

EVALUATION DE PERFORMANCES
ET
APPLICATIONS AUX RESEAUX

PREMIERE PARTIE :

THEORIE DES PROBABILITES
CHAINES DE MARKOV
THEORIE DES FILES D'ATTENTE
SIMULATIONS A EVENEMENTS DISCRETS

Olivier Brun - Urtzi Ayesta

TABLE DES MATIÈRES

1	Introduction	5
1.1	Evaluation des performances des réseaux	5
1.2	Exemple introductif	6
1.3	Méthodes d'évaluation de performances	7
1.4	Plan du cours	8
2	Rappels sur la Théorie des Probabilités	9
2.1	Événement	9
2.2	Probabilité d'un événement	10
2.2.1	Probabilité conditionnelle	10
2.2.2	Loi des probabilités totales	11
2.2.3	Formule de Bayes	11
2.2.4	Indépendance	12
2.3	Variables aléatoires	12
2.3.1	Variables aléatoires continues	13
2.3.2	Variables aléatoires discrètes	13
2.3.3	Cas de plusieurs variables aléatoires	14
2.3.4	Moments d'une variable aléatoire	14
2.4	Quelques distributions de probabilités courantes	15
2.4.1	Distributions discrètes	15
2.4.2	Distributions continues	17
3	Chaînes de Markov	23
3.1	Processus Stochastiques	23
3.2	Processus Markovien	24
3.3	Chaîne de Markov à temps discret	25
3.3.1	Probabilités de transition	25
3.3.2	Probabilités d'état	26

3.3.3	Classification des états	27
3.3.4	Distribution stationnaire	28
3.4	Chaînes de Markov à temps continu	30
3.5	Processus de naissances et de morts	32
3.5.1	Distribution stationnaire	32
3.5.2	Distribution transitoire	33
3.6	Processus de Poisson	34
4	Théorie des files d'attente	39
4.1	Introduction	39
4.2	Les concepts de base des files d'attente	40
4.2.1	La formule de Little	44
4.3	Files d'attente Markoviennes	47
4.3.1	La file $M/M/1$	47
4.3.2	La file $M/M/1/K$	49
4.3.3	La file $M/M/C$	50
4.3.4	La file $M/M/C/C$	52
4.4	Exemples d'application	53
4.4.1	Comparaison de systèmes	53
4.4.2	Etude d'un mécanisme de token bucket	55
4.5	Files d'attente à loi de service générale	56
4.5.1	La file $M/G/1/FCFS$	56
4.5.2	La file $M/G/1/PS$	61
4.6	Réseaux de files d'attente	61
4.6.1	Les réseaux de Jackson	62
4.6.2	Réseaux fermés	64
5	Simulation à Evénements Discrets	69
5.1	Introduction	69
5.2	Génération de nombres aléatoires	69
5.2.1	Générateurs Uniformes	70
5.2.2	Tests sur les générateurs uniformes	73
5.2.3	Générateurs non uniformes	75
5.3	Simulation événementielle	77
5.3.1	Algorithme de simulation	77
5.3.2	Exemple de simulation d'une file $M/M/1$	79
5.3.3	Régime transitoire et régime permanent	82
5.3.4	Estimation des performances en régime permanent	83
5.4	Conclusion	89
5.5	Bibliographie	91

TABLE DES FIGURES

1.1	Point de vue système.	5
1.2	Les trois facteurs clés.	6
1.3	Liaison téléphonique ayant N circuits.	7
2.1	Fonction de répartition de la distribution exponentielle.	18
2.2	Distribution normale.	19
3.1	Convergence de l'espérance et des densités de probabilités vers leurs valeurs stationnaires.	24
3.2	Graphe de transitions d'une chaîne de Markov.	25
3.3	Classification des états d'une chaîne de Markov.	28
3.4	Interprétation de l'équilibre global.	29
3.5	Processus de naissances et de morts.	32
3.6	Processus de Poisson.	34
3.7	Comptage du nombre d'arrivées.	34
4.1	File d'attente.	39
4.2	Représentation graphique d'une file d'attente	40
4.3	Courbe du nombre de clients	43
4.4	Courbe de la charge $W(t)$	45
4.5	Relation entre les courbes $W(t)$ et $N(t)$	46
4.6	Diagramme de transition de la file M/M/1.	47
4.7	File M/M/1/K.	49
4.8	Diagramme de transition de la file M/M/1/K.	49
4.9	File M/M/C.	50
4.10	File M/M/C/C.	52
4.11	Les trois options envisagées.	53
4.12	Principe d'un token bucket.	55

4.13	Temps de service résiduel	57
4.14	Exécution de jobs sur un ordinateur.	63
4.15	Figure pour l'exercice 11	68
5.1	Distribution du khi-carré.	75
5.2	Méthode de la fonction percentile.	76
5.3	Méthode d'acceptation-rejet.	76
5.4	Noyau d'un simulateur à événements discrets.	78
5.5	Gestion du temps de simulation.	79
5.6	Une des trajectoires possibles pour une file M/M/1/∞.	81
5.7	Convergence de l'espérance et des densités de probabilités vers leurs valeurs stationnaires.	83
5.8	Méthode des moyennes par lot.	84
5.9	Méthode des simulations parallèles.	85
5.10	La densité de probabilité de $\bar{X}_n$ se concentre autour de $\bar{X}$ au fur et à mesure que le nombre de trajectoires générées augmente.	86
5.11	Convergence de $\bar{X}_n$ vers $\bar{X}$ et de $Var[\bar{X}_n]$ vers 0.	86

CHAPITRE 1

Introduction

1.1 Evaluation des performances des réseaux

Durant la dernière décennie, les réseaux de télécommunications ont connu un développement sans précédent qui s'est accompagné d'un accroissement important de leur complexité. La maîtrise de ces systèmes complexes nécessite des modèles et des outils permettant de prédire les performances d'un réseau et de les optimiser.

