

La Value-at-Risk

Modèles de la VaR, simulations en Visual Basic (Excel) et autres mesures récentes du risque de marché

François-Éric Racicot*

Département des sciences administratives
Université du Québec, Outaouais

Raymond Théoret

Département Stratégie des Affaires
Université du Québec, Montréal

RePad Working Paper No. 022006

* Adresse postale : François-Éric Racicot, Département des sciences administratives, Université du Québec en Outaouais, Pavillon Lucien Brault, 101 rue Saint Jean Bosco, Gatineau, Québec, Canada, J8Y 3J5.
Correspondance : francoiseric.racicot@uqo.ca.
Raymond Théoret, Département stratégie des affaires, Université du Québec à Montréal, 315 est, Ste-Catherine, Montréal, H2X-3X2. Correspondance : theoret.raymond@uqam.ca.
Ce papier est l'un des chapitres de notre prochain ouvrage intitulé : *Finance computationnelle et gestion des risques*.

Résumé

Depuis la fin des années 90, le Comité de Bâle contraint les banques à calculer périodiquement leur VaR et à détenir un montant de capital suffisant pour essuyer les pertes éventuelles mesurées par la VaR. Il n'existe pas cependant une mesure unique de la VaR. En effet, elle repose sur le concept de volatilité, qui est essentiellement latent. C'est pourquoi les banques se doivent de recourir à plusieurs modèles de VaR de manière à définir la fourchette de leurs pertes éventuelles. Ces calculs sont d'autant plus complexes que la distribution des rendements des titres mesurés à haute fréquence s'éloigne sensiblement de la normale. Cet article propose plusieurs modèles de la VaR ainsi que leur programmation en Visual Basic. Il analyse également d'autres mesures récentes du risque de marché. Nous nous attardons finalement à l'emploi des copules et de la transformée de Fourier au calcul de la VaR.

Abstract

Since the end of the nineties, Basle Committee has required that banks compute periodically their VaR and maintain sufficient capital to pay the eventual losses projected by VaR. Unfortunately, there is not only one measure of VaR because volatility, which is a fundamental component of VaR, is latent. Therefore, banks must use many VaR models to compute the range of their prospective losses. These computations might be complex because the distribution of high frequency returns is not normal. This article analyses many VaR models and produces their programs in Visual Basic. It considers also other new measures of market risk and the use of copulas and Fourier Transform for the computation of VaR.

Mots-clefs : ingénierie financière; simulation de Monte Carlo; banques, copules, transformée de Fourier.

Keywords : financial engineering; Monte Carlo simulation; banks; copulas; Fourier transform.

JEL : G12; G13; G33.

Introduction

Les mesures du risque ont bien évolué depuis que Markowitz a avancé sa célèbre théorie de la diversification de portefeuille à la fin des années 1950, théorie qui devait révolutionner la gestion de portefeuille moderne. L'écart type était alors la mesure du risque d'un portefeuille efficient. Mais pour un titre, cette mesure n'est pas appropriée. En effet, dans le cas d'un titre individuel, le risque est représenté par la covariance de son rendement avec celui des autres titres qui constituent un portefeuille bien diversifié. L'écart type du rendement d'un titre comprend les risques diversifiable et non diversifiable. Or, seul le risque non diversifiable est rémunéré par le marché. Ce risque est représenté par la covariance entre le rendement du titre et les rendements des titres qui constituent un portefeuille hautement diversifié.

Les théories du risque qui ont emboîté le pas à celle de Markowitz se sont attachées aux facteurs qui déterminent le risque d'un titre de même qu'à l'équilibre des marchés financiers. En effet, le modèle de Markowitz exige l'estimation de N variances et de $\frac{N^2 - N}{2}$ covariances si l'on suppose que le nombre de titres qui compose le portefeuille est de N . Quand N devient important, l'estimation de la matrice variance-covariance se présente comme un exercice très laborieux et les risques d'erreurs d'estimations sont loin d'être négligeables, ce qui peut donner lieu à une frontière efficiente pour le moins erronée. Force est donc de simplifier les facteurs de risque. Au lieu de les associer aux rendements des titres, il semble plus approprié d'identifier un nombre limité de facteurs de risque dont dépendent conjointement les rendements. De plus, le modèle de Markowitz ne se donnait pas pour tâche d'expliquer le processus de détermination des niveaux d'équilibre des rendements, les prenant plutôt pour acquis. Le modèle de Markowitz comportait donc de nombreuses failles auxquelles il fallait pallier.

Durant les années 60, Sharpe a proposé le modèle de l'évaluation des actifs financiers, soit le MÉDAF ou le CAPM¹ en anglais. Ce modèle est monofactoriel en ce sens qu'il ne distingue qu'un seul facteur explicatif du risque d'un titre, soit la corrélation entre le rendement de ce titre et celui du portefeuille du marché. C'est ce qu'on appelle le risque systématique² du titre, catégorie de risque qui n'est pas diversifiable. Le risque non systématique, ou risque idiosyncratique, est celui qui est particulier à la compagnie qui émet le titre. Étant diversifiable, il n'est pas rémunéré par le marché. A l'intérieur de la théorie du CAPM, le risque systématique d'un titre équivaut à son bêta qui est une mesure relative du risque établie en comparaison avec le bêta du portefeuille du marché qui, lui, est égal à un.

¹ Acronyme de: *Capital Asset Pricing Model*.

² Dit encore "risque de marché".

Au milieu des années 70 est apparu un autre modèle du risque basé sur l'absence d'arbitrage: l'APT, acronyme de l'expression: *Arbitrage Pricing Theory*. Ce modèle, proposé par Ross, reconnaît que le risque est un phénomène multidimensionnel qui s'explique par plusieurs facteurs. Le modèle APT est donc multifactoriel. Le bêta d'un titre pour un facteur donné est la sensibilité relative du rendement du titre à ce facteur. L'une des faiblesses du modèle APT est qu'il reste muet quant à l'identité des facteurs qui déterminent les rendements des titres.

Au début des années 90, une nouvelle mesure du risque a fait son entrée : la VaR, soit l'acronyme de *Value at Risk*. On reconnaissait en effet de plus en plus les limites des mesures traditionnelles du risque. Il fallait se donner des mesures du risque de baisse de la valeur des actifs. Pour ce faire, il fallait trouver des mesures qui sont davantage reliées à l'ensemble de la distribution des flux monétaires d'un portefeuille. C'est dans ce contexte qu'une mesure nominale du risque a été proposée: la VaR. Cette mesure a d'abord servi à quantifier le risque de marché auquel sont soumis les portefeuilles bancaires. En effet, l'Accord de Bâle a imposé aux banques, en 1997, de détenir un montant de capital réglementaire pour pallier aux risques de marché. Or, ce capital est calculé à partir de la VaR. Cette mesure est ensuite devenue de plus en plus populaire pour évaluer le risque de portefeuilles institutionnels ou individuels. Elle permet entre autres d'évaluer les risques de type asymétrique, comme celui qui est associé aux options, l'écart type et le bêta ne permettant pas de prendre en compte ce risque de façon satisfaisante.

1. VaR et loi normale

Par définition, la VaR est la perte maximale que peut subir un gestionnaire de portefeuille durant une certaine période de temps avec une probabilité donnée. A supposer que cette probabilité soit de 95%, la marge d'erreur ayant trait à cette perte maximale n'est que de 5%. Supposons que la distribution des flux monétaires d'un portefeuille obéisse à une loi normale. Supposons également que la variable aléatoire X représente la valeur du portefeuille, avec $X \sim N(\mu, \sigma^2)$. La variable aléatoire X peut donc être réécrite en termes de la variable normale standard ε , $\varepsilon \sim N(0,1)$:

$$X = \mu + \varepsilon\sigma$$

Soit α le seuil critique associé à la probabilité visée. On peut alors écrire :

$$X = \mu + \alpha\sigma$$

La VaR se calcule alors comme :

$$VaR = E(X) - Q(X, c) = \mu - (\mu + \alpha\sigma) = -\alpha\sigma$$

avec $Q(X,c)$ le quantile associé à la probabilité c . Par exemple, supposons que nous voulions calculer la VaR annuelle d'un portefeuille avec une probabilité de 99%. L'écart-type annuel de ce portefeuille est de 100 millions \$. Pour une probabilité de 99%, $\alpha = -2,326$. La VaR de ce portefeuille pour une période d'un an est donc de : $-\alpha\sigma = 2,326 \times 100 = 232,6 \text{ millions } \$$ ³. Elle est donc ici égale à 232,5 millions \$.

La mesure de VaR que nous venons de donner est une *mesure relative* car elle ne tient pas compte de la moyenne des pertes et profits. Supposons que la VaR soit définie en mesure absolue. La *VaR absolue* est la VaR relative à laquelle on retranche l'espérance du profit au cours de la période considérée⁴. Par exemple, si le profit moyen est de 10 millions \$ dans l'exemple qui vient d'être donné, la VaR absolue est de 222,6 (232,6 – 10) millions \$. Mais comme le profit moyen est généralement quasi-nul sur une courte période de temps, on s'en tient la plupart du temps à la mesure relative de la VaR.

Précisons davantage cette relation entre VaR absolue et VaR relative. On calcule généralement la VaR à partir des rendements d'un titre. Supposons que la période d'observation dt soit d'un mois. Le rendement mensuel espéré pour le titre i est de μ et sa variance mensuelle est de σ^2 . Sa VaR relative au seuil de confiance de 95% est donc de :

$$VaR \text{ relative} = S(-\alpha\sigma\sqrt{dt})$$

avec S , la valeur initiale du portefeuille. On peut calculer cette VaR en recourant la fonction suivante écrite en Visual Basic (Excel) :

Function VaRrel(S, mu, variance, dt, c)

$\alpha = \text{Application.WorksheetFunction.NormSInv}(1 - c)$
 $VaRrel = S * (-\alpha * \text{Sqr}(dt * \text{variance}))$

End Function

avec μ , le rendement mensuel moyen ; V , la variance du rendement ; dt : le pas, ici fixé à un mois et c , le seuil de confiance, ici 95%. Par exemple, supposons que l'on ait investi 10000 \$ dans un titre dont les caractéristiques apparaissent dans le chiffrier suivant :

	A	B
1	S	10000
2	mu	0.01
3	variance	0.005
4	dt	1
5	c	0.95
6		
7	VaR relative	1163.08705

³ N'oublions pas que la VaR est une perte. Elle est ici rapportée en termes absolus, comme cela est l'usage.

⁴ ou à laquelle on ajoute l'espérance de la perte, selon le cas.

La VaR relative de ce portefeuille est alors de 1163,08\$ en utilisant la fonction VaRrel. Ce calcul ne tient pas compte du rendement moyen mensuel espéré pour ce titre, soit 1%. La VaR absolue est obtenue en retranchant ce rendement à la VaR relative, c'est-à-dire :

$$VaR\ absolute = VaR\ relative - (S \times \mu \times dt)$$

On peut calculer la VaR absolue à partir de la fonction VaRabs, qui recourt à la fonction VaRrel:

Function VaRabs(S, mu, variance, dt, c)

alpha = Application.WorksheetFunction.NormSInv(1 - c)

VaRabs = VaRrel(S, mu, variance, dt, c) - (S * mu * dt)

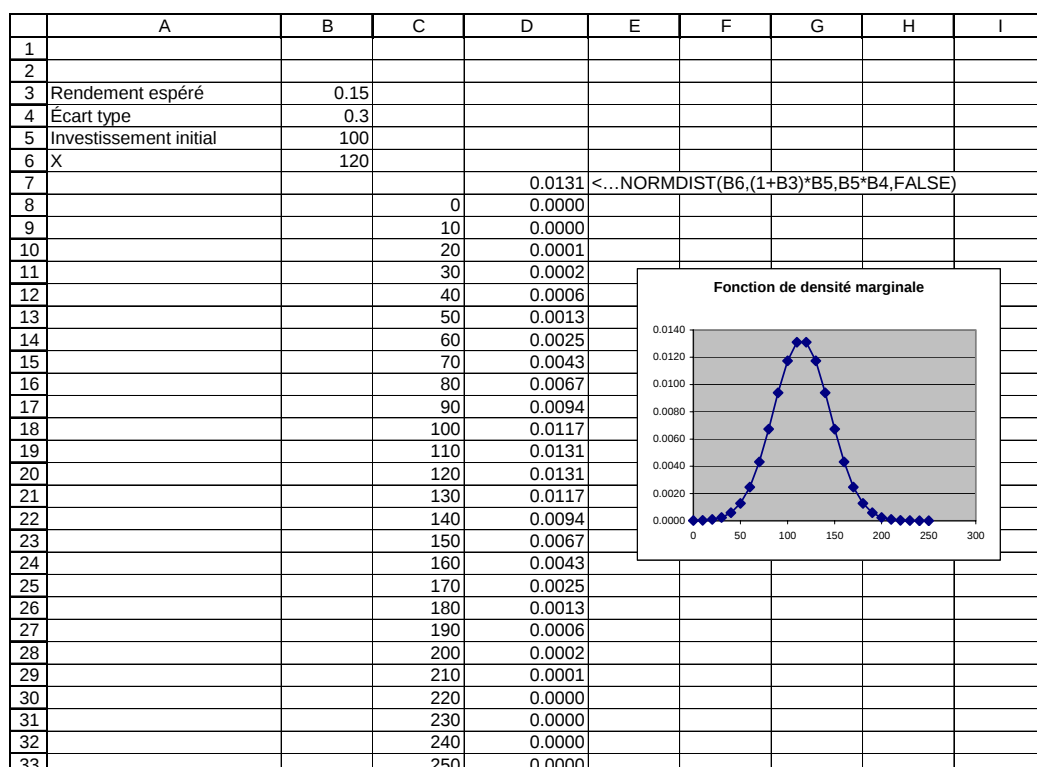
End Function

Dans le cas qui nous intéresse, le profit mensuel espéré sur la détention du titre est de 100\$ (10000 x 0,01). La VaR absolue est donc inférieure à la VaR relative de ce montant. Elle est égale à 1063,08 \$. Dans les sections qui suivent, nous utilisons parfois la VaR relative ou parfois la VaR absolue pour effectuer nos calculs. Le contexte de ces calculs ne devrait pas créer d'ambiguïté.

Pour mieux illustrer le concept de la VaR dans le contexte de la distribution normale, supposons qu'un portefeuille ait un rendement annuel espéré de 15% (μ). L'écart type (σ) de ce rendement est de 30% annuellement. Nous voulons calculer la VaR annuelle avec un probabilité de 95%. L'investissement initial dans ce portefeuille est de 100 \$. Nous supposons que la distribution du rendement du portefeuille est normale. La distribution de ce portefeuille, en dollars, apparaît à la figure 1.

Figure 1

Distribution du portefeuille



Nous nous sommes servis de la fonction Excel *NormDist* pour calculer cette distribution. La formule de cette fonction apparaît à la cellule D7. Elle comprend quatre arguments: 1) la variable aléatoire X qui constitue l'abscisse de la distribution, que l'on appelle aussi *cutoff* en anglais; 2) la moyenne de la distribution, soit: $100(1+\mu)$; 3) l'écart type de la distribution, soit $(100 \times \sigma)$; 4) un argument qui spécifie si l'on veut une distribution cumulative ou non. Ici, nous avons inscrit *FALSE* car nous désirons la distribution marginale. Pour tracer ce graphique, nous avons recouru aux commandes suivantes du menu principal d'Excel: *Data, Table*.

Pour trouver la VaR pour un α de 5%, nous devons calculer la valeur de l'abscisse à gauche de laquelle la surface sous la cloche est de 5%. Il y a deux façons de procéder. Nous pouvons d'abord calculer ce point en recourant à la fonction d'Excel: *Norminv*. Cette fonction comprend trois arguments: i) le seuil (α), ici égal à 5%; ii) l'espérance du portefeuille; iii) l'écart type du portefeuille. Ce calcul apparaît au tableau 1 pour le cas qui nous intéresse.

Tableau 1

Calcul de la VaR pour une distribution normale

	A	B	C	D	E
1					
2					
3	Rendement espéré	0.15			
4	Écart type	0.3			
5	Investissement initial	100			
6	Niveau cible	65.65441	<...=Norminv(0.05,(1+b3)*b5,b4*b5)		

La VaR annuelle pour une probabilité de 95% est donc de:

$$100 - 65.65 = 34.35 \$$$

On aurait pu également recourir à une table normale pour obtenir ce résultat. Pour un α de 5%, la table normale lui associe une valeur égale à 1,655. La perte par dollar est donc de:

$$\text{Perte} = \mu - 1,655\sigma = 0,15 - 1,655(0,30) = -0.3435 \quad (1)$$

Comme le portefeuille se chiffre ici à 100\$, la perte annuelle maximale avec une probabilité de 95%, soit la VaR, est égale à 34,35 \$.

L'autre façon de procéder est de recourir au Solveur d'Excel⁵. Pour ce faire, on efface la formule dans la cellule B6 du tableau 1 et on met un nombre quelconque à la place. Dans la cellule B7, on écrit la formule de la fonction normale de probabilité cumulative comme cela apparaît au tableau 2. L'argument *TRUE* dans la fonction indique que l'on désire la fonction cumulative. Puis on appelle le Solveur. On lui indique que la cellule-cible est B7, que sa valeur doit être de 0,05 et que cette valeur doit être obtenue en modifiant la cellule B6. Le résultat apparaît également au tableau 2.

Tableau 2
Calcul de la VaR à l'aide du Solveur d'Excel

	A	B	C	D	E	F
1						
2						
3	Rendement espéré	0.15				
4	Écart type	0.3				
5	Investissement initial	100				
6	Niveau cible	65.65				
7		0.05	<...=NORMDIST(b6,(1+b3)*b5,b4*b5,TRUE)			

On calcule alors la VaR comme dans le cas précédent. La VaR qui vient d'être calculée couvre une période d'un an et le α correspondant est de 5%. Pour un α de 2,5%, c'est-à-dire pour une probabilité de 97,5%, on remplacerait 1,655 par 1,96 dans l'équation (1) et pour un α de 1%, la valeur critique deviendrait 2,326. Certes, plus on réduit la marge d'erreur, plus la VaR augmente en valeur absolue.

Nous venons de définir la VaR dans le cas le plus simple. Nous avons en effet supposé que la distribution des rendements était normale. La VaR absolue s'exprime alors en termes d'un multiple, disons θ , calculé à partir de la loi normale et relié à la marge d'erreur recherchée, c'est-à-dire⁶:

$$VaR = \text{portefeuille initial}(\mu \Delta t + \theta \sigma \sqrt{\Delta t}) \quad (2)$$

On suppose ici que l'écart type est annualisé; Δt représente la période sur laquelle est calculée la VaR, mesurée en fraction d'année. La VaR dépend donc de trois paramètres: le type de

⁵ Que l'on appelle en cliquant sur *Tools*, puis sur *Solver* dans le menu principal d'Excel.

⁶ N'oublions pas que θ est négatif dans cette formule. La VaR y est donc accompagnée de son signe négatif alors qu'elle était exprimée en termes absolus dans l'exemple antérieur. Il faut toujours bien raisonner une formule pour éviter les ambiguïtés.

distribution auquel obéissent les rendements des titres qui constituent le portefeuille, la période de temps sur laquelle on mesure la VaR⁷, soit l'horizon du calcul, et la marge d'erreur visée. Nous reviendrons ultérieurement sur ces éléments.

Ouvrons ici une parenthèse. En 1996, le Comité de Bâle a décrété qu'il acceptait la «scaling law»⁸. qui dit par exemple que l'écart type mesuré sur 10 jours est égal à l'écart type journalier multiplié par $\sqrt{10}$. Il faut bien se rendre compte que cette loi est associée à la distribution normale. Soit T le nombre de jours et σ , la volatilité journalière du rendement d'un titre. La volatilité sur T jours est alors de $\sigma\sqrt{T}$ selon la «scaling law». Selon McNeil et Frey (2000), pour calculer cette volatilité, il faut plutôt utiliser une formule plus générale, soit σT^{ξ} , où ξ serait supérieur à 0,5 dans la réalité, la distribution des rendements n'étant pas normale. En utilisant la «scaling law» pour calculer la VaR, on sous-estimerait donc systématiquement le risque.

Une question fondamentale se pose ici: à quoi sert la VaR? Mentionnons d'abord qu'elle se révèle d'une grande utilité puisqu'elle est mesurée en termes nominaux en non en pourcentage, tel le bêta. Une fois qu'une institution financière a calculé sa VaR globale, c'est-à-dire la perte maximale qu'elle peut encourir sur l'ensemble de son bilan pour une probabilité prédéterminée, il lui est loisible de se servir de ce montant pour déterminer le capital (avoir propre) minimal qu'elle doit maintenir pour ne pas s'exposer à la faillite. Si en effet elle détient un capital moindre et que la perte maximale se produit, son avoir propre sera négatif et elle devra peut-être déposer son bilan.

La VaR est donc très utile pour une institution financière, car elle lui permet de déterminer le niveau du capital qu'elle doit maintenir pour survivre. Quand la VaR est utilisée à cette fin, on l'appelle plus communément: CaR (*Capital at Risk*), c'est-à-dire que le capital que doit maintenir une institution financière est calculé ou évalué selon les risques auxquels elle est exposée. Plus le risque est important, plus elle devra maintenir un capital élevé. Cela apparaît bien raisonnable, car le capital détenu par une institution financière est d'abord et avant tout un filet de sécurité. Pour une banque, il vise à protéger les dépôts à son passif. La VaR se présente donc comme une mesure appropriée pour définir le capital réglementaire que doit détenir une institution financière. C'est pourquoi le Comité de Bâle, chapeauté par la Banque des règlements internationaux, retenait cette mesure pour calculer le capital réglementaire d'une institution de dépôts en 1995 et qui est devenue effective en janvier 1998. Celles-ci doivent maintenant calculer leur exposition au risque en recourant à la VaR⁹.

⁷ Par exemple, la VaR calculée sur 10 jours est plus élevée que la VaR calculée sur une seule journée.

⁸ Voir à ce sujet: Poon, S.-H. (2005).

⁹ Les règles actuelles concernant le calcul de la VaR pour une institution financière sont les suivantes. La VaR doit être calculée sur un horizon de 10 jours pour un alpha de 1%. On doit recourir à au moins une année

Un autre avantage de la VaR est qu'elle n'est pas assujettie à la distribution normale. Tel n'est pas le cas de l'écart type ou du bêta qui sont des mesures de risque reliées à la loi normale. Il est bien connu que les rendements des titres n'obtempèrent pas à une distribution normale. D'où l'avantage indéniable de la VaR sur les mesures de risque classiques qui sont enchâssées dans la distribution normale.

Nous venons d'étudier la VaR dans un contexte de normalité des rendements. Dans les sections qui suivent, nous considérons d'autres types de distribution des rendements pour calculer la VaR. Nous nous intéresserons successivement aux méthodes de la simulation historique, à la méthode delta, à la simulation de Monte Carlo et à la méthode du *bootstrapping*. Nous nous tournerons finalement vers l'ajustement de Cornish-Fisher.

