- (2) **Dégradation avec les options** : En présence de D options de couverture $\mathbf{Op}_{t_n} \in \mathbb{R}^D$, le FFNN doit allouer ses paramètres entre deux objectifs incompatibles :

$$\delta_{t_n}^{FFNN} = (\hat{\delta}_{t_n}^S, \hat{\delta}_{t_n}^{\mathbf{Op}}) = f_{\theta}(S_{t_n}, V_{t_n}, \mathbf{Op}_{t_n}) \in \mathbb{R}^{1+D}.$$

La position optimale sur les options est :

$$\hat{\delta}_{t_n}^{\mathbf{Op},*} = g^*(t_n, S_{t_n}, V_{t_n}, Z_{t_n}, \mathbf{Op}_{t_n}),$$

qui dépend de Z_{t_n} . Puisque $Z_{t_n} \notin \hat{\mathcal{G}}_{t_n}$, l'estimateur FFNN $\hat{\delta}_{t_n}^{\mathbf{Op}}$ est biaisé :

$$\mathbb{E}^{\mathbb{P}}[\hat{\delta}_{t_n}^{\mathbf{Op}} - \hat{\delta}_{t_n}^{\mathbf{Op},*}] \neq 0,$$

et ce biais est amplifié par la présence des options car leur prix $\mathbf{Op}_{t_n}$ dépend également de Z_{t_n} via le strike effectif de la garantie.

Démonstration : Voir Annexe E

5.6 Reconstruction de la path-dépendance par LSTM

Définition 5.6.1. Soit $\mathcal{X} = \mathbb{R}^d$ l'espace des observations et $\mathcal{T} = \{t_0, t_1, \dots, t_N\}$ une grille temporelle discrète.

- (1) Une *fonctionnelle causale* est une application

$$\Phi : \bigcup_{n=0}^N \mathcal{X}^{n+1} \longrightarrow \mathbb{R}$$

telle que $\Phi(x_0, \dots, x_n)$ ne dépend que des observations passées et présentes $(x_0, \dots, x_n)$, et non des observations futures $(x_{n+1}, \dots, x_N)$.

- (2) Un *système dynamique récurrent* (SDR) de dimension cachée $d_h \in \mathbb{N}$ est un système de la forme :

$$\begin{cases} h_n = \varphi_{\theta}(x_n, h_{n-1}), & n = 1, \dots, N, \\ h_0 = \mathbf{0} \in \mathbb{R}^{d_h}, \\ \hat{y}_n = \psi_{\theta}(h_n), \end{cases}$$

où $\varphi_{\theta} : \mathcal{X} \times \mathbb{R}^{d_h} \rightarrow \mathbb{R}^{d_h}$ est la *fonction de transition d'état* et $\psi_{\theta} : \mathbb{R}^{d_h} \rightarrow \mathbb{R}$ est la *fonction de lecture*, toutes deux paramétrées par $\theta \in \Theta$. La sortie $\hat{y}_n = \psi_{\theta}(h_n)$ est une fonctionnelle causale de $(x_0, \dots, x_n)$.

- (3) Un *LSTM* est un SDR particulier dans lequel φ_{θ} est décomposée en trois portes voir Chapitre 2, avec un état de cellule c_n supplémentaire, satisfaisant les équations.

Définition 5.6.2 (Topologie de la convergence uniforme sur les compacts). On munit l'espace des fonctionnelles causales de la topologie de la convergence uniforme sur les

compacts de trajectoires. Plus précisément, pour tout compact $\mathcal{K} \subset \mathcal{X}^{N+1}$, on définit la semi-norme :

$$\|\Phi\|_{\mathcal{K}} := \sup_{(x_0, \dots, x_N) \in \mathcal{K}} |\Phi(x_0, \dots, x_N)|.$$

Une suite (Φ_k) converge uniformément sur les compacts vers Φ si $\|\Phi_k - \Phi\|_{\mathcal{K}} \rightarrow 0$ pour tout compact $\mathcal{K}$.

Théorème 5.6.1. *Soit $h : \mathbb{R} \rightarrow \mathbb{R}$ une fonction d'activation non polynomiale et mesurable bornée. Soit $\Phi : \mathcal{X}^{N+1} \rightarrow \mathbb{R}$ une fonctionnelle causale continue pour la topologie produit sur $\mathcal{X}^{N+1}$.*

Alors, pour tout compact $\mathcal{K} \subset \mathcal{X}^{N+1}$ et tout $\varepsilon > 0$, il existe un LSTM LSTM_{θ} de dimension cachée $d_h \in \mathbb{N}$ suffisamment grande et de paramètres $\theta \in \Theta$ tels que :

$$\sup_{(x_0, \dots, x_N) \in \mathcal{K}} \left| \Phi(x_0, \dots, x_N) - \psi_{\theta}(h_N^{(d_h)}(x_0, \dots, x_N)) \right| < \varepsilon,$$

où $h_N^{(d_h)}$ désigne l'état caché final du LSTM de dimension d_h construit récursivement par la définition 5.6.1 (2).

En d'autres termes, la famille des LSTM est dense dans l'espace des fonctionnelles causales continues, pour la topologie de la convergence uniforme sur les compacts (définition 5.6.2).

Démonstration : voir Annexe F

Théorème 5.6.2. *Soit $(h_{t_n})_{n \geq 0}$ l'état caché d'un LSTM de dimension d , défini par la récursion :*

$$h_{t_n} = \text{LSTM}_{\theta}(X_{t_n}, h_{t_{n-1}}), \quad h_{t_0} = 0,$$

où $X_{t_n} = (S_{t_n}, V_{t_n})$. Alors, pour tout $p \geq 1$, toute fonction mesurable bornée $g : \mathbb{R} \rightarrow \mathbb{R}$, et tout $\varepsilon > 0$, il existe un LSTM LSTM_{θ} avec état caché de dimension d suffisamment grande tel que :

$$\mathbb{E}^{\mathbb{P}} \left[\left| \mathbb{E}[g(Z_{t_n}) \mid \mathcal{F}_{t_n}] - \psi_{\theta}(h_{t_n}) \right|^p \right]^{1/p} < \varepsilon,$$

où $\psi_{\theta} : \mathbb{R}^d \rightarrow \mathbb{R}$ est une couche linéaire de sortie. Autrement dit, l'état caché du LSTM encode asymptotiquement toute l'information sur Z_{t_n} contenue dans la trajectoire passée $(S_{t_0}, \dots, S_{t_n})$.

Démonstration :

Réprésentation de Z_{t_n} comme fonctionnelle de la trajectoire : Le maximum glissant s'écrit :

$$Z_{t_n} = \max_{0 \leq k \leq n} S_{t_k} = \max(Z_{t_{n-1}}, S_{t_n}).$$

C'est une fonctionnelle *récursive* et *déterministe* de la trajectoire, continue pour la topologie de la convergence uniforme sur les compacts.

Densité des LSTM : Par le théorème toute fonctionnelle causale continue $\Phi : (X_{t_0}, \dots, X_{t_n}) \mapsto \mathbb{R}$ peut être approchée uniformément, sur tout compact de trajectoires, par un LSTM de dimension d suffisamment grande. La fonctionnelle $(X_{t_0}, \dots, X_{t_n}) \mapsto \mathbb{E}[g(Z_{t_n}) \mid \mathcal{F}_{t_n}]$ est une telle fonctionnelle causale, continue par les propriétés de régularité de la mesure de probabilité conditionnelle sous Heston.

Corollaire 5.6.2.1. *Sous les hypothèses du théorème 5.6.2 et de la proposition ??, la stratégie LSTM optimale $\delta^{*,LSTM}$ vérifie :*

$$\rho(H_T - V_T^{\delta^{*,LSTM}}) \xrightarrow{d \rightarrow \infty} \rho(H_T - V_T^{\delta^{*,\Xi}}),$$

où $\delta^{*,\Xi}$ est l'optimum en information complète (définition 5.5.1). En particulier :

$$\rho(H_T - V_T^{\delta^{*,LSTM}}) \leq \rho(H_T - V_T^{\delta^{*,FFNN}}),$$

avec inégalité stricte dès que $\text{Var}^{\mathbb{P}}(Z_{t_n} \mid S_{t_n}, V_{t_n}) > 0$.

