

HICHEM KADDECHE

ANDRÉ-LUC BEYLOT

MONIQUE BECKER

**Approche analytique pour l'étude des performances
de serveurs multimédias multidisques**

RAIRO. Recherche opérationnelle, tome 32, n° 3 (1998),
p. 227-251

http://www.numdam.org/item?id=RO_1998__32_3_227_0

© AFCET, 1998, tous droits réservés.

L'accès aux archives de la revue « RAIRO. Recherche opérationnelle » implique l'accord avec les conditions générales d'utilisation (<http://www.numdam.org/conditions>). Toute utilisation commerciale ou impression systématique est constitutive d'une infraction pénale. Toute copie ou impression de ce fichier doit contenir la présente mention de copyright.

NUMDAM

Article numérisé dans le cadre du programme
Numérisation de documents anciens mathématiques
<http://www.numdam.org/>

APPROCHE ANALYTIQUE POUR L'ÉTUDE DES PERFORMANCES DE SERVEURS MULTIMÉDIAS MULTIDISQUES

par Hichem KADDECHE ⁽¹⁾ ⁽²⁾, André-Luc BEYLOT ⁽³⁾
et Monique BECKER ⁽¹⁾ ⁽²⁾

Résumé. – *Les applications multimédias futures nécessitant des serveurs capables de fournir des séquences vidéo à la demande à plusieurs utilisateurs simultanés géographiquement éloignés sont nombreuses et variées. L'architecture multidisque en grappe semble bien adaptée à ce type d'applications. Cette architecture consiste en un ensemble de nœuds de stockage et de transmission reliés par un réseau d'interconnexion. Chaque nœud de stockage est constitué d'un ou de plusieurs disques. Les fichiers vidéo sont fragmentés en blocs et distribués sur les différents disques.*

Nous présentons dans cet article une approche et un modèle analytique pour l'étude de ces serveurs. Nous présentons en détail la modélisation des différents composants du système et le modèle simplifié final adopté. Le modèle permet de dimensionner rapidement les paramètres de fonctionnement. Il permet d'économiser des simulations coûteuses en temps, éviter les problèmes d'événements rares et d'explorer aussi des domaines de fonctionnement inaccessibles par les simulations. Nous nous sommes intéressés à la qualité de transmission caractérisée par la probabilité de rupture momentanée dans la transmission des blocs. Une rupture momentanée de transmission se produit lorsque le bloc à envoyer n'est pas chargé dans les délais au niveau du nœud de transmission qui s'occupe de son envoi sur le réseau extérieur. Ce bloc se sera pas envoyé sur le réseau extérieur ce qui engendrera une petite interruption au niveau de l'utilisateur. Le critère de performance étudié est alors la probabilité de retard pour un bloc. Cette probabilité dépend de la fonction de répartition du temps de chargement d'un bloc : temps d'attente et d'extraction au niveau du disque et durée de traversée du réseau d'interconnexion. Dans le modèle simplifié final les disques sont modélisés par des files M/D/1 et le réseau d'interconnexion par un délai constant.

Le modèle proposé est validé par de longues simulations réalisées sur un modèle précis sous forme de réseau de files d'attente. L'utilisation du modèle pour le dimensionnement du système est alors présentée. © Elsevier, Paris

Mots clés : Serveurs multimédias, architecture multidisque, modèles analytiques, évaluation de performances, disques, réseaux d'interconnexion.

Abstract. – *Future multimedia applications requiring servers able to provide video sequences on-demand to simultaneous geographically distributed users are numerous and varied. Clustered*

⁽¹⁾ Institut National des Télécommunications, Département informatique, 9 rue Charles-Fourier, 91011 Évry Cedex France. E-mail : kaddeche@etna.int-evry.fr, mbecker@etna.int-evry.fr

⁽²⁾ Laboratoire LIP6, Université de Paris-6, 4, place Jussieu, 75252 Paris Cedex France.

⁽³⁾ Laboratoire PRISM, Université de Versailles Saint-Quentin-en-Yvelines, avenue des États-Unis, 78035 Versailles Cedex France. E-mail : beylot@prism.uvsq.fr

multi-disk architectures seem suitable for designing such servers. They consist in a set of storage nodes and a set of delivery nodes interconnected by a switch. Each storage node is composed of one or several disks. Video files are shared into blocks and distributed over the disks.

This paper presents an analytical model for the performance study of clustered multi-disk multimedia servers. The modeling of the different component of the system is detailed and the final simplified model derived. The model allows a quick dimensioning of the operating parameters. It allows to save expensive simulations, to avoid rare event problems and then to explore values of operation parameters for which performance is not easily evaluated by simulations. The performance evaluation study focuses on the quality of transmission characterized by the probability of occurrence of a break in the delivery of blocks. A block that is not loaded in time in the delivery node that manages its transmission on the external network, will not be transmitted. This will result in a short break at the user end. The performance criterion considered is then the probability of a block to be late. This probability depends on the distribution of the loading time of a block; response time of the disk plus response time of the switch. In the final simplified model, the disks are modeled by M/D/1 queues and the switch is modeled by a deterministic delay.

The proposed model is validated through extensive simulations of an accurate queuing network model. The dimensioning of the system is then derived. © Elsevier, Paris

Keywords: multimedia servers, multi-disk architecture, analytical models, performance evaluation, disk, switch.

1. INTRODUCTION

Les développements récents de l'informatique et des technologies multimédias ont donné naissance à de nombreuses nouvelles applications qui manipulent des données vidéo. La plupart de ces applications sont encore locales sur des machines isolées ou limitées à de petits réseaux locaux. Les applications multimédias à distance disponibles actuellement sont généralement de qualité limitée tant au niveau de l'image qu'au niveau de l'interactivité et du temps de réponse (les services du World Wide Web).

L'arrivée des réseaux à haut débit [Vetter 95] permettra de remédier aux limitations actuelles de transmission et rendra possible le développement de services multimédias distribués de qualité. La réalisation de tels services nécessite le développement de serveurs capables de stocker et de délivrer à la demande des données vidéo pour plusieurs utilisateurs simultanés. Les utilisateurs seront géographiquement distribués et connectés à distance au serveur à travers le réseau. Les charges visées pour ces systèmes sont de l'ordre de centaines voire de milliers d'utilisateurs simultanés. La problématique générale des systèmes multimédias avec ses deux aspects serveur et réseau est présentée dans [Vin 94].

Les données vidéo consistent en une suite d'images qui n'ont de sens que lorsqu'elles sont présentées d'une façon continue contrairement aux données numériques classiques (texte, image) où une simple continuité spatiale suffit. La conséquence de cette continuité temporelle est que ces données nécessitent

des techniques différentes pour leur organisation et leur gestion. Un serveur multimédia doit assurer la distribution des flux vidéo à leur débit temps réel.

La conception des serveurs capables de gérer un grand nombre d'utilisateurs simultanés, est une opération complexe qui lance beaucoup de défis. Aux problèmes de volume de données, de hauts débits et de contraintes de continuité temporelle, s'ajoute la difficulté de la gestion des systèmes multi-utilisateurs. Le système doit être capable de servir le même fichier à différents utilisateurs sans dégradation de ses performances. Des solutions basées sur le parallélisme semblent apporter des réponses à un certain nombre de ces problèmes. La fragmentation des objets multimédias en plusieurs blocs et leur répartition sur plusieurs disques permet d'améliorer considérablement les performances en favorisant un meilleur équilibrage de la charge [Berson 94] [Haskin 95]. Les études des systèmes de stockage multidisques ont donné lieu à plusieurs propositions et à l'évaluation des performances de différents schémas [Livny 87] [Ganger 94].

Les applications futures pour ces serveurs multimédias sont nombreuses et variées. Elles concernent les services d'archives et d'information, la télévision à la demande, les films à la demande, l'enseignement à distance, les informations à la demande... cet ensemble d'applications fait que ces serveurs sont d'une grande importance économique et suscitent l'intérêt d'industriels de l'informatique, des télécommunications et des loisirs. Cet enjeu économique est à la base des nombreux travaux de recherche qui sont menés actuellement dans le domaine.