2. La simulation historique de la VaR¹⁰

La simulation historique est une méthode très simple pour estimer la VaR. Nous l'assimilerons à partir d'un exemple. Nous voulons déterminer la VaR d'un contrat à terme qui permet d'acheter un million \$EU dans trois mois contre livraison de 1,59 million \$CAN. La valeur f_t de ce contrat se calcule comme suit :

$$f_t = S_t \frac{1}{1 + r_t^{EU} \tau} - K \frac{1}{1 + r_t^{CAN} \tau} \quad (3)$$

avec S_t , le prix au comptant du dollar américain en dollars canadiens ; K , le taux de change du dollar américain tel que spécifié dans le contrat ; r^{EU} , le taux d'intérêt sans risque américain, r^{CAN} , le taux d'intérêt sans risque canadien et τ , la durée du contrat, ici trois mois.

d'observations pour calculer cette VaR. L'institution financière doit prendre en compte plusieurs catégories de risques: les risques associés aux instruments financiers non linéaires (produits dérivés), les risques découlant des mouvements de la structure à terme des taux d'intérêt et les risques associés à la base (écart entre le prix au comptant et le prix à terme) pour les matières premières. L'institution financière doit également se livrer au *backtesting*. Mentionnons qu'en réaction à la faillite de la banque Herstatt, les gouverneurs des banques centrales faisant partie du G-10 ont mis sur pied le Comité de Bâle en 1974 dont le rôle est de réglementer et de superviser les pratiques bancaires. Ce Comité est sous la gouverne de la Banque des règlements internationaux. A la suite du krach boursier d'octobre 1987, la Banque des règlements internationaux avait fortement suggéré aux banques en 1988 de détenir un capital réglementaire égal ou supérieur à 8% d'une somme pondérée de leurs actifs risqués. On allouait à chaque actif un coefficient de pondération proportionné à son risque. Cette mesure du risque était statique et ignorait le phénomène de la diversification des portefeuilles. Certes, les banques ne sont pas obligées de suivre les recommandations de la Banque des règlements internationaux. Mais celles qui les négligent risquent de subir une *décote* sur les marchés financiers internationaux. D'où des coûts d'emprunts plus élevés à la clef.

¹⁰ Pour rédiger cette section, nous nous inspirons de : Jorion P. (2003).

Telle qu'elle est écrite, l'équation (3) ressemble beaucoup à celle de Black et Scholes sauf que les probabilités cumulatives $N(d_1)$ et $N(d_2)$ n'apparaissent pas dans cette équation. En effet, un contrat à terme oblige la livraison du sous-jacent, ce qui implique que la probabilité d'exercice, soit $N(d_2)$, est de 1. $N(d_1)$ est par conséquent lui-même égal à 1.

A la date du calcul de la VaR, le taux d'intérêt américain est de 2,75% et le taux d'intérêt canadien s'établit à 3,31%. Le taux de change du dollar américain est alors de 1,588 \$CAN. Selon la formule (3), la valeur du contrat à terme est donc 206\$.

Pour effectuer la simulation historique, on dispose d'une série statistique des 101 jours précédents sur les trois facteurs de risque : le taux d'intérêt américain, le taux d'intérêt canadien et le taux de change du dollar américain. Le tableau 3 reproduit les 25 premières données de cette série. Ces séries ne sont pas en fait des données observées mais plutôt simulées. Les taux d'intérêt sont générés à partir d'une distribution normale alors que le taux de change obéit à la loi de Student qui donne lieu à des sauts sporadiques.

Tableau 3

Séries historiques des trois facteurs de risque

	A	B	C	D
6		r(CAN)	r(EU)	S(\$CAN/\$EU)
7	1	2,4906	2,0732	1,2594
8	2	2,5009	2,0765	1,2868
9	3	2,5011	2,0318	1,3426
10	4	2,5859	2,0859	1,3469
11	5	2,4939	2,0703	1,3504
12	6	2,4902	2,0226	1,3678
13	7	2,5276	2,0833	1,3548
14	8	2,4912	2,0542	1,3765
15	9	2,4127	2,0842	1,3973
16	10	2,4336	2,0798	1,4131
17	11	2,4122	2,0748	1,3514
18	12	2,4109	2,0352	1,3691
19	13	2,4108	1,9962	1,3436
20	14	2,4074	2,0199	1,3297
21	15	2,4138	2,0632	1,2822
22	16	2,3866	2,0030	1,3103
23	17	2,4103	1,9642	1,3730
24	18	2,4504	1,9732	1,3953
25	19	2,4457	1,9295	1,3843
26	20	2,4665	1,9449	1,3203
27	21	2,4858	2,0057	1,3456
28	22	2,5475	1,9707	1,3217
29	23	2,5792	2,0323	1,3434
30	24	2,6175	1,9949	1,3739
31	25	2,6376	1,8744	1,3706

De manière à générer les séries simulées, nous calculons les variations journalières des taux d'intérêt en respectant la séquence historique de même que la variation procentuelle du taux de change de manière à exprimer les données sur la même base. Ces calculs sont répertoriés au tableau 4.

Tableau 4

Variations journalières des données historiques

	F	G	H	I
6		dr(rcan)	dr(rEU)	dS/S
7				
8	1	0,0103	0,0033	0,02177965
9	2	0,0002	-0,0448	0,04338926
10	3	0,0848	0,0542	0,003175
11	4	-0,0920	-0,0157	0,0025894
12	5	-0,0037	-0,0476	0,01289498
13	6	0,0374	0,0606	-0,00952566
14	7	-0,0364	-0,0291	0,01605595
15	8	-0,0785	0,0300	0,01509819
16	9	0,0209	-0,0043	0,01132591
17	10	-0,0214	-0,0051	-0,04369173
18	11	-0,0012	-0,0396	0,01313931
19	12	-0,0001	-0,0390	-0,01861702
20	13	-0,0035	0,0237	-0,01038731
21	14	0,0064	0,0434	-0,0357242
22	15	-0,0272	-0,0602	0,02192705
23	16	0,0238	-0,0388	0,04781492
24	17	0,0400	0,0089	0,01624869
25	18	-0,0047	-0,0436	-0,00784783
26	19	0,0208	0,0154	-0,04623512
27	20	0,0193	0,0608	0,01913887
28	21	0,0616	-0,0350	-0,01773081
29	22	0,0318	0,0617	0,01642149
30	23	0,0383	-0,0374	0,02271151
31	24	0,0201	-0,1205	-0,00239577
32	25	-0,0460	-0,0051	-0,01498118

Nous en sommes maintenant à l'étape de la simulation historique. L'input de la simulation historique est constitué des valeurs observées des trois facteurs de risque le jour du calcul de la VaR et des variations du tableau 4. Au jour du calcul de la VaR, le taux d'intérêt canadien se situe à 3,31%, le taux d'intérêt américain à 2,75% et le taux de change du dollar américain en termes du dollar canadien à 1,588. La simulation historique consiste à appliquer les variations du tableau 4 à ces données de façon à obtenir les valeurs simulées du contrat. Par exemple, le jour 1, on calcule les valeurs révisées des trois facteurs de risque. Le taux d'intérêt canadien est alors de :

$$3,31 + 0,0103 = 3,3141$$

Le taux d'intérêt américain est pour sa part de :

$$2,75 + 0,0033 = 2,8789$$

Et finalement, le taux de change du dollar américain est de :

$$1,588 (1 + 0,0217) = 1,6153$$

En reportant ces valeurs dans la formule (3), on obtient une première valeur simulée pour le contrat à terme :

$$f_t = \left[1,6153 \frac{1}{1 + (0,028789 \times 0,25)} - 1,59 \frac{1}{1 + (0,033141 \times 0,25)} \right] \times 10^6 = 26859$$

Pour le jour 2, on procède de la même façon, c'est-à-dire que l'on applique les variations du jour 2 encore une fois aux valeurs actuelles des trois facteurs de risque qui sont de : 3,31% (taux d'intérêt canadien), 2,75% (taux d'intérêt américain) et 1,588 (taux de change du dollar américain). Les données simulées de la valeur du contrat pour les 25 premiers jours se retrouvent au tableau 5.

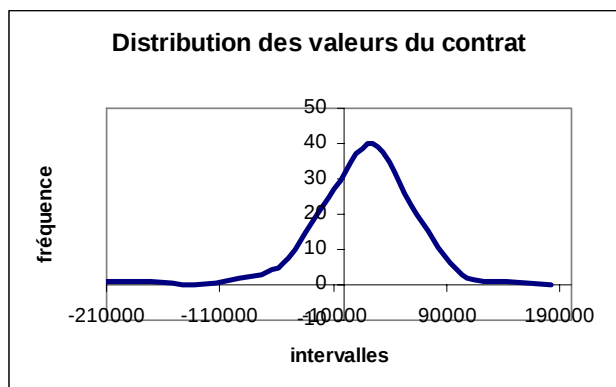
Tableau 5

Facteurs de risque simulés et valeurs correspondantes du contrat

	K	L	M	N	O	P
5	Facteurs de risque simulés					
6	r(1\$CAN)	r(1\$EU)	S	PV(1\$CAN)	PV(1\$EU)	Valeur du contrat
7						
8	3,3141	2,8789	1,61533663	0,9917829	0,99285413	26859
9	3,3039	2,8308	1,64949937	0,99180784	0,9929727	60933
10	3,3886	2,9297	1,58592444	0,99159977	0,9927289	-2251
11	3,2117	2,8599	1,58499867	0,99203461	0,99290102	-3588
12	3,3000	2,8279	1,60129081	0,99181742	0,99297978	13060
13	3,3411	2,9362	1,56584591	0,99171643	0,99271293	-22394
14	3,2673	2,8465	1,60628799	0,99189791	0,99293403	17820
15	3,2252	2,9055	1,60477387	0,99200137	0,99278855	15919
16	3,3246	2,8712	1,59881025	0,99175695	0,99287307	10522
17	3,2823	2,8705	1,51183259	0,99186099	0,99287482	-75998
18	3,3025	2,8360	1,60167706	0,99181131	0,99295987	13421
19	3,3036	2,8366	1,55147333	0,99180858	0,9929585	-36427
20	3,3003	2,8992	1,56448372	0,99181683	0,99280406	-23763
21	3,3102	2,9189	1,5244285	0,99179252	0,99275551	-63565
22	3,2765	2,8153	1,61556964	0,9918753	0,99301085	27196
23	3,3275	2,8368	1,65649592	0,99174988	0,99295791	67948
24	3,3438	2,8845	1,6065927	0,99170992	0,99284037	18271
25	3,2991	2,8320	1,56849839	0,99181198	0,99296988	-19522
26	3,3246	2,8909	1,50781173	0,99175711	0,9928245	-79901
27	3,3230	2,9363	1,61116181	0,99176086	0,99271266	22521
28	3,3654	2,8406	1,55287434	0,99165675	0,99294869	-34810
29	3,3355	2,9373	1,60686588	0,99173021	0,99271039	18301
30	3,3420	2,8382	1,61680981	0,99171419	0,99295453	28593
31	3,3239	2,7550	1,57711759	0,99175881	0,9931595	-10567
32	3,2577	2,8705	1,55722124	0,99192152	0,99287491	-31029

Une fois calculées les 100 valeurs simulées du contrat, on les agence par ordre croissant et on construit l'histogramme. Cet histogramme apparaît à la figure 2.

Figure 2



On constate à la figure 2 que la distribution des valeurs du contrat n'est pas normale puisque par construction, le taux de change obéit à la loi de Student. Le coefficient d'asymétrie de cette distribution est de -0,7 et le leptokurtisme se chiffre à 7,85.

La VaR pour un seuil de confiance de 95% correspond à la 5^{ième} donnée de la distribution de la VaR lorsque les données sont organisées par ordre croissant¹¹. Elle est ici de 80107 \$.

3. La méthode delta du calcul de la VaR¹²

Reprenons l'équation (3) en la simplifiant quelque peu :

$$f_t = S_t P_t^* - K P_t \quad (4)$$

avec P_t^* le facteur d'escompte américain et P_t , le facteur d'escompte canadien. La méthode delta du calcul de la VaR est basée sur l'expansion de Taylor du premier degré de cette équation¹³, soit :

$$df = \frac{\partial f}{\partial S} dS + \frac{\partial f}{\partial P} dP + \frac{\partial f}{\partial P^*} dP^*$$

En calculant les dérivées à partir de l'équation (4), on obtient :

$$df = P^* dS + S dP^* - K dP$$

ce qui s'écrit, en termes d'accroissements :

$$df = (SP^*) \frac{dS}{S} + (SP^*) \frac{dP^*}{P^*} - (KP) \frac{dP}{P} \quad (5)$$

Le contrat à terme équivaut donc à un portefeuille composé des trois éléments suivants : i) une position en compte égale à (SP^*) dans le taux de change ; ii) une position en compte également de (SP^*) en dépôts américains ; iii) une position à découvert à hauteur de KP (emprunt) en dépôts canadiens. Le degré d'exposition du portefeuille dans le facteur de risque représenté par S est de (SP^*) , lequel s'applique également au facteur de risque représenté par P^* . Quant au facteur de risque P , son degré d'exposition est négatif et égal à (KP) .

On peut réécrire l'équation (5) comme suit :

$$df = \chi_1 dz_1 + \chi_2 dz_2 + \chi_3 dz_3$$

où les χ_i représentent les degrés d'exposition aux divers facteurs de risque et les dz_i , les facteurs de risque censés être des variables normales.

¹¹ Puisque la distribution comporte 100 données.

¹² Cette section s'inspire également de Jorion, P. (2003).

¹³ Nous supposons ici qu'il n'existe aucune interaction entre les facteurs de risque.

Soit Ω la matrice variance-covariance des facteurs de risque et χ , le vecteur des degrés d'exposition. Nous calculons la VaR à partir de la variance de df qui est égale à :

$$\sigma^2(df) = \chi' \Omega \chi$$

Nous calculons le vecteur des expositions à partir des données actuelles de S, P et P*, qui sont respectivement de 1,588, 0,9917 et 0,9931¹⁴. Le vecteur des expositions aux trois facteurs de risque est donc de :

	(K*P)	S*P*	S*P*
expositions	-1576803	1577043	1577043

On calcule la matrice variance-covariance à partir des écarts-types historiques et des corrélations des trois facteurs de risque. Nous recourons à notre échantillon de 100 données pour y arriver. Le vecteur des écarts-types est le suivant :

	dP/P	dP*/P*	dS/S
écart-type	1,0637E-05	1,11E-05	0,03230684

tandis que la matrice des corrélations des trois facteurs de risque se lit comme suit :

	dP/P	dP*/P*	dS/S
dP/P	1		
dP*/P*	0,092007741	1	
dS/S	0,038434537	0,106729337	1

Comme nous savons que :

$$Cov(X, Y) = \sigma_x \sigma_y \rho_{xy}$$

nous en déduisons la matrice variance-covariance des facteurs de risque :

	dP/P	dP*/P*	dS/S
dP/P	1,12021E-10	1,07548E-11	1,3076E-08
dP*/P*	1,07548E-11	1,21972E-10	3,789E-08
dS/S	1,30763E-08	3,78901E-08	0,00103329

Nous pouvons donc calculer l'écart-type de df :

$$\sigma(df) = \sqrt{\chi' \Omega \chi} = 50695$$

Pour un seuil de confiance de 95%, le multiple correspondant à la distribution normale se situe à 1,645. La VaR à 95% est donc de :

¹⁴ Le taux d'intérêt canadien est en effet égal à 3,31% et le taux américain, à 2,75%. D'où les facteurs d'escompte P et P*, calculés sur trois mois.

$$50695 \times 1,645 = 83393$$

Ce nombre est relativement rapproché de la VaR calculée en recourant à la simulation historique, soit 80107 \$¹⁵.

4. La simulation de Monte Carlo

Nous voulons toujours calculer la VaR d'un contrat à terme qui permet d'acheter 1 million \$ ÉU dans trois mois contre livraison de 1,59 million \$ CAN. La simulation de Monte Carlo a pour but de calculer la distribution des profits et pertes du contrat, qui nous servira à calculer la VaR du contrat. Pour ce faire, nous devons dans un premier temps spécifier les processus stochastiques du taux de change du dollar canadien, du taux d'intérêt canadien et du taux d'intérêt américain. Cela exige l'estimation des paramètres de ces équations qui relève du calibrage des processus stochastiques, lequel fera l'objet d'un autre chapitre. Puis nous devons imaginer des scénarios pour les trois variables du contrat. A chacun de ces scénarios, nous calculons le profit ou la perte associés au contrat. La simulation d'un nombre approprié de scénarios nous permet de dégager la distribution des profits et pertes du contrat, nécessaire pour le calcul de la VaR.

Nous supposons que le taux de change du dollar canadien suit le mouvement brownien géométrique suivant :

$$d\text{tc} = (\text{tc} \times 0,5 \times dt) + (\text{tc} \times 0,2 \times \varepsilon \sqrt{dt})$$

Le terme aléatoire du taux de change obtempère à une distribution de Student avec 4 degrés de liberté. Cela permet au dollar canadien d'enregistrer des sauts épisodiques. La simulation des scénarios du taux de change du dollar canadien apparaît au tableau 6.

Tableau 6 Programme Visual Basic des scénarios du taux de change du dollar canadien

Sub CAN()

$T = 0.25$

$N = 100$

$dt = T / N$

$\mu = 0.5$

$\sigma = 0.2$

For j = 1 To 1000

$\text{tc} = 1.588$

¹⁵ N'oublions pas que le taux de change obéit à la loi de Student. Le multiple utilisé pour calculer la VaR qui repose sur la loi de Student n'est donc qu'approximatif.

```

For i = 1 To N
Randomize
eps = Application.WorksheetFunction.TInv(Rnd, 4)

If Rnd <= 0.5 Then
eps = -eps
Else
eps = eps
End If

tc = tc + (tc * mu * dt) + (tc * sigma * eps * Sqr(dt))

Next i
Range("tc").Offset(j, 0) = tc
Next j

End Sub

```

Comme on peut le constater au tableau 6, chaque scénario comporte 100 pas. Autrement dit, la durée de trois mois du contrat a été divisée en 100 sous-périodes. Nous inscrivons le résultat de chaque scénario dans le chiffrier, de manière à calculer la distribution des profits et pertes du contrat.

Les taux d'intérêt canadien et américain obéissent au même mouvement brownien arithmétique, soit le suivant :

$$dr = dt + 0,4 \times \varepsilon \sqrt{dt}$$

Seul leur niveau de départ diffère. Le taux d'intérêt canadien est fixé initialement à 3,31% et le taux américain, à 2,75%. Le programme Visual Basic qui effectue la simulation du taux canadien se retrouve au tableau 7.

Tableau 7 Programme Visual Basic de la simulation du taux d'intérêt canadien

```

Sub rcan()

a = 1
sigma = 0.4
T = 0.25
N = 100
dt = T / N

For j = 1 To 1000

rc = 3.31

For i = 1 To N

Randomize

```

$eps = Application.WorksheetFunction.NormSInv(Rnd)$

$rc = rc + (a * dt) + (sigma * eps * Sqr(dt))$

$Range("tauxc").Offset(j, 0) = rc$

Next i

$Range("tauxc").Offset(j, 0) = rc$

Next j

End Sub

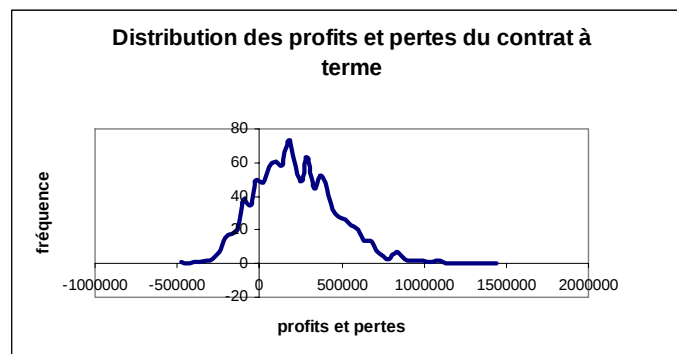
Les dix premiers scénarios des profits et pertes du contrat sont reproduits au tableau 8. Nous avons généré 1000 scénarios de la sorte. Le profit ou la perte du contrat est calculé à partir d'une valeur initiale du contrat de 206\$.

Tableau 8 Scénarios des profits et pertes du contrat

	A	B	C	D	E	F	G	H	I
4		Taux de change	r _{can}	re _u	VP(1\$CAN)	VP(1\$EU)		Valeur contrat	Profits et pertes
5	1	1,648220947	3,23052858	3,23338822	0,99198838	0,99198135		57742,90825	57536,90825
6	2	1,78092442	3,41157263	3,22230085	0,9915432	0,99200862		190138,703	189932,703
7	3	1,883836503	3,48269711	2,82110402	0,99136841	0,99299663		294367,5323	294161,5323
8	4	2,288224234	3,58502268	3,08617825	0,99111706	0,99234363		694828,6143	694622,6143
9	5	1,512627687	3,67379769	2,91593785	0,99089909	0,99276291		-73848,88989	-74054,88989
10	6	1,313164459	3,2467163	3,06882564	0,99194856	0,99238635		-274031,7304	-274237,7304
11	7	1,78732394	3,61169468	3,24261614	0,99105156	0,99195865		197179,4553	196973,4553
12	8	1,689759688	3,44782282	3,14129189	0,9914541	0,99220796		100180,991	99974,99101
13	9	2,422211869	3,44768009	2,73576649	0,99145446	0,99320704		829345,3052	829139,3052
14	10	1,547418464	3,63950249	2,89610453	0,99098328	0,99281178		-39368,13758	-39574,13758

L'histogramme des profits et pertes est tracé à la figure 3. La VaR du contrat, qui correspond à la cinquantième donnée de la série ordonnée des profits et pertes, est ici estimée à 168 937 \$.

Figure 3



5. La technique du bootstrapping¹⁶

Les techniques de rééchantillonnage (*resampling*) se sont vu donner des noms exotiques en anglais: *jackknife* et *bootstrap* (ou bootstrapping) sont de ceux-là. Dans cette section, nous nous intéressons à la technique du bootstrapping. Son instigateur est Efron (1979)¹⁷.

Le concept de la méthode du bootstrapping est simple à saisir. A partir d'un échantillon donné, telle une série statistique sur le prix d'un titre pour une période donnée, on effectue des réaménagements aléatoires de cette série, avec ou sans remise, de façon à décrypter la distribution de la variable qui fait l'objet de l'étude. Sous des conditions très générales, la distribution qui a fait l'objet de bootstrapping converge vers la véritable distribution quand le nombre d'observations tend vers l'infini.¹⁸

5.1 Bootstrapping d'un seul titre

Pour fixer les idées, considérons la situation hypothétique suivante. Nous effectuons un placement dans l'indice boursier TSE300¹⁹ et nous voulons calculer la VaR de notre placement pour une période donnée. Pour ce faire, nous pourrions supposer que la distribution du rendement du TSE300 est normale et calculer sa VaR en recourant à l'équation (2). Mais, selon la figure 4, la distribution du rendement du TSE300 n'est pas normale. Alors comment procéder?