Démonstration : voir annexe G

Proposition 5.6.1. *En présence de D options de couverture, le LSTM apprend la règle de décision :*

$$\hat{\delta}_{t_n}^{\text{Op},LSTM} \approx g^*(t_n, S_{t_n}, V_{t_n}, Z_{t_n}^{LSTM}, \text{Op}_{t_n}),$$

où $Z_{t_n}^{LSTM} = \psi_{\theta}(h_{t_n})$ est l'estimé de Z_{t_n} par l'état caché. La réduction du CVaR99 de 32.87 (stock seul) à -2.80 (stock + puts) s'explique par l'utilisation de deux mécanismes complémentaires :

- (1) **Couverture du risque de queue :** les puts de strikes $K > S_{t_n}$ ont un payoff $(K - S_T)_+$ positivement corrélé à $(Z_T - S_T)_+$ quand $Z_T \approx K$, ce qui réduit directement les pertes extrêmes.

$$\text{Cov}^{\mathbb{P}}((Z_T - S_T)_+, (K - S_T)_+) > 0 \quad \text{pour } K \asymp \mathbb{E}[Z_T],$$

car les deux payoffs sont des fonctions croissantes du maximum Z_T .

- (2) **Conditionnement sur l'état du ratchet :** le LSTM détecte, via h_{t_n} , si $S_{t_n} \approx Z_{t_n}$ qui est le moment critique de nouveau maximum potentiel. Dans ce cas, il augmente la position en puts pour anticiper une future chute ; sinon, il réduit la protection. Cette stratégie conditionnelle n'est pas accessible au FFNN.

5.7 Les perte entropique et CVaR

Théorème 5.7.1. *Considérons deux mesures de risque : la perte entropique $\rho_{\lambda}^{\text{ent}}$ et le CVaR_{α} . Alors :*

- (1) **Cohérence temporelle :** $\rho_{\lambda}^{\text{ent}}$ est cohérent temporel, i.e., pour tous $0 \leq m \leq n \leq$

N et toute variable aléatoire $X \in L^1(\Omega, \mathcal{F}, \mathbb{P})$:

$$\rho_\lambda^{\text{ent}}(X) = \rho_\lambda^{\text{ent}}\left(-\frac{1}{\lambda} \log \mathbb{E}[e^{\lambda X} \mid \mathcal{F}_{t_n}]\right).$$

Le CVaR_α n'est pas time-consistent pour $\alpha > 0$ (Artzner et al., 1999; Riedel, 2004).

- (2) **Sensibilité à toute la distribution** : Pour deux PnL X, Y avec même CVaR_α mais distributions différentes sur $(-\infty, \text{VaR}_\alpha)$, on a en général $\rho_\lambda^{\text{ent}}(X) \neq \rho_\lambda^{\text{ent}}(Y)$.
Le développement de Taylor :

$$\rho_\lambda^{\text{ent}}(X) = \mathbb{E}[X] + \frac{\lambda}{2} \text{Var}(X) + \frac{\lambda^2}{6} \mu_3(X) + O(\lambda^3),$$

où $\mu_3(X) = \mathbb{E}[(X - \mathbb{E}[X])^3]$ est le moment centré d'ordre 3, montre que $\rho_\lambda^{\text{ent}}$ pénalise simultanément la moyenne, la variance et l'asymétrie du PnL.

- (3) **Optimalité uniforme** : La stratégie $\delta^{*, \text{ent}}$ solution de $\min_\delta \rho_\lambda^{\text{ent}}(H_T - V_T^\delta)$ est l'unique stratégie qui maximise l'utilité espérée sous utilité exponentielle $U(x) = -e^{-\lambda x}$:

$$\delta^{*, \text{ent}} = \arg \max_{\delta \in \mathcal{H}} \mathbb{E}^\mathbb{P}[U(V_T^\delta - H_T)],$$

ce qui constitue un critère cohérent avec la théorie de l'utilité espérée sur l'horizon complet $[0, T]$.

Démonstration : voir Annexe [H](#)

Proposition 5.7.1. Pour $N = 180$ périodes, l'optimisation sous CVaR_α avec $\alpha = 0.99$ génère une stratégie qui :

- (1) ignore le comportement du PnL sur les 99 premiers centiles, créant des pertes modérées mais fréquentes;
- (2) n'est pas cohérent temporel, entraînant des incohérences dynamiques : la stratégie optimale à $t = 0$ n'est plus optimale à $t = t_n$ pour $n > 0$.

Ces deux phénomènes expliquent les résultats observés : sous CVaR_{99} , le LSTM affiche un $\text{CVaR}_{99} = 0.89 > -5.88$ pour la perte entropique, malgré un taux de succès comparable.

Démonstration :

La représentation du CVaR donne :

$$\text{CVaR}_\alpha(X) = \inf_{w \in \mathbb{R}} \left\{ w + \frac{1}{1 - \alpha} \mathbb{E}[(X - w)_+] \right\}.$$

La condition du premier ordre du problème d'optimisation de couverture sous CVaR est :

$$\mathbb{E}^\mathbb{P}[\mathbf{1}_{\{H_T - V_T^\delta > w^*\}} \cdot (S_{t_{n+1}} - e^{r\Delta t} S_{t_n})] = 0,$$

qui ne fait intervenir que les scénarios où le PnL dépasse le seuil w^* . Les 99% restants de la distribution n'influencent pas les positions optimales. La stratégie résultante peut donc laisser des pertes modérées non couvertes sur la majorité des scénarios.

La non cohérence temporelle implique que :

$$\text{CVaR}_\alpha\left(\text{CVaR}_\alpha(H_T - V_T^\delta \mid \mathcal{F}_{t_n})\right) \neq \text{CVaR}_\alpha(H_T - V_T^\delta),$$

Une stratégie entraînée pour minimiser $\text{CVaR}_{0.99}$ sur l'horizon $[0, T]$ n'est plus optimale pour $\text{CVaR}_{0.99}$ sur $[t_n, T]$ pour un n intermédiaire. Ceci crée des positions sous-optimales à chaque rebalancement, qui s'accumulent sur les 180 pas de temps, dégradant la performance globale par rapport à la perte entropique, qui est cohérent.

Théorème 5.7.2. *Sous les hypothèses du modèle de Heston avec $\sigma_V > 0$, $\rho < 0$, et pour le passif GMAB lookback sur $N = 180$ périodes, les stratégies vérifient la hiérarchie suivante pour toute mesure de risque convexe ρ satisfaisant les conditions des théorèmes 5.4.4 , 5.7.1, la proposition 5.7.1 et le corollaire 5.6.2.1 :*

$$\rho\left(H_T - V_T^{\delta^*, \text{ent}, LSTM}\right) \leq \rho\left(H_T - V_T^{\delta^*, CVaR, LSTM}\right) \leq \rho\left(H_T - V_T^{\delta^*, FFNN}\right) \leq \rho\left(H_T - V_T^{\delta^\Delta}\right),$$

Conclusion

Ce mémoire s’est inscrit dans un contexte de complexification croissante des produits financiers et assurantiels, notamment ceux intégrant des garanties, pour lesquels les méthodes classiques de couverture montrent certaines limites opérationnelles. Face aux difficultés liées à la spécification et à la calibration des modèles paramétriques traditionnels, ce travail a exploré une approche alternative fondée sur l’apprentissage profond, à travers le cadre du **deep hedging**.

L’intérêt principal de cette approche réside dans sa capacité à exploiter directement les données historiques de marché afin d’apprendre des stratégies de couverture adaptées, sans recourir à une modélisation explicite des dynamiques sous-jacentes. En s’appuyant sur une démarche **data-driven**, le modèle est en mesure de capturer des comportements complexes, non linéaires et dépendants du chemin, tout en optimisant directement des critères de performance pertinents du point de vue du gestionnaire de risque. Cette flexibilité constitue un atout majeur dans le contexte des produits d’assurance à long terme, où les hypothèses classiques peuvent s’avérer insuffisantes.

L’application menée sur des produits de type GMAB a permis de mettre en évidence la pertinence de cette approche, en particulier à travers l’utilisation d’architectures récurrentes de type LSTM, capables de modéliser efficacement les dépendances temporelles inhérentes à ces contrats. Les résultats obtenus confirment que l’apprentissage à partir de données historiques peut constituer une alternative crédible aux approches traditionnelles de couverture, notamment pour des produits fortement dépendants du chemin.

Enfin, une extension naturelle de ce travail consisterait à intégrer explicitement le risque de mortalité dans le cadre du deep hedging. En effet, les produits d’assurance vie reposent sur une interaction étroite entre risques financiers et risques biométriques, dont la prise en compte conjointe reste encore limitée dans les approches actuelles. L’intégration de modèles de mortalité au sein de ce cadre permettrait de mieux refléter la réalité des portefeuilles d’assurance et d’ouvrir la voie à des stratégies de couverture globales, tenant compte simultanément des dynamiques de marché et des évolutions démographiques.