Comme illustré sur la figure 1.1, on considérera ici un réseau de télécommunication comme un système servant le trafic entrant. Ce trafic entrant est généré par les utilisateurs du réseau. On ne sait jamais qui communique avec qui, quand et pour combien de temps. Par conséquent, une caractéristique essentielle des modèles réseaux que nous allons étudier est que ce sont des modèles stochastiques, c'est à dire qu'ils intègrent l'incertitude liée d'une part aux arrivées des communications dans le réseau et d'autre part à leurs durées.

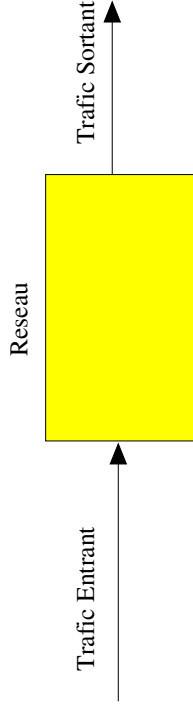


FIGURE 1.1 – Point de vue système.

Les modèles réseaux doivent nous permettre de répondre à trois grandes

questions :

- Etant donné la demande (le trafic) et les ressources disponibles, quelle est la qualité de service offerte aux utilisateurs ?
- Etant donné la demande et une qualité de service cible, comment dimensionner les ressources du système ?
- Etant donné les ressources disponibles et une qualité de service cible, quelle est la demande maximale ?

Autrement dit, comme illustré sur la figure 1.2, ces modèles doivent nous permettre d'étudier les relations entre d'une part la demande (le trafic) et la qualité de service requise et d'autre part la capacité du système (quantité de ressources disponibles) et donc son coût.

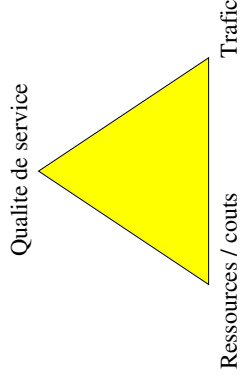


FIGURE 1.2 – Les trois facteurs clés.

De manière générale, les techniques d'évaluation des performances doivent nous permettre :

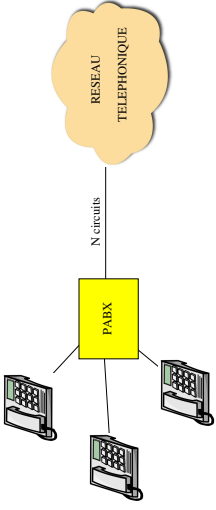
- de prédire les performances (QoS, utilisation des ressources) avant la construction ou l'extension d'un réseau.
- d'évaluer et comparer plusieurs scénarii : modification des équipements, configuration des protocoles de routage, etc.
- d'anticiper le comportement du réseau suite à une prévision d'augmentation du trafic ou suite à une panne d'équipement,
- etc.

1.2 Exemple introductif

Pour illustrer ceci, considérons tout de suite un exemple. Une société est raccordée au réseau téléphonique par une liaison supportant N appels simultanément (cf. figure 1.3). Un appel arrivant alors que les N circuits sont pris est bloqué et rejeté.

On sait que :

- La société compte 100 employés. En heure de pointe, chacun téléphone en moyenne 4 fois, soit 400 appels/heure (un appel toutes les 9 secondes).
- Chaque appel dure en moyenne 3 minutes (180 secondes).

FIGURE 1.3 – Liaison téléphonique ayant N circuits.

Dans cet exemple, la capacité du système correspond aux N circuits. La demande correspond aux appels téléphoniques des employés de la société. Enfin, la qualité de service correspond au taux de blocage des appels. On voudrait un taux de blocage inférieur à 1%.

La première question que l'on peut se poser est la suivante. S'il n'y avait pas de blocage, combien y aurait-il en moyenne d'appels simultanés sur le lien reliant l'entreprise au réseau téléphonique ? Nous verrons plus tard, que si la capacité $N \rightarrow \infty$, il y aurait en moyenne $Y = (1/9) * 180 = 20$ appels sur le lien.

Sachant que le coût de la liaison dépend directement de sa capacité, il est clair que prendre par exemple $N = 200$ si on ne doit avoir que 20 appels simultanés en moyenne n'est pas rentable.

Intuitivement, on peut penser que, vu que le trafic offert est de 20, une capacité $N = 20$ doit convenir. En fait, comme on le verra, ce choix conduit à ce que 16% des appels soient rejetés (64 appels sur les 400). Ceci s'explique en fait par la variabilité liée à l'arrivée des appels et à leurs durées.

Si on veut un taux de blocage de 1%, il faut en fait comme on le verra une capacité de $N = 30$ circuits.