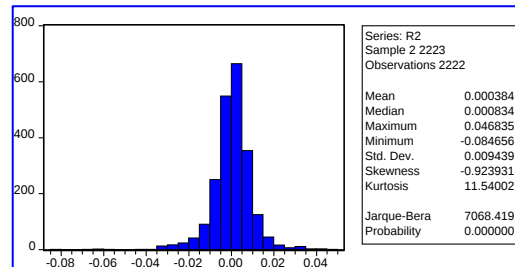
¹⁶ Pour cette section, nous nous référons à: Stuart A., Ord K. et Arnold S., Kendall's Advanced Theory of Statistics. Volume 2A: Classical Inference and the Linear Model, Arnold, 1999; Judge, G.G. et al., Introduction to the Theory and Practice of Econometrics, 2^{ième} édition, Wiley, New York, 1988, chap.9; Benninga, S.(2000), op.cit..

¹⁷ Efron, B. (1979), *Bootstrap methods: Another look at the jackknife*, Ann. Statis., 7. Ce chercheur a publié par la suite plusieurs articles sur le sujet dont l'un des plus récents est: Efron, B. (1994), *Missing Data, Imputation and the bootstrap*, J. Amer. Statis. Ass., 89.

¹⁸ Cependant, lorsque les observations ne sont pas indépendantes, la méthode semble moins bien performer. A ce sujet, voir: Kendall's Advanced Theory of Statistics, volume 2A (1999), p.7.

¹⁹ Qui est devenu le S&PTSX en 2002.

Figure 4
Distribution des rendements du TSE300 1992-2001



Source: EViews, version 5.1.

Benninga (2000) suggère la simulation suivante. L'input de la simulation sera ici une série des rendements journaliers du TSE300 du début de 1999 au début de 2001. On génère d'abord des nombres aléatoires, un pour chaque observation. Ces nombres ont une distribution uniforme et non normale pour ne pas biaiser les résultats. Puis on classe les observations sur le TSE300 par ordre croissant des nombres aléatoires et on obtient une série réaménagée du TSE300. On associe un rendement à cette série. C'est la différence entre le TSE300 de la fin de la série et celui du début, exprimée en pourcentage du TSE300 du début de la série²⁰. Cela constitue un premier élément pour calculer la VaR. On répète cette procédure un grand nombre de fois de façon à générer la distribution empirique du rendement du TSE 300. On peut alors calculer la VaR du portefeuille en trouvant le rendement associé au seuil α recherché.

Pour fixer les idées, nous empruntons ici un programme Visual Basic à Benninga (2000), que nous aurons l'occasion de modifier à souhait par la suite. Pour classer une série, nous recourons à la commande *Sort* d'Excel 2000 (version anglaise), qui classe une série selon différents ordres. On accède à cette commande en cliquant sur la commande *Data* du menu principal, puis sur *Sort*. Mais comme nous voulons ici intégrer cette fonction dans le cadre d'un programme Visual Basic de façon à calculer la VaR du TSE300 par la méthode du bootstrapping, nous devons connaître la syntaxe de la commande *Sort*. La meilleure façon d'y parvenir est d'enregistrer la macro pertinente dans Visual Basic.

²⁰ On pourrait également calculer le rendement en utilisant la formule suivante: $\ln(P_t/P_{t-1})$, avec P_{t-1} l'indice initial et P_t l'indice final.

Tableau 9
Série à classer par ordre croissant de nombres aléatoires

	A	B
1	X	
2	7	
3	8	
4	10	
5	12	
6	15	

Nous voulons classer la série X qui apparaît au tableau 9 sur un chiffrier Excel 2000 (version anglaise) par ordre croissant d'une variable aléatoire. Pour ce faire, nous insérons la fonction *Rand*⁽²¹⁾ dans la cellule B2 et nous la recopions dans les cellules B3 à B6. Le résultat apparaît au tableau 10.

Tableau 10
Génération d'une variable aléatoire pour classer la série X

	A	B
1	X	
2	7	0.399099
3	8	0.787238
4	10	0.678131
5	12	0.136128
6	15	0.122104

Puisque nous avons inséré des variables aléatoires dans la colonne B, elles se recalculent régulièrement. Pour éviter ceci, on les copie dans la même colonne, puis on effectue un collage spécial dans lequel on spécifie que l'on veut des valeurs (fixes) dans cette colonne.

Puis on enregistre la macro correspondant à la commande *Sort*. On clique sur *Tools* dans le menu principal, puis sur *macro*. On clique ensuite sur *Record New Macro*. Une boîte apparaît dans laquelle on demande un nom pour la macro et une clef d'exécution. Nous retenons le nom suggéré et nous choisissons la clef *ctrl+a* pour l'exécution. Nous cliquons

²¹ Cette fonction fournit un nombre aléatoire dont la distribution est uniforme dans l'intervalle [0,1].

sur *OK* et nous sommes alors en mode d'enregistrement de la macro. Une petite boîte pour arrêter l'exécution au moment voulu apparaît à l'écran.

Nous sélectionnons alors les cellules A2 à B6 et nous cliquons sur *Data* dans le menu principal, puis sur *Sort*. Apparaît alors une boîte à l'écran qui nous demande certaines spécifications. On doit d'abord choisir la colonne B comme série servant au classement de la colonne A. On spécifie que l'on veut un classement par ordre croissant des nombres qui sont dans la colonne B. Puis on lui indique qu'il n'y a pas d'entête au niveau de la sélection des séries. On passe ensuite aux options pour indiquer que l'on veut un classement normal, c'est-à-dire de haut en bas. On quitte par la suite la commande *Sort* en cliquant sur *OK* et on arrête l'enregistrement de la macro. Le résultat du classement apparaît au tableau 11.

Tableau 11
Variable X classée par ordre croissant d'une variable aléatoire

	A	B
2	15	0.122104
3	12	0.136128
4	7	0.399099
5	10	0.678131
6	8	0.787238

Les nombres aléatoires sont bien dans un ordre croissant dans la colonne B et la série X a été classée par ordre croissant de ce nombre. Pour visualiser l'enregistrement de la macro, on clique simultanément sur les touches *ALT-F11*, puis sur module puisque la macro a été insérée dans un module, puis sur module 1. La macro qui apparaît alors à l'écran se retrouve au tableau 12.

Tableau 12
La syntaxe de la commande *Sort* en langage Visual Basic

```
Sub Macro1()
'
' Macro1 Macro
' Macro recorded 6/30/2001 by dsa
'
' Keyboard Shortcut: Ctrl+a
'
    Range("A2:B6").Select
    Application.CutCopyMode = False
```

```
Selection.Sort Key1:=Range("A2"), Order1:=xlAscending, Header:=xlNo, _
```

```
OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom
```

```
End Sub
```

Quelques mots sur le programme qui se retrouve au tableau 12. La formule: *Range("A2:B6").Select* permet de sélectionner le rectangle délimité par les cellules A2 à B6. La syntaxe de la commande *Sort* apparaît sur les deux lignes qui commencent par *Selection.Sort*. On voit que cette syntaxe relève d'un langage sibyllin qui ne peut être déchiffré que par l'enregistrement de la macro qui effectue cette opération. Lorsque l'on ne connaît pas la syntaxe d'une commande dans Visual Basic, il est donc recommandé d'enregistrer tout bêtement la macro qui correspond à cette commande. La commande *Record New Macro* enregistre en effet tout ce que vous faites tant que vous ne la stoppez pas!

Nous voulons tout d'abord calculer la VaR d'un placement hypothétique dans le TSE300 à partir des cotes journalières de cette série débutant en janvier 1999 et se terminant en mars 2001. Nous recourons à la méthode du bootstrapping pour y arriver. La mise en forme du chiffrier qui a été utilisé pour effectuer cette simulation se retrouve au tableau 13.

Tableau 13
Chiffrier Excel pour *bootstrapper* le TSE300

	A	B	C	D	E	F	G	H	I	J	K
1		TSE300					Rendemen	-0.0643444		rMax	0.728978
2	1/4/99	10235.39	4.53E-05							rMin	-0.44128
3	1/5/99	6889.77	0.000666								
4	1/6/99	9557.65	0.001353				starttime	4:48:57 PM			
5	1/7/99	7665.62	0.006276				elapsed	05:26.0			
6	1/8/99	7292.58	0.006767								
7	1/11/99	7063.29	0.007993				iterations	1000			
8	1/12/99	7271.3	0.00949								
9	1/13/99	9461.57	0.01059								
10	1/14/99	7042.46	0.014284								
11	1/15/99	7193.21	0.014477								
12	1/18/99	8899.1	0.015301								
13	1/19/99	8937.8	0.016483								
14	1/20/99	8822.48	0.01672								
15	1/21/99	7693.11	0.016735								
16	1/22/99	7023.62	0.019041								
17	1/25/99	6970.81	0.019404								
18	1/26/99	10701.4	0.020311								
19	1/27/99	9836.5	0.02048								

Nous n'avons pas inclus toutes les cotes du TSE300 dans le chiffrier du tableau 13 du fait de la longueur de la série. Plusieurs cellules ont été nommées, de façon à ce que le programme Visual Basic qui effectuera la simulation puisse les reconnaître. Pour nommer une cellule,

vous cliquez dans la case réservée au nom de la cellule. Vous effacez ce nom, disons C2, et vous inscrivez le nom désiré à la place. Par exemple, la cellule C2 porte le nom TSE en plus de C2. Le programme Visual Basic l'identifiera comme TSE.

Rappelons la procédure pour la simulation. Des nombres aléatoires sont générés dans la colonne C du tableau 13, autant que la série du TSE comporte d'observations. La série du TSE est ensuite classée par ordre croissant des nombres aléatoires. Cela constitue une itération. On retient le rendement obtenu sur le TSE au cours de cette itération en faisant la différence entre la cote finale et la cote initiale de la nouvelle série du TSE, rapportée à la cote initiale. La formule qui apparaît dans la cellule H1, qui comptabilise le rendement moyen de l'itération, est donc la suivante:

$$=(B555/B2)-1$$

la dernière observation de la série simulée étant B555. On reporte ensuite le résultat dans la colonne O (tableau 16), qui enregistre le résultat de chaque itération. La première itération est alors terminée.

Comme nous le disions, nous avons nommé des cellules dans le chiffrier pour que le programme Visual Basic puisse les reconnaître. Celles-ci apparaissent au tableau 14.

Tableau 14
Noms de certaines cellules du chiffrier

<i>Cellule</i>	<i>Nom</i>
C2	TSE
H1	Rmoyen
H4	Starttime
H5	Elapsed
H7	Iterations
O2	Returndata

Les cellules H4 et H5 enregistrent l'heure à laquelle a débuté la simulation et le temps qu'elle a duré²². La cellule H7 spécifie le nombre d'itérations désiré. La cellule O2 spécifie l'endroit où le programme enregistre les résultats, ici le rendement correspondant à chaque itération.

²² Ces cellules doivent donc être *formatées* en mode temps.

Fort de cette mise en forme, nous pouvons écrire le programme de la simulation, qui est retranscrit au tableau 15. Pour y arriver, nous touchons simultanément les touches du clavier *ALT* et *F11* pour appeler Visual Basic. Puis, dans le menu principal de Visual Basic, nous cliquons sur *Insert*, puis sur *module*. Nous sommes alors prêts à enregistrer notre programme.

Tableau 15
Programme Visual Basic pour *bootstrapper* le TSE300

```
Sub Varin()

Range("starttime") = Time
Range("O1:O15000").ClearContents
Application.ScreenUpdating = False
For Iteration = 1 To Range("iterations")
For Row = 1 To 554
Range("TSE").Cells(Row, 1) = Rnd
Next Row
Range("B2:C555").Select
Selection.Sort Key1:=Range("C2"), Order1:=xlAscending, Header:=xlNo, _
OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom

Range("returndata").Cells(Iteration, 1) = Range("rmoyen")
Next Iteration

Range("elapsed") = Time - Range("starttime")

End Sub
```

Quelques commentaires sur ce programme²³. La première ligne indique d'insérer l'heure du début de la simulation dans la cellule nommée *starttime* dans le chiffrier. La deuxième ligne

²³ Il y a deux façons de mettre en branle ce programme. La première est de lui définir une clef. Pour ce faire, on clique sur la commande *Tools* du menu principal, puis sur *Macro* et encore une fois sur *Macro*. On clique alors sur le nom de la macro à laquelle on veut donner une clef, puis sur *options*. Dans la section *short key* de la boîte, on introduit une lettre dans la case *ctrl+*, disons *a*. A chaque fois que l'on pressera simultanément sur les touches du clavier: *ctrl* et *a*, la macro se mettra en branle. Une autre façon d'enclencher la macro est de cliquer sur l'icône

indique d'effacer le contenu des cellules O1 à O15000. C'est en effet dans cette colonne que seront répertoriés les résultats de chaque itération. La commande:

Application.Screenupdating=False est là pour éviter qu'Excel ne fasse état de toutes les étapes de la simulation.

Le programme comporte une première boucle:

```
For Iteration = 1 To Range("iterations")  
    ...  
Next Iteration
```

Cette boucle indique au programme d'effectuer le nombre d'itérations contenu dans la cellule nommée "iterations" du chiffrier.

Le programme comporte une seconde boucle:

```
For Row = 1 To 554  
    ...  
Next Row
```

Cette boucle indique d'insérer un nombre aléatoire dans la cellule nommée TSE, qui est la cellule C2, et de remplir les cellules suivantes de cette colonne jusqu'à la ligne 554, comme le commande la boucle. La commande suivante permet cette insertion:

```
Range("TSE").Cells(Row, 1) = Rnd
```

où *Rnd* désigne "variable aléatoire" dans Visual Basic.

Une fois la colonne C remplie de nombres aléatoires, on peut alors classer la série du TSE par ordre croissant des variables aléatoires comme cela a été indiqué auparavant. C'est ce qu'effectue la commande *Selection.Sort*, dont nous avons découvert la syntaxe par simple enregistrement de la macro pertinente. A remarquer que le *Range* spécifié, soit C2, est celui du début de la série de variables aléatoires.

"Run Sub" du menu principal de Visual Basic. Cet icône reproduit la touche *Play* d'un système de sons ou d'un magnétoscope.

La commande:

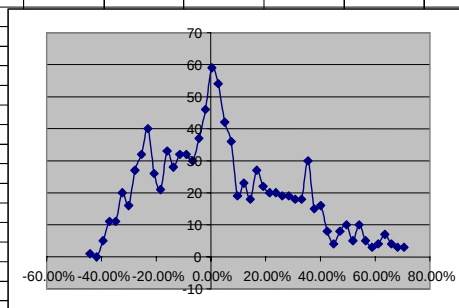
`Range("returndata").Cells(Iteration, 1) = Range("rmoyen")`

indique de reporter le résultat de l'itération, qui se trouve dans la cellule nommée: "rmoyen", dans la cellule nommée: "returndata". La syntaxe *Cells* indique que la cellule O2, nommée "returndata", sera d'abord remplie, soit la *Cells(1,1)* puisque l'on est encore à la première itération. Les autres cellules seront remplies dans l'ordre des itérations.

La première itération est maintenant terminée. La procédure se répète en vertu de la boucle qui commande les itérations et génère 1000 rendements moyens dont il est temps de tracer la distribution. C'était en effet là l'objectif de la procédure du bootstrapping. Les résultats de la simulation apparaissent dans la colonne O du tableau 16.

Tableau 16
Chiffrier final du bootstrapping du TSE300

	O	P	Q	R	S	T	U	V	W	X	Y	Z
1			Bins	Frequence	Cumul.							
2	-0.1519377	1	-44.13%	1	0.10%							
3	0.13334252	2	-41.79%	0	0.10%							
4	0.24828427	3	-39.45%	5	0.60%		Distribution du TSE300: 1999-2001					
5	-0.25333122	4	-37.11%	11	1.70%							
6	-0.0989036	5	-34.77%	11	2.80%							
7	-0.05821718	6	-32.43%	20	4.80%							
8	0.14001748	7	-30.08%	16	6.40%							
9	-0.35540171	8	-27.74%	27	9.10%							
10	0.23027556	9	-25.40%	32	12.30%							
11	0.55162304	10	-23.06%	40	16.30%							
12	-0.0013449	11	-20.72%	26	18.90%							
13	0.48547081	12	-18.38%	21	21.00%							
14	0.41338192	13	-16.04%	33	24.30%							
15	-0.17090294	14	-13.70%	28	27.10%							
16	0.11316556	15	-11.36%	32	30.30%							
17	0.28256049	16	-9.02%	32	33.50%							
18	0.01874663	17	-6.68%	30	36.50%							
19	-0.32817803	18	-4.34%	37	40.20%							
20	0.32164276	19	-2.00%	46	44.80%							
21	-0.28148972	20	0.34%	59	50.70%							
22	-0.09365043	21	2.68%	54	56.10%							
23	0.06552027	22	5.02%	42	60.30%							
24	-0.18251872	23	7.36%	36	63.90%							
25	0.13089187	24	9.70%	19	65.80%							
26	-0.11678212	25	12.04%	23	68.10%							
27	0.00721073	26	14.38%	18	69.90%							
28	-0.30545654	27	16.73%	27	72.60%							
29	0.12920488	28	19.07%	22	74.80%							
30	0.14738985	29	21.41%	20	76.80%							
31	0.00838192	30	23.75%	20	78.80%							
32	-0.2052579	31	26.09%	19	80.70%							
33	0.45701048	32	28.43%	19	82.60%							
34	0.34637657	33	30.77%	18	84.40%							
35	-0.2702949	34	33.11%	18	86.20%							
36	0.16044986	35	35.45%	30	89.20%							
37	-0.2822838	36	37.79%	15	90.70%							
38	0.05121386	37	40.13%	16	92.30%							
39	0.15291653	38	42.47%	8	93.10%							
40	0.14666799	39	44.81%	4	93.50%							
41	-0.11049552	40	47.15%	8	94.30%							
42	0.09096978	41	49.49%	10	95.30%							
43	0.02834678	42	51.83%	5	95.80%							
44	0.19978831	43	54.17%	10	96.80%							
45	-0.23803734	44	56.51%	5	97.30%							
46	0.02166637	45	58.85%	3	97.60%							
47	0.08516526	46	61.20%	4	98.00%							
48	-0.01094322	47	63.54%	7	98.70%							
49	0.03800681	48	65.88%	4	99.10%							
50	-0.11927685	49	68.22%	3	99.40%							
51	0.39145659	50	70.56%	3	99.70%							



Pour calculer la VaR, nous devons d'abord calculer la distribution des rendements du TSE300. De façon à établir les subdivisions (*bins*) de la distribution, nous avons prévu deux formules dans notre chiffrier (tableau 13). D'abord, une formule qui calcule le maximum des rendements dans la colonne O (MAX O:O) et une autre qui en calcule le minimum (MIN O:O). Le premier *bin* que nous introduisons dans la cellule Q2 est K2, soit le minimum des rendements. Nous fixons ici le nombre de bins à 50. Nous entrons donc la formule suivante dans la cellule Q3:

$$=(\$K\$1-\$K\$2)/50 + Q2$$

l'expression entre parenthèses étant la différence entre le rendement maximal et le rendement minimal. Puis nous recopions cette formule des cellules Q4 à Q51. Pour chacun de ces *bins*, nous devons calculer la fréquence des observations qui lui correspond. Pour ce faire, nous recourons à la fonction *Frequency* d'Excel. Nous sélectionnons les cellules R2 à R51 et nous écrivons la formule suivante:

$$=Frequency(O:O, Q2:Q51)$$

Puis nous pressons simultanément les touches au clavier *CRTL-SHIFT-ENTER*²⁴ et les fréquences apparaissent dans la colonne R.

Dans la colonne S, nous calculons la probabilité cumulative associée aux *bins* puisque cette colonne nous servira à calculer la VaR. Dans la cellule S2, nous avons inscrit:

$$=R2/iiterations$$

C'est-à-dire la fréquence marginale du premier *bin* divisée par le nombre d'itérations, ici 1000. Puis nous cumulons les fréquences. A la cellule S3, nous avons inscrit:

$$=R3/iiterations + S2$$

puis nous recopions cette formule jusqu'au *bin* final, c'est-à-dire jusqu'à la cellule S51.

Au tableau 16 apparaît également la distribution des rendements du TSE300 qui a été obtenue par la méthode du bootstrapping. Pour la construire, nous sélectionnons les cellules appropriées des colonnes Q et R puis nous cliquons sur la commande *Chart* dans le menu principal. Nous choisissons la catégorie de graphique *XY-Scatter*. Le reste n'est qu'une simple réponse aux options que présente le menu de construction de graphiques d'Excel.

²⁴ La fonction *Frequency* est en effet une fonction de type *Array* (Rangée) dans le langage Excel comme le sont d'ailleurs les matrices, c'est-à-dire qu'elle lie ensemble plusieurs cellules. Il faut alors presser les touches *CRTL-SHIFT-ENTER* pour la mettre à exécution et non pas seulement sur la touche *ENTER*.

On est maintenant à même de constater que la distribution des rendements journaliers du TSE300 dévie beaucoup de la normale. Il serait donc très hasardeux de calculer la VaR du TSE300 en recourant à la loi normale. Elle est même bimodale. Les rendements nuls sont par ailleurs les plus fréquents.

On peut calculer la VaR du TSE300 en interpolant la fonction de probabilité cumulative qui apparaît à la colonne S. Pour un α de 1%, le rendement que lui associe la fréquence cumulative serait de -38,6% et pour un α de 5%, de -32,1%. Ce résultat est obtenu sur une période de 554 jours, soit le nombre d'observations de la simulation. Pour un α de 1%, le rendement associé à la VaR serait d'environ -0,07% par jour. Sur 10 jours, ce qui est un intervalle courant pour calculer la VaR, la perte serait de 7000 \$ sur un portefeuille d'un million \$.

5.2 Bootstrapping d'un portefeuille de titres

Nous nous attaquons maintenant au bootstrapping des rendements d'un portefeuille de titres. Nous considérons ici trois entreprises canadiennes oeuvrant dans le secteur de la biotechnologie et cotées à la Bourse de Toronto: Axcan Pharma Inc., Glyco Biomedical Inc. et Theratechnologies Inc. Nous les désignerons par leur cote boursière respective soit: AXP, GBL et TH. Nous disposons de données journalières sur les cotes de ces titres pour la période s'étirant du début de janvier 1999 au mois de mars 2001, soit un total de 574 observations. Ces actions ne versent aucun dividende.

A l'instar de la procédure utilisée pour transposer la méthode du bootstrapping au calcul de la VaR du TSE300, nous utilisons dans un premier temps une variable aléatoire différente pour classer chacun des trois titres. Le programme Visual Basic utilisé à cette fin apparaît au tableau 17.