En définitive, ce mémoire met en évidence le potentiel des approches d’apprentissage profond pour renouveler les méthodes de couverture en finance et en assurance, en proposant un cadre flexible, directement fondé sur les données, et adapté aux défis posés par les produits complexes.

Bibliographie

- [1] Buehler, H., Gonon, L., Teichmann, J., & Wood, B. (2019). Deep Hedging. *Quantitative Finance*, 19(8), 1271–1291.
- [2] Limmer, Y., & Horvath, B. (2024). Robust Hedging GANs : Towards Automated Robustification of Hedging Strategies. *Applied Mathematical Finance*, 31(3), 164–201.
- [3] Carbonneau, A. (2021). Deep hedging of long-term financial derivatives. *Insurance : Mathematics and Economics*, 98, 115–133.
- [4] Herrmann, S., Muhle-Karbe, J., & Seifried, F. T. (2017). Hedging with small uncertainty aversion. *Finance and Stochastics*, 21(1), 1–64.
- [5] Perez, D. (2022). *Étude des Variable Annuities*. Mémoire de Master, ISUP.
- [6] Krah, A. (2020). LSMC Methods in the Life Insurance Sector. *ASTIN Bulletin*, 50(3), 737–771.
- [7] EL QALLI, Y. (2020). *Théorie des Options*. INSEA.
- [8] Devolder, P. (2016). *Stochastic Finance in Insurance*. UCLouvain.
- [9] Picha, A. K. (2022). *Deep hedging in financial markets with a large trader*. Diplomarbeit, TU Wien.
- [10] Hainaut, D., & Dupret, J. L. (2024). Option pricing with model-constrained Gaussian process regressions. *UCLouvain Discussion Paper*.
- [11] Gan, G., & Valdez, E. A. (2019). Computational Problems in Variable Annuities. *Risks*, 7(3), 91.
- [12] Coleman, T., Kim, Y., Li, Y., & Patron, M. (2007). Robustly hedging variable annuities with guarantees under jump and volatility risks. *Journal of Risk and Insurance*, 74(2), 347–376.
- [13] Moukodouma, D.-F. B., Denis, C., Mbourou, D. R. R., & Nkoulembene, C. A. (2023). Comparison of LSTM and Wavelet Transform-LSTM Models for Temperature Prediction in a part of Congo Basin. *Preprint*.
- [14] Lai, T. L., & Xing, H. (2025). *Data Science and Risk Analytics in Finance and Insurance*. Chapman & Hall/CRC Financial Mathematics Series.
- [15] Privault, N. (2013). *Notes on Stochastic Finance*. World Scientific.
- [16] Fine, T. L. (1999). *Feedforward Neural Network Methodology*. Springer-Verlag New York.

- [17] Achdou, Y., & Pironneau, O. (2009). *Computational Methods for Option Pricing*. SIAM.
- [18] Ben Slama, E. (2021). *Deep pricing and deep calibration : applications in finance and insurance*. Mémoire de fin d'études, ENSAE Paris – Institut des Actuaire.
- [19] Haggouch, M. *Modélisation et prédiction de la courbe des taux à l'aide des modèles : Nelson-Siegel Svensson et CNN-LSTM*. Projet de fin d'études, INSEA, Royaume du Maroc.
- [20] Sazli, M. H. (2006). A Brief Review of Feed-Forward Neural Networks. *Communications Faculty of Sciences University of Ankara Series A2-A3*.
- [21] Li, T., Zhang, Y., Wang, H., & Liu, M. (2020). A Hybrid CNN-LSTM Model for Forecasting Particulate Matter (PM2.5) Concentration. *IEEE Access*, 8, 27164–3984.
- [22] Brigo, D., & Mercurio, F. (2006). *Interest Rate Models — Theory and Practice* (2nd ed.). Springer Finance.
- [23] Université catholique de Louvain. (2025–2026). *LACTU2240 – Processus avancés et ingénierie de l'assurance vie*. Notes de cours, Master en sciences actuarielles.
- [24] Bielecki, T. R., Cialenco, I., & Pitera, M. (2017). *A survey of time consistency of dynamic risk measures and dynamic performance measures in discrete time : LM-measure perspective*. Probability, Uncertainty and Quantitative Risk, 2, 3. <https://doi.org/10.1186/s41546-017-0012-9>
- [25] DIENDERE, W. K. O. (2026). *Deep Hedging for Variable Annuities and Universal Life Contracts*. GitHub repository. Disponible sur : <https://github.com/olisehStar>

Annexe A. Convexité des mesures de risque

A.1 Preuve : CVaR convexe

La représentation duale (définition 1.2.2) et la convexité de $x \mapsto (x)_+$ donnent, pour $\lambda \in [0, 1]$ et $a = X - w$, $b = Y - w$: $(\lambda a + (1 - \lambda)b)_+ \leq \lambda(a)_+ + (1 - \lambda)(b)_+$. Donc $\text{CVaR}_\alpha(\lambda X + (1 - \lambda)Y) \leq \lambda \text{CVaR}_\alpha(X) + (1 - \lambda) \text{CVaR}_\alpha(Y)$. $\square$

A.2 Preuve : risque entropique convexe

Pour $\mu \in [0, 1]$, l'inégalité de Holder avec $p = 1/\mu$, $q = 1/(1 - \mu)$: $\mathbb{E}[e^{\lambda\mu X} e^{\lambda(1-\mu)Y}] \leq (\mathbb{E}[e^{\lambda X}]^\mu (\mathbb{E}[e^{\lambda Y}])^{1-\mu})$. En prenant $\log / \lambda$: $\rho_\lambda^{\text{ent}}(\mu X + (1 - \mu)Y) \leq \mu \rho_\lambda^{\text{ent}}(X) + (1 - \mu) \rho_\lambda^{\text{ent}}(Y)$.

Ces convexités garantissent que le problème (1.1) est bien posé et compatible avec l'algorithme 1.

Annexe B. Portefeuille auto-financé en temps discret

On considère un marché composé de $d + 1$ actifs :

Un **actif sans risque** de prix

$$S_t^{(0)} = B(t) = \exp\left(\int_0^t r_s ds\right),$$

où r_t est le taux d'intérêt instantané (taux court).

d **actifs risqués** de prix

$$\mathbf{S}_t = (S_t^{(1)}, \dots, S_t^{(d)})^\top.$$

Soit $\delta_t = (\delta_t^{(1)}, \dots, \delta_t^{(d)})^\top$ le vecteur des quantités détenues dans les actifs risqués à la date t . La quantité détenue dans l'actif sans risque est notée $\delta_t^{(0)}$.

Valeur du portefeuille

La valeur totale du portefeuille à la date t est :

$$V_t^\delta = \delta_t^{(0)} S_t^{(0)} + \delta_t^\top \mathbf{S}_t.$$

Stratégie auto-financée

Une stratégie est *auto-financée* si toute variation de composition est financée par les ressources internes du portefeuille. Cela se traduit par :

$$dV_t^\delta = \delta_t^{(0)} dS_t^{(0)} + \delta_t^\top d\mathbf{S}_t.$$

Comme $dS_t^{(0)} = r_t S_t^{(0)} dt$, et en substituant

$$\delta_t^{(0)} = \frac{V_t^\delta - \delta_t^\top \mathbf{S}_t}{S_t^{(0)}},$$

on obtient la dynamique fondamentale :

$$dV_t^\delta = \delta_t^\top d\mathbf{S}_t + r_t (V_t^\delta - \delta_t^\top \mathbf{S}_t) dt.$$

Discrétisation temporelle

Considérons une grille de temps $0 = t_0 < t_1 < \dots < t_N = T$ avec $\Delta t_n = t_{n+1} - t_n$. En supposant que les positions $\delta_{t_{n+1}}$ sont choisies à la date t_n (**stratégie prévisible**, et restent constantes sur $[t_n, t_{n+1})$), on intègre l'équation continue :

$$V_{t_{n+1}}^\delta - V_{t_n}^\delta = \int_{t_n}^{t_{n+1}} \delta_s^\top d\mathbf{S}_s + \int_{t_n}^{t_{n+1}} r_s (V_s^\delta - \delta_s^\top \mathbf{S}_s) ds.$$

Sous l'hypothèse de **positions constantes** ($\delta_s = \delta_{t_{n+1}}$ pour $s \in [t_n, t_{n+1})$) et en négligeant la variation intra-période de V_s^δ , on obtient :

$$V_{t_{n+1}}^\delta - V_{t_n}^\delta \approx \delta_{t_{n+1}}^\top (\mathbf{S}_{t_{n+1}} - \mathbf{S}_{t_n}) + r_{t_n} (V_{t_n}^\delta - \delta_{t_{n+1}}^\top \mathbf{S}_{t_n}) \Delta t_n.$$