1.3 Méthodes d'évaluation de performances

Les approches traditionnelles dans ce domaine sont le prototypage, la modélisation analytique et la simulation événementielle.

Le prototypage est en général hors de coût même pour des réseaux de taille modeste. La métrologie permet éventuellement de mesurer certains indices de performances, quand la quantité de données à collecter n'est pas monumentale. Mais il peut être difficile d'effectuer ces mesures sans perturber le réseau.

opérationnel. Surtout, cette approche ne permet pas de prédire les performances du réseau en cas de panne, d'évolution de la matrice de trafic ou de modification de la configuration du réseau.

La modélisation analytique, basée sur la théorie des files d'attente, est une approche puissante mais dont la capacité à représenter finement certains mécanismes complexes, et notamment le comportement réactif des protocoles, reste limitée. Par conséquent, une modélisation analytique exacte ayant toutes les propriétés précédentes semble aujourd'hui complètement inenvisageable compte tenu de nos connaissances théoriques. En particulier, il est bien connu qu'une modélisation analytique des phénomènes transitoires dans les réseaux de files d'attente est très complexe et peu de résultats généraux existent dans ce domaine.

Enfin, les modèles de simulation à événements discrets n'ont pas ces limitations, bien qu'ils requièrent souvent des temps de calcul prohibitifs pour des réseaux de taille réelle. De par sa grande flexibilité, la simulation constitue donc un outil souvent incontournable pour évaluer les performances des réseaux de communication.

1.4 Plan du cours

Le cours est organisé de la manière suivante :

- Rappels sur l'évaluation des performances :
 - Théorie des probabilités et chaînes de Markov (O. Brun)
 - Files d'attente élémentaires (U. Ayesta)
- Simulation événementielle (O. Brun)
 - Réseaux à commutation de circuits (O. Brun)
 - Réseaux à commutation de paquets (U. Ayesta)

CHAPITRE 2

Rappels sur la Théorie des Probabilités

L'intérêt principal de la théorie des probabilités est de nous permettre de modéliser une expérience (ou un système) sans disposer de l'information nécessaire pour la décrire exactement. Supposons par exemple que l'on lance une pièce de monnaie en l'air et qu'on observe sur quelle face elle va tomber. Pour prédire si le résultat sera pile ou face, il faut disposer d'une modélisation parfaite de la pièce, de l'état de l'air, de la vitesse à laquelle elle est lancée, etc. De même, si on veut étudier le trafic dans un réseau téléphonique, il faudrait pouvoir prédire, pour chaque abonné au téléphone, quand-est qu'il va téléphoner, à qui, et pendant combien de temps. C'est parce que nous sommes incapables de faire cette modélisation, ou que nous y renonçons, qu'il est nécessaire de faire appel aux probabilités.

Ce chapitre présente un rappel (très) synthétique des principales notions de la théorie des probabilités.

2.1 Événement

On considère un ensemble $\mathcal{S}$ qui représente l'ensemble des issues possibles d'une expérience. Par exemple, si je jette un dés, l'ensemble des issues possibles est $\mathcal{S} = \{1, 2, 3, 4, 5, 6\}$. C'est un ensemble discret. Au contraire, si mon expérience consiste à mesurer très précisément le temps jusqu'à ce que quelqu'un entre dans mon bureau, l'ensemble $\mathcal{S}$ est continu : $\mathcal{S} = \{x \in \mathbb{R} | x > 0\}$.

Un événement A est un sous-ensemble de l'ensemble $\mathcal{S}$ des issues possibles pour l'expérience. Si le résultat de l'expérience appartient à A , on dit que

l'événement A s'est réalisé. Par exemple, je peux m'intéresser à l'événement $A = \{2, 4, 6\} \subset \mathcal{S}$ que le jet d'un dés donne un nombre pair. Ou bien m'intéresser à l'événement $B = \{x \in \mathbb{R} | x > 10.4\} \subset \mathcal{S}$ que personne ne rentre dans mon bureau avant 10.4 secondes.

L'ensemble $\mathcal{S}$ tout entier est appelé l'événement certain, alors que le sous-ensemble vide $\emptyset$ de $\mathcal{S}$ est appelé l'événement impossible.

Etant donnés deux événements A et B , si le résultat de l'expérience appartient à $A \cup B$, on dira que l'événement A **ou** l'événement B s'est produit. Si le résultat de l'expérience appartient à $A \cap B$, on dira que A **et** B se sont produits. Si $A \cap B = \emptyset$, les deux événements ne peuvent se produire simultanément : ils sont dits mutuellement exclusifs. Enfin, si le résultat de l'expérience n'appartient pas à A , c'est qu'il appartient à son complémentaire $\bar{A} = \{e \in \mathcal{S} | e \notin A\}$ dans $\mathcal{S}$ (non A).

Un ensemble d'événements $A_1, A_2, \dots$ forment une partition de $\mathcal{S}$ s'ils sont mutuellement exclusifs ($A_i \cap A_j = \emptyset$ si $i \neq j$) et couvrent l'ensemble des issues possibles de l'expérience ($\cup_i A_i = \mathcal{S}$).

2.2 Probabilité d'un événement

Jusqu'à maintenant, on a considéré une expérience réalisée une seule fois. Supposons qu'elle soit répétée N fois. On peut alors, pour chaque événement A définir le nombre $N(A)$ de fois où il se produit (i.e le résultat appartient à A).