Tableau 17

**Programme Visual Basic pour *boostrapper* 3 titres
sans prise en compte de leur corrélation historique**

```
Sub ValueatRisk1()  
Range("starttime") = Time  
Range("O1:O15000").ClearContents  
Application.ScreenUpdating = False  
For Iteration = 1 To Range("iterations")  
For Row = 1 To 573  
Range("AXPRAND").Cells(Row, 1) = Rnd  
Next Row  
For Row = 1 To 573  
Range("GBLRAND").Cells(Row, 1) = Rnd  
Next Row  
For Row = 1 To 573  
Range("THRAND").Cells(Row, 1) = Rnd  
Next Row  
Range("B2:C574").Select  
Selection.Sort Key1:=Range("C2"), Order1:=xlAscending, Header:=xlNo, _  
    OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom  
Range("D2:E574").Select  
Selection.Sort Key1:=Range("E2"), Order1:=xlAscending, Header:=xlNo, _  
    OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom  
Range("F2:G574").Select  
Selection.Sort Key1:=Range("G2"), Order1:=xlAscending, Header:=xlNo, _  
    OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom  
  
Range("returndata").Cells(Iteration, 1) = Range("rmoyen")  
Next Iteration  
Range("elapsed") = Time - Range("starttime")  
  
End Sub
```

Le tableau 18 donne l'allure générale du chiffrier dans ce cas.

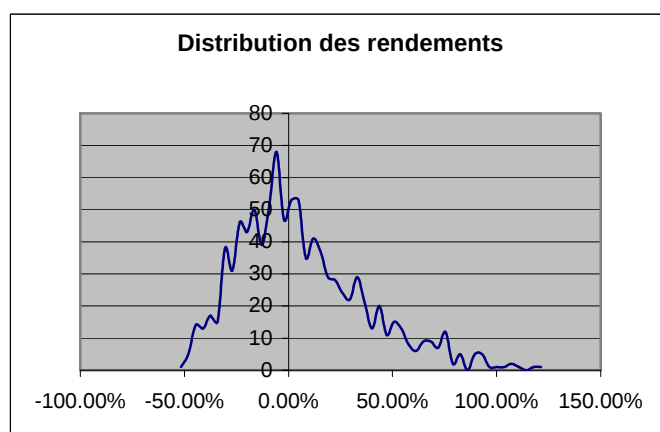
Tableau 18
Chiffrier du bootstrapping de 3 titres
sans prise en compte de leur corrélation historique

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
1		AXP		GBL		TH		Portefeuille			Rendement	-14.29%	rMAX	124.77%	-38.84%
2	1/1/99	9.2	0.000167	6.8	0.001091	5	0.007113099	21					rMin	-51.60%	-46.95%
3	1/4/99	9.8	0.006143	6.05	0.00191	5.2	0.00854224	21.05							32.48%
4	1/5/99	11.7	0.007217	7.5	0.003407	9.25	0.017462015	28.45			starttime	3:23:56 PM			1.07%
5	1/6/99	7	0.008459	6.2	0.006903	12	0.017480493	25.2			elapsed	10:56.0			13.00%
6	1/7/99	15	0.00879	9.2	0.009418	8.3	0.017623723	32.5							-22.99%
7	1/8/99	10.15	0.009539	6.65	0.010701	10.75	0.018807471	27.55			iterations	1000			41.37%
8	1/11/99	6.2	0.013214	6.8	0.011409	11.95	0.020773172	24.95							-32.85%
9	1/12/99	7.45	0.01469	6.5	0.01323	3.8	0.022529185	17.75							-13.26%
10	1/13/99	15.2	0.020979	5.55	0.014458	12.25	0.024957716	33							8.12%
11	1/14/99	9.95	0.024017	5.45	0.015619	4.75	0.028709233	20.15							-7.37%
12	1/15/99	15.2	0.024357	8.35	0.018087	9	0.029308498	32.55							-0.17%
13	1/18/99	18.5	0.024634	5.2	0.020659	4.9	0.030032635	28.6							7.82%
14	1/19/99	8.4	0.025367	6.6	0.021486	8.65	0.032485604	23.65							2.91%
15	1/20/99	10.8	0.025905	6.2	0.02232	4.7	0.03264457	21.7							-26.07%
16	1/21/99	10.95	0.026744	6.65	0.023269	4.6	0.033358872	22.2							42.24%
17	1/22/99	9.65	0.026872	6	0.025281	5.15	0.033763409	20.8							-14.74%
18	1/25/99	9.9	0.029379	6	0.02566	5	0.034745574	20.9							12.32%

La cellule C2 a été nommée AXPRAND; la cellule E2, GBLRAND et la cellule G2, THRAND. Ce sont ces noms qui apparaissent dans le programme Visual Basic (tableau 17) et qui servent à insérer les variables aléatoires dans les cellules.

La distribution des rendements du portefeuille qui découle de l'application de la technique du bootstrapping pour 1000 itérations apparaît à la figure 5.

Figure 5
Distribution des rendements du portefeuille
résultant du programme du tableau 17

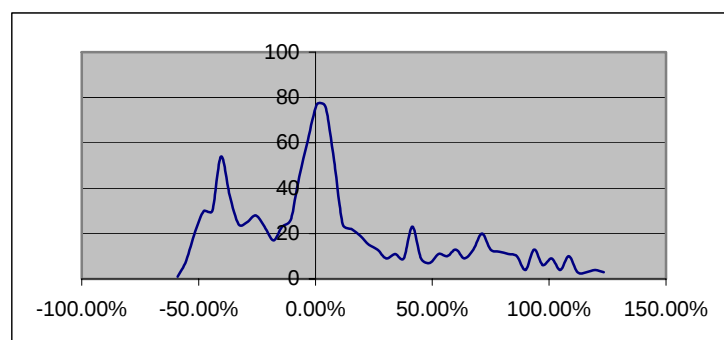


Comme on peut le constater à la figure 5, la distribution des rendements du portefeuille des trois titres a un coefficient d'asymétrie nettement positif, ce qui diminue de beaucoup son lien de parenté avec la distribution normale. Selon la fonction de distribution cumulative, le rendement associé à la VaR pour un alpha de 1% serait de -47,4%, soit -0,08% par jour.

La méthode que nous avons utilisée pour *bootstrapper* le portefeuille des trois titres n'est pas satisfaisante car elle ne prend pas en compte toute l'information qui est incorporée dans l'échantillon. Elle néglige en effet la corrélation qui existe entre les rendements des trois titres. Comme nous avons utilisé une variable aléatoire différente pour les trois titres, nous avons supposé implicitement que la corrélation entre les rendements des trois titres était quasi nulle, ce qui semble être la situation la plus favorable au plan de la diversification du portefeuille car, habituellement, la corrélation entre les rendements des titres est positive.

Une façon rapide de pallier à ce problème est d'appliquer la méthode du bootstrapping à l'ensemble du portefeuille plutôt qu'à chacun des trois titres. La distribution des rendements du portefeuille qui en résulte se retrouve à la figure 6.

Figure 6
Distribution des rendements
Bootstrapping de l'ensemble du portefeuille



La distribution qui apparaît à la figure 6, obtenue par un bootstrapping de l'ensemble du portefeuille en faisant abstraction des titres qui le constituent, diffère nettement de celle qui apparaît à la figure 5, qui elle résulte d'un bootstrapping effectué sur chaque titre. La distribution de la figure 6 ressemble plutôt à celle du TSE300 (figure 4) qui elle-même fut construite sur l'ensemble de l'indice en ne prenant pas en compte les titres qui le constituent. Elle est bimodale, comme cela est le cas pour le TSE300. La distribution des rendements du portefeuille comporte en effet un mode au niveau des rendements nuls et un autre dans la zone des rendements négatifs importants, ce qui, certes, se révèle très défavorable s'agissant du calcul de la VaR. Cela atteste des risques substantiels liés à l'investissement dans le secteur de la biotechnologie.

Pour le portefeuille analysé, le rendement associé à un alpha de 1% est ici de -55,1% sur l'ensemble de la période, ce qui est nettement supérieur au taux obtenu lors du bootstrapping de chacun des titres. Cela était anticipé puisque ce dernier supposait une corrélation nulle entre les rendements des titres, ce qui n'est pas le cas.

Une façon qui nous semble plus satisfaisante pour effectuer le bootstrapping du portefeuille parce qu'elle prélève davantage d'informations dans l'échantillon est de préserver la corrélation historique entre les rendements des trois titres à chaque itération. La matrice de corrélation des cotes boursières des trois titres pour la période étudiée se retrouve au tableau 19.

Tableau 19
Matrice de corrélation des cours
AXP, GBL et TH

	AXP	GBL	TH
AXP	1	-0,023	0,868
GBL	-0,023	1	0,244
TH	0,868	0,244	1

La façon de procéder est la suivante. Il suffit de classer les titres avec la même variable aléatoire. Ici, pour chaque itération, on introduit une variable aléatoire dans la colonne à droite de AXP et on la recopie dans les colonnes limitrophes aux deux autres titres. On s'assure de la sorte que la corrélation des rendements que renferme le tableau 19 est préservée d'une itération à l'autre. Le tableau 20 fait état du programme Visual Basic que nous avons créé à cette fin.

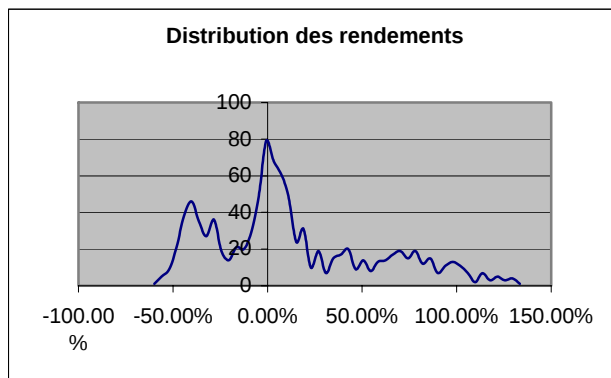
Tableau 20

**Programme Visual Basic pour *bootstrapper* 3 titres
avec prise en compte de leur corrélation historique**

```
Sub ValueatRisk1()  
Range("starttime") = Time  
Range("O1:O15000").ClearContents  
Application.ScreenUpdating = False  
For Iteration = 1 To Range("iterations")  
For Row = 1 To 573  
Range("AXPRAND").Cells(Row, 1) = Rnd  
Range("GBLRAND").Cells(Row, 1) = Range("AXPRAND").Cells(Row, 1)  
Range("THRAND").Cells(Row, 1) = Range("AXPRAND").Cells(Row, 1)  
Next Row  
Range("B2:C574").Select  
Selection.Sort Key1:=Range("C2"), Order1:=xlAscending, Header:=xlNo, _  
OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom  
Range("D2:E574").Select  
Selection.Sort Key1:=Range("E2"), Order1:=xlAscending, Header:=xlNo, _  
OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom  
Range("F2:G574").Select  
Selection.Sort Key1:=Range("G2"), Order1:=xlAscending, Header:=xlNo, _  
OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom  
  
Range("returndata").Cells(Iteration, 1) = Range("rmoyen")  
Next Iteration  
Range("elapsed") = Time - Range("starttime")  
  
End Sub
```

Le chiffrier a la même apparence que celui du tableau 18. La distribution des rendements qui résulte de cette simulation se lit à la figure 7.

Figure 7
Distribution des rendements du portefeuille
avec prise en compte de la corrélation historique des rendements



Cette distribution ressemble évidemment beaucoup à celle de la figure 6 bien que certaines différences soient observables, notamment au niveau du premier mode qui est moins prononcé. Pour un alpha de 1%, le rendement est de -54,2%, donc très rapproché de celui établi par la méthode précédente, qui se situe à -55,1%. Pour un alpha de 5%, les résultats diffèrent cependant davantage. Pour la méthode qui prend en compte la corrélation des rendements, le rendement est de -46,1% alors qu'il est de -49,1% dans le cas de la méthode qui consiste à effectuer un bootstrapping sur l'ensemble du portefeuille.

Avant de quitter cette section, nous pouvons émettre une faiblesse des programmes antérieurs inspirés de Benninga (2000). Comme on aura pu le constater, à chaque itération, nous effectuons le bootstrap de la série qui a fait l'objet du bootstrap antérieur. On pourrait penser qu'il est plus logique de bootstrapper la même série à chaque itération, soit la série historique du titre ou du portefeuille. Bien qu'asymptotiquement ces deux façons de procéder seraient équivalentes, elles ne le seront sans doute pas dans une procédure de tirage sans remise. Dans l'annexe de cet article, nous corrigeons le programme antérieur de manière à bootstrapper la même série à chaque itération.

5.3 Autres façons de *bootstrapper* des séries dans Excel

La version 2000 d'Excel comporte un menu appelé *Poptools* qui permet, entre autres, de *bootstrapper* des séries de deux autres façons. La procédure de bootstrapping que nous venons d'utiliser en est une sans remise. Pour comprendre ce dont il s'agit, considérons un boulier rempli de boules, qui sont les prix de nos actions. Lors de nos simulations, nous avons tiré des prix sans remise. La fonction *Resample*²⁵ de *Poptools* permet d'effectuer des tirages avec remise tandis que la fonction *Shuffle*, également intégrée au menu de *Poptools*, permet, à l'instar de la fonction *Sort* d'Excel, d'effectuer des tirages sans remise. Le tableau 21 prend acte du programme Visual Basic que nous avons conçu pour *bootstrapper* l'ensemble de notre portefeuille de trois titres. Nous savons maintenant en effet qu'il est préférable de *bootstrapper* l'ensemble du portefeuille plutôt que chaque titre séparément dont les résultats sont par la suite additionnés pour former le portefeuille. La fonction *Resample* ne permet pas toutefois de prendre en compte directement la corrélation historique entre les prix des titres du portefeuille, méthode que nous favorisons.

Tableau 21
Programme Visual Basic du bootstrapping avec remise
de l'ensemble du portefeuille

```
Sub ValueatRisk1()  
Range("starttime") = Time  
Range("O1:O15000").ClearContents  
Application.ScreenUpdating = False  
For Iteration = 1 To Range("iterations")  
  
Range("h2:h574").Select  
Selection.FormulaArray = "=Resample(g2:g574)"  
  
Range("returndata").Cells(Iteration, 1) = Range("rmoyen")  
Next Iteration  
Range("elapsed") = Time - Range("starttime")  
  
End Sub
```

Le chiffrier ayant servi à la simulation se retrouve au tableau 22.

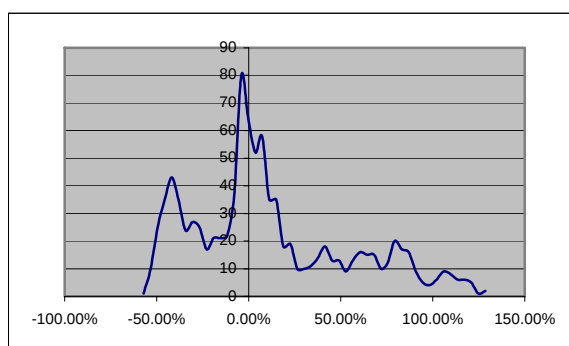
²⁵ En fait, Efron (1979) favorisait cette méthode.

Tableau 22
Chiffrier du bootstrapping avec remise
de l'ensemble du portefeuille

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O
1		AXP	GBL	TH				Portefeuille			Rendement	50.98%	rMAX	132.46%	-5.90%
2	1/1/99	10.75	5.55	2.9			19.2	18.81					rMin	-57.09%	11.13%
3	1/4/99	10.55	5.5	2.85			18.9	32.35							28.71%
4	1/5/99	10.2	5.55	2.8			18.55	17.8			starttime	2:46:10 PM			3.13%
5	1/6/99	10.1	5.55	2.8			18.45	17.8			elapsed	00:23.0			-4.09%
6	1/7/99	10.3	5.5	2.8			18.6	17.9							106.63%
7	1/8/99	10.75	5.6	2.8			19.15	18.05			iterations	1000			81.92%
8	1/11/99	10.7	6	2.8			19.5	25.25							-36.23%
9	1/12/99	10.7	6	2.9			19.6	19.35							58.47%
10	1/13/99	10.75	6.1	2.9			19.75	32.5							-36.25%
11	1/14/99	10.5	6.05	2.7			19.25	18.25							-17.49%
12	1/15/99	10.2	6	2.7			18.9	20.05							10.36%
13	1/18/99	10.05	5.95	2.85			18.85	28.05							71.81%
14	1/19/99	10.15	6	2.8			18.95	30.1							-32.04%
15	1/20/99	10.15	6.1	2.9			19.15	18.65							-48.80%
16	1/21/99	9.85	6.05	2.75			18.65	36.5							-22.89%
17	1/22/99	10	6.1	2.8			18.9	18.65							23.90%

Comme l'indique le programme du tableau 21, à chaque itération, on sélectionne le *range* H2 à H574 de façon à insérer les résultats de l'itération et on appelle la fonction *Resample* pour *bootstrapper* avec remise la série des données de notre portefeuille qui se trouve dans les cellules G2 à G574. A noter que ces dernières cellules doivent avoir été converties en valeurs avant d'effectuer la simulation. Les autres lignes du programme ont déjà été expliquées. La distribution des rendements du portefeuille qui résulte de ces itérations se retrouve à la figure 8.

Figure 8
Distribution des rendements du portefeuille:
Bootstrapping avec remise (1000 itérations)

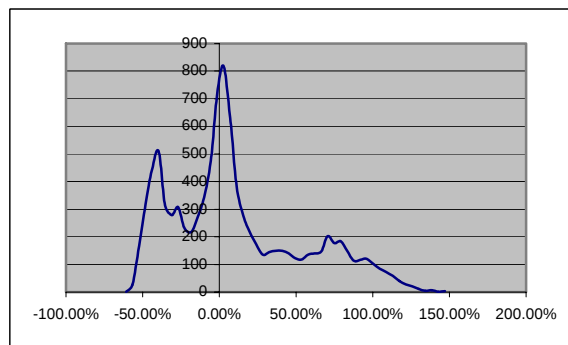


La figure 8 ressemble à la figure 6 bien qu'elle comporte moins de soubresauts. Son premier mode est moins accentué, ce qui la rapproche davantage de la figure 7.

L'un des avantages de la fonction *Resample* est qu'elle permet d'effectuer rapidement un très grand nombre d'itérations, ce qui n'est pas le cas pour la fonction *Sort*. A la figure 9,

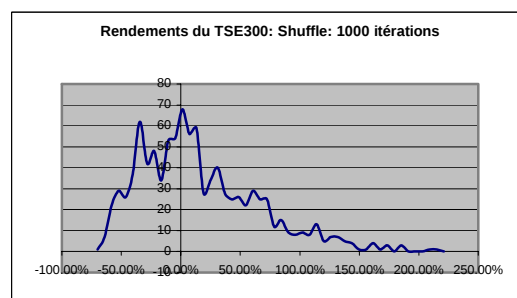
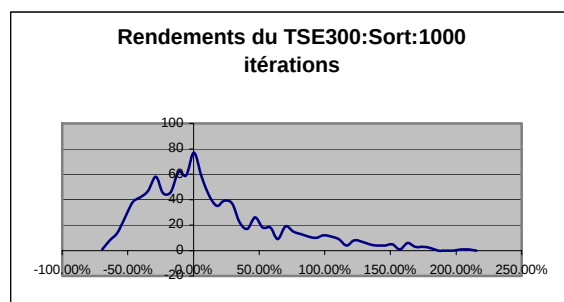
10000 itérations ont été effectuées pour obtenir la distribution plutôt que 1000 comme à la figure 8.

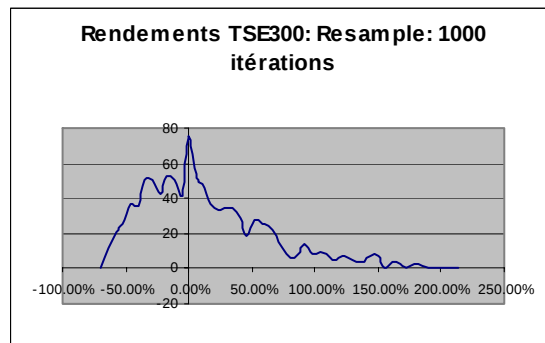
Figure 9
Bootstrapping avec remise:
10000 itérations



Après 10000 itérations, l'aspect de la distribution ressemble beaucoup à celui de la figure 8 mais celle-ci est devenue plus continue, comme il fallait s'y attendre. Il n'y aurait donc guère de gain additionnel à *bootstrapper* une série au delà de 1000 itérations.

Figure 10
Comparaison des trois techniques de bootstrapping





A la figure 10, de façon à comparer les trois techniques de bootstrapping, nous *bootstrapons* les rendements journaliers du TSE300 pour la période 1992-2001 selon les trois techniques dont il a été question dans ce chapitre. D'abord, par la technique *Sort* puis par les deux fonctions de *Poptools*: *shuffle* (bootstrapping sans remise) et *resample* (bootstrapping avec remise). Le programme qui a servi à créer la distribution avec la fonction *Shuffle*, qui n'a pas encore été fourni, apparaît au tableau 23. On est à même de constater que ces trois méthodes donnent des résultats très comparables. La distribution des rendements journaliers du TSE300 présente un coefficient d'asymétrie positive important. On remarque cependant que le bootstrapping avec remise altère quelque peu l'allure de la distribution.

Tableau 23
Programme Visual Basic:
Bootstrapping du portefeuille avec la fonction *Shuffle*

```
Sub Varin()

Range("starttime") = Time
Range("O1:O15000").ClearContents
Application.ScreenUpdating = False
For Iteration = 1 To Range("iterations")
Range("C2:C2224").Select
' Les résultats seront compilés dans les cellules C2 à C224
Selection.FormulaArray = "=Shuffle(b2:b2224)"
' La série du TSE se trouve dans les cellules b2:b2224. On demande de bootstrapper cette série _
' par la fonction Shuffle
Range("returndata").Cells(Iteration, 1) = Range("rmoyen")
```

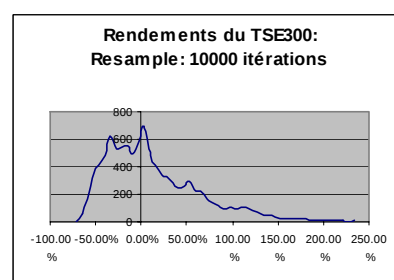
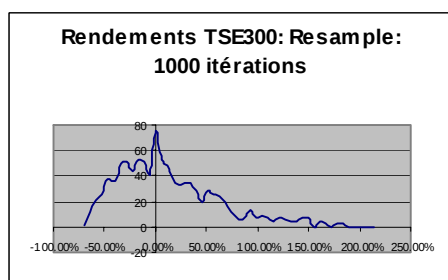
```
Next Iteration
```

```
Range("elapsed") = Time - Range("starttime")
```

```
End Sub
```

N.B. Visual Basic ne prend pas en compte les lignes qui s'ouvrent par : '. Ce sont de simples explications ou aide-mémoire.

Figure 11
Bootstrapping du TSE300 avec la fonction *Resample*:
1000 et 10000 itérations



Finalement, à la figure 11, on peut comparer un *resample* qui comporte 1000 itérations avec un autre qui en comprend 10000. Comme nous l'avons allégué auparavant, l'augmentation du nombre d'itérations au delà de 1000 n'altère guère les résultats, sinon qu'elle adoucit les fluctuations de la distribution.