Un certain nombre d'architectures ont été proposées pour les serveurs multimédias ces dernières années. Un état de l'art de la problématique et des solutions possibles est présenté dans [Gemmell 95]. Nous nous intéressons dans cette étude à une architecture qui semble bien adaptée pour ce genre de système [Damm 94] [Kaddeche 96] [Tewari 96]. Il s'agit d'une architecture qui comporte un ensemble de nœuds de stockage reliés à un ensemble de nœuds de transmission par un réseau d'interconnexion. Chaque nœud de stockage est composé d'un ou de plusieurs disques. On parle alors de serveurs multidisques ou serveurs en grappe (*clustered servers*). L'étude que nous présentons traite des performances du système indépendamment du réseau extérieur.

Nous avons concentré principalement notre étude sur la qualité de la transmission qui dépend des ruptures momentanées dans la transmission des blocs. Nous avons estimé la probabilité pour qu'un bloc provoque une rupture

(c'est-à-dire la probabilité pour qu'un bloc d'information arrive après sa date prévue de transmission).

Nous présentons dans cet article une approche et un modèle analytique simplifié que nous validons par des simulations et exploitons pour le dimensionnement ainsi que pour l'étude de l'effet de certains paramètres internes de fonctionnement du système.

L'article est structuré comme suit. Le paragraphe 2 décrit l'architecture et le principe de fonctionnement du système. Dans le paragraphe 3, nous présentons le principe de l'approche analytique. Les paragraphes 4 et 5 présentent l'étude détaillée des disques et du réseau d'interconnexion. Dans le paragraphe 6, nous décrivons le modèle de simulation utilisé pour la validation de l'approche analytique. La validation et l'exploitation du modèle analytique sont alors présentées dans la paragraphe 7.

2. ARCHITECTURE

Les séquences vidéo considérées sont numérisées et compressées suivant le format MPEG-2 (*Motion Picture Expert Group*) [ISO 93]. Le format de compression MPEG-2 utilisé induit un débit constant (*Constant Bit Rate, CBR*) de 0,5 Mo/s. La fonction du serveur est de stocker et de transmettre à la demande ces fichiers vidéo en respectant le débit requis.

Chaque fichier vidéo est fragmenté en blocs de même taille correspondant à des fragments de séquences de même durée. Les blocs sont distribués sur les différents disques suivant une politique statique de placement. Lors de la demande d'une séquence vidéo par un utilisateur, on transmettra dans le bon ordre et d'une manière isochrone les différents blocs de la séquence en respectant le débit MPEG imposé.

L'architecture comprend des nœuds de stockage et des nœuds de transmission, reliés par un réseau d'interconnexion. Les nœuds de transmission assurent la collecte des blocs et leur transmission sur le réseau extérieur au travers de processeurs d'interfaçage. Les nœuds de stockage assurent la gestion des disques ou des groupes de disques qui leur sont connectés. Le réseau d'interconnexion assure l'acheminement des différentes entités échangées.

Quand une requête arrive au système, elle est affectée à un nœud de transmission qui gèrera son service. Le nœud assurera la collecte et la transmission des différents blocs de la séquence demandée. L'envoi des blocs sur le réseau extérieur commence après une étape d'anticipation qui consiste

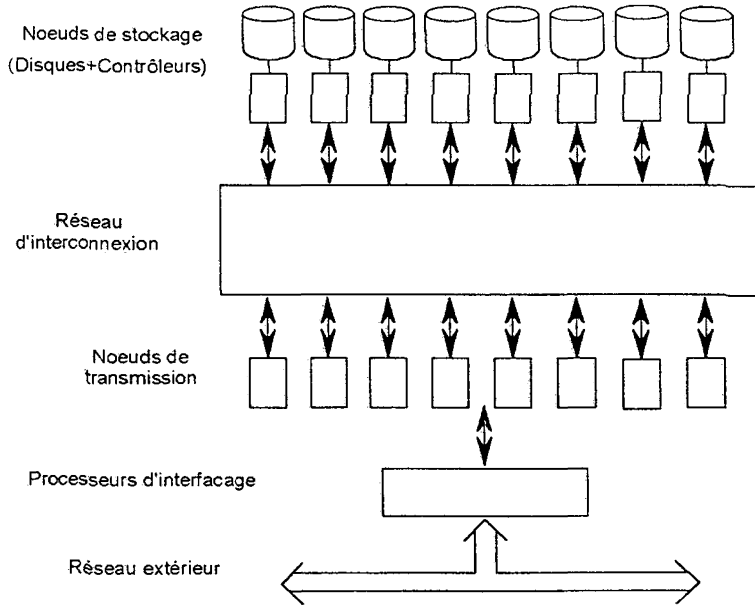


Figure 1. - Architecture du serveur.

à précharger simultanément les A premiers blocs de la séquence considérée. Le nombre d'anticipation A est un paramètre interne de fonctionnement du système qui est fixé au préalable. Après cette étape d'anticipation commence l'étape de transmission. Les blocs sont envoyés sur le réseau extérieur d'une manière cyclique isochrone suivant une période égale à la période $\Delta_{\text{émission}}$ d'émission d'un bloc (taille du bloc/débit MPEG). La collecte des blocs se fait en parallèle en suivant les mêmes cycles. A chaque envoi de bloc sur le réseau extérieur un ordre de chargement du bloc à extraire est envoyé au disque concerné. L'ordre d'extraction généré par le nœud de transmission traverse le réseau d'interconnexion pour atteindre le nœud de stockage puis le disque qui contient le bloc demandé. Une fois le bloc extrait, il est envoyé au nœud de transmission demandeur via le réseau d'interconnexion en empruntant le chemin inverse de l'ordre. Le bloc est alors stocké dans la mémoire du nœud de transmission jusqu'à l'arrivée de sa date d'échéance d'envoi sur le réseau extérieur. Si le bloc arrive après cette date d'échéance, il sera considéré comme perdu. Cela induit une petite interruption au niveau de l'utilisateur de durée égale à la durée d'émission du bloc. Ces cycles sont répétés jusqu'à la fin de la séquence ou la demande d'interruption par l'utilisateur.

Deux politiques de placement des blocs sont principalement utilisées : le placement rotatoire (*round-robin*) et le placement aléatoire (*random*). Dans le placement aléatoire, on fait toutefois en sorte de ne pas placer deux blocs successifs sur le même disque. Les avantages et inconvénients de ces deux politiques sont exposés dans [Tewari 96]. Pour cette étude nous considérons le placement aléatoire qui conduit à des résultats très comparables à ceux obtenus par un placement rotatoire pour le système étudié.

Dans la suite, le critère de performance (probabilité de retard) sera estimé en fonction des deux paramètres internes de fonctionnement : taille des blocs suivant laquelle seront fragmentés les fichiers et nombre de blocs d'anticipation (A).

3. APPROCHE ANALYTIQUE

Nous présentons un modèle analytique simplifié qui permet d'étudier l'effet de ces deux paramètres. Les simulations du fonctionnement exact du système sont très coûteuses en temps de calcul. Le modèle proposé permettra de réaliser rapidement un premier dimensionnement de ces paramètres et d'explorer des domaines de fonctionnement inaccessibles par les simulations.

La symétrie dans l'architecture et le bon équilibrage de la charge entraînent une forte symétrie dans le fonctionnement du système. Les nœuds de stockage peuvent alors être considérés comme équivalents sur le plan opérationnel. Il en est de même pour les nœuds de transmission. Cette observation permet de n'étudier les résultats que pour un couple de nœuds. Nous ferons une approximation d'indépendance et nous étudierons un couple isolé de nœuds en simplifiant les interactions avec les autres couples.

Comme nous nous intéressons à la qualité de la transmission, nous allons essayer d'exprimer analytiquement la probabilité pour que la transmission d'un bloc cause une rupture momentanée de séquence. Pour ce faire nous allons analyser les différentes opérations impliquées dans la transmission d'une séquence.