Empiriquement, on va définir la probabilité $P[A]$ de l'événement A comme la fréquence relative à laquelle se produit cet événement au cours d'une expérience répétée un nombre infini de fois :

$$P[A] = \lim_{N \rightarrow \infty} N(A)/N$$

On en déduit un certain nombre de propriétés :

1. $0 \leq P[A] \leq 1$
2. $P[\mathcal{S}] = 1$ et $P[\emptyset] = 0$
3. $P[A \cup B] = P[A] + P[B] - P[A \cap B]$
4. $P[A \cup B] = P[A] + P[B]$ si $A \cap B = \emptyset$
5. $P[\bar{A}] = 1 - P[A]$
6. $P[A] \leq P[B]$ si $A \subset B$

2.2.1 Probabilité conditionnelle

L'occurrence d'un événement A peut dépendre du fait que l'événement B s'est produit. La probabilité de l'événement A sachant que B s'est produit est donnée par,

$$P[A|B] = \frac{P[A \cap B]}{P[B]}$$

On a donc : $P[A \cap B] = P[A|B] \times P[B]$.

2.2.2 Loi des probabilités totales

Soit $B_1, B_2, \dots, B_n$ des événements formant une partition de $\mathcal{S}$. Tout événement A peut alors s'écrire,

$$A = A \cap \mathcal{S} = A \cap (\cup_i B_i) = \cup_i (A \cap B_i)$$

d'où l'on déduit que,

$$P[A] = \sum_{i=1}^n P[A \cap B_i] = \sum_{i=1}^n P[A|B_i] P[B_i]$$

Cela permet de calculer la probabilité de A en conditionnant par les événements B_i . Autrement dit, on décompose le problème en sous-problèmes.

2.2.3 Formule de Bayes

De nouveau, $B_1, B_2, \dots, B_n$ sont des événements formant une partition de $\mathcal{S}$. On veut calculer la probabilité de l'événement B_i sachant que l'événement A s'est produit,

$$P[B_i|A] = \frac{P[A \cap B_i]}{P[A]} = \frac{P[A|B_i] P[B_i]}{\sum_{j=1}^n P[A|B_j] P[B_j]}$$

La formule de Bayes permet de calculer une probabilité conditionnelle quand on connaît les probabilités conditionnelles inverses. Prenons un exemple. Soit 3 cartes. La première a ses deux faces rouges (rr). La deuxième a ses deux faces bleues (bb). Enfin la troisième a une face rouge et une face bleu (rb). On tire une carte au hasard et on ne regarde que la face de dessus. Elle est rouge. Quelle est la probabilité que l'autre face soit bleue ?

$$\begin{aligned} P[\text{rb}|\text{rouge}] &= \frac{P[\text{rouge}|\text{rb}] P[\text{rb}]}{P[\text{rouge}|\text{rr}] P[\text{rr}] + P[\text{rouge}|\text{bb}] P[\text{bb}] + P[\text{rouge}|\text{rb}] P[\text{rb}]} \\ &= \frac{1/2 \times 1/3}{1 \times 1/3 + 0 \times 1/3 + 1/2 \times 1/3} \\ &= 1/3 \end{aligned}$$

2.2.4 Indépendance

Deux événements A et B sont dits indépendants si et seulement si $P[A \cap B] = P[A] P[B]$. Dans ce cas, on a :

$$P[A|B] = \frac{P[A \cap B]}{P[B]} = P[A]$$

Autrement, savoir que B s'est produit n'apporte aucune information sur la probabilité que A se produise.

Exercice : on jette deux dés. Soit A l'événement "on a un 6 avec le premier dés" et B l'événement "on a un 1 avec le second dés". Montrer que les deux événements sont indépendants.

Exercice : on jette deux dés. Soit A l'événement "on a un 6 avec le premier dés" et B l'événement "la somme des deux dés fait 9". Montrer que les deux événements ne sont pas indépendants.

2.3 Variables aléatoires

Si on tire 10 fois à pile ou face, ce qui va nous intéresser c'est plus le nombre de fois où la pièce est tombée sur face que la séquence de piles et de faces obtenue. Ainsi, si on obtient le résultat pppppfff, ce qui nous intéresse c'est que la pièce est tombée 4 fois sur face.

De manière générale, on est souvent plus intéressé par une mesure (un nombre) associée à l'expérience que par le résultat proprement dit. Cette mesure est appelée variable aléatoire (v.a.). Ainsi, une variable aléatoire à valeurs réelles est une fonction $X : \mathcal{S} \rightarrow \mathbb{R}$ qui associe à chaque issue e possible de l'expérience un nombre réel $X(e)$.

L'image d'une variable aléatoire X est l'ensemble,

$$\mathcal{S}_X = \{x \in \mathbb{R} | \exists e \in \mathcal{S}, X(e) = x\}$$

L'image de X est donc l'ensemble des valeurs que cette v.a. peut prendre. Cet ensemble peut être fini ou dénombrable : on parlera alors de v.a. discrète. Sinon, on parlera de v.a. continue.

Etant donné un réel x , on peut définir les événements suivants,

$$\begin{aligned} [X = x] &= \{e \in \mathcal{S} | X(e) = x\} \\ [X \leq x] &= \{e \in \mathcal{S} | X(e) \leq x\} \end{aligned}$$

Pour des facilités d'écriture, ces événements sont notés de manière spéciale, mais ce sont des événements bien définis. Le premier correspond à l'ensemble des résultats possibles de l'expérience dont la valeur (par X) est égale à x . Le

second aux issues possibles de l'expérience dont la valeur est inférieure à x .