En utilisant l'approximation exponentielle $e^{r_{t_n} \Delta t_n} \approx 1 + r_{t_n} \Delta t_n$, on obtient la forme discrète standard :

$$V_{t_{n+1}}^\delta = e^{r_{t_n} \Delta t_n} V_{t_n}^\delta + \delta_{t_{n+1}}^\top (\mathbf{S}_{t_{n+1}} - e^{r_{t_n} \Delta t_n} \mathbf{S}_{t_n})$$

En présence de coûts de transaction $C_{t_{n+1}}(\delta)$, la dynamique devient :

$$V_{t_{n+1}}^\delta = e^{r_{t_n} \Delta t_n} V_{t_n}^\delta + \delta_{t_{n+1}}^\top (\mathbf{S}_{t_{n+1}} - e^{r_{t_n} \Delta t_n} \mathbf{S}_{t_n}) - C_{t_{n+1}}(\delta).$$

Annexe C. Preuve du théorème 5.4.1

Soit $\tilde{P}(t_n) := e^{-rt_n} P(t_n, S_{t_n}, V_{t_n})$ le prix actualisé du passif résiduel sous Q , calculé par FFT (voir partie 4.2.1). On pose également $\tilde{S}_{t_n} := e^{-rt_n} S_{t_n}$ et $\tilde{H}_T := e^{-rT} H_T$.

Dynamique de $\tilde{S}_t$. Par la règle d'Itô appliquée à $\tilde{S}_t = e^{-rt} S_t$:

$$d\tilde{S}_t = e^{-rt} dS_t - r e^{-rt} S_t dt = (\mu - r) \tilde{S}_t dt + \sqrt{V_t} \tilde{S}_t dW_t^{(1), \mathbb{P}}.$$

Dynamique du portefeuille actualisé. La dynamique auto-financée (définition 3.2.1) s'écrit en valeurs nominales :

$$V_{t_{n+1}} = e^{r\Delta t} V_{t_n} + \delta_{t_n}^\Delta (S_{t_{n+1}} - e^{r\Delta t} S_{t_n}).$$

En multipliant par $e^{-rt_{n+1}}$:

$$\tilde{V}_{t_{n+1}} = e^{-rt_{n+1}} V_{t_{n+1}} = e^{-rt_n} V_{t_n} + \delta_{t_n}^\Delta (e^{-rt_{n+1}} S_{t_{n+1}} - e^{-rt_n} S_{t_n}) = \tilde{V}_{t_n} + \delta_{t_n}^\Delta (\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}).$$

En sommant sur $n = 0, \dots, N-1$, la télescopie est exacte car $\tilde{V}_{t_{n+1}} - \tilde{V}_{t_n}$ se somme terme à terme :

$$\tilde{V}_T - \tilde{V}_0 = \sum_{n=0}^{N-1} \delta_{t_n}^\Delta (\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}).$$

Décomposition télescopique de l'erreur. Puisque $\tilde{V}_0 = \tilde{P}(t_0)$ (même capital initial) et $\tilde{P}(t_N) = \tilde{H}_T$ (prix du passif à maturité), on écrit :

$$\begin{aligned} \tilde{H}_T - \tilde{V}_T &= \tilde{P}(t_N) - \tilde{V}_0 - \sum_{n=0}^{N-1} \delta_{t_n}^\Delta (\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}) \\ &= \sum_{n=0}^{N-1} (\tilde{P}(t_{n+1}) - \tilde{P}(t_n)) - \sum_{n=0}^{N-1} \delta_{t_n}^\Delta (\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}) \\ &= \sum_{n=0}^{N-1} \underbrace{\left[\tilde{P}(t_{n+1}) - \tilde{P}(t_n) - \delta_{t_n}^\Delta (\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}) \right]}_{\tilde{\varepsilon}_{t_n}}. \end{aligned}$$

La deuxième égalité utilise la télescopie exacte $\sum_{n=0}^{N-1} (\tilde{P}(t_{n+1}) - \tilde{P}(t_n)) = \tilde{P}(t_N) - \tilde{P}(t_0)$.

Développement de $\tilde{P}(t_{n+1}) - \tilde{P}(t_n)$ par le lemme d'Itô. On applique le lemme d'Itô multidimensionnel (lemme 1.4.1) à $\tilde{P}(t, \tilde{S}_t, V_t)$ sur $[t_n, t_{n+1}]$:

$$d\tilde{P} = \frac{\partial \tilde{P}}{\partial t} dt + \frac{\partial \tilde{P}}{\partial \tilde{S}} d\tilde{S}_t + \frac{\partial \tilde{P}}{\partial V} dV_t + \frac{1}{2} \frac{\partial^2 \tilde{P}}{\partial \tilde{S}^2} (d\tilde{S}_t)^2 + \frac{1}{2} \frac{\partial^2 \tilde{P}}{\partial V^2} (dV_t)^2 + \frac{\partial^2 \tilde{P}}{\partial \tilde{S} \partial V} d\tilde{S}_t dV_t.$$

Les termes de variation quadratique valent :

$$(d\tilde{S}_t)^2 = V_t \tilde{S}_t^2 dt, \quad (dV_t)^2 = \sigma^2 V_t dt, \quad d\tilde{S}_t dV_t = \rho \sigma V_t \tilde{S}_t dt.$$

$\tilde{P}$ satisfait l'EDP de Heston actualisée sous Q :

$$\frac{\partial \tilde{P}}{\partial t} + \frac{1}{2} V_t \tilde{S}_t^2 \frac{\partial^2 \tilde{P}}{\partial \tilde{S}^2} + \rho \sigma V_t \tilde{S}_t \frac{\partial^2 \tilde{P}}{\partial \tilde{S} \partial V} + \frac{1}{2} \sigma^2 V_t \frac{\partial^2 \tilde{P}}{\partial V^2} + \kappa(\tilde{\theta} - V_t) \frac{\partial \tilde{P}}{\partial V} = 0.$$

Tous les termes déterministes s'annulent via cette EDP. En intégrant sur $[t_n, t_{n+1}]$, il reste :

$$\begin{aligned}\tilde{P}(t_{n+1}) - \tilde{P}(t_n) &= \frac{\partial \tilde{P}}{\partial \tilde{S}}(t_n) (\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}) \\ &\quad + \int_{t_n}^{t_{n+1}} \frac{\partial \tilde{P}}{\partial V} [\kappa(\theta - V_s) - \tilde{\kappa}(\tilde{\theta} - V_s)] ds \\ &\quad + \int_{t_n}^{t_{n+1}} \sigma \sqrt{V_s} \frac{\partial \tilde{P}}{\partial V}(s, \tilde{S}_s, V_s) dW_s^{(2), \mathbb{P}} + \underbrace{R_n}_{\tilde{\varepsilon}_{t_n}^{(1)}},\end{aligned}$$

où $R_n = O((\Delta t)^{3/2})$ est le résidu de discrétisation d'Euler.

En effet, le schéma d'Euler-Maruyama appliqué à $\tilde{S}_t$ sur $[t_n, t_{n+1}]$ donne :

$$\tilde{S}_{t_{n+1}} \approx \tilde{S}_{t_n} + (\mu - r)\tilde{S}_{t_n}\Delta t + \sqrt{V_{t_n}}\tilde{S}_{t_n}\Delta W_n^{(1)}.$$

La valeur exacte de $\tilde{S}_{t_{n+1}}$ est donnée par l'intégrale stochastique :

$$\tilde{S}_{t_{n+1}} = \tilde{S}_{t_n} + \int_{t_n}^{t_{n+1}} (\mu - r)\tilde{S}_s ds + \int_{t_n}^{t_{n+1}} \sqrt{V_s}\tilde{S}_s dW_s^{(1), \mathbb{P}}.$$

L'erreur de discrétisation provient du remplacement des intégrandes évalués en s par leur valeur en t_n :

$$\int_{t_n}^{t_{n+1}} (\mu - r)\tilde{S}_s ds - (\mu - r)\tilde{S}_{t_n}\Delta t = \int_{t_n}^{t_{n+1}} (\mu - r)(\tilde{S}_s - \tilde{S}_{t_n}) ds.$$

Or $\tilde{S}_s - \tilde{S}_{t_n} = O(\sqrt{\Delta t})$ car $\tilde{S}$ est un processus de diffusion. Donc :

$$\int_{t_n}^{t_{n+1}} (\mu - r)(\tilde{S}_s - \tilde{S}_{t_n}) ds = O(\sqrt{\Delta t}) \cdot O(\Delta t) = O((\Delta t)^{3/2}).$$

Identification des trois termes d'erreur :

En substituant dans $\tilde{\varepsilon}_{t_n}$:

$$\begin{aligned}\tilde{\varepsilon}_{t_n} &= \tilde{P}(t_{n+1}) - \tilde{P}(t_n) - \delta_{t_n}^\Delta(\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}) \\ &= \left(\frac{\partial \tilde{P}}{\partial \tilde{S}}(t_n) - \delta_{t_n}^\Delta \right) (\tilde{S}_{t_{n+1}} - \tilde{S}_{t_n}) + \tilde{\varepsilon}_{t_n}^{(3)} + \tilde{\varepsilon}_{t_n}^{(1)}.\end{aligned}$$

Terme $\tilde{\varepsilon}_{t_n}^{(1)}$ erreur de discrétisation :

$$\tilde{\varepsilon}_{t_n}^{(1)} = R_n = O((\Delta t)^{3/2}),$$

résidu du schéma d'Euler sur $[t_n, t_{n+1}]$.