Suite à la phase d'anticipation, la transmission de la séquence se fait d'une manière cyclique. Un cycle de transmission correspond à l'extraction d'un bloc et son envoi sur le réseau extérieur. Les opérations du cycle de transmission se produisent entre le nœud de transmission qui gère la séquence, le disque contenant le bloc à extraire et le nœud de stockage auquel est associé ce disque. La première opération consiste en la génération et l'envoi de l'ordre d'extraction du nœud de transmission vers le nœud de stockage adéquat. C'est l'étape de traversée du réseau d'interconnexion par

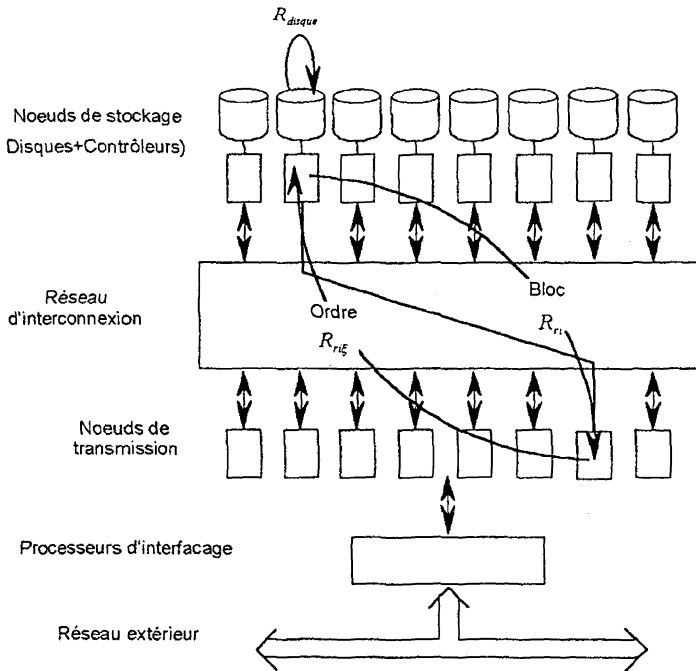


Figure 2. - Cycle de transmission.

l'ordre d'extraction (durée $R_{ri\xi}$). La deuxième opération est l'extraction du bloc du disque (durée R_{disque}). La troisième opération est la traversée du réseau d'interconnexion par le bloc pour aller du nœud de stockage vers le nœud de transmission demandeur (durée R_{ri}). La durée totale du cycle est alors la somme de ces trois v.a. temps de réponse : $R_{ri\xi}$, R_{disque} et R_{ri}).

Il y a une rupture momentanée de séquence quand le bloc n'arrive pas au nœud de transmission à la date prévue pour son envoi sur le réseau extérieur. Ceci se produit lorsque le cycle de transmission dure plus que la période qui lui est allouée Δ_{cycle} . La durée allouée au cycle correspond à l'intervalle de temps qui sépare la date d'envoi de l'ordre d'extraction du bloc de la date prévue pour son envoi sur le réseau extérieur. Cette période n'est autre que le produit de la durée de transmission d'un bloc par le nombre d'anticipation :

$$\Delta_{cycle} = A \times \Delta_{émission} \quad (1)$$

Cette période est constante pour une taille de bloc et un nombre d'anticipation donnés.

Ainsi il y a une rupture momentanée de séquence chaque fois que pour un bloc : $R_{ri\xi} + R_{disque} + R_{ri} > \Delta_{cycle}$. La probabilité π pour que la transmission d'un bloc occasionne une rupture momentanée de séquence s'écrit alors : $\pi = \Pr(R_{ri\xi} + R_{disque} + R_{ri} > \Delta_{cycle})$.

Sachant que la durée de traversée du réseau d'interconnexion est fortement liée à la taille de l'entité traitée et vu la différence de taille entre un ordre d'extraction (64 octets) et celle d'un bloc de donnée (+ 128 Koctets), la durée de traversée de l'ordre d'extraction $R_{ri\xi}$ sera négligée dans la suite de l'étude.

La probabilité de retard s'écrit alors :

$$\pi = \Pr(R_{disque} + R_{ri} > \Delta_{cycle}) \quad (2)$$

Une résolution exacte de l'équation 2 même en faisant l'approximation de l'indépendance de ces deux v.a. est très compliquée. En effet, les densités de probabilité des temps de réponse sont très difficiles à obtenir et particulièrement celle du réseau d'interconnexion.

Nous avons approché le temps de réponse du réseau d'interconnexion par une constante égale à sa moyenne $\bar{R}_{ri}$. L'équation (2) devient : $\pi = \Pr(R_{disque} > \Delta_{cycle} - \bar{R}_{ri})$.

Le problème est alors ramené au calcul de la fonction de répartition du temps de réponse d'un disque, $F_{R_{disque}}$, au point $(\Delta_{cycle} - \bar{R}_{ri})$:

$$\pi = 1 - F_{R_{disque}}(\Delta_{cycle} - \bar{R}_{ri}) \quad (3)$$

Il nous faudra donc obtenir une bonne approximation de la fonction de répartition du temps de réponse disque. Le modèle analytique final du système est donné sur la figure 3. Il comprend deux files en tandem : une file M/D/1 modélisant les disques et un délai $\bar{R}_{ri}$ modélisant le réseau d'interconnexion. La résolution de cette équation ainsi que la justification et validation des hypothèses simplificatrices sont présentées dans l'étude détaillée qui suit.

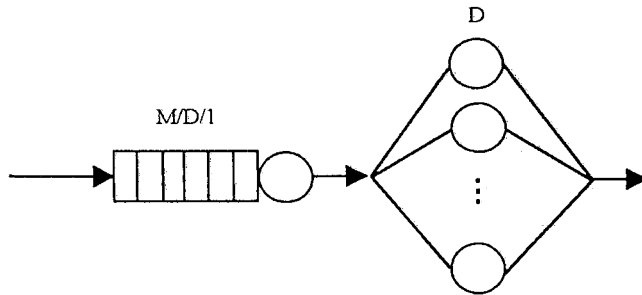


Figure 3. – Modèle analytique final.

4. ÉTUDE DES DISQUES

Présentons dans cette partie l'étude menée sur les composants disques de notre système. Nous commencerons par la présentation de l'architecture des disques, de leur fonctionnement et des différents modèles proposés dans la littérature. Nous présenterons ensuite la démarche adoptée pour notre modélisation, le modèle final simplifié, sa résolution et sa validation.

4.1. Architecture et fonctionnement

Un disque est composé de plusieurs plateaux tournant continûment autour d'un axe, à une vitesse angulaire constante. Les données sont inscrites sur la surface des plateaux et organisées en cercles équidistants ayant l'axe du disque comme centre commun (Fig. 4).

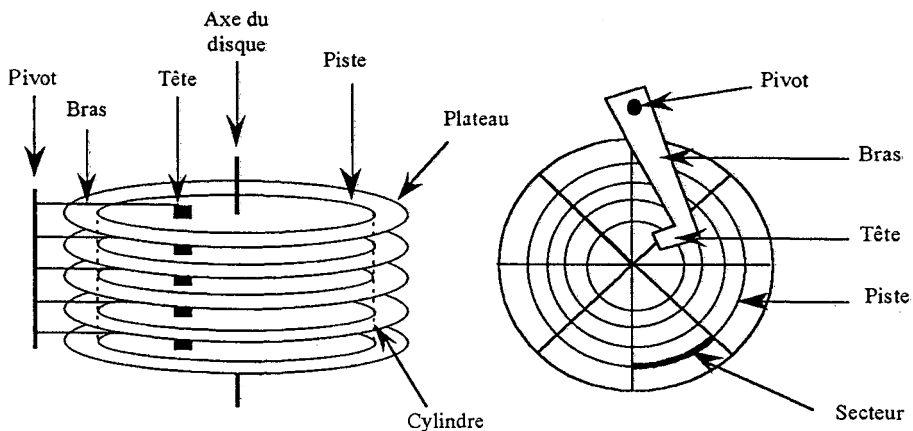


Figure 4. – Architecture d'un disque.