5.4 Autre façon de calculer la VaR d'un portefeuille de titres

Soit un portefeuille constitué de deux titres, libellés 1 et 2. Nous disposons d'une estimation des VaR respectives de ces titres, soit VaR_1 et VaR_2 . Qui plus est, nous connaissons la matrice de corrélation des rendements de ces deux titres, donnée par :

$$\begin{bmatrix} 1 & \rho_{12} \\ \rho_{21} & 1 \end{bmatrix}$$

La VaR du portefeuille des deux titres, notée par VaR_p , est alors égale à :

$$VaR_p = \sqrt{\begin{bmatrix} VaR_1 & VaR_2 \end{bmatrix} \begin{bmatrix} 1 & \rho_{12} \\ \rho_{21} & 1 \end{bmatrix} \begin{bmatrix} VaR_1 \\ VaR_2 \end{bmatrix}} \quad (6)$$

Pour analyser la formule (6), on peut distinguer trois cas. D'abord celui pour lequel la corrélation entre les rendements des deux titres est de 1. Selon la formule (6), la VaR du portefeuille est alors de :

$$VaR_p = VaR_1 + VaR_2$$

La VaR du portefeuille est donc dans ce cas égale à la somme des VaR des deux titres. On retrouve ici l'un des principes de la diversification de Markowitz. Aucune diversification des portefeuilles n'est en effet possible lorsque la corrélation entre les rendements des titres regroupés dans un portefeuille est de 1. Il n'y a alors aucun avantage à diversifier un portefeuille.

Passons maintenant au cas pour lequel la corrélation entre les rendements des deux titres est de -1. L'application de la formule (6) donne alors :

$$VaR_p = |VaR_1 - VaR_2|$$

Selon Markowitz, c'est là le cas idéal de la diversification. On peut en arriver à une couverture parfaite si $VaR_1 = VaR_2$, c'est-à-dire que VaR_p est alors nul.

Soit un troisième cas, celui pour lequel la corrélation entre les rendements des deux titres est nulle. On parle alors de «pooling» des risques. Selon la formule (6), la VaR du portefeuille est alors de :

$$VaR_p = \sqrt{VaR_1^2 + VaR_2^2}$$

6. L'expansion de Cornish-Fisher et la VaR²⁶

La VaR n'est certes qu'une approximation. C'est pourquoi plusieurs méthodes valent mieux qu'une pour la calculer. L'expansion de Cornish-Fisher²⁷ est l'une de celles-là. Celle-ci est une relation approximative entre les percentiles d'une distribution et ses moments. Au dire de Stuart et al. (1999), un très grand nombre de distributions que l'on retrouve en Statistique tendent vers la normale quand n (le nombre d'observations) se dirige vers l'infini, mais dans des échantillons de moindre envergure, la distribution normale peut laisser beaucoup à désirer. C'est pourquoi il faut recourir à l'expansion de Cornish-Fisher dans ce cas pour approximer les percentiles d'une distribution. Cette approximation, basée sur la série de Taylor, recourt aux moments d'une distribution qui dévie de la normale pour calculer

²⁶ Pour cette section, nous nous référons à: Stuart A., Ord K. et Arnold S.(1999), Kendall's Advanced Theory of Statistics. Volume 1: Distribution Theory, Arnold, pp. 236-240; Hull J.C. (2000), Options, Futures and Other Derivatives, Prentice Hall, New Jersey, chap. 14.

²⁷ Cornish, E.A et Fisher, R.A. (1937), *Moments and Cumulants in the Specification of Distributions*, Rev. Int. Statist. Inst., 307.

ses percentiles. Hull (2000) fournit cette approximation jusqu'au troisième moment d'une distribution. L'approximation de Cornish-Fisher s'écrit alors comme suit:

$$w_{\alpha} \cong z_{\alpha} + \frac{1}{6}(z_{\alpha}^2 - 1)AS \quad (7)$$

Dans cette expression, w_{α} est le percentile corrigé de la distribution au seuil α ; z_{α} est le percentile correspondant à une $N(0,1)$ et AS est le coefficient d'asymétrie. Pour fixer les idées, reprenons l'équation (2) dans laquelle nous supposons que l'espérance du rendement d'un portefeuille se situe à 0,15, son écart-type annuel, à 0,30 et où nous supposons également que les rendements ont un coefficient d'asymétrie nul ($AS = 0$). Nous voulons calculer la VaR annuelle par dollar pour un alpha de 5%. L'équation (7) s'écrit alors: $w_{5\%} = z_{5\%} = -1,655$, soit le percentile ou valeur critique de la normale au seuil de 5%. La VaR annuelle par dollar est, sous ces hypothèses:

$$Perte = \mu - 1,655\sigma = 0,15 - 1,655(0,30) = -0.3435$$

Supposons maintenant que le coefficient d'asymétrie des rendements soit de -0,5. L'expansion de Cornish-Fisher s'écrit alors:

$$w_{\alpha} = z_{\alpha} + \frac{1}{6}(z_{\alpha}^2 - 1)AS = -1,655 + \frac{1}{6}[(-1,655)^2 - 1] \times -0,5 = -1,80$$

Corrigée de l'asymétrie des rendements du portefeuille, la VaR par dollar devient:

$$0,15 - 1,80(0,30) = -0,3900$$

La VaR se voit rehaussée de 13,5% à la suite du changement du coefficient d'asymétrie de 0 à -0,5. Si la distribution des rendements présente un coefficient d'asymétrie négatif, cela augmente la VaR, donc le risque, comme nous le soulignons précédemment. Un coefficient d'asymétrie positif diminue la VaR, donc le risque.

Il est malheureux que Hull (2000) ne soit pas allé plus loin que le troisième moment d'une distribution pour écrire l'expansion de Cornish-Fisher --, qui, rappelons-le, met en cause tous les moments d'une distribution, -- car le principal problème que présentent les distributions de rendements est leur caractère leptocurtique, ou leur excédent de *kurtosis* si l'on veut, problème qui concerne le quatrième moment d'une distribution. On rappelle que le quatrième moment d'une normale est de 3. Or, les distributions de rendements présentent généralement un coefficient plus élevé que 3, d'où leur caractère leptocurtique ou excédent de *kurtosis*. En

prenant en compte l'excédent de kurtosis (EKUR), l'expansion de Cornish-Fisher devient, en négligeant les termes peu significatifs²⁸:

$$w_\alpha \cong z_\alpha + \frac{1}{6}(z_\alpha^2 - 1)AS + \frac{1}{24}(z_\alpha^3 - 3z_\alpha)EKUR - \frac{1}{36}(2z_\alpha^3 - 5z_\alpha)AS^2 \quad (8)$$

où EKUR représente l'excès de kurtosis. L'expansion de Cornish-Fisher se complique donc sensiblement lorsque l'on introduit le quatrième moment d'une distribution. A titre d'exemple, reprenons le cas précédent où, cette fois-ci, $AS = 0$ mais $EKUR = 4$, c'est-à-dire que le coefficient d'aplatissement des rendements est de 7 plutôt que de 3 comme c'est le cas pour la loi normale. L'expansion de Cornish-Fisher devient dans ce cas:

$$w_\alpha = z_\alpha + \frac{1}{24}(z_\alpha^3 - 3z_\alpha)EKUR = -1,655 + \frac{1}{24}[(-1,655)^3 - 3(-1,655)]4 = -1,5830$$

La VaR par dollar, pour un alpha de 5%, est donc de:

$$0,15 - 1,583(0,30) = -0,3249$$

L'excès de kurtosis a donc réduit la VaR pour un alpha de 5%. L'excès de kurtosis n'a cependant pas un effet à sens unique sur la Var. Son effet dépend de la marge d'erreur que l'on recherche²⁹. Calculons en effet la VaR pour un alpha de 1%. z_α est alors égal à -2,33 et w_α , le multiple corrigé pour l'excédent de kurtosis, est égal à -3,27 en vertu de l'équation (8). La VaR est égale à -0,83\$, chiffre sensiblement supérieur à la situation sans excès de kurtosis, soit -0,3435. Du fait que l'on vise habituellement des marges d'erreur très faibles lorsque l'on calcule la VaR, de l'ordre de 1%, on se rend compte que l'excédent de kurtosis au chapitre de la distribution des rendements peut augmenter sensiblement le risque d'un portefeuille.

Au tableau 24 apparaissent le degré d'asymétrie, l'excédent de kurtosis et la statistique w_α calculée à partir de l'équation (8) pour un seuil de 1%, cela pour les trois titres étudiés dans ce chapitre ainsi que pour l'indice TSE300 et l'indice des titres biotechnologiques. Les coefficients d'asymétrie ainsi que l'excès de kurtosis ont été obtenus à partir du logiciel EViews pour la période s'étirant du début de 1999 au début de 2001. Comme on peut le constater au tableau 24, l'excès de kurtosis domine nettement l'asymétrie dans le calcul de l'expansion de Cornish-Fisher. En utilisant -2,33 comme multiple dans l'équation (2), multiple représenté par θ , on sous-estime donc beaucoup le risque de ces cinq titres ou indices. C'est bien souvent le double de ce multiple, voire davantage, qu'il faut utiliser pour

²⁸ On retrouvera l'expansion de Cornish-Fisher jusqu'au sixième moment dans: Stuart A., Ord K. et Arnold S., Kendall's Advanced Theory of Statistics. Volume 1: Distribution Theory, Arnold, 1999, pp. 238.

²⁹ En fait le point mort, soit le point où l'excédent de kurtosis commence à jouer négativement sur la VaR, se situe ici à $z_\alpha = -1,73$.

calculer la VaR, toujours à partir de l'équation (2). Le titre AXP a même un multiple de -6,49. Bien que son asymétrie positive ait tendance à diminuer son multiple, son excès de kurtosis qui est très important joue fort défavorablement au chapitre de son w_α .

Tableau 24
Statistiques w_α au seuil de 1%
Titres sous étude
1999 - 2001

	AS	EKUR	w_α
TH	1,3	8,5	-5,85
GBL	0,2	3,7	-4,13
AXP	1,5	10,0	-6,49
TSE300	-0,6	3,7	-4,72
Indice Bio.	-0,3	2,9	-4,08

Finalement, le tableau 25 fournit la VaR journalière pour les titres ou indices du tableau 24 calculée à partir de l'équation (2), les multiples respectifs étant ceux qui apparaissent au tableau 24. Les indices ont une VaR plus faible, comme il se doit, car l'écart type de leur rendement est inférieur à celui des titres individuels du fait du caractère diversifié des indices. Selon le tableau 25, c'est le TSE300 qui présenterait le moins de risque, comme il se doit. Les moyennes et écarts types respectifs des rendements ont été calculés à partir du logiciel Eviews.

Tableau 25
VaR journalière au seuil de 1% ajustée par la méthode Cornish-Fisher
Titres sous-étude

	μ	σ	VaR
TH	0,002	0,0472	-0,2741
GBL	0,000	0,0458	-0,1892
AXP	0,000	0,0349	-0,2265
TSE300	0,000	0,0139	-0,0656
Indice Bio.	0,000	0,0221	-0,0902

7. Méthodes du calcul de la VaR utilisant une distribution autre que la loi normale mais qui restent basées sur l'emploi d'un multiple

L'approche au calcul de la VaR basée sur l'expansion de Cornish-Fisher vise à modifier le multiple associé à la loi normale de manière à intégrer les troisième et quatrième moments de la distribution des rendements. Disons que le multiple modifié au seuil de $\alpha\%$ soit égal à $\theta_{cf,\alpha}$. Si l'investisseur détient une position en compte dans un titre (« long position »), on peut alors calculer la VaR du titre à partir de la limite à gauche de l'intervalle de confiance du rendement de ce titre, c'est-à-dire :

$$\mu_r + \theta_{cf,\alpha} \sigma_r$$

avec μ_r et σ_r , respectivement l'espérance et l'écart-type (volatilité) du rendement. En multipliant ce résultat par la valeur du portefeuille, on obtient la VaR en dollars.

Plutôt que de calculer le multiple associé à l'expansion de Cornish-Fisher, plusieurs institutions financières gonflent le multiple associé à la loi normale de manière à prendre en compte les troisième et quatrième moments de la distribution des rendements. Disons qu'elles veulent calculer le multiple associé à un α de 5%. La loi normale associe un multiple de -1,65 à ce seuil. Pour intégrer le leptokurtisme de la distribution des rendements à l'intérieur de leurs calculs de la VaR, plusieurs institutions financières gonflent ce multiple à -2, voire à -3, de manière à en arriver à un calcul plus conservateur. Il va sans dire qu'une telle procédure est peu orthodoxe.

L'expansion de Cornish-Fisher est valable lorsque la distribution du rendement d'un titre ne s'éloigne pas trop de la loi normale. Mais comment procéder si cette distribution dévie sensiblement de la normale et que l'on veuille conserver quand même l'approche par le multiple. Et bien, il suffit d'adopter la méthode suivante : i) prévoir la volatilité des rendements en recourant à une distribution autre que la normale et qui semble bien représenter la distribution du rendement du titre étudié; ii) se servir de la même distribution pour calculer le multiple.

Précisons davantage cette procédure. Situons-nous dans le cadre de la distribution de Student. Le coefficient de leptokurtisme de cette distribution peut être supérieur à 3, soit le coefficient associé à la loi normale, en autant que le nombre de degrés de liberté (ν) de ladite distribution soit peu élevé mais supérieur à 4, car il n'est pas défini en deçà de ce nombre de degré de liberté. Supposons que nous voulons estimer la VaR journalière d'un titre en recourant à cette distribution. Nous estimons la volatilité du titre en recourant à un processus GARCH(1,1) et

en utilisant pour ce faire une distribution de Student avec ν degrés de liberté. Le modèle comporte une équation pour le rendement du titre, donnée par :

$$r_t = c + \mu_t$$

où $\mu_t = \varepsilon_t \sqrt{h_t}$, ε_t obtempérant à une distribution de Student standardisée, et où $\text{var}(\mu_t) = h_t$. Le processus GARCH(1,1) auquel obéit la variance conditionnelle, processus bien connu, s'écrit :

$$h_t = \beta_0 + \beta_1 h_{t-1} + \beta_2 \mu_{t-1}^2$$

On estime ces deux équations sur un échantillon journalier qui se termine le jour auquel la prévision est effectuée. La loi de Student avec ν degrés de liberté est utilisée pour estimer l'équation GARCH(1,1). La prévision de la variance qui suit le jour de la dernière observation est égal à :

$$\sigma_{t+1}^2 = E(h_{t+1}) = \beta_0 + \beta_1 h_t + \beta_2 \mu_t^2$$

La racine carrée de cette variance est l'écart-type qui est introduit dans l'équation du multiple pour calculer la VaR, ce qui montre que la VaR est de nature prévisionnelle.

Une fois la prévision de la volatilité effectuée, il est relativement simple de calculer la VaR à partir de la distribution utilisée pour estimer le modèle GARCH. Supposons que pour estimer le modèle GARCH, nous ayons eu recours à la loi t de Student standardisée comportant ν degrés de liberté. L'équation du multiple devient alors :

$$\mu_{r,t+1} + st_{\alpha,\nu} \sigma_{r,t+1}$$

avec $t_{\alpha,\nu}$, le multiple associé à la loi de Student standardisée pour un seuil de $\alpha\%$ et un nombre de degrés de liberté égal à ν . A titre d'exemple, si l'on fixe ν à 5, ce qui se traduit par une distribution de Student plutôt leptokurtique, $st_{\alpha,\nu}$ est alors égal au ratio du quantile de la distribution de Student (non standardisée) et de $\sqrt{\frac{\nu}{\nu-2}}$. Si l'on fixe ν à 5 et α à 1%, on obtient

le multiple suivant :

$$\frac{-3,3649}{\sqrt{\frac{5}{3}}} = -2,6064$$

en regard du multiple de la loi normale qui est égal, pour un α de 1%, à -2,3263. Pour des α faibles, les multiples de la Student sont donc plus importants que ceux de la loi normale, ce qui se traduit par des VaR plus élevées du côté de la Student.

Certes, il existe bien d'autres façons de prévoir la volatilité. Giot et Laurent³⁰ (2003) suggèrent d'utiliser le modèle APARCH³¹(1,1) plutôt que le modèle GARCH(1,1) pour effectuer cette prévision. Le premier modèle donnerait en effet des résultats supérieurs au second. Le modèle APARCH(1,1) est issu d'une transformation Box-Cox et s'écrit :

$$\sigma_t^\delta = \omega + \alpha_1 \left[|\mu_{t-1}| - \alpha_n \mu_{t-1} \right]^\delta + \beta_1 \sigma_{t-1}^\delta$$

Le paramètre δ représente la transformation Box-Cox de la volatilité conditionnelle. Par ailleurs, on sait que la volatilité des rendements est sujette à un effet de levier. En effet, les chocs négatifs exerceraient un effet plus important sur la volatilité conditionnelle que les chocs positifs. Cet effet sera observé si le coefficient α_n est positif. En utilisant ce modèle de volatilité, la VaR est donc donnée par :

$$\mu_{r,t+1} + st_{\alpha,v} \sigma_{r,t+1}$$

si l'on utilise encore une fois la loi de Student pour caractériser la distribution des rendements. Giot et Laurent proposent même d'utiliser la loi de Student asymétrique qui donnerait de meilleurs résultats.

On peut également recourir à la volatilité réalisée pour calculer la volatilité journalière. Ce concept a gagné en popularité ces dernières années et fait appel aux données intra-journalières. Disons que pour une journée donnée, nous disposons de rendements d'un titre mesurés toutes les cinq minutes. La volatilité réalisée se définit alors comme suit :

$$\sigma_i^2 = \sum_{j=1}^n r_{i,j}^2$$

avec σ_i^2 , la variance réalisée pour la journée i et n , le nombre d'intervalles de cinq minutes dans une journée de transactions boursières. On parle presque de volatilité observée dans ce cas même si en principe la volatilité est un concept latent.

8. Mesures du risque : une généralisation

La VaR est une mesure du risque qui fait appel à la probabilité cumulative de la distribution des rendements. Or, c'est en se basant sur cette approche que les chercheurs ont pu récemment découvrir des mesures du risque plus satisfaisantes que le simple écart-type des rendements ou le bêta d'un titre. La probabilité cumulative est en effet susceptible de capter les moments supérieurs à deux d'une distribution. Or, ces moments représentent des dimensions additionnelles du risque.

³⁰ Giot, P. et Laurent, S. (2003), Value-at-Risk for long and short trading positions, *Journal of Applied Econometrics*, 18, 641-664.

³¹ Soit le Asymetric-Power ARCH

Pour le montrer, désignons l'utilité du rendement du portefeuille par $U(R_p)$. Cette fonction peut être approximée par une série de Taylor, soit :

$$dU \approx \frac{\partial U}{\partial R_p} dR_p + \frac{1}{2!} \frac{\partial^2 U}{\partial R_p^2} dR_p^2 + \frac{1}{3!} \frac{\partial^3 U}{\partial R_p^3} dR_p^3 + \dots = \sum_{n=1}^{\infty} \frac{1}{n!} \frac{\partial^n U}{\partial R_p^n} \partial R_p^n$$

Cette fonction peut être écrite autour de la moyenne $\bar{R}_p$. On a :

$$U(R_p) \approx U(\bar{R}_p) + \sum_{n=1}^{\infty} \frac{1}{n!} \frac{\partial^n U(\bar{R}_p)}{\partial R_p^n} (R_p - \bar{R}_p)^n$$

En prenant l'espérance des deux côtés et en réarrangeant, on obtient :

$$\begin{aligned} E(U(R_p)) \approx & U(\bar{R}_p) + \frac{1}{2} U'' E[\underbrace{(R_p - \bar{R}_p)^2}_{\text{variance}}] + \frac{1}{3!} U''' E[\underbrace{(R_p - \bar{R}_p)^3}_{\text{skewness}}] + \frac{1}{4!} U'''' E[\underbrace{(R_p - \bar{R}_p)^4}_{\text{kurtosis}}] + \\ & \sum_{n=5}^{\infty} \frac{1}{n!} \frac{\partial^n U}{\partial R_p^n} \underbrace{\left[E(R_p - \bar{R}_p)^n \right]}_{\text{moments supérieurs}} \end{aligned} \quad (9)$$

On peut réécrire cette fonction d'utilité de manière équivalente de la façon suivante. Soit w la richesse initiale d'un investisseur et $\tilde{x}$, un accroissement aléatoire de la richesse. Sa fonction d'utilité s'écrit : $U = U(\tilde{x} + w)$. Le retour sur l'investissement est alors de : $\tilde{r} = \frac{\tilde{x}}{w}$. En substituant cette valeur dans U , on obtient : $U = U(rw + w)$. En définissant : $\mu = E(w + rw)$ et en faisant une expansion de Taylor de la fonction d'utilité, on a :

$$E(U) \approx U(\mu) + \frac{U^2(\mu)}{2} \sigma^2 + \sum_{i=3}^{\infty} \frac{\mu_i}{i!} U^i(\mu)$$

avec μ_i : ième moment central et $U^i = \frac{\partial^i U}{\partial \mu^i}$.

Les moments supérieurs à 1 représentent les diverses dimensions du risque. On admet généralement que les dérivées impaires de la fonction d'utilité sont positives et les dérivées paires, négatives. Ainsi, un investisseur préfère-t-il une distribution à asymétrie positive à une autre à asymétrie négative. Il éprouve par ailleurs une aversion pour une distribution à leptokurtisme positif.

Comme nous le disions, la probabilité cumulative est susceptible de capter les moments d'ordre supérieurs à deux de la distribution des rendements. C'est dans ce sens que la VaR représente une amélioration sur les mesures classiques. L'approche par la probabilité cumulative permet également de récupérer l'apport de la théorie de la dominance stochastique.

La théorie de la dominance stochastique est basée sur l'expansion de Taylor de la fonction d'utilité espérée. La règle de la dominance stochastique du premier degré suppose tout simplement que la fonction d'utilité est croissante dans la richesse, c'est-à-dire $U' \geq 0$. Cela revient à dire que l'investisseur préfère une espérance de rendement (premier moment de la distribution) supérieure à une espérance de rendement inférieure. Selon la règle de la dominance stochastique du premier degré, l'actif x est préféré à l'actif y si x génère davantage de richesse dans tous les états de la nature. En termes de la fonction de probabilité cumulative de la richesse, cela s'écrit :

$$F_x(W) \leq G_y(W) \quad \forall W$$

avec W , la richesse, et F et G , respectivement les fonctions de probabilité cumulative de F et de G . Une inégalité stricte doit tenir pour au moins un niveau de W . Cette condition revient à dire que la fonction de probabilité cumulative de x se situe à droite de celle de y .