Terme $\tilde{\varepsilon}_{t_n}^{(2)}$ erreur de path-dépendance :

Le delta FFT traite Z_{t_n} comme un strike fixe, approchant $\frac{\partial \tilde{P}}{\partial \tilde{S}}$ par $e^{-rt_n} \Delta_P^{\text{euro}}(t_n, Z_{t_n})$. Or la vraie dérivée du passif conditionnel contient un terme supplémentaire dû à la sensibilité de Z_T à S_{t_n} :

$$\frac{\partial \tilde{P}}{\partial \tilde{S}}(t_n) - \delta_{t_n}^\Delta = e^{-rt_n} Q(Z_T = S_{t_n} \mid \mathcal{F}_{t_n}) \cdot \frac{\partial Z_{t_n}}{\partial S_{t_n}} =: \tilde{\varepsilon}_{t_n}^{(2)}.$$

Ce terme est non nul uniquement lorsque $S_{t_n} = Z_{t_n}$ (nouveau maximum atteint), et nul sinon.

Terme $\tilde{\varepsilon}_{t_n}^{(3)}$ risque de volatilité non couvert :

$$\tilde{\varepsilon}_{t_n}^{(3)} = \int_{t_n}^{t_{n+1}} \sigma \sqrt{V_s} \frac{\partial \tilde{P}}{\partial V}(s, \tilde{S}_s, V_s) dW_s^{(2), \mathbb{P}} + \int_{t_n}^{t_{n+1}} \frac{\partial \tilde{P}}{\partial V} [\kappa(\theta - V_s) - \tilde{\kappa}(\tilde{\theta} - V_s)] ds$$

Ce terme est une intégrale stochastique par rapport au brownien de volatilité, non couvert par δ^Δ (qui ne détient que $\tilde{S}$). Comme $\sigma > 0$ et $\frac{\partial \tilde{P}}{\partial V} \neq 0$, on a $\mathbb{E}^\mathbb{P}[(\tilde{\varepsilon}_{t_n}^{(3)})^2] > 0$. En conclusion, en sommant sur $n = 0, \dots, N-1$, on obtient la décomposition :

$$\tilde{H}_T - \tilde{V}_T^{\delta^\Delta} = \sum_{n=0}^{N-1} \tilde{\varepsilon}_{t_n}^{(1)} + \sum_{n=0}^{N-1} \tilde{\varepsilon}_{t_n}^{(2)} + \sum_{n=0}^{N-1} \tilde{\varepsilon}_{t_n}^{(3)}.$$

Les trois sommes sont $\mathcal{F}_T$ -mesurables et d'espérance non nulle sous $\mathbb{P}$ dès que $\sigma > 0$ et $\rho < 0$, ce qui établit la sous-optimalité structurelle du delta-hedging (corollaire 5.4.1.1).

Annexe D. Preuve du proposition 5.4.1

Convexité uniforme du risque entropique sur L^2 : Pour $\rho = \rho_\lambda^{\text{ent}}$, on utilise le développement de Taylor de la fonction génératrice des cumulants. Pour tout $X \in L^2$ centré (sans perte de généralité) :

$$\rho_\lambda^{\text{ent}}(X) = \frac{1}{\lambda} \log \mathbb{E}[e^{\lambda X}] = \mathbb{E}[X] + \frac{\lambda}{2} \text{Var}(X) + O(\lambda^2).$$

Pour deux variables $X, Y \in L^2$, on calcule :

$$\begin{aligned} & \frac{\rho_\lambda^{\text{ent}}(X) + \rho_\lambda^{\text{ent}}(Y)}{2} - \rho_\lambda^{\text{ent}}\left(\frac{X+Y}{2}\right) \\ &= \frac{1}{2} \left(\mathbb{E}[X] + \frac{\lambda}{2} \text{Var}(X) \right) + \frac{1}{2} \left(\mathbb{E}[Y] + \frac{\lambda}{2} \text{Var}(Y) \right) - \mathbb{E}\left[\frac{X+Y}{2}\right] - \frac{\lambda}{2} \text{Var}\left(\frac{X+Y}{2}\right) + O(\lambda^2). \end{aligned}$$

Les termes en espérance s'annulent. Il reste :

$$= \frac{\lambda}{4} \text{Var}(X) + \frac{\lambda}{4} \text{Var}(Y) - \frac{\lambda}{2} \text{Var}\left(\frac{X+Y}{2}\right).$$

Or par la formule de la variance :

$$\text{Var}\left(\frac{X+Y}{2}\right) = \frac{1}{4} (\text{Var}(X) + \text{Var}(Y) + 2\text{Cov}(X, Y)).$$

En substituant :

$$\begin{aligned} & \frac{\lambda}{4} \text{Var}(X) + \frac{\lambda}{4} \text{Var}(Y) - \frac{\lambda}{8} \text{Var}(X) - \frac{\lambda}{8} \text{Var}(Y) - \frac{\lambda}{4} \text{Cov}(X, Y) \\ &= \frac{\lambda}{8} (\text{Var}(X) + \text{Var}(Y) - 2\text{Cov}(X, Y)) = \frac{\lambda}{8} \text{Var}(X - Y) = \frac{\lambda}{8} \|X - Y\|_{L^2}^2. \end{aligned}$$

Donc le risque entropique est uniformément convexe sur L^2 avec constante $\kappa_\rho = \lambda$:

$$\frac{\rho_\lambda^{\text{ent}}(X) + \rho_\lambda^{\text{ent}}(Y)}{2} - \rho_\lambda^{\text{ent}}\left(\frac{X+Y}{2}\right) \geq \frac{\lambda}{8} \|X - Y\|_{L^2}^2.$$

Convexité uniforme du CVaR sur L^2 : Pour $\rho = \text{CVaR}_\alpha$, on utilise la représentation variationnelle :

$$\text{CVaR}_\alpha(X) = \inf_{w \in \mathbb{R}} \left\{ w + \frac{1}{1-\alpha} \mathbb{E}[(X - w)^+] \right\}.$$

On note $\phi_w(X) := w + \frac{1}{1-\alpha} \mathbb{E}[(X - w)^+]$.

Pour tout $w \in \mathbb{R}$, la fonction $x \mapsto (x - w)^+$ est convexe. De plus, sous l'hypothèse que la loi de X est absolument continue (sans atome), la fonction $X \mapsto \mathbb{E}[(X - w)^+]$ est strictement convexe sur L^2 . Plus précisément, pour tout $X \neq Y$ dans L^2 et $\lambda \in (0, 1)$:

$$\mathbb{E}\left[\left(\frac{X+Y}{2} - w\right)^+\right] \leq \frac{1}{2} \mathbb{E}[(X - w)^+] + \frac{1}{2} \mathbb{E}[(Y - w)^+] - \frac{1}{8} \mathbb{E}\left[\mathbf{1}_{\{X > w\} \cap \{Y > w\}} (X - Y)^2\right].$$

En notant $p_w := \mathbb{P}(X > w, Y > w) > 0$ pour w au niveau du VaR, on obtient :

$$\frac{\text{CVaR}_\alpha(X) + \text{CVaR}_\alpha(Y)}{2} - \text{CVaR}_\alpha\left(\frac{X+Y}{2}\right) \geq \frac{p_w}{8(1-\alpha)} \|X - Y\|_{L^2}^2.$$

Donc le CVaR est uniformément convexe sur L^2 avec constante $\kappa_\rho = \frac{p_w}{1-\alpha} > 0$.