L'opération de lecture (ou d'écriture) de données se fait en trois étapes. La première étape consiste à déplacer la tête de lecture jusqu'à la bonne piste. Elle engendre une latence d'accès T_{seek} (*seek latency*). Ensuite, pour commencer la lecture il faut attendre que la donnée passe sous la tête. Cette étape de positionnement engendre une latence de rotation T_{pos} (*rotation latency*). La troisième étape est l'opération de lecture proprement dite. Cette étape a une durée T_{lect} . La durée totale T_{tot} de l'opération de lecture d'un bloc sur un disque s'écrit alors comme la somme de ces trois durées, $T_{tot} = T_{seek} + T_{pos} + T_{lect}$.

4.2. Modèles de disques

Suivant le degré de précision recherché plusieurs modèles sont proposés dans la littérature. Notons qu'en raison de leur complexité la plupart des modèles qui se veulent précis sont uniquement étudiés par simulation. La modélisation du fonctionnement du disque en lecture consiste en la modélisation des trois étapes (accès, positionnement et lecture) décrites ci-dessus et qui constituent l'opération élémentaire de lecture d'un bloc.

La modélisation de l'opérateur d'accès consiste en la modélisation du mouvement de déplacement du bras entraînant les têtes de lecture. Le mouvement du bras est composé : d'une phase d'accélération, d'une phase uniforme si la vitesse maximale est atteinte et d'une phase de ralentissement. Suivant les caractéristiques du disque (accélération du bras et rayon du disque) et le nombre de cylindres parcourus, ces étapes sont plus ou moins importantes les unes par rapport aux autres. La durée T_{seek} de cette étape d'accès sera le cumul des durées de ces trois phases. Les phases accélérée et décélérée font apparaître un terme en $\sqrt{D}$, la phase uniforme fait apparaître un terme en D où D représente le nombre de cylindres parcourus pendant chacune de ces phases. Il existe dans la littérature principalement trois modèles précis : $T_{seek} = \alpha\sqrt{D} + \beta$ [Hennesy 90] [Scranton 83], $T_{seek} = \alpha\sqrt{D} + \beta D + \gamma$ [Lee Edward 93] et $T_{seek} = \begin{cases} \alpha_1\sqrt{D} + \beta_1, & D < D_0 \\ \alpha_2\sqrt{D} + \beta_2, & D \geq D_0 \end{cases}$ [Ruemmler 94]. Il existe aussi des modèles simples pour lesquels cette étape est remplacée par un temps constant égal à la moyenne de cette opération observée sur plusieurs déplacements. Là aussi plusieurs méthodes sont proposées pour estimer cette moyenne.

La durée de l'opération de positionnement est généralement modélisée par une loi uniforme entre 0 et la durée d'une rotation complète du disque ou simplement par la moyenne de cette loi, soit la moitié de la durée d'une rotation complète.

L'opération de lecture a une durée T_{lec} constante qui dépend du débit de lecture D_{lec} du disque et de la taille du bloc t_{bloc} , $T_{lec} = t_{bloc}/D_{lec}$.

On voit que les modèles précis aboutissent à une partie variable (accès + positionnement) et une partie constante (lecture). Les modèles simples consistent à approcher la durée totale de l'opération par une constante, somme des moyennes des trois opérations.

4.3. Modélisation

Pour notre étude, vu l'importance des disques (goulets d'étranglement) dans les performances du système, nous avons tout d'abord cherché à utiliser pour le temps de service disque un modèle assez précis. Par souci de simplification de la résolution analytique nous avons ramené ensuite ce modèle à un modèle simple à temps de lecture constant. Nous présentons ci-dessous le modèle de départ et nous exposons la démarche de simplification entreprise tout en la justifiant.

4.3.1. Modèle adopté

Chaque disque est modélisé par une file d'attente. Nous supposons que les arrivées d'ordres d'extraction sont poissonniennes. Cette hypothèse est justifiée par le grand nombre de requêtes émises et par le placement aléatoire des blocs sur les disques. Cette hypothèse sera validée par la suite.

Pour le temps de service, le modèle précis est le suivant :

$$T_{tot} = T_{seek} + T_{pos} + T_{lec} \text{ avec } \begin{cases} T_{seek} = \alpha\sqrt{D} + \beta, D > 0 \\ T_{pos} = \text{unif}[0, \gamma] \\ T_{lec} = \chi = t_{bloc}/D_{lec} \end{cases}$$

où D est le nombre de cylindres parcourus par la tête de lecture. α et β sont des constantes déterminées de sorte à retrouver T_{seek} minimal (déplacement d'un cylindre) et T_{seek} maximal (déplacement de tous les cylindres) donnés par le constructeur. γ est la durée d'une rotation entière du disque. Cette durée est directement déduite de la vitesse de rotation donnée par le constructeur. Les données techniques des disques utilisés sont résumées dans le tableau 1. Nous avons considéré les Deskstar3 d'IBM, disques d'actualité adaptés à ce type d'applications [IBM].

TABLE 1
Données techniques des disques

Capacité	2 Go
Rotation	5400 trs/min
D_{tec}	5,2 Mo/s
accès minimal	0,6 ms
accès maximal	17,4 ms
nb de cylindres	2870

D correspond à la différence entre la position du bras au moment de deux lectures consécutives. $D = \Delta X = |X_n - X_{n-1}|$ où X représente le numéro du cylindre sur lequel se trouve le bloc considéré. X est une variable discrète, mais vu le grand nombre de cylindres (~ 3000) nous allons relaxer cette contrainte afin de faciliter les calculs par la suite. Si l'on considère que la densité d'information est la même sur tout le disque, on peut alors prendre pour X une loi uniforme entre 0 et N , $N + 1$ étant le nombre total de cylindres. Les ΔX successifs vont être corrélés. On peut néanmoins considérer qu'ils sont indépendants et supposer que X_n et X_{n-1} suivent deux lois uniformes indépendantes entre 0 et N .

Nous aboutissons ainsi à une file M/G/1 avec un service $B = T_{tot}$. Rappelons que le but de cette étape est de déterminer la densité de probabilité du temps de réponse des disques afin de calculer le critère de performance : probabilité de retard (cf. équation 3).

4.3.2. Étude du temps de service

Dans ce qui suit, étudions la densité de probabilité $b(x)$ du temps de service et ses premiers moments. Le temps d'accès peut se mettre sous la forme $T_{seek} = \Theta_s + \beta$ avec $\Theta_s = \alpha\sqrt{\Delta X}$. Il vient alors : $f_{\Theta_s}(y) = \frac{2y}{\alpha^2} f_{\Delta X}\left(\left(\frac{y}{\alpha}\right)^2\right)$. Sachant que : $f_{X_n}(x) = \frac{1}{N} 1_{\{0 \leq x \leq N\}}$ on trouve : $f_{\Delta X}(x) = \frac{2(N-x)}{N^2}$. D'où : $f_{\Theta_s}(y) = \frac{2y}{\alpha^2} \left\{ \frac{2}{N} - \frac{2}{N^2} \left(\frac{y}{\alpha}\right)^2 \right\}$ et $F_{\Theta_s}(y) = \frac{2}{\alpha^4 N} y^2 - \frac{1}{\alpha^2 N^2} y^4$. On en déduit : $E[\theta_s] = \frac{8}{15} \alpha \sqrt{N}$, $\text{Var}[\theta_s] = \frac{11}{225} \alpha^2 N$.

Pour le temps de positionnement T_{pos} , on a immédiatement :
 $f_{T_{pos}}(x) = \frac{1}{\gamma} 1_{\{0 \leq x \leq \gamma\}}$, ainsi que les premiers moments : $E[T_{pos}] = \frac{\gamma}{2}$,
 $\text{Var}[T_{pos}] = \frac{\gamma^2}{12}$.