La règle de la dominance stochastique du second degré suppose que la fonction d'utilité est concave, c'est-à-dire que $U'' \leq 0$. Selon la fonction d'utilité espérée, elle concerne donc la variance de la distribution. Un investisseur préfère donc un placement à variance faible à un placement à variance élevée, toutes choses égales d'ailleurs. Ce principe définit le degré d'aversion au risque d'un investisseur. Selon la règle de la dominance stochastique du second degré, l'investissement x sera préféré à l'investissement y si :

$$\int_{-\infty}^{W_i} F_x(W) \leq \int_{-\infty}^{W_i} G_y(W) \quad \forall W_i$$

Cette inégalité signifie que la surface cumulative sous $G(\cdot)$ doit être supérieure à celle sous $F(\cdot)$, ceci quel que soit le niveau de W et avec une inégalité stricte pour au moins une valeur de W . Par rapport à la dominance stochastique du premier degré, les fonctions de probabilité cumulative peuvent se croiser dans ce cas mais il reste que la règle de surface doit être respectée quel que soit le niveau de W .

La règle de la dominance stochastique du troisième degré concerne la dérivée troisième de la fonction d'utilité. On postule que $U''' \geq 0$, c'est-à-dire que le degré d'aversion absolue au risque décroît avec la richesse. Selon la fonction d'utilité espérée, cette règle se rapporte au troisième moment de la distribution, une distribution à asymétrie positive étant préférée à une autre à asymétrie négative. La règle de la dominance stochastique du troisième degré fait appel à l'intégrale double. Définissons :

$$F_1(W_i) = \int_{-\infty}^{W_i} F(t) dt$$

et,

$$G_1(W_i) = \int_{-\infty}^{W_i} G(t) dt$$

et de la même manière,

$$F_2(W_i) = \int_{-\infty}^{W_i} F_1(t) dt$$

$$G_2(W_i) = \int_{-\infty}^{W_i} G_1(t) dt$$

En vertu de la règle de la dominance stochastique du troisième degré, étant donné les deux options F et G, F est préféré à G si et seulement si :

$$F_2 \leq G_2 \quad \forall W_i$$

et avec une inégalité stricte pour au moins un niveau de W. Pour définir des règles de dominance stochastique d'ordre supérieur, il suffit d'intégrer davantage la probabilité cumulative. La règle de dominance stochastique du n^{ième} degré s'énonce comme suit. La fonction de probabilité f(x) domine g(x) selon cette règle si : $F_{n-1}(x) \leq G_{n-1}(x)$, avec une inégalité stricte pour au moins une valeur de x. La règle de dominance stochastique du quatrième degré concerne le quatrième moment d'une distribution, soit le leptokurtisme. Règle générale, les dérivées d'ordre impair de la fonction d'utilité sont positives et les dérivées d'ordre pair sont négatives. Le leptokurtisme se rapportant à la dérivée quatrième de la fonction d'utilité, les investisseurs préfèrent une situation dans laquelle le leptokurtisme est plus faible à une autre dans laquelle il est plus important.

Avant d'examiner davantage l'approche par la probabilité cumulative, qui a permis de définir, entre autres, le ratio oméga, nous rappelons deux mesures du risque de seconde génération, soit la semi-variance et l'indice de Sortino.

La variance ne permet pas de déterminer si les écarts en termes de la moyenne se sont produits en-dessous ou au-dessus de la moyenne puisque ceux-ci sont exprimés au carré. La semi-variance vient remédier à cette lacune en ne retenant que les rendements sous la moyenne. Elle se définit comme suit :

$$\frac{1}{T} \sum_{\substack{0 \leq t \leq T \\ R_{Pt} < \bar{R}_P}} (R_{Pt} - \bar{R}_P)^2$$

Markowitz avait envisagé cette mesure au début des années 50 mais il ne l'afficha pas comme mesure du risque. D'abord parce qu'à l'époque le temps de calcul de cette statistique était beaucoup plus important que celui de la variance. Ensuite, parce que cette statistique se prêtait mal au développement d'une théorie de la diversification du portefeuille alors que la variance autorisait une approche beaucoup plus élégante. Et finalement, comme Markowitz prenait pour acquis que la

distribution des rendements obéissait à une loi normale, la variance équivalait dans ce contexte à la semi-variance étant donné la symétrie de la distribution normale.

La mesure de Sortino est une mesure qui fait appel à une variante du concept de semi-variance. Au lieu de ne retenir que les rendements situés sous la moyenne, elle ne retient que ceux qui se situent sous un seuil jugé acceptable, désigné par MAR (rentabilité minimale acceptable). L'indice de Sortino se définit comme suit :

$$\frac{E(R_P) - MAR}{\sqrt{\frac{1}{T} \sum_{\substack{t=0 \\ R_{pt} < MAR}}^T (R_{pt} - MAR)^2}}$$

C'est donc une mesure du rendement excédentaire par unité de risque. Ce ratio est donc très apparenté à ceux de Sharpe et de Jensen, sauf que les mesures de rendement excédentaire et de risque diffèrent d'un ratio à l'autre.

Les deux ratios précédents, soit la semi-variance et le ratio de Sortino, représentent des améliorations en regard de la variance comme mesures du risque mais il reste que ce sont de simples mesures de dispersion. Elles ne peuvent mettre en lumière les non-linéarités qui sont inhérentes aux moments d'ordre supérieur à 2 et qui influencent la fonction d'utilité des investisseurs.

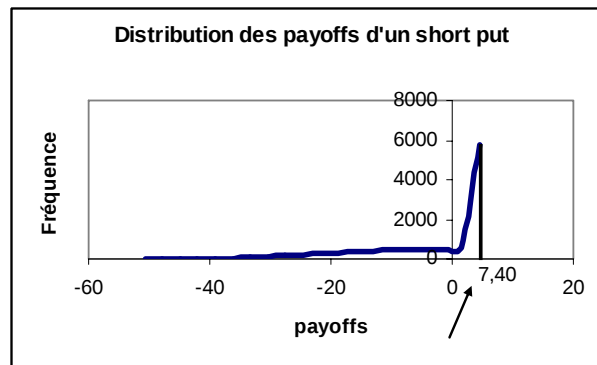
Les moments d'ordre supérieur à 2 sont susceptibles de conférer aux payoffs d'un portefeuille une structure semblable à ceux d'une position à découvert dans un put. C'est d'ailleurs dans cette perspective que Merton avait originellement envisagé le risque de crédit : une obligation risquée équivaut à une obligation sans risque à laquelle s'ajoute une position à découvert dans un put défini sur la valeur nominale de la dette de l'entreprise émettrice. Un événement rare défavorable à un investisseur peut donc être analysé par une position à découvert dans un put. Or, les événements rares relèvent du quatrième moment d'une distribution. Il existe donc une parenté étroite, sur le plan analytique, entre le quatrième moment d'une distribution et une position à découvert dans un put.

Nous avons représenté à la figure 12 les payoffs, à l'échéance, d'un put européen à découvert aux caractéristiques suivantes : i) prix de l'action sous-jacente : 100 \$; ii) prix d'exercice : 95\$; iii) taux sans risque : 0% ; iv) volatilité du sous-jacent : 0,5 ; v) durée de l'option : 0,25 an. Nous avons simulé 10000 payoffs de cette position à découvert en représentant l'évolution du prix de l'action comme un mouvement brownien géométrique, soit le suivant :

$$dS = 0,5 \times S \times dz$$

La dérive («drift») est ici nulle puisque l'on suppose le taux d'intérêt nul. Les payoffs incluent la prime de l'option, soit le prix du put, qui est égal à 7,40\$. La moyenne des payoffs est très près de 0, comme il se doit³², avec une médiane et un mode de 7,40\$. Comme le révèle la figure 12, la distribution des payoffs comporte un fort degré d'asymétrie négative. En fait, le coefficient d'asymétrie de cette distribution s'établit à -1,44. Le leptokurtisme, à hauteur de 4,24, est lui-même important. Décidément, l'investissement dans un «short put» se révèle une opération pour le moins risquée, et les troisième et quatrième moments sont les véritables moteurs de son risque.

Figure 12



C'est dans cet esprit qu'est défini l'indicateur de risque prénommé «oméga», qui constitue une synthèse entre la théorie de la dominance stochastique, qui fait appel à la notion de probabilité cumulative, et la théorie des options. Cet indicateur est donc très avant-gardiste. La mesure dite «oméga», symbolisée par $\Omega(L)$, se définit comme suit :

$$\Omega(L) = \frac{\int_a^b [1 - F(x)] dx}{\int_a^b F(x) dx} \quad (10)$$

Dans cette équation, x représente le rendement aléatoire sur l'investissement d'une période ; L est le seuil de rendement sélectionné par l'investisseur ; (a,b) représente les bornes inférieure et supérieure du rendement ; $F(y) = P\{x \leq y\}$, soit la probabilité cumulative du rendement d'une période.

³² En effet, comme les payoffs incluent ici la prime, leur moyenne simulée doit être nulle. Sans la prime, la moyenne des payoffs simulés doit être la valeur négative de la prime, ou du prix du short put, puisque le prix du put est la valeur espérée, dans un univers risque-neutre, des payoffs finaux. En fait, la moyenne calculée des payoffs sur les 10000 simulations est de -7,44\$, ce qui est très rapproché de la valeur négative du prix du put (-7,40).

Le numérateur de l'équation (10) est une mesure probabiliste du rendement excédentaire de l'investissement alors que le dénominateur est une mesure probabiliste de son risque. Pour mieux comprendre l'indicateur oméga, on peut le réécrire comme suit :

$$\Omega(L) = \frac{C(L)}{P(L)}$$

avec $C(L)$, le prix d'un call européen sur l'investissement et $P(L)$, le prix d'un put européen sur l'investissement. L'échéance des options est d'une période (1 mois) et le prix d'exercice des options est de L . Le numérateur de l'indicateur oméga est égal à :

$$\int_L^b [1 - F(x)] dx = \int_L^b (x - L) f(x) dx = E(x - L)^+$$

soit l'espérance du payoff d'un call définie sous la mesure objective et non sous la mesure risque-neutre. N'oublions ici que nous sommes en gestion des risques et non en «pricing». Dans la théorie de la valorisation des options, c'est la mesure risque-neutre qui tient le haut du pavé. C'est une mesure artificielle qui facilite le calcul du prix d'une option. Mais en matière de gestion des risques, on doit s'en tenir à la mesure objective. Cela deviendra clair quand on analysera le risque de crédit.

Il suit :

$$C(L) = e^{-r} E(x - L)^+$$

où r est le taux d'intérêt instantané sous la mesure objective. Par ailleurs, le dénominateur de l'indicateur oméga est égal à :

$$\int_a^L F(x) dx = \int_a^L (L - x) f(x) dx = E(L - x)^+$$

Il suit :

$$P(L) = e^{-r} E(L - x)^+$$

Le numérateur de l'oméga définit donc le rendement excédentaire de façon probabiliste, soit comme la probabilité cumulative d'excéder le rendement minimal fixé par l'investisseur. C'est le coût d'une option d'achat qui assure que le rendement se situera au-dessus de L . Par contre, le risque qui apparaît au dénominateur de l'expression est mesuré par le prix du put qui assure que le rendement ne se situera pas en deçà du rendement minimal recherché. Le risque est donc défini en termes d'une police d'assurance.

L'une des variantes des mesures oméga est celui de Sharpe. L'oméga de Sharpe se définit comme suit :

$$\frac{\bar{x} - L}{P(L)}$$

avec $\bar{x}$ le rendement espéré de l'investissement. Le numérateur représente le rendement excédentaire comme un simple écart arithmétique entre le rendement attendu et le rendement minimal tandis que la formule générale définit le rendement excédentaire sur une base probabiliste ou prospective. La formule générale de l'oméga définit donc le rendement excédentaire de façon plus satisfaisante.

On peut calculer $P(L)$ comme suit :

$$P(L) = e^{L-r_f} N(-d_2) - e^{\bar{x}-r_f} N(-d_1)$$

avec $d_1 = \frac{\bar{x} - L + 0,5\sigma^2}{\sigma}$ et $d_2 = d_1 - \sigma$. Dans la formule de B&S, on a remplacé le taux sans risque r_f par le rendement moyen sur l'investissement.

L'oméga de Sharpe constitue cependant une amélioration en regard du ratio de Sharpe S_p qui est défini comme ceci :

$$S_p = \frac{\bar{x} - r_f}{\sigma_p}$$

avec r_f le taux sans risque et σ_p l'écart-type du rendement du portefeuille. L'oméga de Sharpe remplace donc le taux sans risque par le rendement minimal recherché, ce qui semble plus approprié. De même substitue-t-il à l'écart-type du rendement le coût d'un put protecteur. Cette dernière mesure du risque est plus satisfaisante car elle prend en compte la forme de la distribution des rendements. Comme nous l'avons dit antérieurement, la distribution des payoffs d'un put incorpore les troisième et quatrième moments de la distribution des rendements.

Les mesures du risque que nous venons d'analyser relèvent des nouvelles théories dites du «risk of shortfall», du «downside risk» ou plus communément du risque baissier. Le risque n'est donc plus un concept symétrique comme c'était le cas pour l'écart-type. Ces théories sont très préoccupées par les moments d'ordre supérieur à 2, et notamment par le quatrième moment qui représente le risque associé à des événements extrêmes. La VaR relève de la première génération de ces théories. Elle mesure en effet le risque par la probabilité cumulative. Mais elle comporte plusieurs défauts. D'abord, elle n'est pas une mesure cohérente du risque. La VaR d'un portefeuille de titres n'est pas, dans certains cas, plus faible que la somme des VaR de ses composantes. Elle escamote alors le principe de la diversification, un principe de base en finance. Qui plus est, elle ne tient pas compte de la forme de la distribution des rendements et dès lors ne renseigne guère sur les pertes extrêmes que peut subir un portefeuille.

C'est dans ce contexte qu'a été développée la VaR conditionnelle, soit la CVaR. On peut la définir comme suit :

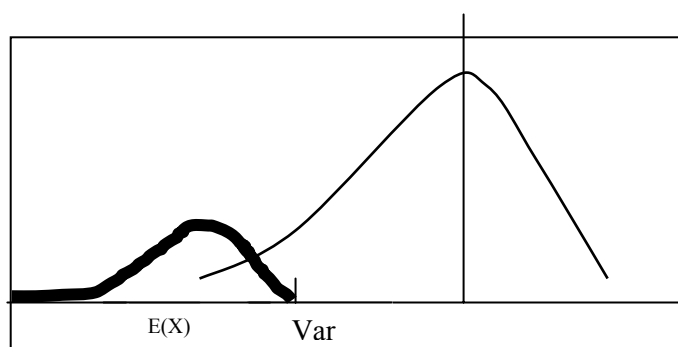
$$CVaR = E[X | X \leq VaR]$$

avec X, une variable aléatoire qui tient lieu des pertes enregistrées par le portefeuille. Cette mesure est représentée à la figure 13.

Figure 13

La CVaR

Distribution des profits et pertes d'un portefeuille



Comme l'indique la figure 13, la CVaR est calculée comme l'espérance des pertes à gauche de la VaR. Elle est calculée à partir de la distribution des réalisations comprises dans la queue comme l'indique la figure. Elle tient donc compte de l'asymétrie de cette région. Cette distribution est également influencée par l'épaisseur de la queue, qui est liée à l'intensité des événements extrêmes.

Les développements récents en matière d'analyse du risque ont donc fait litière de la distribution normale et de la fonction d'utilité quadratique qui conditionnaient auparavant le choix des indicateurs de risque. En effet, on peut se limiter à l'analyse des deux premiers moments d'une distribution que si l'une ou l'autre des deux conditions suivantes est valable. D'abord, si la distribution des rendements est normale. Cette distribution est alors complètement caractérisée par ses deux premiers moments. Ensuite, si la fonction d'utilité des investisseurs est quadratique. L'investisseur ne considère alors comme mesure de risque que le deuxième moment et non les moments d'ordre supérieur. Mais, dans les faits, aucune de ces deux conditions n'est valide. Les théories de la gestion du risque ont donc évolué, comme on l'a vu, vers des mesures qui intègrent les moments d'ordre supérieur à 2.

Ces nouvelles théories ont permis, entre autres, une analyse plus étoffée des activités des fonds spéculatifs (*hedged funds*) dont les payoffs s'apparentent à ceux des options. Elles ont donné lieu au développement d'un CAPM à quatre facteurs qui intègre les moments d'ordre supérieur à deux. Rubinstein (1973) a donné la forme suivante au CAPM à trois moments :

$$E(R_i) - R_F = \beta_i \left[\frac{(E(R_i) - R_F)(E(R_m) - R_F)}{(E(R_m) - R_F)^2} \right] + \left[\frac{(E(R_i) - R_F)(E(R_m) - R_F)^2}{(E(R_m) - R_F)^3} \right]$$

Le premier terme à droite représente la prime de risque pour la volatilité du marché et le second, la prime pour la co-asymétrie du marché. Ce nouveau terme est introduit pour indiquer que le marché ne rémunère pas l'asymétrie des rendements comme telle, de la même façon qu'il ne rémunère pas toute la volatilité d'un titre. Il ne rémunère que l'asymétrie du rendement d'un titre en relation avec le rendement du portefeuille du marché, ce qu'indique le second terme de l'équation. Pour des fins d'estimation, cette équation peut être simplifiée et intégrer d'autres moments comme le quatrième. Voici une forme cubique du CAPM qui incorpore le quatrième moment :

$$R_i - R_F = \alpha_i + \beta_{i1}[R_m - R_F] + \beta_{i2}[R_m - R_F]^2 + \beta_{i3}[R_m - R_F]^3 + \varepsilon_i$$

Dans cette équation, le carré de la prime de risque prend en compte de troisième moment et le cube, le quatrième.

La réforme de l'analyse du risque met à mal les modèles qui établissent une relation linéaire entre le rendement espéré et les facteurs de risque, tels le CAPM et l'APT, l'introduction des moments d'ordre supérieur à 2 dans l'analyse ayant pour conséquence d'introduire des non-linéarités importantes dans l'analyse rendement-risque. Un modèle qui a gagné en popularité au cours de la décennie 90, soit celui de Fama et French (1992), se voit lui aussi remis en question. Fama et French croient avoir identifié aux moins deux autres facteurs, en sus de la prime de risque du marché, qui expliquent le rendement espéré d'un portefeuille. D'abord, les actions émises par les entreprises à faible capitalisation semblent rapporter davantage que celles des entreprises à forte capitalisation. Ensuite, plus le ratio de la valeur aux livres à la valeur marchande d'une action est important, plus son rendement est susceptible de l'être également, toutes choses égales d'ailleurs.

L'équation de Fama et French s'écrit comme suit :

$$R_i - R_f = \alpha_i + \beta_{i1}(R_m - R_f) + \beta_{i2}SMB + \beta_{i3}HML + \varepsilon_i$$

Dans cette équation, R_i représente le rendement d'un titre ou d'un portefeuille et $(R_m - R_f)$, la prime de risque du marché. La variable SMB représente le rendement d'un portefeuille qui est en compte («long») dans les actions de firmes à faible capitalisation et à découvert («short») dans les actions de firmes à forte capitalisation. La variable HML, qui concerne le ratio : (valeur aux livres/valeur marchande), est construite de la même façon.

Or, le modèle de Fama et French fait maintenant l'objet de plus en plus de critiques. En effet, lorsqu'on introduit les troisième et quatrième moments dans l'équation du rendement d'un titre, l'influence des facteurs mis de l'avant par Fama et French s'estompe. Ces facteurs seraient donc des succédanés aux troisième et quatrième moments de la distribution des rendements. On peut donc désormais identifier les facteurs de Fama et French à des facteurs de risque, et non à des sources d'inefficiences de marché qui donneraient lieu à des opérations d'arbitrage couronnées d'un gain sans risque et sans mise de fonds additionnelle.

Les auteurs ne sont guère allés au-delà du quatrième moment de la distribution des rendements dans leur analyse du risque. Il est en effet difficile de trouver une signification a priori à ces moments. Mais un avant-gardiste, Taleb (1996), intègre dans son analyse du risque les moments d'ordre supérieur à 4. D'emblée, il classe les moments pairs d'une distribution comme des facteurs de convexité (concavité), et les moments impairs, comme des facteurs d'asymétrie. Selon Taleb, le cinquième moment mesure la «sensibilité asymétrique» du quatrième moment. Par exemple, Taleb envisage le cas d'une option barrière américaine qui est couverte par une option classique. Un tel portefeuille est fort sensible au 5^{ème} moment. Sa concavité en regard du marché augmente en effet beaucoup lorsqu'on se rapproche de la barrière. Par ailleurs, sa relation au marché devient plus linéaire lorsqu'on s'en éloigne. Un autre moment qui s'avère important est le 7^{ème}. Il correspond au signe du changement de la convexité quand le sous-jacent augmente ou diminue. Il va sans dire que les moments d'ordre supérieur à 4 sont surtout importants pour l'analyse des portefeuilles d'options exotiques qui incorporent de fortes non-linéarités.

9. Frontière efficiente, moments supérieurs et cumulants

La frontière efficiente est une relation qui associe à un rendement espéré d'un portefeuille de titres l'écart-type minimal pour ce rendement. Pour un rendement espéré donné, le problème consiste donc à trouver les pondérations optimales des titres du portefeuille qui minimisent la variance du rendement dudit portefeuille sous la contrainte d'un rendement espéré donné. De plus, la somme des coefficients de pondération des titres qui composent le portefeuille doit être égale à l'unité. Le problème classique de la construction de la frontière efficiente se présente donc comme suit :

$$\inf_{w_i \in [0,1]} \text{var}(\{w_i\})$$

$$\sum w_i E(R_i) = E^*$$

$$\sum w_i = 1$$

avec $\text{var}(\cdot)$, la variance du rendement du portefeuille; w_i , la pondération du titre i dans le portefeuille; $E(R_i)$, le rendement espéré du titre i et E^* , le niveau-cible du rendement espéré du portefeuille.

Compte tenu des nouvelles mesures du risque, cette approche apparaît insatisfaisante car elle ne prend en compte que les deux premiers moments de la distribution des rendements d'un titre. Elle néglige les moments d'ordre supérieur qui représentent d'autres dimensions du risque.

Comme le montrent Malevergne et Sornette (2005), on peut redéfinir le concept de frontière efficiente en recourant aux cumulants et prendre ainsi en compte les moments d'ordre supérieur. Les cumulants d'une distribution sont définis à partir de la fonction caractéristique d'une variable aléatoire, disons X . Cette fonction se définit comme suit :

$$\varphi_X(t) = E(e^{itX}), \quad -\infty < t < \infty$$

avec $i = \sqrt{-1}$. La fonction caractéristique est donc basée sur les nombres complexes. Selon James et Webber (2000), la fonction caractéristique représente la transformation de Fourier de ladite variable aléatoire.