2.3.1 Variables aléatoires continues

La fonction de répartition (CDF) de la v.a. X est définie par,

$$F_X(x) = P[X \leq x]$$

$F_X(x)$ représente ainsi la probabilité que l'événement $[X \leq x]$ se produise. Évidemment on a $\lim_{x \rightarrow -\infty} F_X(x) = 0$. De plus, si X est à valeurs positives, $F_X(0) = 0$. La fonction de répartition d'une v.a. positive est donc croissante entre 0 et 1.

La probabilité d'un intervalle (i.e. $x_1 \leq X(e) \leq x_2$) est :

$$P[x_1 \leq X \leq x_2] = F_X(x_2) - F_X(x_1)$$

La densité de probabilité (pdf) d'une v.a. continue est définie par,

$$f_X(x) = \frac{dF_X(x)}{dx} = \lim_{dx \rightarrow 0} \frac{F_X(x+dx) - F_X(x)}{dx}$$

On a évidemment $F_X(x) = \int_0^x f_X(x) dx$ et $\int_0^\infty f_X(x) dx = 1$. Pour dx "petit", $f_X(x) dx$ correspond à $P[x \leq X \leq x+dx]$.

2.3.2 Variables aléatoires discrètes

Si X est une v.a. discrète prenant ses valeurs dans $x_1, x_2, \dots$, alors on peut définir sa fonction de répartition par,

$$F_X(x) = \sum_{x_i \leq x} p_i$$

où $p_i = P[X = x_i]$. $F_X(x)$ est donc une fonction en escalier ayant des sauts de hauteurs p_i aux points x_i et telle que $F_X(x) = 0$ pour $x < x_1$ et $F_X(x) = 1$ pour $x > x_n$.

L'équivalent de la densité de probabilité est la fonction de masse de probabilité définie par,

$$p(x) = P[X = x] = \begin{cases} p_i & \text{si } x = x_i \\ 0 & \text{sinon} \end{cases}$$

2.3.3 Cas de plusieurs variables aléatoires

Si X et Y sont deux v.a., on peut définir leur CDF jointe et leur pdf jointe par,

$$F_{X,Y}(x, y) = P[X \leq x, Y \leq y] \quad \text{et} \quad f_{X,Y}(x, y) = \frac{\partial}{\partial x} \frac{\partial}{\partial y} F_{X,Y}(x, y)$$

Ainsi si dx et dy sont "petits", $f(x, y) dx dy$ correspond à $P[x \leq X \leq x+dx, y \leq Y \leq y+dy]$.

La densité marginale de Y est donnée par,

$$f_Y(y) = \int_0^\infty f_{X,Y}(x, y) dx$$

Les v.a. X et Y sont dites indépendantes si les événements $[X \leq x]$ et $[Y \leq y]$ sont indépendants quels que soient x et y . Cela se traduit par,

$$F_{X,Y}(x, y) = F_X(x) F_Y(y) \quad \text{et} \quad f_{X,Y}(x, y) = f_X(x) f_Y(y)$$

De la même façon, on peut définir la pdf conditionnelle de X sachant Y ,

$$f_{X|Y}(x, y) = \frac{f_{X,Y}(x, y)}{f_Y(y)}$$

Si les deux v.a. sont indépendantes, on a $f_{X|Y}(x, y) = f_X(x)$.

2.3.4 Moments d'une variable aléatoire

Soit un entier $k \geq 1$. Le moment d'ordre k d'une v.a. continue est donné par,

$$E[X^k] = \int_0^\infty x^k f_X(x) dx$$

et pour une v.a. discrète,

$$E[X^k] = \sum_{i=1}^\infty x_i^k p_i$$

Les deux premiers moments ($k = 1, 2$) jouent un rôle particulier. Le premier est appelé l'espérance mathématique (ou moyenne) :

$$E[X] = \int_0^\infty x f_X(x) dx \quad (\text{v.a. continue}) \quad \text{ou} \quad E[X] = \sum_{i=1}^\infty x_i p_i \quad (\text{v.a. discrète})$$

L'espérance est linéaire : $E[aX + bY] = aE[X] + bE[Y]$ pour a et b réels. Par contre, $E[XY] = E[X]E[Y]$ n'est vérifiée que si X et Y sont des v.a.

indépendantes.

Le second moment intervient dans le calcul de la variance :

$$V[X] = E[(X - E[X])^2] = E[X^2] - E[X]^2$$

La variance vérifie $V[cX] = c^2 V[X]$. Attention : $V[X + Y] = V[X] + V[Y]$ n'est vrai que si les v.a. X et Y sont indépendantes ! Enfin, signalons qu'on utilise souvent l'écart type d'une v.a. X , qui correspond à la racine carrée de la variance.

2.4 Quelques distributions de probabilités courantes

Jusqu'ici, nous avons considéré une expérience, ses résultats possibles et les événements correspondants. Nous avons défini la probabilité d'un événement comme la fréquence à laquelle il se produirait si l'expérience était renouvelée un nombre infini de fois. La notion de v.a. nous a permis d'associer une valeur à chaque résultat possible. La fonction de répartition, la densité de probabilité, l'espérance mathématique et la variance que nous avons vues ensuite, sont simplement des outils commodes pour caractériser la probabilité d'obtenir telle ou telle valeur.