On obtient alors les premiers moments de la distribution du temps de service :

$$E[B] = \frac{8}{15} \alpha \sqrt{N} + \beta + \frac{\gamma}{2} + \chi \quad (4)$$

$$\text{Var}[B] = \frac{11}{225} \alpha^2 N + \frac{\gamma^2}{12} \quad (5)$$

La densité de probabilité de $\Theta_s + T_{pos}$ se calcule par convolution :
 $f_{\Theta_s + T_{pos}}(x) = f_{\Theta_s}(x) \otimes f_{T_{pos}}(x)$. Pour les valeurs numériques considérées, on a $\gamma \leq \alpha \sqrt{N}$, il vient alors :

$$f_{\Theta_s + T_{pos}}(x) = \begin{cases} \frac{1}{\gamma} \left[\frac{2}{\alpha^4 N} x^2 - \frac{1}{\alpha^2 N^2} x^4 \right], & 0 \leq x \leq \gamma \\ \frac{1}{\gamma} \left[\left\{ \frac{2}{\alpha^2 N} x^2 - \frac{1}{\alpha^4 N^2} x^4 \right\} - \left\{ \frac{2}{\alpha^2 N} (x - \gamma)^2 - \frac{1}{\alpha^4 N^2} (x - \gamma)^4 \right\} \right], & \gamma \leq x \leq \alpha \sqrt{N} \\ \frac{1}{\gamma} \left[1 - \left\{ \frac{2}{\alpha^2 N} (x - \gamma)^2 - \frac{1}{\alpha^4 N^2} (x - \gamma)^4 \right\} \right], & \alpha \sqrt{N} \leq x \leq \gamma + \alpha \sqrt{N} \end{cases}$$

La densité de probabilité du temps de service s'obtient par :

$$b(x) = f_{\Theta_s + T_{pos}}(x - \beta - \chi) \quad \text{pour } \beta + \chi \leq x \leq \alpha \sqrt{N} + \beta + \gamma + \chi \quad (6)$$

4.3.3. Résolution du modèle

Soient λ le taux d'arrivée, $b = E[B]$ le temps moyen de service et $\rho = \lambda b$ la charge de la file. Nous noterons par R et W respectivement le temps de réponse et le temps d'attente de la file. pour alléger les équations on posera systématiquement $Y(x)$ la fonction de répartition de la variable aléatoire Y , $y(x)$ sa densité de probabilité et $Y^*(s)$ sa transformée de Laplace. On notera les différents moments de la distribution de la manière suivante : $y_i = E[Y^i]$.

On a alors (formule de Pollaczek-Kinchin) : $R^*(s) = B^*(s)W^*(s) = \frac{(1 - \rho)sB^*(s)}{s - \lambda + \lambda B^*(s)}$. Cette transformée de Laplace est le plus souvent impossible

à inverser formellement. C'est le cas, en l'occurrence pour le service décrit par le système (6). Pour obtenir la probabilité de retard nous allons approcher le temps de service par un temps de service déterministe égal à sa moyenne (cf. équation 4) et de là calculer la fonction de répartition du temps de réponse. La probabilité de retard est alors donnée par l'équation (3). Le modèle du disque est ainsi ramené à une file M/D/1. On validera l'approximation en comparant la fonction de répartition obtenue à celle donnée par simulation sur le modèle exact.

Dans le cas de la file M/D/1, on peut facilement calculer la valeur de la fonction de répartition du temps de réponse aux points kb . Pour cela, constatons que si la file est vide quand un client arrive, son temps de réponse sera égal à b . S'il y a k clients à son arrivée, son temps de réponse vaudra : $kb + \Omega$ où Ω est le temps résiduel de service.

Comme Ω est compris strictement entre 0 et b , le temps de réponse sera compris strictement entre kb et $(k+1)b$. par conséquent, si l'on note π_k , la probabilité limite stationnaire qu'il y ait k clients à l'arrivée du client, on aura :

$$\begin{cases} R((k+1)b) - R(kb) = \pi_k, & k \geq 1 \\ R(b) = \pi_0 \end{cases}$$

D'où :

$$R(kb) = \sum_{j=0}^{k-1} \pi_j, \quad k \geq 1 \quad (7)$$

Les π_k sont données par [Jain 92] :

$$\begin{cases} \pi_0 = 1 - \rho \\ \pi_1 = (1 - \rho)(e^\rho - 1) \\ \pi_k = (1 - \rho) \sum_{j=0}^k \frac{(-1)^{k-j} (j\rho)^{k-j-1} (j\rho + k - j) e^{j\rho}}{(k-j)!}, & k \geq 2. \end{cases}$$

Pour avoir la fonction de répartition en un point autre que les multiples de b , procédons par interpolation :

$$\text{Log}(R(kb + v)) = \frac{b-v}{b} \text{Log}(R(kb)) + \frac{v}{b} \text{Log}(R((k+1)b)), \quad 0 \leq v \leq b.$$

Avant de passer à l'évaluation du modèle retenu, nous présentons les arguments qui justifient *a priori* la simplification effectuée.

4.3.4. Justification de la simplification

Pour bien utiliser un disque il faut choisir une taille de bloc assez grande car le débit global du disque croît avec la taille des blocs manipulés. En

effet, la durée de l'opération de lecture croît avec la taille du bloc alors que celles des opérations d'accès et de positionnement n'en dépendent pas. Par conséquent, la proportion de la durée de l'opération de lecture p_{lec} dans le temps de service disque croît avec la taille des blocs. Le disque est alors mieux utilisé et le débit global augmente.

Le débit global du disque, $D_g = t_{bloc}/E[T_{seek} + T_{pos} + T_{lec}]$, s'écrit : $D_g = p_{lec} \times D_{lec}$ avec $p_{lec} = E[T_{lec}]/E[T_{seek} + T_{pos} + T_{lec}]$. En développant cette expression on trouve que p_{lec} s'écrit comme une fonction hyperbolique croissante de t_{bloc} :

$$p_{lec} = t_{bloc}/(t_{bloc} + \nu), \nu \text{ constante} \quad (8)$$

p_{lec} admet 1 comme limite supérieure, ce qui correspond à un débit global maximum égal au débit de lecture D_{lec} , pour une taille de bloc très élevée.

Une borne inférieure de la taille de bloc peut être définie en se fixant un seuil minimum de la proportion de lecture du disque. Pour notre application un seuil de 0,6, qui correspond à une taille de bloc de 128 Ko, semble raisonnable. Dans la suite de l'étude nous nous intéresserons à des tailles supérieures à ce seuil.

Pour évaluer l'effet de la partie variable dans le temps total de service du disque, nous avons tracé son coefficient de variation $\sqrt{\text{Var}[B]}/E[B]$ en fonction de la taille des blocs. Cette variance a été calculée grâce aux équations (4) et (5). Ce coefficient décroît avec la taille des blocs et est inférieur à 0,12 pour les tailles qui nous intéressent (supérieures à 128 Ko). Ceci appuie l'hypothèse simplificatrice qui consiste à prendre le temps de service constant.

Dans le même souci de justification de l'approximation effectuée, nous présentons ci-dessous un tableau comparatif des quatre premiers moments du temps de réponse des disques pour chacun des deux modèles : modèle précis et modèle simplifié déterministe sous l'hypothèse d'arrivées poissonniennes. Les moments de la distribution du temps de réponse sont obtenus en dérivant successivement la transformée de Laplace suivante, formule de

$$\text{Takacs [Takacs 62]} : \begin{cases} w_k = \frac{\lambda}{1-\rho} \sum_{j=1}^k \binom{k}{j} \frac{b_{j+1}}{j+1} w_{k-j} \\ w_0 = 1 \end{cases}$$

Les différents moments du temps de réponse valent alors :

$$r_k = \sum_{j=0}^k \binom{k}{j} b_j w_{k-j}$$

Pour un temps de service déterministe, nous avons : $b_k = b_1^k = b^k$

Pour le modèle précis, notons ν_k les moments de la distribution de $T_{pos} + \theta_s$, γ_k ceux de T_{pos} et η_k ceux de θ_s . Il vient :

$$b_k = \sum_{j=0}^k \binom{k}{j} \nu_j (\beta + \chi)^{k-j} \text{ où } \nu_k = \sum_{j=0}^k \binom{k}{j} \gamma_j \eta_{k-j}$$

avec

$$\gamma_k = \frac{\gamma^k}{k+1} \text{ et } \eta_k = 4 \left\{ \frac{1}{k+2} - \frac{1}{k+4} \right\} (\alpha\sqrt{N})^k$$

On voit sur le tableau 2 que les moments pondérés du temps de réponse pour le modèle précis et pour le modèle déterministe sont assez proches; ce qui vient encore appuyer l'approximation.