Application à la proposition : Dans les deux cas, il existe $\kappa_\rho > 0$ tel que :

$$\frac{\rho(X) + \rho(Y)}{2} - \rho\left(\frac{X+Y}{2}\right) \geq \frac{\kappa_\rho}{8} \|X - Y\|_{L^2}^2.$$

On pose $X = H_T - V_T^{\delta^\Delta}$ et $Y = H_T - V_T^{\delta^*, \text{Markov}}$.

Par optimalité de δ^*, Markov et linéarité de V_T^δ en δ :

$$\rho(Y) \leq \rho\left(\frac{X+Y}{2}\right).$$

En combinant avec l'inégalité de convexité uniforme :

$$\frac{\rho(X) - \rho(Y)}{2} \geq \frac{\kappa_\rho}{8} \|X - Y\|_{L^2}^2.$$

Donc :

$$\rho(X) - \rho(Y) \geq \frac{\kappa_\rho}{4} \|X - Y\|_{L^2}^2.$$

Rélation $\|X - Y\|_{L^2}$ aux stratégies : On calcule :

$$X - Y = V_T^{\delta^*, \text{Markov}} - V_T^{\delta^\Delta} = \sum_{n=0}^{N-1} (\delta_{t_n}^{*, \text{Markov}} - \delta_{t_n}^\Delta) \underbrace{(S_{t_{n+1}} - e^{r\Delta t} S_{t_n})}_{=: A_{t_n}}.$$

En prenant la norme L^2 et en appliquant l'inégalité de Cauchy-Schwarz :

$$\|X - Y\|_{L^2}^2 = \mathbb{E} \left[\left(\sum_{n=0}^{N-1} (\delta_{t_n}^{*, \text{Markov}} - \delta_{t_n}^\Delta) A_{t_n} \right)^2 \right].$$

Par la non-dégénérescence des incréments sous $\mathbb{P}$ (modèle de Heston avec $\sigma_V > 0$), il existe une constante $c > 0$ telle que $\mathbb{E}[A_{t_n}^2] \geq c > 0$ pour tout n . Par l'inégalité de Cauchy-Schwarz discrète :

$$\|X - Y\|_{L^2}^2 \geq c \sum_{n=0}^{N-1} \mathbb{E}[(\delta_{t_n}^{*, \text{Markov}} - \delta_{t_n}^\Delta)^2] = c \mathbb{E}^\mathbb{P}[\|\delta^\Delta - \delta^{*, \text{Markov}}\|_{\ell^2([0, T])}^2].$$

En substituant dans l'inégalité de l'étape précédente et en posant $\tilde{\kappa}_\rho = \frac{\kappa_\rho c}{4}$:

$$\Delta\rho = \rho(H_T - V_T^{\delta^\Delta}) - \rho(H_T - V_T^{\delta^{*, \text{Markov}}}) \geq \frac{\kappa_\rho}{2} \mathbb{E}^\mathbb{P}[\|\delta^\Delta - \delta^{*, \text{Markov}}\|_{\ell^2([0, T])}^2] > 0.$$

L'inégalité est stricte dès que $\delta^\Delta \neq \delta^{*, \text{Markov}}$ $\mathbb{P}$ -p.s. et qui notre cas car $\sigma_V > 0$ et $\rho \neq \text{MSE}$, car le risque de volatilité $\varepsilon_{t_n}^{(3)}$ (théorème 5.4.1) est non nul sous ces conditions.

Annexe E. Preuve du Théorème 5.5.1

Démonstration du gap d'information :

(1) Définition des espérances conditionnelles : On définit les deux espérances conditionnelles de H_T sous les deux niveaux d'information :

$$\mu_n := \mathbb{E}^{\mathbb{P}}[H_T \mid \mathcal{F}_{t_n}], \quad \mu_n^{\text{FFNN}} := \mathbb{E}^{\mathbb{P}}[H_T \mid \hat{\mathcal{G}}_{t_n}],$$

où $\hat{\mathcal{G}}_{t_n} = \sigma(S_{t_n}, V_{t_n}) \subsetneq \mathcal{F}_{t_n} = \sigma(S_{t_0}, \dots, S_{t_n}, V_{t_n})$.

Par la loi des espérances itérées :

$$\mathbb{E}^{\mathbb{P}}[\mu_n] = \mathbb{E}^{\mathbb{P}}[\mu_n^{\text{FFNN}}] = \mathbb{E}^{\mathbb{P}}[H_T].$$

Les deux espérances conditionnelles ont donc la même moyenne.

(2) Décomposition de la variance : Puisque $\hat{\mathcal{G}}_{t_n} \subsetneq \mathcal{F}_{t_n}$, on applique la loi de la variance totale à μ_n conditionnellement à $\hat{\mathcal{G}}_{t_n}$:

$$\text{Var}^{\mathbb{P}}(\mu_n) = \underbrace{\text{Var}^{\mathbb{P}}(\mathbb{E}[\mu_n \mid \hat{\mathcal{G}}_{t_n}])}_{(a)} + \underbrace{\mathbb{E}^{\mathbb{P}}[\text{Var}(\mu_n \mid \hat{\mathcal{G}}_{t_n})]}_{(b) \geq 0}.$$

Or par la loi des espérances itérées, puisque $\hat{\mathcal{G}}_{t_n} \subset \mathcal{F}_{t_n}$:

$$\mathbb{E}[\mu_n \mid \hat{\mathcal{G}}_{t_n}] = \mathbb{E}[\mathbb{E}[H_T \mid \mathcal{F}_{t_n}] \mid \hat{\mathcal{G}}_{t_n}] = \mathbb{E}[H_T \mid \hat{\mathcal{G}}_{t_n}] = \mu_n^{\text{FFNN}}.$$

Donc le terme (a) vaut exactement $\text{Var}^{\mathbb{P}}(\mu_n^{\text{FFNN}})$, et on obtient :

$$\text{Var}^{\mathbb{P}}(\mu_n) = \text{Var}^{\mathbb{P}}(\mu_n^{\text{FFNN}}) + \mathbb{E}^{\mathbb{P}}[\text{Var}(\mu_n \mid \hat{\mathcal{G}}_{t_n})].$$

Puisque le terme $\mathbb{E}^{\mathbb{P}}[\text{Var}(\mu_n \mid \hat{\mathcal{G}}_{t_n})] \geq 0$, on conclut :

$$\text{Var}^{\mathbb{P}}(\mu_n) \geq \text{Var}^{\mathbb{P}}(\mu_n^{\text{FFNN}}).$$

(3) Stricte positivité de la variance résiduelle : Le terme résiduel est strictement positif :

$$\mathbb{E}^{\mathbb{P}}[\text{Var}(\mu_n \mid \hat{\mathcal{G}}_{t_n})] > 0.$$

En effet, $\mu_n = \mathbb{E}[H_T \mid \mathcal{F}_{t_n}]$ dépend de $Z_{t_n} = \max_{k \leq n} S_{t_k}$ via le passif $H_T = (Z_T - S_T)^+$. Or Z_{t_n} n'est pas $\hat{\mathcal{G}}_{t_n}$ -mesurable : la loi conditionnelle de Z_{t_n} sachant (S_{t_n}, V_{t_n}) est non dégénérée sous le modèle de Heston pour tout $n \geq 1$, c'est-à-dire :

$$\text{Var}^{\mathbb{P}}(Z_{t_n} \mid S_{t_n}, V_{t_n}) > 0.$$

Donc μ_n n'est pas $\hat{\mathcal{G}}_{t_n}$ -mesurable, ce qui implique la stricte positivité annoncée. On peut donc écrire :

$$\text{Var}^{\mathbb{P}}(\mu_n) - \text{Var}^{\mathbb{P}}(\mu_n^{\text{FFNN}}) = \mathbb{E}^{\mathbb{P}}[\text{Var}(\mu_n | \hat{\mathcal{G}}_{t_n})] = \text{Var}^{\mathbb{P}}(\mu_n - \mu_n^{\text{FFNN}}) > 0.$$

(4) Passage aux valeurs de risque via la convexité uniforme : Par la convexité uniforme de ρ sur L^2 (proposition ??), il existe $\kappa_\rho > 0$ tel que pour tout $X, Y \in L^2$:

$$\rho(X) - \rho(Y) \geq \frac{\kappa_\rho}{4} \|X - Y\|_{L^2}^2.$$

On pose :

$$X = H_T - V_T^{\delta^*, \text{FFNN}}, \quad Y = H_T - V_T^{\delta^*, \Xi}.$$

Par optimalité de δ^*, Ξ sous information complète $\mathcal{F}_{t_n}$, et par linéarité de V_T^δ en δ :

$$\rho(Y) \leq \rho\left(\frac{X + Y}{2}\right),$$

ce qui, combiné à la convexité uniforme, donne :

$$\rho(X) - \rho(Y) \geq \frac{\kappa_\rho}{4} \|V_T^{\delta^*, \Xi} - V_T^{\delta^*, \text{FFNN}}\|_{L^2}^2.$$

(5) Minoration de $\|V_T^{\delta^*, \Xi} - V_T^{\delta^*, \text{FFNN}}\|_{L^2}^2$ par la variance résiduelle. La différence des portefeuilles vaut :

$$V_T^{\delta^*, \Xi} - V_T^{\delta^*, \text{FFNN}} = \sum_{n=0}^{N-1} (\delta_{t_n}^{\delta^*, \Xi} - \delta_{t_n}^{\delta^*, \text{FFNN}}) A_{t_n},$$

où $A_{t_n} = S_{t_{n+1}} - e^{r\Delta t} S_{t_n}$.