Dans la pratique, quand on doit modéliser un problème, la démarche suivie est souvent inverse. On commence par identifier les v.a. intervenants dans le problème. Puis on suppose qu'elles suivent une certaine distribution de probabilités. Cette hypothèse devra ensuite être validée par des expérimentations, une fois l'analyse effectuée.

Dans cette optique de modélisation, certaines distributions de probabilités jouent un rôle important. Tout simplement parce qu'on les rencontre souvent dans les systèmes réels. Nous en présentons ici quelques unes.

2.4.1 Distributions discrètes

Distribution de Bernoulli

Elle modélise une expérience avec deux résultats possibles : succès (avec probabilité p) ou échec (avec probabilité $q = 1 - p$). Une v.a. X suivant la distribution de Bernoulli est définie par,

$$X = \begin{cases} 1 & \text{avec probabilité } p \\ 0 & \text{avec probabilité } q = 1 - p \end{cases}$$

On obtient immédiatement $E[X] = p$ et $V[X] = pq$. La v.a. X peut par exemple modéliser un système avec deux états : ON et OFF.

Distribution Binomiale

Soit X une v.a. représentant le nombre de succès après n expériences de Bernoulli indépendantes,

$$X = \sum_{i=1}^n Y_i \quad \text{où} \quad Y_i \sim \text{Bernoulli}(p)$$

On a immédiatement $E[X] = nE[Y_i] = np$ et $V[X] = nV[Y_i] = np(1 - p)$. La probabilité d'avoir i succès est de plus donnée par :

$$p_i = C_i^n p^i (1 - p)^{n-i}$$

Distribution Géométrique

Soit X une v.a. représentant le nombre de tentatives avant le premiers succès dans une série d'expériences de Bernoulli indépendantes :

$$p_i = P[X = i] = (1 - p)^{i-1} p$$

On en déduit le nombre moyen de tentatives $E[X] = 1/p$ et sa variance $V[X] = (1 - p)/p^2$. La probabilité qu'il faille plus de n tentatives avant le premier succès est donnée par,

$$P[X > n] = \sum_{i=n+1}^{\infty} p_i = (1 - p)^n$$

La distribution géométrique est sans mémoire :

$$\begin{aligned} P[X > i + j \mid X > i] &= \frac{P[X > i + j, X > i]}{P[X > i]} \\ &= \frac{P[X > i + j]}{P[X > i]} \\ &= \frac{(1 - p)^{i+j}}{(1 - p)^i} \\ &= (1 - p)^j = P[X > j] \end{aligned}$$

Autrement dit, le fait de savoir qu'il y a déjà eu i tentatives infructueuses, ne change pas la probabilité qu'il en faille encore plus de j pour avoir un succès. Les tentatives passées n'ont pas d'influence sur les tentatives futures.

Distribution de Poisson

Une v.a. discrète X à valeurs entières positives suit la loi de Poisson de paramètre λt (deux réels positifs fixés) si sa distribution est donnée par,

$$p_i = P[X = i] = \frac{(\lambda t)^i}{i!} \exp(-\lambda t)$$

On peut montrer que la loi de Poisson correspond à la limite de la loi binomiale $B(n, p)$ quand on fixe le nombre moyen de succès $np = \lambda t$ et qu'on fait tendre n vers l'infini. En conséquence, on peut voir le paramètre t comme une durée et λ comme un taux de succès. Ainsi, une v.a. $X \sim \text{Poisson}(\lambda t)$ permet de représenter le nombre d'événements sur un intervalle de longueur t pour un processus de Poisson d'intensité λ . Le nombre moyen d'événements est donné par $E[X] = \lambda t$ et sa variance est $V[X] = \lambda t$. De plus,

- La probabilité d'un événement (un succès) durant dt est λdt ,
- La probabilité de plus d'un événement durant dt est négligeable,
- Le nombre d'événements se produisant sur deux intervalles de temps dis-joints sont indépendants.

Enfin, on peut signaler que le processus de Poisson est stable par agrégation, c'est à dire si $X_1 \sim \text{Poisson}(\lambda_1 t)$ et $X_2 \sim \text{Poisson}(\lambda_2 t)$, alors $X_1 + X_2 \sim \text{Poisson}((\lambda_1 + \lambda_2) t)$.

La loi de Poisson va nous servir par exemple à représenter le nombre d'appels téléphoniques sur un intervalle de temps.

2.4.2 Distributions continues

Distribution uniforme

Une v.a. $X \sim U(a, b)$ a une pdf constante sur l'intervalle $[a, b]$,

$$f_X(x) = \begin{cases} 1/(b-a) & a < x < b \\ 0 & \text{sinon} \end{cases}$$

Sa moyenne est $E[X] = (a+b)/2$ et sa variance est $V[X] = (b-a)^2/12$.