TABLE 2
Premiers « moments » r_k/b^k du temps de réponse

	$\rho = 0.5$		$\rho = 0.7$		$\rho = 0.9$	
Moment	Précis	Détermi.	Précis	Détermi.	Précis	Détermi.
1	1.507	1.500	2.183	2.167	5.562	5.500
2	2.889	2.833	6.988	6.833	54.894	53.500
3	7.275	7.000	31.774	30.555	806.955	775.000
4	23.713	22.367	191.433	181.004	15813.0	14695.3

4.5. Évaluation du modèle simplifié

Nous passons maintenant à l'évaluation du modèle simplifié en comparant le temps de réponse qu'il engendre au temps de réponse obtenu par le modèle précis. Pour ce faire nous comparons la fonction de répartition du temps de réponse calculée analytiquement pour M/D/1 (cf. équation 7) à celle obtenue par simulation sur le modèle précis avec arrivées poissonniennes.

Nous avons dressé plusieurs familles de courbes pour différentes charges. La figure 5 présente deux de ces familles. Ces courbes montrent que le modèle simplifié donne des résultats assez proches de ceux obtenus sur le modèle précis. Nous nous sommes limités aux valeurs élevées de la fonction de répartition du temps de réponse du disque puisque nous nous intéressons à un fonctionnement avec de faibles probabilités de retard de blocs (cf. équation 3).

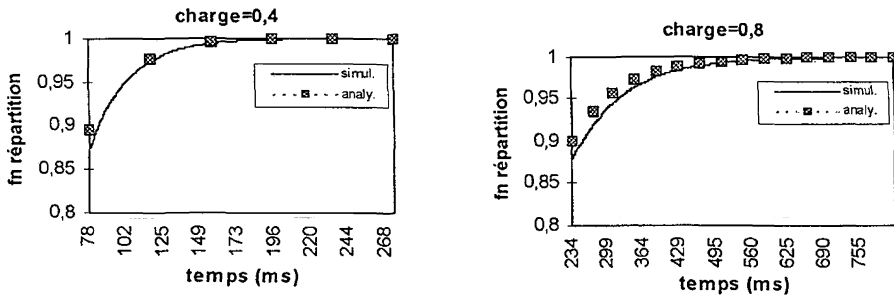


Figure 5. – Fonction de répartition du temps de réponse.

5. ÉTUDE DU RÉSEAU D'INTERCONNEXION

Le développement des architectures parallèles a beaucoup contribué au développement des réseaux d'interconnexion. Les réseaux d'interconnexion permettent de connecter les différents composants d'une machine parallèle : processeurs, mémoires, disques, etc. Nous nous intéressons dans cette étude à un réseau multiétage de type Oméga. Une étude bibliographique sur les réseaux multiétages est présentée dans [Siegle 87].

5.1. Architecture et fonctionnement

Le réseau multiétage considéré est un réseau Oméga constitué par des éléments de commutation (commutateurs) interconnectés et organisés en n étages. C'est un réseau monochemin, entre une entrée et une sortie quelconques il existe un chemin unique. Les éléments de commutation considérés ont 2 entrées et 2 sorties. Le nombre d'entrées est égal au nombre de sorties et égal à 2^n .

Nous supposons un fonctionnement avec un routage de messages (store-and-forward). Les blocs avancent dans le réseau vers leurs destinations en transitant dans les commutateurs intermédiaires. Lorsqu'un bloc arrive à un commutateur, il est stocké dans sa mémoire jusqu'à libération du lien à emprunter. A chaque étape le lien emprunté est aussitôt libéré. La taille des blocs implique que le délai de traitement induit dans la traversée d'un étage est principalement dû au temps de transmission sur le lien.

5.2. Modélisation

Les critères de performances généralement recherchés lors de l'étude de tels réseaux d'interconnexion sont : le temps de traversée et le débit

maximum du système. Pour l'architecture retenue le débit maximal est égal au débit des liens multiplié par le nombre d'entrées si nous considérons qu'il n'y a pas de perte causée par la limitation des mémoires des commutateurs.

Le système peut être modélisé par un réseau de files d'attente. Chaque commutateur est modélisé par deux files à la sortie. Le temps de service de ces files correspond alors au temps de traversée sur un lien. Pour une taille de bloc donnée t_{bloc} ce service est constant égal à $T_{lien} = t_{bloc}/D_{lien}$, D_{lien} étant le débit du lien.

La résolution analytique d'un tel modèle est rendue difficile à cause de la dépendance entre les étages. Alors que nous trouvons beaucoup de travaux sur les réseaux à fonctionnement synchrone à cause de leurs utilisation dans les télécommunications (ATM), peu de travaux sont disponibles pour les réseaux d'interconnexion à temps continu qui nous intéressent. Les études disponibles sont généralement limitées à la détermination du temps moyen de traversée. [Mohapatra 96] présente un bref état de l'art sur le sujet.

Pour cette étude nous avons considéré des commutateurs avec des débits bidirectionnels de 20 Mo/s. On trouve actuellement sur le marché des commutateurs avec des débits bidirectionnels de 50 Mo/s (l'Elite de Pci) voire 200 Mo/s (Paragon Mesh Routing Chip d'Intel) [Intel 91].

Notons qu'avec les données techniques d'actualité choisies pour les disques et pour les commutateurs, la charge du réseau d'interconnexion reste inférieure à 0,2. Le débit global maximum d'un disque est égal à $D_{lec} = 5,2 \text{ Mo/s}$ ou D_{lec}/t_{bloc} (blocs/s) (cf. équation 8). Le débit maximum de l'ensemble des disques est alors : $nb. \text{ disques} \times D_{lec}/t_{bloc}$ (blocs/s). Le débit d'un lien de commutateur est $D_{lien} = 20 \text{ Mo/s}$. Le débit du réseau en entier est égal au nombre de ports d'entrée multiplié par ce débit ou encore $nb. \text{ nœuds de stockage} \times D_{lien}/t_{bloc}$ (blocs/s). Si nous considérons qu'il y a un disque par nœud de stockage on aura : $nb. \text{ disques} = nb. \text{ nœuds de stockage}$ et le goulet d'étranglement est alors situé au niveau des disques. Pour une charge de disque inférieure à 0,8 la charge du réseau d'interconnexion est inférieure à 0,2.

Si on veut utiliser le réseau d'interconnexion avec une charge supérieure, on peut augmenter le nombre de disques par nœud de stockage. Un tel fonctionnement induirait des conflits au niveau du nœud de stockage et une dégradation du temps de traversée du réseau d'interconnexion. L'utilisation du réseau avec une charge forte induit des temps de traversée élevés et dispersés à cause des problèmes de contention qui détériorent les performances du système [Lee Gyungho 89].

Compte tenu du matériel disponible actuellement (performances/coût), l'architecture choisie est telle que le goulet d'étranglement de ces systèmes se trouve au niveau des disques. Pour optimiser l'utilisation de l'ensemble du serveur, on est amené à utiliser le réseau d'interconnexion avec une charge faible. Les simplifications qui suivent sont alors généralement applicables pour tous les serveurs ayant cette architecture.

L'étude du réseau d'interconnexion dans les conditions de fonctionnement qui nous intéressent (faible charge et confrontation avec les disques) montre que nous pouvons réduire le réseau en entier à un simple délai égal à la moyenne de son temps de traversée. En effet pour les charges considérées nous avons un coefficient de variation qui est inférieur à 10^{-1} . La comparaison de la variance de la durée de traversée du réseau d'interconnexion et de celle du temps de réponse des disques montre que la part de variation induite par le réseau dans la durée totale de traversée (cycle de transmission) est négligeable par rapport à celle causée par les disques. Le rapport de ces variances varie entre 10^{-3} et 10^{-2} pour les valeurs de charge qui nous intéressent. Cette variation dans le temps de traversée du réseau d'interconnexion peut alors être négligée. Ces observations nous ont amenés à modéliser le réseau d'interconnexion par un délai constant égal à la moyenne de son temps de traversée.