L'écart $\delta_{t_n}^{\delta^*, \Xi} - \delta_{t_n}^{\delta^*, \text{FFNN}}$ provient du fait que $\delta_{t_n}^{\delta^*, \Xi}$ conditionne sur $\mathcal{F}_{t_n}$ (donc utilise Z_{t_n}) tandis que $\delta_{t_n}^{\delta^*, \text{FFNN}}$ ne conditionne que sur $\hat{\mathcal{G}}_{t_n}$. Par les conditions du premier ordre du problème de couverture optimale, cet écart est contrôlé par la variance résiduelle de μ_n :

$$\mathbb{E}^{\mathbb{P}}\left[(\delta_{t_n}^{\delta^*, \Xi} - \delta_{t_n}^{\delta^*, \text{FFNN}})^2\right] \geq c \cdot \text{Var}^{\mathbb{P}}(\mu_n - \mu_n^{\text{FFNN}}),$$

où $c > 0$ est une constante liée à la non-dégénérescence des incréments A_{t_n} sous Heston ($\sigma_V > 0$). En sommant sur n :

$$\|V_T^{\delta^*, \Xi} - V_T^{\delta^*, \text{FFNN}}\|_{L^2}^2 \geq c \sum_{n=0}^{N-1} \text{Var}^{\mathbb{P}}(\mu_n - \mu_n^{\text{FFNN}}).$$

En combinant les étapes 4 et 5, et en posant $\tilde{\kappa}_\rho = \frac{\kappa_\rho \cdot c}{4}$:

$\rho(H_T - V_T^{\delta^*, \text{FFNN}}) - \rho(H_T - V_T^{\delta^*, \Xi}) \geq \kappa_\rho \text{Var}^\mathbb{P}(\mathbb{E}[H_T | \mathcal{F}_{t_n}] - \mathbb{E}[H_T | \hat{\mathcal{G}}_{t_n}]) > 0$. L'égalité est stricte car $\text{Var}^\mathbb{P}(Z_{t_n} | S_{t_n}, V_{t_n}) > 0$ sous Heston pour tout $n \geq 1$ (étape 3).

Biais des positions sur options : La stratégie optimale sur les options, pour le passif lookback, est solution du système de conditions du premier ordre :

$$\mathbb{E}^\mathbb{P}[\rho'(H_T - V_T^{\delta^*}) \cdot (\mathbf{Op}_{t_{n+1}} - e^{r\Delta t} \mathbf{Op}_{t_n})] = 0.$$

Le terme $\rho'(H_T - V_T^{\delta^*})$ dépend de Z_T , et donc de Z_{t_n} . L'espérance conditionnelle de ρ' sur $\hat{\mathcal{G}}_{t_n}$ vaut :

$$\mathbb{E}^\mathbb{P}[\rho'(H_T - V_T^{\delta^*}) | \hat{\mathcal{G}}_{t_n}] = \int \rho'(\cdot) \mu_\mathbb{P}(dz | S_{t_n}, V_{t_n}),$$

où $\mu_\mathbb{P}(\cdot | S_{t_n}, V_{t_n})$ est la loi conditionnelle de Z_{t_n} sachant (S_{t_n}, V_{t_n}) , qui est non dégénérée. Marginaliser sur Z_{t_n} introduit un biais systématique dans l'estimation de la position optimale sur les options, d'autant plus prononcé que $\mathbf{Op}_{t_n}$ et Z_{t_n} sont corrélés.

Annexe F. Preuve du Théorème 5.6.1

Les RNN sont denses dans les SDR continus :

Par le théorème d'approximation universelle, 2.1.2, pour tout compact $K \subset \mathcal{X} \times \mathbb{R}^{d_h}$ et tout $\eta > 0$, la fonction de transition φ_θ d'un RNN à une couche cachée avec activation non polynomiale h peut approcher uniformément n'importe quelle fonction continue $\varphi^* : K \rightarrow \mathbb{R}^{d_h}$ à η près.

Formellement, la composition récursive

$$h_n = \varphi_\theta(x_n, h_{n-1}), \quad n = 1, \dots, N,$$

définit un opérateur sur les trajectoires qui est continu pour la topologie produit sur $\mathcal{X}^{N+1}$. La densité de la famille $\{\varphi_\theta : \theta \in \Theta\}$ dans $\mathcal{C}(K, \mathbb{R}^{d_h})$ implique, par continuité de la composition et induction sur n , que les états cachés $(h_1, \dots, h_N)$ peuvent approcher arbitrairement les états de tout autre SDR continu sur les compacts.

Les LSTM sont denses dans les RNN :

Un RNN standard est un cas particulier de LSTM avec la porte d'oubli fixée à $f_n \equiv \mathbf{1}$ et la porte d'entrée à $i_n \equiv \mathbf{1}$. Réciproquement, par le théorème 2.1.2 appliqué à chacune des fonctions de porte, un LSTM avec activations sigmoïde et tanh peut approcher uniformément sur tout compact n'importe quel RNN continu. Les LSTM héritent donc de la propriété de densité des RNN.

En conclusion, la composition des deux étapes donne : pour toute fonctionnelle causale continue Φ et tout $\varepsilon > 0$, il existe un LSTM LSTM_θ de dimension d_h assez grande tel que

$$\|\psi_\theta(h_N^{(d_h)}(\cdot)) - \Phi(\cdot)\|_{\mathcal{K}} < \varepsilon,$$

Remarque. : Dans notre cadre section 5.6, la fonctionnelle causale cible est :

$$\Phi_g(x_0, \dots, x_n) := \mathbb{E}^{\mathbb{P}}[g(Z_{t_n}) \mid \mathcal{F}_{t_n}],$$

où $x_k = (S_{t_k}, V_{t_k})$ et $Z_{t_n} = \max_{k \leq n} S_{t_k}$. Cette fonctionnelle est causale au sens de la définition 5.6.1 (1) car Z_{t_n} ne dépend que de $(S_{t_0}, \dots, S_{t_n})$.

La continuité de Φ_g pour la topologie produit sur $\mathcal{X}^{n+1}$ résulte de la régularité de la mesure conditionnelle sous le modèle de Heston (définition 1.4.1) : le semigroupe de transition $(t, s, v) \mapsto \mathbb{E}^{\mathbb{P}}[g(Z_{t_n}) \mid S_t = s, V_t = v]$ est continu par les propriétés de régularité parabolique des EDP associées .

Le théorème 5.6.1 garantit donc l'existence d'un LSTM approchant Φ_g uniformément sur tout compact, ce qui fonde le théorème 5.6.2.

Annexe G. Preuve du Corrolaire 5.6.2.1

Convergence vers l'optimum en information complète :

Par le théorème 5.6.2, pour tout $\varepsilon > 0$, il existe un LSTM de dimension cachée d suffisamment grande et de paramètres $\theta \in \Theta$ tel que :

$$\mathbb{E}^{\mathbb{P}} \left[\left\| \delta_{t_n}^{*, \Xi} - \delta_{t_n}^{\text{LSTM}, \theta} \right\|^2 \right]^{1/2} \leq \varepsilon \quad \text{pour tout } n = 0, \dots, N-1.$$

Puisque V_T^δ est linéaire en δ , cette proximité des stratégies se traduit en proximité des PnL :

$$\left\| (H_T - V_T^{\delta^{*, \Xi}}) - (H_T - V_T^{\delta^{*, \text{LSTM}}}) \right\|_{L^2} = \left\| V_T^{\delta^{*, \Xi}} - V_T^{\delta^{*, \text{LSTM}}} \right\|_{L^2} \leq \varepsilon \cdot C_A,$$

où $C_A = \left(\mathbb{E} \left[\sum_{n=0}^{N-1} A_{t_n}^2 \right] \right)^{1/2} < \infty$ sous Heston.