Distribution exponentielle

Une v.a. continue X suit une loi exponentielle de paramètre μ si sa CDF est donnée par (cf. figure 2.1),

$$F_X(x) = \begin{cases} 1 - \exp(-\mu x) & x \geq 0 \\ 0 & x < 0 \end{cases}$$

La pdf est alors,

$$f_X(x) = \begin{cases} \mu \exp(-\mu x) & x \geq 0 \\ 0 & x < 0 \end{cases}$$

La loi exponentielle va par exemple nous servir à modéliser la durée d'un appel téléphonique, ou encore la durée entre l'arrivée de deux appels. Sa moyenne et sa variance sont données par,

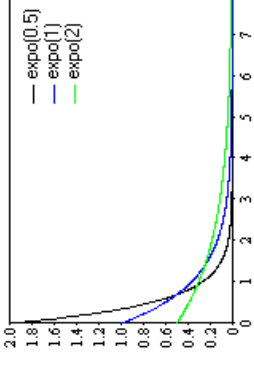


FIGURE 2.1 – Fonction de répartition de la distribution exponentielle.

$$E[X] = 1/\mu \quad \text{et} \quad V[X] = 1/\mu^2$$

Comme la loi géométrique dans le cas discret, la loi exponentielle a la propriété d'être sans mémoire. Supposons que X représente la durée aléatoire d'un appel téléphonique. Quelle est la probabilité que l'appel dure au moins x secondes de plus, sachant qu'il a déjà duré t secondes.

$$\begin{aligned} P[X > t + x \mid X > t] &= \frac{P[X > t + x, X > t]}{P[X > t]} \\ &= \frac{P[X > t]}{P[X > t]} \\ &= \frac{\exp(-\mu(t+x))}{\exp(-\mu t)} \\ &= \exp(-\mu x) = P[X > x] \end{aligned}$$

Autrement dit, le fait de savoir que l'appel a déjà duré t secondes, ne change pas la probabilité qu'il en faille encore plus de x .

Remarquons enfin que le paramètre μ s'interprète comme un taux de terminaison. En effet, si on sait qu'un appel a déjà duré t secondes, quelle est la probabilité qu'il se termine en dt secondes supplémentaires :

$$P[X < t+dt \mid X > t] = P[X < dt] = 1 - \exp(-\mu dt) = 1 - (1 - \mu dt + (\mu dt)^2/2 \dots) \approx \mu dt$$

Ainsi μ est le taux de terminaison de l'appel par unité de temps, $1/\mu$ étant la durée moyenne de l'appel.

Distribution normale

La pdf d'une v.a. gaussienne de paramètres μ et σ^2 est (cf. figure 2.2),

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} \exp\left(-\frac{1}{2}(x - \mu)^2/\sigma^2\right)$$

Les paramètres μ et σ^2 correspondent à la moyenne et à la variance : $E[X] = \mu$ et $V[X] = \sigma^2$.

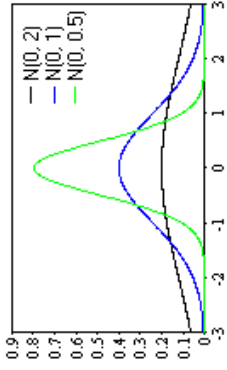


FIGURE 2.2 – Distribution normale.

Outre ses nombreuses applications dans le domaine de la physique, cette distribution est surtout importante car elle intervient dans le théorème central limite. Illustrons ceci par un exemple. Supposons que la durée d'un appel soit une v.a. de moyenne μ et de variance σ^2 . On mesure n fois la durée d'un appel pour obtenir n v.a. $X_1, \dots, X_n$ indépendantes et identiquement distribuées (i.i.d.). On s'intéresse à la moyenne empirique de ces mesures,

$$W_n = \frac{1}{n} \sum_{i=1}^n X_i$$

W_n est une v.a. dont l'espérance et la variance sont données par $E[W_n] = \mu$ et $V[W_n] = \sigma^2/n$. Intuitivement, on voit donc que plus on va avoir de mesures, plus la moyenne arithmétique sera "proche" de l'espérance de la durée d'un appel, avec une variance de mesure qui décroît linéairement. La loi faible des grands nombres précise cette intuition en affirmant que, pour une précision ϵ donnée, on a,

$$\lim_{n \rightarrow \infty} P[W_n - \mu \geq \epsilon] = 0$$

et la loi forte des grands nombres nous dit que W_n converge vers μ avec probabilité 1. Le théorème central limite nous permet de mesurer cette convergence. Il affirme que pour n suffisamment grand,

$$\frac{\sqrt{n}}{\sigma} (W_n - \mu) \sim N(0, 1)$$

Si on connaît la moyenne μ et la variance σ de la durée d'un appel, on peut ainsi évaluer l'erreur commise en utilisant la moyenne empirique pour estimer la

moyenne. On verra que ce résultat nous sera utile pour établir un critère d'arrêt pour des simulations événementielles.

Exercices

Exercice 1 Soit X une variable aléatoire positive telle que $E[X^2] < \infty$ et $f(\cdot)$ sa densité de probabilité. Montrer que :

$$E[X] = \int_0^\infty P(X > u) du.$$

et,

$$E[X^2] = \int_0^\infty 2uP(X > u) du.$$

Exercice 2 Soit X_i , $i = 1, 2$, des v.a. exponentiellement distribuées de paramètres λ_i .

- Calculer $P(X_1 > u + v | X_1 > u)$
- Quelle est la probabilité que $X_1 \leq X_2$?
- Quelle est la distribution de la v.a. $\min\{X_1, X_2\}$?