5.3. Calcul du délai de traversée et validation

L'étude détaillée du réseau d'interconnexion montre que le temps moyen de traversée d'un étage se stabilise pour devenir constant à partir d'un certain rang [Mohapatra 96]. L'ensemble du réseau d'interconnexion est modélisé par un délai de valeur constante égale au nombre d'étages multiplié par le temps de réponse moyen d'une file M/D/1 modélisant un étage. Cette approximation est validée par comparaison avec la simulation. Les résultats présentés sur la figure 6 sont relatifs à un réseau à 64 entrées/sorties, soit 6 étages de commutateurs, pour des blocs de 128 Ko.

On constate (*Fig. 6*) que l'erreur commise dans l'estimation du temps de traversée du réseau d'interconnexion en utilisant l'approximation M/D/1 pour la détermination du temps de traversée de chaque étage est faible dans les conditions de fonctionnement considérées (erreur relative maximale = 4%).

6. MODÈLE DE SIMULATION

Pour valider cette approche analytique nous avons développé un modèle précis du système en entier sous forme de réseau de files d'attente. Le

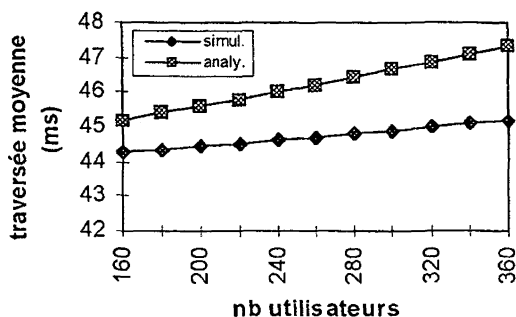


Figure 6. – Temps moyen de traversée global du réseau d'interconnexion.

modèle tient compte des principales opérations qui interviennent dans le fonctionnement du système. L'approche de modélisation a été validée dans une étude antérieure [kaddeche 96] sur un prototype développé à cet effet. Les simulations ont été réalisées avec le simulateur QNAP2 (*Queueing Networks Analysis Package*) [POT 84].

La résolution est effectuée en réseau fermé : nous avons un nombre fixe d'utilisateurs simultanés constamment actifs. Dès que le système a terminé avec la transmission d'une séquence, fin de la séquence ou demande d'interruption par l'utilisateur, une nouvelle requête est générée. Cette requête peut provenir de ce même utilisateur ou d'un utilisateur nouvellement connecté. Le choix de la séquence à visualiser est effectué suivant une loi uniforme modélisant la demande des utilisateurs. Pour avoir une bibliothèque de séquences assez variée, nous avons considéré 300 séquences dont les durées sont fixées au départ suivant une loi exponentielle de moyenne 10 minutes. Pour tenir compte du comportement de l'utilisateur nous avons considéré qu'il interrompt la visualisation de la séquence, si celle-ci n'est pas terminée, au bout d'un temps qui suit une loi exponentielle de moyenne 6 minutes. Les simulations sont effectuées sur un serveur de 64 disques à raison d'un disque par nœud de stockage.

La probabilité de retard est obtenue en divisant le nombre de blocs arrivés en retard par le nombre total de blocs arrivés.

Les durées des simulations sont fixées de manière à avoir un intervalle de confiance relatif, pour la probabilité de retard, de moins de 25% pour un niveau de confiance de 95% pour les simulations avec la taille de bloc de 128 Ko et de moins de 75% pour les simulations avec les tailles de blocs supérieures.

7. VALIDATION ET UTILISATION DU MODÈLE ANALYTIQUE

La figure 7 présente les probabilités de retard de blocs obtenues par simulation (courbes en pointillé) et par la méthode analytique approchée (courbes en continu). La taille des blocs est de 128 Ko. Le nombre d'anticipation varie entre 1 et 3. Les nombres d'utilisateurs considérés varient entre 160 et 340 ce qui pour la taille de blocs considérée représente une charge du disque comprise entre 0.4 et 0.85.

La comparaison entre la simulation et le modèle analytique approché peut être effectuée lorsque les simulations sont possibles, c'est à dire lorsque les retards ne sont pas trop rares. Cette comparaison montre que le modèle analytique approché donne des résultats valides.

L'analyse de ces courbes montre que les probabilités de retard augmentent avec la charge et diminuent lorsque le nombre d'anticipation augmente. Lorsque la charge augmente les temps de réponse des disques et du réseau d'interconnexion augmentent. La durée du cycle de transmission (extraction sur disque + traversée du réseau d'interconnexion) du bloc augmente, ce qui provoque l'augmentation du nombre de blocs retardés. L'anticipation permet d'augmenter la durée allouée à l'opération de transmission d'un bloc qui est directement proportionnelle à ce nombre (cf. équation 1). Ainsi l'augmentation du nombre d'anticipation favorise le bon déroulement de l'opération de transmission.

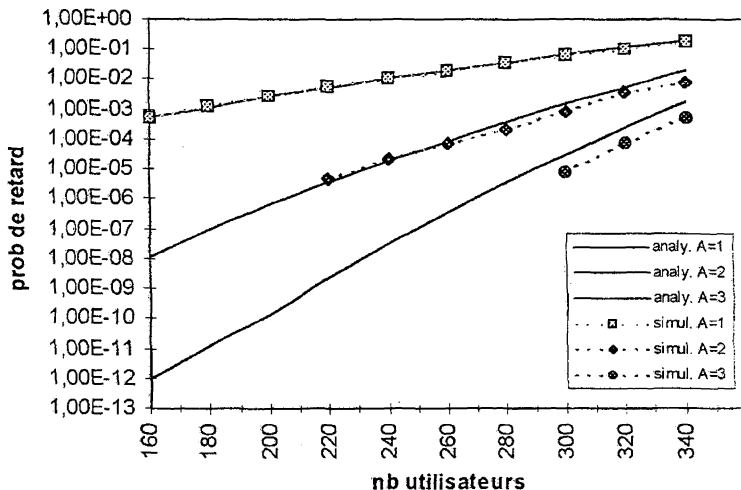


Figure 7. - Comparaison méthode analytique - Simulations.

Pour une taille de bloc de 128 Ko on peut fixer la probabilité de retard maximum à 10^{-4} , ce qui se traduit au niveau de l'utilisateur par au maximum une interruption de 250 ms toutes les 41 minutes, interruption non sensible pour ce type d'application. On voit sur la figure 7 que de telles probabilités de retard sont obtenues pour une charge inférieure à 260 et 300 utilisateurs respectivement pour 2 et 3 anticipations. Avec un seul bloc anticipé les probabilités de retard sont toujours supérieures à ce seuil pour le domaine d'utilisation étudié. Pour une taille de bloc donnée, on choisit ainsi le nombre d'anticipation qui permet d'atteindre une charge (en nombre d'utilisateurs) avec une probabilité de retard seuil fixée. Notons qu'on peut calculer une borne supérieure théorique de la charge par le rapport du débit total de tous les disques par le débit requis pour une séquence. Pour les données techniques considérées cette borne vaut 400 utilisateurs.

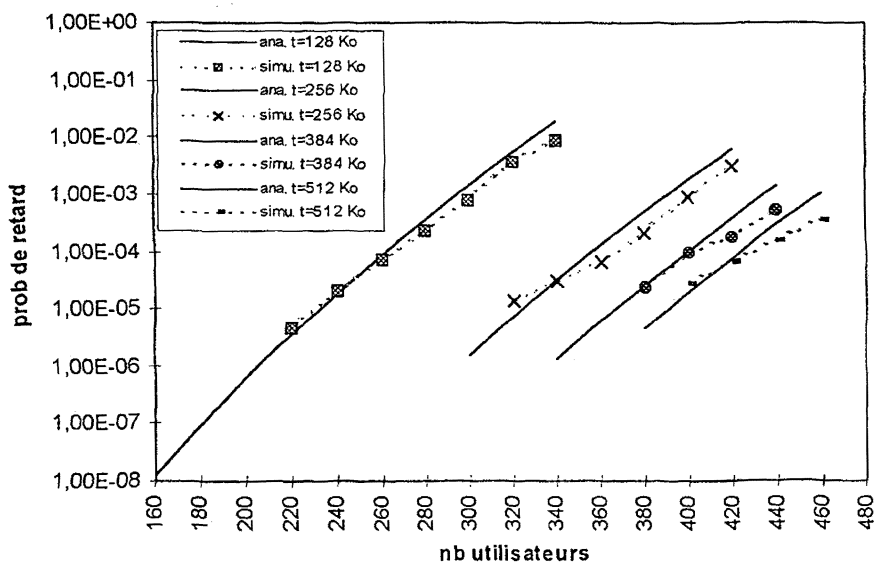


Figure 8. – Effet de la taille des blocs sur la probabilité de retard.