Par la propriété Lipschitz de ρ sur L^2 (proposition 5.4.1) avec constante $C_\rho > 0$:

$$\left| \rho(H_T - V_T^{\delta^{*, \text{LSTM}}}) - \rho(H_T - V_T^{\delta^{*, \Xi}}) \right| \leq C_\rho \cdot \varepsilon \cdot C_A.$$

Comme $\varepsilon \rightarrow 0$ quand $d \rightarrow \infty$, on conclut :

$$\rho(H_T - V_T^{\delta^{*, \text{LSTM}}}) \xrightarrow{d \rightarrow \infty} \rho(H_T - V_T^{\delta^{*, \Xi}}).$$

Inégalité stricte LSTM $\leq$ FFNN. On utilise la décomposition suivante. Par le théorème 5.5.1 , appliqué cette fois en comparant le LSTM au FFNN plutôt qu'à l'optimum en information complète :

$$\rho(H_T - V_T^{\delta^{*, \text{FFNN}}}) - \rho(H_T - V_T^{\delta^{*, \text{LSTM}}}) \geq \frac{\kappa_\rho}{4} \left\| V_T^{\delta^{*, \text{LSTM}}} - V_T^{\delta^{*, \text{FFNN}}} \right\|_{L^2}^2.$$

Par la convexité uniforme de ρ sur L^2 (proposition 5.4.1), on minore la norme L^2 par la variance résiduelle :

$$\left\| V_T^{\delta^*, \text{LSTM}} - V_T^{\delta^*, \text{FFNN}} \right\|_{L^2}^2 \geq c \cdot \text{Var}^{\mathbb{P}} \left(\mathbb{E}[H_T \mid \mathcal{F}_{t_n}] - \mathbb{E}[H_T \mid \hat{\mathcal{G}}_{t_n}] \right),$$

où $c > 0$ provient de la non-dégénérescence des incréments A_{t_n} sous Heston ($\sigma_V > 0$).

En combinant et en posant $\tilde{\kappa}_\rho = \frac{\kappa_\rho \cdot c}{4}$:

$$\rho(H_T - V_T^{\delta^*, \text{FFNN}}) - \rho(H_T - V_T^{\delta^*, \text{LSTM}}) \geq \tilde{\kappa}_\rho \text{Var}^{\mathbb{P}} \left(\mathbb{E}[H_T \mid \mathcal{F}_{t_n}] - \mathbb{E}[H_T \mid \hat{\mathcal{G}}_{t_n}] \right).$$

Stricte positivité : Le membre droit est strictement positif si et seulement si $\mathbb{E}[H_T \mid \mathcal{F}_{t_n}]$ n'est pas $\hat{\mathcal{G}}_{t_n}$ -mesurable. Or :

$$\mathbb{E}[H_T \mid \mathcal{F}_{t_n}] \text{ dépend de } Z_{t_n} \notin \hat{\mathcal{G}}_{t_n} \iff \text{Var}^{\mathbb{P}}(Z_{t_n} \mid S_{t_n}, V_{t_n}) > 0,$$

condition vérifiée sous Heston pour tout $n \geq 1$ par les propriétés de régularité de la mesure conditionnelle (remarque en annexe F). On conclut :

$$\rho(H_T - V_T^{\delta^*, \text{LSTM}}) < \rho(H_T - V_T^{\delta^*, \text{FFNN}})$$

Annexe H. Preuve du Théorème 5.7.1

Time-consistency de $\rho_\lambda^{\text{ent}}$:

Par la loi des espérances itérées :

$$\mathbb{E}^{\mathbb{P}}[e^{\lambda X}] = \mathbb{E}^{\mathbb{P}}[\mathbb{E}[e^{\lambda X} \mid \mathcal{F}_{t_n}]].$$

On pose :

$$Y_n := \frac{1}{\lambda} \log \mathbb{E}[e^{\lambda X} \mid \mathcal{F}_{t_n}],$$

de sorte que $\mathbb{E}[e^{\lambda X} \mid \mathcal{F}_{t_n}] = e^{\lambda Y_n}$.

En substituant :

$$\mathbb{E}^{\mathbb{P}}[e^{\lambda X}] = \mathbb{E}^{\mathbb{P}}[e^{\lambda Y_n}].$$

En prenant $\frac{1}{\lambda} \log(\cdot)$ des deux membres :

$$\frac{1}{\lambda} \log \mathbb{E}^{\mathbb{P}}[e^{\lambda X}] = \frac{1}{\lambda} \log \mathbb{E}^{\mathbb{P}}[e^{\lambda Y_n}],$$

c'est-à-dire :

$$\rho_\lambda^{\text{ent}}(X) = \rho_\lambda^{\text{ent}}(Y_n) = \rho_\lambda^{\text{ent}}\left(\frac{1}{\lambda} \log \mathbb{E}[e^{\lambda X} \mid \mathcal{F}_{t_n}]\right),$$

.

Non time-consistency du CVaR : Pour le CVaR, on exhibe un contre-exemple. Soient $X, Y \in L^1$ avec $X \geq Y$ p.s. La time-consistency exigerait :

$$\text{CVaR}_\alpha(\text{CVaR}_\alpha(X \mid \mathcal{F}_{t_n})) = \text{CVaR}_\alpha(X).$$

Dans le document [24], page 34, Exemple 04 démontre que le CVaR n'est pas time-consistent.

Développement de Taylor : La fonction génératrice des cumulants de X est :

$$K(\lambda) = \log \mathbb{E}[e^{\lambda X}].$$

Par définition des cumulants $\kappa_r = K^{(r)}(0)$, on a :

$$\kappa_1 = \mathbb{E}[X], \quad \kappa_2 = \text{Var}(X), \quad \kappa_3 = \mu_3(X) = \mathbb{E}[(X - \mathbb{E}[X])^3].$$

Le développement en série de $K(\lambda)$ autour de $\lambda = 0$ donne :

$$K(\lambda) = \sum_{r=1}^{\infty} \frac{\kappa_r \lambda^r}{r!} = \kappa_1 \lambda + \frac{\kappa_2}{2} \lambda^2 + \frac{\kappa_3}{6} \lambda^3 + O(\lambda^4).$$

En divisant par λ :

$$\rho_\lambda^{\text{ent}}(X) = \frac{K(\lambda)}{\lambda} = \kappa_1 + \frac{\kappa_2}{2} \lambda + \frac{\kappa_3}{6} \lambda^2 + O(\lambda^3) = \mathbb{E}[X] + \frac{\lambda}{2} \text{Var}(X) + \frac{\lambda^2}{6} \mu_3(X) + O(\lambda^3).$$

Équivalence avec l'utilité exponentielle : Pour $U(x) = -e^{-\lambda x}$ avec $\lambda > 0$:

$$\begin{aligned} \arg \max_{\delta \in \mathcal{H}} \mathbb{E}^\mathbb{P}[U(V_T^\delta - H_T)] &= \arg \max_{\delta \in \mathcal{H}} \mathbb{E}^\mathbb{P}[-e^{-\lambda(V_T^\delta - H_T)}] \\ &= \arg \min_{\delta \in \mathcal{H}} \mathbb{E}^\mathbb{P}[e^{\lambda(H_T - V_T^\delta)}]. \end{aligned}$$

Puisque $x \mapsto \frac{1}{\lambda} \log(x)$ est strictement croissante sur $\mathbb{R}_{++}$ et que $\mathbb{E}^\mathbb{P}[e^{\lambda(H_T - V_T^\delta)}] > 0$, on peut appliquer cette transformation sans changer l'argmin :

$$\begin{aligned} \arg \min_{\delta \in \mathcal{H}} \mathbb{E}^\mathbb{P}[e^{\lambda(H_T - V_T^\delta)}] &= \arg \min_{\delta \in \mathcal{H}} \frac{1}{\lambda} \log \mathbb{E}^\mathbb{P}[e^{\lambda(H_T - V_T^\delta)}] \\ &= \arg \min_{\delta \in \mathcal{H}} \rho_\lambda^{\text{ent}}(H_T - V_T^\delta). \end{aligned}$$

On conclut :

$$\delta^{*, \text{ent}} = \arg \max_{\delta \in \mathcal{H}} \mathbb{E}^\mathbb{P}[U(V_T^\delta - H_T)] = \arg \min_{\delta \in \mathcal{H}} \rho_\lambda^{\text{ent}}(H_T - V_T^\delta)$$