Les cumulants d'ordre j sont obtenus en dérivant la fonction caractéristique par rapport à t et en évaluant cette dérivée au point $t = 0$. Ainsi le cumulant d'ordre j , désigné par κ_j , est égal à :

$$\kappa_j = \frac{1}{i^j} \left[\frac{\partial^j \varphi_X(t)}{\partial t^j} \right]_{t=0}$$

Pour illustrer l'utilisation de cette formule, supposons que la variable aléatoire X ait une distribution normale d'espérance nulle et de variance σ^2 . Sa fonction caractéristique est alors de³³ :

$$\varphi_X(t) = e^{-\frac{\sigma^2 t^2}{2}}$$

Le cumulant d'ordre 1 est donc égal à :

$$\frac{1}{i} \left[\frac{\partial}{\partial t} e^{-\frac{\sigma^2 t^2}{2}} \right]_{t=0} = \frac{1}{i} \left[-\sigma^2 t e^{-\frac{\sigma^2 t^2}{2}} \right]_{t=0} = 0$$

Le cumulant d'ordre 2 est pour sa part égal à :

$$\frac{1}{i^2} \left[\frac{\partial^2}{\partial t^2} e^{-\frac{\sigma^2 t^2}{2}} \right]_{t=0} = \frac{1}{i^2} \left[(-\sigma^2 + \sigma^4 t^2) e^{-\frac{\sigma^2 t^2}{2}} \right]_{t=0} = \sigma^2$$

On constate que dans ce cas, les deux premiers cumulants de cette distribution normale sont égaux aux deux premiers moments centrés. Pour mieux établir la relation entre moments et cumulants, définissons le moment centré d'ordre j comme :

$$\mu_j = E(X - \mu)^j$$

On peut alors écrire :

$$k_1 = \mu$$

$$k_2 = \mu_2$$

$$k_3 = \mu_3$$

$$k_4 = \mu_4 - 3\mu_2^2$$

On constate donc que les trois premiers cumulants d'une distribution sont identiques aux trois premiers moments centrés de cette distribution. Comme le soulignent Malevergne et Sornette (2005), le recours aux cumulants comme mesures du risque différentes des moments centrés devient intéressant au-delà du troisième ordre. Au dire de ces auteurs, les cumulants ont donc une

³³ Pour la dérivation de cette formule, voir : Hoel, P.G., Port, S.C. et Stone, C.J.(1971), Introduction to Probability Theory, Houghton Mifflin Company.

interprétation fort importante : ils mesurent le degré avec lequel la distribution d'un risque donné s'éloigne de la distribution normale. De plus, ils constituent des mesures cohérentes du risque.

Fort de ces développements, Malevergne et Sornette (2005) généralisent le problème d'optimisation relié à la construction de la frontière efficiente en modifiant la fonction objective comme suit :

$$\inf_{w_i \in [0,1]} \rho_n(\{w_i\})$$

ρ_n étant une mesure quelconque du risque. Par exemple, ce peut être un cumulant d'ordre n . Supposons que ρ_n soit le cumulant d'ordre 4 d'une distribution. La fonction objective du problème de la construction de la frontière efficiente devient alors :

$$\inf_{w_i \in [0,1]} \kappa_4 = \inf_{w_i \in [0,1]} (\mu_4 - 3\mu_2^2)$$

Un investisseur qui dispose de cette fonction objective manifeste une aversion pour les fluctuations qui sont représentées par le quatrième moment mais est attiré par les fluctuations mesurées par la variance. Selon Malevergne et Sornette (2005), un tel comportement n'est pas irrationnel car l'investisseur demeure globalement averse au risque. Comme les risques reliés au quatrième moment (leptokurtisme) sont beaucoup plus élevés que ceux qui sont associés au deuxième (variance)³⁴, l'investisseur essaie alors d'éviter les risques importants mais est prêt à assumer les risques moindres.

Une frontière efficiente basée sur la minimisation du quatrième cumulant, et par conséquent qui pénalise le leptokurtisme de la distribution des rendements, donnera probablement lieu à des portefeuilles optimaux fort différents d'une frontière classique qui se fonde strictement sur la minimisation de la variance des rendements. Supposons en effet que la distribution du rendement d'un titre présente un fort degré de leptokurtisme mais une variance relativement faible. Une frontière efficiente classique le surpondérera en regard d'une frontière qui minimise le quatrième cumulant. En adoptant cette dernière frontière, on pourrait augmenter le rendement espéré du

³⁴ En effet, le quatrième moment mesure les risques associés aux événements rares.

portefeuille tout en réduisant les sources importantes de risque qui sont emmagasinées dans le degré de leptokurtisme d'une distribution.

Malevergne et Sornette (2005) en sont donc arrivés à définir des portefeuilles « beaucoup plus efficaces » que ceux qui résultent de la frontière efficace classique. Leur approche à la construction de la frontière efficace s'avère fort prometteuse car en recourant aux cumulants pour définir le risque, on en arrive à prendre en compte plusieurs dimensions du risque. A l'instar des composantes principales, les cumulants s'avèrent une technique de réduction des dimensions d'un phénomène, soit le risque dans le cas de la construction de la frontière efficace.

10. Les copules (*copulas*)³⁵, la transformée de Fourier et le calcul de la VaR

Les copules sont une mesure alternative de la dépendance entre deux variables aléatoires, au même titre que la corrélation de Pearson sert à mesurer le lien linéaire statistique entre deux variables aléatoires.

Afin de pouvoir comprendre la provenance et l'utilité des copules, nous allons débiter cette section par quelques rappels statistiques pour ensuite les relier avec la théorie des copules. Enfin, nous terminerons la section par une application des copules à la VaR du crédit.

Une autre technique qui est depuis peu utilisée pour calculer la VaR est la transformée de Fourier. Elle a également d'autres applications en finance comme la résolution rapide d'un arbre binomial. C'est pourquoi nous fournissons une introduction des séries de Fourier dans cet article avec applications Excel et Matlab.

10.1. La corrélation

Considérons deux variables aléatoires X et Y . Ces deux variables ont pour distribution marginale, respectivement : $F_x(x) = P[X \leq x]$ et $F_y(y) = P[Y \leq y]$. Ces distributions donnent la probabilité

³⁵ Pour rédiger cette section, nous nous sommes fortement inspirés des documents suivants : Dowd, K. (2002, 2005), Measuring market risk, 1^e et 2^e ed., Wiley; Hull, J.C. (2006), Options, Futures and Other Derivatives, 6^e ed., Pearson Prentice Hall; Jäckel, P. (2002), Monte Carlo Methods in Finance, Wiley. Pour d'autres applications des copules en finance, on consultera : Alexander, C. (2001), Market Models, Wiley; Wilmott, P. (2006), The Best of Wilmott 2, Wiley.

cumulative de chaque variable. La distribution conjointe de la matrice de la variables aléatoire (X, Y) se représente comme suit: $F(x, y) = P[X \leq x, Y \leq y]$. Elle donne la probabilité que X soit inférieure ou égal à x et que simultanément Y n'est pas supérieur à y . Les distributions marginales considèrent donc chaque variable séparément et la distribution conjointe prend en compte la structure de dépendance qui existe entre elles.

La façon la plus simple de mesurer cette dépendance est d'utiliser la corrélation linéaire de Pearson³⁶ définie par : $\rho_{XY} = \frac{Cov(X, Y)}{\sqrt{V(X)V(Y)}}$. La question est : dans quelle circonstance cette mesure

est-elle fiable ? La corrélation de Pearson est une mauvaise mesure de dépendance linéaire lorsque la distribution multivariée de nos variables aléatoires est non elliptique, c'est-à-dire des distributions du type Lévy³⁷ avec queues épaisses ou dans le cas pour lequel les séries ont un trend (drift) et sont non cointégrées. Mais dans le cas où nous sommes en présence d'une distribution multivariée de variables aléatoires elliptiques, c'est-à-dire de distribution normale ou t de Student, la corrélation standard s'avère une bonne mesure de dépendance. La corrélation traduit tout ce qu'il faut savoir sur la dépendance entre les variables aléatoires.

Revenons au cas non elliptique. Le problème le plus sérieux dans ce cas est que les distributions marginales et les corrélations ne suffisent plus pour déterminer la distribution conjointe multivariée. La corrélation ne nous fournit pas toute l'information sur la dépendance entre les variables aléatoires, et encore moins dans les queues de la distribution. Nous avons donc besoin d'une mesure alternative de la dépendance entre nos variables aléatoires.

10.2. La théorie des copules³⁸

Le terme *Copula* provient du latin et signifie : joindre ensemble ou connecter. Ce terme est très rapproché du mot français ou anglais : couple. En ce qui nous concerne, nous sommes intéressés à sa signification économétrique. Dans ce contexte, une copule est une fonction qui relie une fonction de distribution multivariée à un ensemble de fonctions de distributions marginales. Les copules nous permettent d'extraire la structure de dépendance d'une fonction de distribution conjointe. Ce faisant, on arrive à séparer la structure de dépendance à partir des fonctions de distributions marginales.

³⁶ Pour d'autres mesures de corrélation, on consultera : Racicot, F.E. et R. Théoret (2001), Traité d'économétrie financière, Presses de l'Université du Québec (PUQ).

³⁷ À noter qu'il y a un lien entre un tirage d'une distribution de Lévy (ou d'une t) et la distribution de Fréchet (ou la distribution Gumbel), respectivement. En effet, les tirages de la Lévy convergent vers une distribution de Fréchet à mesure que la taille de l'échantillon (n) augmente. De même pour les tirages d'une normale ou lognormale, ils obéiront à une distribution Gumbel pour un n important.

³⁸ Pour de plus amples informations sur la théorie des copules, on consultera : Nelson (1999) et Cherubini et al. (2004).

Un résultat clé dans ce domaine est celui de Sklar (1959)³⁹. Afin d'illustrer le concept de copule, considérons encore une fois nos deux variables aléatoires X et Y. En supposant que $F(x,y)$ est une fonction de distribution conjointe avec pour distributions marginales (continues) $F_x(x) = u$ et $F_y(y) = v$, alors $F(x,y)$ peut être écrite en terme d'une fonction unique $C(u,v)$ telle que : $F(x,y) = C(u,v)$ où $C(u,v)$ est connue sous le nom de copule de $F(x,y)$. Ainsi, la fonction copule décrit comment la fonction multivariée $F(x,y)$ est dérivée ou couplée à des fonctions de distributions marginales $F_x(x)$ et $F_y(y)$. On peut comprendre la copule comme étant celle qui donne la structure de dépendance de la fonction de distribution multivariée $F(x,y)$.

Ce résultat est donc très important car il nous permet de construire des fonctions de distributions conjointes à partir des fonctions de distributions marginales de façon à tenir compte de la structure de dépendance de nos variables aléatoires. Pour modéliser la distribution conjointe, tout ce dont nous avons besoin de faire est de spécifier nos distributions marginales, choisir une copule pour représenter la structure de dépendance, estimer les paramètres impliqués et enfin, appliquer la fonction de copule à nos marginales. Une fois que nous sommes en mesure de modéliser la fonction de distribution conjointe, nous pouvons en principe l'utiliser pour estimer n'importe quelle mesure de risque comme celle qui nous intéresse, c'est-à-dire la VaR⁴⁰.

³⁹ Sklar, A. (1959), Fonctions de répartition à n dimensions et leur merges, Publ. Inst. Stat. Univ., Paris 8 : 229-231.

⁴⁰ La VaR et la ES : *Expected Shortfall* (synonyme de ETL ou C-VaR) sont en fait des cas particuliers d'une fonction connue sous le nom de fonction d'aversion au risque et également connue sous le nom de spectre de risque (risk-spectrum). Cette mesure du risque est aussi considérée comme étant plus cohérente que la VaR ou l'ES car elle est plus générale que l'une ou l'autre de ces mesures. Elle est donnée par : $M_\phi = \int_0^1 \phi(p) q_p dp$, où

$\phi(p)$ est une fonction de poids qui est généralement supposée définie par : $\phi_\gamma(p) = \frac{e^{-(1-p)/\gamma}}{\gamma(1 - e^{-1/\gamma})}$ qui est, on le

reconnaîtra, la fonction d'aversion pour le risque exponentielle tant utilisée en microéconomie. γ est le degré d'aversion au risque (plus γ est faible, plus le degré d'aversion au risque est grand), q_p est le quantile p du portefeuille prospectif des profits/pertes et p est le degré de confiance ou probabilité cumulative. Dans ce contexte, la VaR est simplement : $VaR = -q_p$. Par exemple, si la distribution de la série est normale, alors $q_p = 1.645$ pour une VaR à 95%. En pratique, on définit la VaR d'un portefeuille comme suit :

$VaR = -\mu_{p/L} + \sigma_{p/L} Z_\alpha$ pour des rendements normaux. Baumol (1963) est l'un des premiers à avoir discuté de cette mesure du risque mais c'est la firme JP Morgan qui a consacré cette mesure. L'ES quant à elle se définit

comme étant la moyenne des pires pertes au niveau α et est donnée par : $ES_\alpha = \frac{1}{1-\alpha} \int_\alpha^1 q_p dp$. Si la distribution

des pertes est discrète alors l'ES est donné par : $ES_\alpha = \frac{1}{1-\alpha} \sum_{p=0}^\alpha (p^{\text{ième}} \text{ grande perte}) \times (\text{prob. de } p^{\text{ième}} \text{ grand perte})$. En

substituant la fonction d'aversion au risque exponentielle dans la fonction spectre de risque, on obtient :

$M_\phi = \int_0^1 \phi(p) q_p dp = \int_0^1 \frac{e^{-(1-p)/\gamma}}{\gamma(1 - e^{-1/\gamma})} q_p dp$ qui se nomme : la mesure spectrale exponentielle du risque et est donc une moyenne pondérée des quantiles avec des poids définis par la fonction d'aversion pour le risque exponentielle.

10.2.1. Les copules les plus utilisées dans les applications financières

Parmi les copules les plus populaires, on compte les copules gaussiens, copules t , les copules logistiques ou de Gumbel. Dans les formes les plus simples, on compte les copules dites d'indépendance (ou produit), les copules minimales (ou comonotoniques) et les copules du maximum (contra-monotonique). Notre présentation ne se veut aucunement exhaustive mais seulement une présentation permettant au lecteur d'avoir une idée des copules. Commençons par présenter l'allure que prend la copule gaussienne :

$$C_{\rho}^{\text{Ga}}(x, y) = \int_{-\infty}^{\Phi^{-1}(x)} \int_{-\infty}^{\Phi^{-1}(y)} \frac{1}{2\pi(1-\rho^2)^{1/2}} e^{-\left(\frac{s^2 - 2\rho st + t^2}{2(1-\rho^2)}\right)} ds dt$$

où $-1 \leq \rho \leq 1$ et $\Phi^{-1}(\cdot)$ représente l'inverse de la fonction de distribution normale univariée. Il est important ici de noter que cette copule dépend de la corrélation ρ . Cela confirme que ρ est suffisant pour représenter toute la structure de dépendance. A partir de distributions marginales normales et de cette structure de dépendance, nous pouvons obtenir une loi normale multivariée. Malheureusement, les copules gaussiennes n'ont pas de solution analytique : il faudra avoir recours à des méthodes numériques pour les estimer comme par exemple la méthode FFT.

La copule t n'est autre qu'une généralisation de la copule gaussienne et est donnée par :

$$C_{v,\rho}^t(x, y) = \int_{-\infty}^{t_v^{-1}(x)} \int_{-\infty}^{t_v^{-1}(y)} \frac{1}{2\pi(1-\rho^2)^{1/2}} e^{-\left(1 + \frac{s^2 - 2\rho st + t^2}{v(1-\rho^2)}\right)^{-\frac{v+2}{v}}} ds dt$$

où $t_v^{-1}(\cdot)$ est l'inverse de la distribution de Student univariée à v degrés de liberté.

Une autre copule habituellement rencontrée dans la littérature financière est la copule logistique, dite copule de Gumbel, qui a l'allure suivante :

$$C_{\beta}^{\text{Gu}}(x, y) = e^{-\left(-[\log x]^{\beta/\beta} + [-\log y]^{\beta/\beta}\right)^{\beta}}$$

où $0 \leq \beta \leq 1$. Ce paramètre détermine le niveau de dépendance entre nos variables. En effet, quand $\beta = 1$, cela signifie que les variables sont indépendantes ; $\beta > 0$ signifie que nos variables ont un niveau de dépendance limité et quand $\beta = 0$, cela signifie que nos variables sont parfaitement dépendantes. Une caractéristique importante de cette copule est qu'elle est cohérente avec la théorie

Elle peut être mise en œuvre en utilisant des méthodes numériques pour calculer la moyenne pondérée car M_{ϕ} n'est rien d'autre qu'une moyenne pondérée. D'ailleurs une méthode similaire peut être appliquée au calcul de l'ES car l'ES n'est en fait que la VaR moyenne, soit la moyenne des VaR des différents niveaux de confiance postulés. Voir Artzner et al. (1997, 1999) pour plus de détails sur la théorie d'une mesure du risque cohérente.

des valeurs extrêmes (*EV theory*), ce qui n'est pas le cas de la copule gaussienne que nous venons de discuter.

Notons que les copules peuvent être classifiées en familles distinctes. Elles incluent, par exemple, les copules elliptiques et les copules archimédiennes. Les copules archimédiennes peuvent s'écrire sous la forme suivante :

$$C(u, v) = \varphi^{-1}(\varphi(u) + \varphi(v))$$

où $\varphi(\cdot)$ est une fonction strictement décroissante et convexe et $\varphi(1) = 0$. Les copules archimédiennes sont considérées comme étant faciles d'utilisation et peuvent s'ajuster à une variété de comportements de dépendance.

Une autre famille de copules fréquemment rencontrées en finance est celle des copules à valeurs extrêmes (EV-copulas). Elles sont issues d'une version multivariée du théorème de Fisher-Tippett (1928). L'application de celui-ci nous dit que si $G(x,y)$ est une distribution à valeurs extrêmes généralisée (GEV distribution), alors $G(x,y)$ doit avoir une copule qui satisfait la condition suivante :

$$C(u^t, v^t) = (C(u, v))^t$$

Donc si nous sommes en présence d'extrêmes multivariés, nous devrions choisir des copules qui satisfont cette condition. Les copules EV de base dans cette classe sont connues sous les noms de : copule produit (copule d'indépendance), copule minimum, copule Gumbel, copule Gumbel II et copule Galambos et sont données par :

$$C_{\text{ind}}(u, v) = u \times v$$

$$C_{\text{min}}(u, v) = \min(u, v)$$

$$C_{\beta}^{\text{Gu II}}(x, y) = uv e^{\left(\frac{\beta uv}{u+v}\right)}$$

$$C_{\beta}^{\text{Gal}}(x, y) = uv e^{\left(u^{-\beta} + v^{-\beta}\right)^{-1/\beta}}$$

où $\beta \in [0,1]$ pour $C_{\beta}^{\text{Gu II}}(\cdot)$ et $\beta \in [0,\infty]$ pour $C_{\beta}^{\text{Gal}}(\cdot)$. Mentionnons en terminant que la copule maximale (ou copule contra-monotone), comme la copule minimale (ou copule comonotone), est considérée comme la copule la plus simple qui existe. Elle est donnée par :

$$C_{\text{max}}(u, v) = \min(u + v - 1, 0).$$

10.3. Estimation des copules

Pour effectuer l'estimation d'une Copule, il faut choisir sa forme fonctionnelle et ensuite estimer les paramètres associés. On peut estimer ses paramètres par une méthode paramétrique classique comme par exemple, la méthode du maximum de vraisemblance. Pour ce faire, il suffit de construire la

fonction de vraisemblance de la forme fonctionnelle choisie et ensuite maximiser cette fonction en utilisant un algorithme d'optimisation comme par exemple, ceux qui sont disponibles dans le logiciel *EViews 5.1*. On peut également utiliser les méthodes non paramétriques qui sont généralement considérées simples de mise en œuvre et une variété de méthodes sont disponibles⁴¹.

10.4. Un exemple d'application des copules à la VaR du risque de crédit

La VaR du risque de crédit peut être définie de manière analogique à la VaR pour le risque de marché. En effet, la VaR de crédit à un niveau de confiance de 99,9% pour un horizon d'un an est la certitude à 99,9% que la perte de crédit n'excèdera pas ce niveau.

La VaR définie à l'aide de copules est donnée par :

$$\text{VaR} = L(1 - R)V(X, T)$$

où

1) L : valeur du portefeuille de prêt ;

2) R : taux de recouvrement de la dette ;

3) $V(X, T) = N\left(\frac{N^{-1}(Q(T)) + \sqrt{\rho}N^{-1}(X)}{\sqrt{1 - \rho}}\right)$: définit la certitude à $X\%$ que le pourcentage de perte sur

les années T pour un grand portefeuille sera plus petite que cette valeur $V(X, T)$ ⁴² ;

4) $Q(T)$: probabilité cumulée qu'une entreprise aura fait défaut au temps T ;

5) ρ : corrélation définie en terme de copules gaussiennes ;

6) $N(\cdot)$: distribution normale cumulée et l'inverse de cette distribution au niveau 5%, par exemple, est : $N(1,645) = 5\% \Rightarrow 1,645 = N^{-1}(0,05)$.

Clarifions ici certains concepts avant de passer à l'exemple concret du calcul de la VaR de crédit. Notons que la probabilité de défaut est reliée à un autre concept souvent utilisé en économétrie, c'est-à-dire la fonction de hasard⁴³ (taux de hasard), aussi appelé intensité (ou intensité de défaut, en finance). Considérons une courte période de temps notée Δt . L'intensité de défaut, notée $\lambda(t)$, est la probabilité conditionnelle de défaut et est définie de telle sorte que $\lambda(t)\Delta t$ est la probabilité de défaut pour l'intervalle t et $t + \Delta t$ conditionnellement à ce qu'il n'y ait eu aucun défaut pour cet intervalle de temps.

⁴¹ Pour plus de détails sur l'estimation des copules, on consultera : Bouyé et al. (2000), Scaillet (2000a).

⁴² Ce résultat est dû à Vasicek, O. (1987), Probability of Loss on Loan Portfolio, Working Paper, KMV.

⁴³ Pour d'autres applications financières ou économiques, on consultera : Racicot, F.E. (2003), Trois essais sur l'analyse des données économiques et financières, chap. 2, thèse de doctorat, ESG-UQAM.