Pour étudier l'effet de la taille des blocs, nous considérons un nombre d'anticipation égal à 2. La figure 8 présente la probabilité de retard en fonction du nombre d'utilisateurs pour des tailles de blocs de 128, 256, 384, 512 Ko. Les courbes en pointillé correspondent aux résultats des simulations, les courbes en trait plein correspondent aux résultats obtenus par le modèle analytique approché.

Sur les diverses courbes, on peut estimer la précision de la méthode analytique. Cette dernière nous permet dans tous les cas d'obtenir rapidement un ordre de grandeur de la probabilité de retard.

Lorsque la taille des blocs augmente, le débit global des disques augmente et, ainsi, la probabilité de retard diminue. On remarque que le gain en charge est plus important pour les premières variations de la taille (premières courbes plus espacées). Ceci s'explique en partie par le fait que le débit des disques croît d'une façon hyperbolique en fonction de la taille des blocs (cf. équation 8). Il n'est donc pas intéressant d'accroître indéfiniment la taille des blocs d'autant que l'on est limité par le coût des mémoires dans les disques et dans les autres composants du système. Une trop grande taille de blocs pourrait aussi remettre en cause les propriétés de l'architecture à savoir le bon équilibrage de la charge sur les disques et sur le réseau d'interconnexion. Pour les valeurs des paramètres que nous avons considérées, une taille de blocs de 128Ko semble raisonnable.

8. CONCLUSION

Nous avons présenté dans cette étude un modèle analytique simple pour l'analyse des performances d'une architecture générale de serveurs multimédias. Il s'agit d'une architecture multidisque formée par un ensemble de nœuds de stockage reliés à un ensemble de nœuds de transmission par un réseau d'interconnexion.

Nous nous sommes intéressés à la qualité de transmission exprimée en terme de probabilité de retard des blocs de données qui induisent des petites interruptions au niveau de l'utilisateur. Vu la complexité du système étudié, nous avons essayé de trouver un compromis entre la complexité du modèle et sa précision. Pour ce faire, nous avons étudié en détail chacun des composants en essayant de simplifier au maximum sa modélisation. A chaque étape, les approximations effectuées sont justifiées et validées. Le modèle final du système complet est à son tour validé. Les validations sont effectuées par de longues simulations réalisées sur des modèles précis sous forme de réseau de files d'attente. Notons que la validation a été effectuée avec les données techniques de matériel performant actuellement disponible sur le marché. Dans le modèle final les disques sont représentés par des files M/D/1 et le réseau d'interconnexion par un délai constant.

Le modèle analytique est alors utilisé pour dimensionner les paramètres internes de fonctionnement du système. Nous avons montré que l'architecture présente un bon comportement multi-utilisateur et que la technique

d'anticipation est intéressante. L'augmentation de la taille des blocs améliore les performances du système jusqu'à un seuil donné au dessus duquel la faible amélioration ne justifie plus le coût des mémoires nécessaires aux composants du système.

La méthode analytique présentée est d'un intérêt assez général et s'adapte à beaucoup de variantes de l'architecture proposée. Nous travaillons actuellement sur des approximations moins fortes pour le temps de service disque. Des études pourraient être faites pour améliorer le modèle du réseau d'interconnexion.

BIBLIOGRAPHIE

- [Berson 94] S. BERSON, S. GHANDEHARIZADEH, R. MUNTZ, et X. JU., Staggered striping in multimedia information systems. *Proceedings of the Fifth Intl. Conf. On management of data*, May 1994.
- [Damm 94] G. DAMM, G. BABONNEAU, A. L. BEYLOT et M. BECKER, Performance evaluation of a multimedia server for ATM networks, *Proceedings of the 27th Annual Simulation Symposium*, La Jolla, April 1994, p. 41-50.
- [Ganger 94] Gregory R. GANGER, Bruce L. WORTHINGTON, Robert Y. HOU, et Yale N. PATT. Disk Arrays, High-Performance, High-Reliability Storage Subsystems. *IEEE Computer*, March 1994, p. 30-36.
- [Gemmell 95] D. James GEMMELL, Harrick M. VIN, Dilip D. KANDLUR, P. VENKAT RANGAN, Multimedia Storage Servers : A Tutorial and Survey. *IEEE Computer*, May 1995.
- [Haskin 95] Roger HASKIN et Frank L. STEIN, A system for delivery of interactive television programming. *COMPCON*, Spring 1995.
- [Hennesy 90] J. L. HENNESY et PATTERSON, Computer architecture a quantitative approach. Morgan Kauffman, 1990.
- [IBM] <http://www.ibm.com>, <http://www.storage.ibm.com/oem/tinfo.html>.
- [Intel 91] Intel Corporation. PARAGON XP/S, product overview, 1991.
- [ISO 93] ISO/IEC JTC1/SC29/WG11, N0601. Coding of Moving Pictures and Associated Audio. Nov., 1993.
- [Jain 92] Raj JAIN, The art of computer systems performance analysis. Wiley, wiley professional computing. 1992.
- [Kaddeche 96] H. KADDECHE, G. DAMM, G. BABONNEAU et M. BECKER, Design and performance analysis of a multimedia server for high speed networks. In *Proceedings of TDP96*, p. 359-374, La Londe les Maures, France, June 1996.
- [Lee Edward 93] Edward K. LEE, Randy H. KATZ, An analytic performance model of disk arrays. In *Proceedings of the ACM Sigmetrics*, May 1993, p. 98-109.
- [Lee Gyungho 89] Gyungho LEE, A performance bound of multistage combining networks. *IEEE Transactions on computers*, 38, n° 10, October 1989.

- [Livny 87] M. LIVNY, S. KHOSHAFIANS et H. BORAL, Multi-disk management algorithms, *Proceedings of the 1987 ACM SIGMETRICS*, May 1987, p. 69-77.
- [Mohapatra 96] Parsant MOHAPATRA et Chita R. DAS, Performance analysis of finite-buffered asynchronous multistage interconnection networks. *IEEE Transactions on parallel and distributed systems*, 7, n° 1, January 1996.
- [Potier 84] D. POTIER et M. VERAN, QNAP2 : a Portable Environment for Queuing Systems Modelling, *International Conference on Modelling and Performance Analysis Tools*, May 1984.
- [Ruemmler 94] Chris RUEMMLER et John WILKES, An introduction to disk drive modeling. *IEEE Computer*, March 1994, p. 17-28.
- [Scranton 83] R. A. SCRANTON, D. A. THOMPSON, et D. W. HUNTER. The access time myth. IBM Res. Rep. RC10197, Sept. 1983.
- [Siegle 87] H. J. SIEGLE, T. SCHWEDERSKI, D. G. MEYER, Hsu TSUN-YUNK, Large scale parallel processing systems. Microprocessors and microsystems. 11, n° 1, Jan./Feb. 1987, p. 3-20.
- [Takacs 62] L. TAKACS. Introduction to the theory of queues. Oxford University Press, New York, 1962.
- [Tewari 96] Renu TEWARI, Rajat MUKHERJEE, Daniel M. DIAS, Harrik M. VIN, Design and Performance tradeoffs in clustered video servers. In *The proceedings of IEEE-ICMCS-Tokyo*, June 1996.
- [Vetter 95] Ronald VETTER, ATM concept, architecture and protocols. CACM, February 1995.
- [Vin 94] Harrick M. VIN, Multimedia systems architecture, In *Defining the Global Information Infrastructure: Infrastructure, Systems, and Services*, Ed. Stephen F. Lundstrom, CR56, SPIE Press, November 1994, p.287-296.