Définissons $V(t)$ par la probabilité cumulative qu'une entreprise survive jusqu'au temps t , c'est-à-dire qu'il n'y ait eu aucun défaut jusqu'au temps t , alors la relation entre $V(t)$ et $\lambda(t)$ est donnée par :

$$V(t + \Delta t) - V(t) = -\lambda(t)V(t)\Delta t$$

En prenant la limite de cette expression pour $\Delta t \rightarrow 0$, on obtient :

$$\begin{aligned}\lim_{\Delta t \rightarrow 0} (V(t + \Delta t) - V(t)) &= \lim_{\Delta t \rightarrow 0} (-\lambda(t)V(t)\Delta t) \\ \Rightarrow dV(t) &= -\lambda(t)V(t)dt \\ \Rightarrow \frac{dV(t)}{dt} &= -\lambda(t)V(t)\end{aligned}$$

Cette équation n'est qu'une simple équation différentielle habituelle⁴⁴ dont la solution pour $V(t)$ est :

$$V(t) = e^{-\int_0^t \lambda(\tau) d\tau}$$

En définissant $Q(t)$ comme la probabilité de défaut cumulée jusqu'au temps t et sachant que l'aire totale sous la distribution de probabilité de $-\infty$ à ∞ est 1, alors $Q(t)$ est donnée par :

$$Q(t) = 1 - V(t) = 1 - e^{-\int_0^t \lambda(\tau) d\tau}$$

Mais une intégrale définie n'est en fait qu'une somme de petits rectangles, alors on peut redéfinir

$\int_0^t \lambda(\tau) d\tau$ par la moyenne de $\lambda(t)$ multipliée par t : $t\bar{\lambda}(t)$. On obtient alors que $Q(t)$ est définie en terme d'intensité moyenne de défaut entre 0 et t (ou probabilité conditionnelle de défaut), c'est-à-dire ;

$$Q(t) = 1 - e^{-\bar{\lambda}(t)t}$$

Donc, pour trouver $Q(t)$, tout ce dont nous avons besoin n'est qu'une probabilité moyenne conditionnelle de défaut. Moody's, par exemple, effectue le calcul d'intensité moyenne de défaut pour différentes cotes d'obligations corporatives. Une obligation cotée Aaa de 4 ans a une probabilité de défaut cumulée : $Q(4) = 1 - e^{-\bar{\lambda}(4)4} = 0,04\%$.

L'intensité moyenne de défaut, sur une période de 4 ans, est donc de :

$$\bar{\lambda}(4) = -\frac{1}{4} \ln(1 - Q(4)) = -\frac{1}{4} \ln(1 - 0,0004) = 0.0010\%.$$

Il est important ici de noter que ces probabilités de défaut ne sont pas risque neutres. En effet, ces probabilités sont des probabilités physiques définies dans un monde réel et sont obtenues à partir de données historiques. Elles sont généralement très inférieures aux probabilités risque neutres. On obtient les probabilités risque neutres à partir des prix d'obligations en appliquant la formule : $h = s/(1 - R)$ où h est la probabilité conditionnelle (par année) à aucun défaut préalable ; s est le spread entre le rendement d'une obligation corporative et le taux sans risque et R est le taux espéré de recouvrement de dette. On observe que les probabilités réelles convergent vers les probabilités risque neutres à mesure que la cote de l'obligation diminue. On peut donc se demander

⁴⁴ Voir notre annexe sur les rappels mathématiques et les équations différentielles dans notre prochain livre : Finance computationnelle et gestion des risques.

quelle est la meilleure procédure pour évaluer la probabilité de défaut. La réponse dépend de l'analyse effectuée, c'est-à-dire que lorsqu'on tente d'évaluer l'impact du risque de défaut sur le pricing de produits dérivés du crédit, il est préférable d'utiliser les probabilités risque neutres. Cela est dû au fait que lorsque l'on effectue ce type d'analyse, on est confronté à des calculs qui requièrent une actualisation de cash-flows espérés qui implique implicitement ou explicitement l'utilisation de l'évaluation risque neutre. Mais il est préférable d'utiliser les probabilités objectives dans le cas où l'on est confronté à l'analyse de scénarios de pertes potentielles futures causées par des défauts de paiements.

Revenons à notre exemple de copule gaussienne et du calcul de la VaR de crédit. Supposons qu'une banque ait fait des prêts pour un total de $L = 150$ millions \$. La probabilité de défaut moyenne pour une année est de 2%, le taux de recouvrement moyen est de 61%, la corrélation en terme de la copule estimée est de 15%. En effectuant nos calculs dans Excel, on obtient le tableau 26.

Tableau 26
Calcul de la VaR de crédit à l'aide de la copule gaussienne

	A	B	C	D	E	F	G	H
3	Loi normale inverse au niveau 5% :	-1.64485363	=LOI.NORMALE.INVERSE(0.05;0;1)					
4	Loi normale inverse au niveau 2% :	-2.05374891	=LOI.NORMALE.INVERSE(0.02;0;1)					
5	Loi normale inverse au niveau 99.9% :	3.09023231	=LOI.NORMALE.INVERSE(0.999;0;1)					
6	L =	150						
7	R =	61%						
8	corrélation copule =	0.15						
9	Q(T) =	2%						
10	$N^{-1}(0,02)$	-2.05374891						
11	$N^{-1}(0,99)$	3.09023231						
12	V(0.999,1)	0.17632894	=LOI.NORMALE.STANDARD(((\$B\$8+RACINE(\$B\$6)*\$B\$9)/RACINE(1-\$B\$6)))					
13								
14	VaR =	10.3152429	=\$B\$6*(1-\$B\$7)*\$B\$12					
15								
16								

La VaR est donc dans ce cas de : $VaR = 10,31$ millions de dollars.

Notre en terminant que l'accord de *Basel II* propose des changements nouveaux au calcul de la VaR de crédit. En effet, ces nouvelles façons de calculer la VaR impliquent l'utilisation de copules et cela afin d'obtenir une méthode raisonnable cherchant à tenir compte de la structure de dépendance. Comme nous l'avons vu précédemment, la théorie des copules propose une méthode de calcul de la corrélation qui est robuste, d'où les changements éventuels au calcul de la VaR proposés par le nouvel accord de *Basel II*.

10.5. La transformée de Fourier

Soit $f(t)$ une série temporelle continue non périodique. Sa transformée de Fourier⁴⁵ est égale à:

$$\mathfrak{F}(\alpha) = \frac{1}{\sqrt{2\pi}} \int_{t=-\infty}^{\infty} f(t) e^{-i\alpha t} dt$$

$\mathfrak{F}(\alpha)$ est également appelé: fonction de densité spectrale de $f(t)$. Pour sa part, la transformée inverse de Fourier s'écrit:

$$\mathfrak{F}^{-1} = f(x) = \frac{1}{2\pi} \int_{-\pi}^{\pi} \mathfrak{F}(\alpha) e^{i\alpha x} d\alpha$$

avec $i = \sqrt{-1}$. La transformée de Fourier⁴⁶ permet, entre autres, d'accélérer les calculs effectués à l'aide d'un arbre binomial dans le but de valoriser une call européen. La formule utilisée est la suivante:

$$C_0 = \mathfrak{F}\left\{\mathfrak{F}^{-1}(C_T) \times [\mathfrak{F}(b)]^N\right\} \quad (1)$$

avec C_0 , le prix du call européen; $\mathfrak{F}(\cdot)$, l'opérateur de la transformée de Fourier; $\mathfrak{F}^{-1}(\cdot)$, l'opérateur de la transformée inverse de Fourier; C_T , le vecteur des payoffs de l'option à son échéance dont la dimension est de $N+1$, N étant le nombre de pas de l'arbre binomial.

Le vecteur $\mathbf{b}$ est un vecteur de dimension $((N+1) \times 1)$. La première entrée du vecteur est le ratio des deux termes suivants: i) la probabilité de hausse du prix de l'action dans l'arbre binomial; ii) $(1 + r_f)$, r_f étant le taux sans risque. La deuxième entrée contient les mêmes calculs pour la probabilité de baisse. Les autres cellules du vecteur $\mathbf{b}$ sont nulles.

Pour illustrer la formule (1), nous recourrons tour à tour au logiciel Excel puis au logiciel Matlab. Nous considérons le cas de la construction d'un arbre binomial à 3 pas ($N=3$) pour un call européen aux caractéristiques suivantes: $S = X = 100$ \$; $T = 0,25$; $r = 0\%$ et $\sigma = 0,1$.

⁴⁵ Pour une excellente présentation des séries de Fourier, voir : Piskounov, N., Calcul différentiel et intégral, tome II, éditions MIR, Moscou, chap. 17.

⁴⁶ Sur les applications de la transformée de Fourier en finance, on consultera, entre autres: Cerny, A. (2004), Mathematical Techniques in Finance: Tools for Incomplete Markets, Princeton University Press, chap. 7. Il est à noter que nous avons modifié la transformée de Fourier qui apparaît dans ce manuel pour calculer le prix d'un call européen. De même, nous utilisons les langages Excel et Matlab et non le langage Gauss pour calculer cette fonction.

Nous voulons solutionner cet arbre de façon accélérée en recourant à l'équation (1). On calcule d'abord les vecteurs C_T et b sous les données du problème:

$$C_T = \begin{bmatrix} 9,04 \\ 2,92 \\ 0 \\ 0 \end{bmatrix}$$

$$b = \begin{bmatrix} 0,4928 \\ 0,5072 \\ 0 \\ 0 \end{bmatrix}$$

Puis nous appliquons la formule (1) pour calculer le prix du call européen. L'utilitaire d'analyse d'Excel permet d'effectuer les transformées de Fourier d'un vecteur de même que sa transformée inverse. La transformée inverse de C_T , soit $\mathfrak{F}^{-1}(C_T)$, est égale à:

$$\begin{bmatrix} 2.99 \\ 2.26+0.73i \\ 1.53 \\ 2.26-0.73i \end{bmatrix}$$

avec $i = \sqrt{-1}$, soit un nombre imaginaire.

Tandis que la transformée de Fourier du vecteur b est de:

$$\begin{bmatrix} 1 \\ 0.4928-0.5072i \\ -1.44E-002 \\ 0.4928+0.5072i \end{bmatrix}$$

L'équation (1) nous indique qu'il faut élever ce vecteur à l'exposant N , ici 3. Nous nous servons donc de la fonction Excel suivante pour effectuer cette opération:

$$= \text{complexe.exposant}(\text{nombre 1}, \text{exposant})$$

À la suite de cette opération, nous obtenons le vecteur suivant, soit $\mathfrak{F}(b)^N$:

$$\begin{bmatrix} 1 \\ -0.260643733504-0.239045226496i \\ -2.985984E-006+1.09699649364359E-021i \\ -0.260643733504+0.239045226496i \end{bmatrix}$$

L'équation (1) indique qu'il faut multiplier les deux vecteurs $\mathfrak{F}^{-1}(C_T)$ et $\mathfrak{F}(b)^N$ et effectuer par la suite une autre transformée de Fourier sur le produit de ces deux vecteurs pour en arriver au résultat souhaité. Pour effectuer le produit de deux nombres complexes, nous nous servons de la fonction Excel suivante:

= *complexe.produit (nombre 1, nombre 2)*

Le produit des deux vecteurs donne le résultat suivant:

```
2.99  
-0.41455182237696-0.73051213733888i  
-4.56855552E-006+1.67840463527469E-021i  
-0.41455182237696+0.73051213733888i
```

Pour calculer le prix de l'action, il faut effectuer une dernière transformée de Fourier sur ce vecteur, selon l'équation (1). On obtient:

```
2.16089178669056  
1.52898029387776  
3.8190990761984  
4.45102884323328
```

Le prix du call européen est la première cellule de ce vecteur, soit 2,16 \$. Ce résultat est identique à celui que l'on obtiendrait en solutionnant directement l'arbre binomial sans recourir à la transformée de Fourier. On voit que les nombres complexes ne sont que subsidiaires aux calculs. Ils disparaissent lors de la solution finale. Le prix de ce call est de 1,99 \$ selon l'équation de Black et Scholes.

Nous recourons maintenant au logiciel Matlab pour calculer l'équation (1). Dans ce logiciel, la fonction `fftn` effectue la transformée de Fourier d'un vecteur de dimension `n`. Par exemple, `fftn(c)` commande la transformée du vecteur `c`. Par ailleurs, la fonction `ifftn` effectue la transformée inverse de Fourier d'un vecteur.

Nous écrivons d'abord le vecteur **c** dans le logiciel Matlab:

```
>> c=[9.04 2.92 0 0]  
  
c =  
  
9.0400 2.9200 0 0
```

Puis le vecteur **b**:

```
>> b=[.4927 .5073 0 0]  
  
b =  
  
0.4927 0.5073 0 0
```

Puis nous écrivons la formule (1) dans le langage Matlab:

```
>> pro=fftn(ifftn(c).*((1)*fftn(b)).^3)
```

```

pro =
    2.1600    1.5295    3.8200    4.4505

```

Nous obtenons donc le même vecteur que celui du logiciel Excel. Il est à noter que nous avons écrit l'équation (1) d'un seul trait dans Matlab. Nous avons fait suivre les vecteurs de points, par exemple nous avons écrit: `ifftn(c)`., pour bien signifier que nous voulons effectuer des opérations terme par terme sur les vecteurs et non des opérations matricielles classiques. Matlab estime donc le prix du call à 2,16 \$ en recourant à la transformée de Fourier, soit un prix identique à celui que nous avons obtenu avec le logiciel Excel.

Pour se rapprocher davantage de la solution donnée par B&S, on peut fixer le nombre de pas à 10. Le vecteur C_T des payoffs de l'option à l'échéance est alors de:

```

>> c=[17.12 13.48 9.95 6.52 3.21 0 0 0 0 0]

c =

Columns 1 through 7
    17.1200    13.4800    9.9500    6.5200    3.2100         0         0

Columns 8 through 11
         0         0         0         0

```

Tandis que le vecteur b est de:

```

>> b=[.4960 .5040 0 0 0 0 0 0 0 0]

b =

Columns 1 through 7
    0.4960    0.5040         0         0         0         0         0

Columns 8 through 11
         0         0         0         0

```

La solution pour le prix du call est donc de:

```

>> pro=fftn(ifftn(c).*((1)*fftn(b)).^10)

pro =

Columns 1 through 7
    1.9429    0.7586    0.3955    0.9752    2.8055    5.7268    8.5574

Columns 8 through 11
    9.7798    8.8565    6.5265    3.9553

```

Avec 10 pas, la valeur du call est donc de 1,94 \$, chiffre plus rapproché de la valeur théorique du call, à hauteur de 1,99\$, qu'avec seulement 3 pas.

Conclusion

Le risque est un concept multidimensionnel. C'est pour cette raison qu'il est très difficile à mesurer. Ces dernières années, la théorie du risque s'est ralliée à celle de la dominance stochastique, qui semble une voie très prometteuse pour en arriver éventuellement à une approche intégrée du risque. En privilégiant la distribution cumulative, les nouvelles théories du risque s'attachent en effet à la forme de la distribution des rendements, ce que ne faisaient que très sommairement les théories traditionnelles du risque car elles prenaient pour acquise la normalité de la distribution des rendements. Or, celle-ci ne l'est pas.

La théorie moderne du risque intègre progressivement les enseignements de la théorie des options. Par exemple, les payoffs causés par un événement rare défavorable sont assimilables à ceux d'une position à découvert dans une option de vente. On se sert donc maintenant de plus en plus des payoffs d'une option pour mesurer le risque. A cet égard, l'indicateur oméga recourt au prix d'un put à court terme pour mesurer le risque d'un portefeuille. Ce prix mesure alors le coût de protéger ce portefeuille, coût que l'on peut donc assimiler au risque dudit portefeuille. Certes, les perspectives qu'ouvre la théorie des options pour la mesure du risque s'avèrent très prometteuses, les options pouvant en effet représenter tous les moments d'une distribution. L'approche au risque par les cumulants d'une distribution est une autre voie d'avenir.

Annexe

Modification du programme de bootstrapping

Au tableau A1, on retrouve le programme Visual Basic qui nous permet de bootstrapper la série temporelle originale à chaque itération et non la série résultant de la dernière itération, comme cela est effectué dans l'article, à la section 5.2. Nous avons bootstrappé le portefeuille antérieur de trois actions du secteur de la biotechnologie. Pour ce faire, nous recopions à chaque itération dans une colonne du chiffrier la série originale, que nous bootstrapons par la suite.

Tableau A1 Programme Visual Basic du bootstrapping d'un portefeuille composé de trois titres biotechnologiques, le bootstrap s'effectuant constamment sur la série originale

```

Sub Varin()

Range("starttime") = Time
Range("O1:O15000").ClearContents

For Iteration = 1 To Range("iterations")
For Row = 1 To 574
Range("TSE").Cells(Row, 1) = Rnd
Range("tsecopie").Cells(Row, 1) = Range("tse300").Cells(Row, 1)
Next Row
Range("c2:d574").Select
Selection.Sort Key1:=Range("d2"), Order1:=xlAscending, Header:=xlNo, _
    OrderCustom:=1, MatchCase:=False, Orientation:=xlTopToBottom

Range("returndata").Cells(Iteration, 1) = Range("rmoyen")
Next Iteration

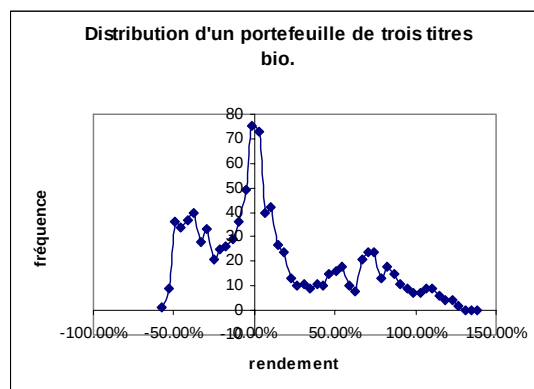
Range("elapsed") = Time - Range("starttime")

End Sub

```

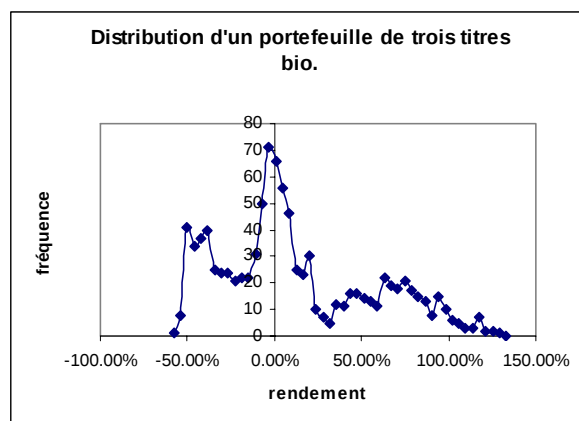
La distribution des rendements du portefeuille qui résulte du bootstrapping se retrouve à la figure A1.

Figure A-1



À la figure A-2 se retrouve le même exercice sauf que l'on utilise la technique du bootstrapping qui apparaît dans l'article et qui effectue, à chaque itération, le bootstrapping de la série obtenue lors de l'itération antérieure. On constate que les différences entre les graphiques A-1 et A-2 sont mineures. Comme nous le disions antérieurement, des différences risquent d'apparaître dans un exercice de bootstrapping sans remise.

Figure A-2



Bibliographie

- Alexander, C.(1996), The Handbook of Risk Management and Analysis, Wiley, New York.
- Benninga, S.(2000), Financial Modeling, The MIT Press, Cambridge.
- Bera, A.K. et C.M. Jarque (1981), *An Efficient Large-Sample Test for Normality of Observations and Regression Residuals*, Australian National University Working Papers in Econometrics, 40, Canberra.
- Bessis, J.(1998), Risk Management in Banking, John Wiley and Sons, New York.
- Cooley, P.L. (1977), A multidimensional Analysis of Institutional Investor Perception of Risk, *Journal of Finance*, 32, p. 67-78.
- Copeland, T.E., J.F. Weston et K. Shastri (2005), Financial Theory and Corporate Policy, Pearson.
- Cornish, E.A et Fisher, R.A. (1937), *Moments and Cumulants in the Specification of Distributions*, *Rev. Int. Statist. Inst.*, 307
- Das, S.(1997), Risk Management and Financial Derivatives, McGraw Hill, New York.
- Dowd, K. (2002, 2005), Measuring market risk, 1^e et 2^e ed., Wiley.
- Efron, B. (1994), *Missing Data, Imputation and the bootstrap*, *J. Amer. Statis. Ass.*, 89
- Efron, B. (1979), *Bootstrap methods: Another look at the jackknife*, *Ann. Statis.*, 7.
- Esch, K. et al.(1997), Value at Risk: vers un risk management moderne, DeBoeck, Bruxelles.
- Fama, E.F. et French, K.R. (1992), The cross-section of expected stock returns, *Journal of Finance*, 47, 427-465.
- Fung, W et D.A. Hsieh (1997), Empirical Characteristics of Dynamic trading strategies : the case of hedge funds, *The Review of financial studies*, 10, 1997, p. 275-303.
- Giot, P. et Laurent, S. (2003), Value-at-Risk for long and short trading positions, *Journal of Applied Econometrics*, 18, 641-664.
- Hoel, P.G., Port, S.C. et Stone, C.J.(1971), Introduction to Probability Theory, Houghton Mifflin Company.
- Hull, J.C. (2006), Options, Futures and Other Derivatives, 6^e ed., Pearson Prentice Hall
- Hull J.C. (2000), Options, Futures and Other Derivatives, Prentice Hall, New Jersey
- James, J. et Webber, N. (2000), Interest Rate Modelling, Wiley.
- Jorion, P.(2003), Financial Risk Manager Handbook, Wiley.
- Judge, G.G. et al.(1988), Introduction to the Theory and Prattice of Econometrics, 2^{ième} édition, Wiley, New York.

Kazemi, H., T. Schneeweis et R. Gupta (2003), Omega as a Performance Measure, working paper, 2003.

Kraus, A. et R.H. Litzenberger(1976), Skewness Preference and the valuation of Risk assets, *Journal of Finance*, 31, p. 1085-1100.

Levy, H. et M. Sarnat (1984), Portfolio and Investment Selection : Theory and Practice, Prentice Hall.

Lhabitant, F.S.(2004), Hedge Funds : Quantitative Insights, John Wiley & Sons.

Malevergne, Y. et Sornette, D. (2005), Higher-Moment Portfolio Theory, Capitalizing on behavioral anomalies of stock markets, *Journal of Portfolio Management*, 31, 4, 49-55.

McNeil, A.J. et Frey, R. (2000), Estimation of tail-related risk measures for heteroscedastic financial time series: An extreme Value Approach, *Journal of Empirical Finance*, 7, 271-300.

Poon, S.-H.(2005), A Practical Guide to Forecasting Financial Market Volatility, Wiley.

Racicot, F.-É. et R. Théoret (2001), Traité d'économétrie financière: modélisation financière, Les Presses de l'Université du Québec, Ste-Foy.

Scott, R.C. et P.A. Hovarth (1980), On the direction of preference for moments of higher order than the variance, *Journal of Finance*, 35, pp. 915-919.

Stuart A., Ord K. et Arnold S. (1999), Kendall's Advanced Theory of Statistics. Volume 1: Distribution Theory, Arnold.

Stuart A., Ord K. et Arnold S.(1999), Kendall's Advanced Theory of Statistics. Volume 2A: Classical Inference and the Linear Model, Arnold.

Taleb, N. (1997), Dynamic Hedging. Managing Vanilla and Exotic Options, Wiley.

Tehrani, H. (1980), Empirical Studies in Portfolio Performance Using Higher Degrees of Stochastic Dominance, *Journal of Finance*, 35, p. 159-171.

Théoret, R. et P. Rostan (2000), *Empirical Comparative Study of Monte Carlo Simulation Methods versus Historical Methods to estimate Value at Risk*, Centre de recherche en gestion, École des sciences de la gestion, UQAM, CRG 19-2000.

Théoret, R.(2000), Traité de gestion bancaire, Presses de l'Université du Québec, Ste-Foy.

Wilmott, P. (2006), The Best of Wilmott 2, Wiley.