Exercice 3 (distribution Hyper-Exponentielle) Soit X_i , $i = 1, 2$, des v.a. exponentiellement distribuées de paramètres λ_i . Soit B une v.a. telle que $P(B = 1) = p$ et $P(B = 2) = 1 - p$, avec $0 < p < 1$. Définissons alors la v.a. Y de la façon suivante : $Y = X_1$ si $B = 1$ et $Y = X_2$ si $B = 2$. Calculer $P(Y > x)$ en conditionnant sur la valeur de B . En déduire la densité de Y .

Exercice 4 Une connexion est constitué de 4 liens consécutifs non fiables. Sur chacun de ces liens, la probabilité qu'un bit (0 ou 1) transmis soit reçu correctement est 90%, alors que dans 10% des cas le bit reçu est erroné. Quelle est la probabilité qu'un bit transmis sur la connexion soit reçu correctement à l'autre bout.

Exercice 5 Soit X_i , $i = 1, 2, 3$, des v.a. indépendantes exponentiellement distribuées de paramètres λ_i , $i = 1, 2, 3$. Déterminer $P(X_1 < X_2 < X_3)$.

Indication. En conditionnant $P(X_1 < X_2 < X_3) = \int_0^\infty P(X_1 < t < X_3 | X_2 = t) \lambda_2 e^{-\lambda_2 t} dt$.

Exercice 6 On considère 2 dés. On note n_i le résultat du dé $i = 1, 2$. Nous avons vu en cours que les événements $A = \{n_1 = 6\}$ et $B = \{n_2 = 1\}$ sont indépendants. Considérons maintenant les événements $A = \{(n_1, n_2) \text{ t.q. } n_1 = 6\}$ et $B = \{(n_1, n_2) \text{ t.q. } n_1 + n_2 = 9\}$.

- a) Ecrire explicitement les ensembles A, B et $A \cap B$.
- b) Calculer $P(A)$, $P(B)$ et $P(A \cap B)$.
- c) A et B sont-ils indépendants ?

Exercice 7 Dans une année quelconque, le nombre d'élèves inscrits à l'Université de Toulouse suit une loi de Poisson de paramètre λ . L'université propose

aux élèves N filières en 4ème année. Un élève choisit la filière i avec une probabilité p_i indépendamment des autres élèves. Quelle est la loi de probabilité du nombre d'élèves qui s'inscrivent en filière i , $i = 1, 2, \dots, N$?

Exercice 8 Le centre d'appels d'une compagnie aérienne sert deux types de clients : les premiers parlent uniquement l'anglais et les autres uniquement le français. Les nombres d'appels anglophones et francophones respectivement d'une journée suivent une loi de Poisson de paramètres λ_a et λ_f respectivement. Au bout du premier jour, vous êtes informé qu'un total de N appels ont été traités par le centre. Quel est la probabilité qu'il y ait eu k appels anglophones pendant la journée ? Même question si le nombre de appels des deux langues sont distribués géométriquement de paramètres p_a et p_f , respectivement ?

Exercice 9 On vient d'installer N ampoules de marques différentes. La durée de fonctionnement d'une ampoule de marque i est exponentiellement distribuée de taux λ_i . Quelle est la probabilité que l'ampoule de marque i soit la première à tomber en panne. Quelle est cette probabilité si la durée de fonctionnement est uniforme entre $[0, 2/\lambda_i]$?

Exercice 10 Il y a une maladie contagieuse qui circule. Vous êtes un médecin et lors d'un contrôle vous constatez qu'une personne a attrapé cette maladie. Cependant lors du dernier contrôle il y a J jours, la même personne a été contrôlée négative. Afin de pouvoir prescrire le bon traitement, il faut savoir quand la maladie a été rattrapée. Le patient n'est pas en mesure de vous renseigner. Néanmoins vous êtes au courant d'une étude qui montre qu'une personne après avoir été contrôlée négative peut attraper cette maladie dans une durée aléatoire, et cette étude précise la loi de probabilité associée. Vous cherchez donc le moment entre les deux contrôles qui maximise la probabilité d'attraper cette maladie. Quel est ce moment si selon l'étude :

1. la distribution est exponentielle de moyenne D jours ?
2. la distribution est uniforme de $D \geq J$ jours ?

Exercice 11 Vous voulez garer votre voiture dans une parking souterrain de N places au centre ville. Quand vous y arrivez, celui-ci affiche complet. Néanmoins vous y entrez en espérant qu'une place se libérera entre temps. Dans le parking, les voitures sont garées le long du côté droit de la voie de circulation, et vous-êtes obligé de rouler à une vitesse telle que vous passez devant une place de parking en δ secondes. Supposons qu'à partir de votre entrée dans le parking, il vous faille δ secondes pour arriver jusqu'à la première place de parking, et que la durée de stationnement d'une voiture soit exponentielle avec moyenne $1/\mu$ secondes. Si vous ne trouvez pas de place, vous êtes obligé de sortir.

1. Supposons que vous ne puissiez pas faire marche-arrière. Quelle est la probabilité que vous trouviez la i ème place libre ? Quelle est la probabilité

que vous sortiez sans trouver de place ? Et si les voitures étaient garées des deux côtés de la voie de circulation ?

2. Supposons qu'il n'y ait personne d'autre qui cherche une place et que vous puissiez faire marche-arrière uniquement si une place se libère derrière vous ? Quelle est la probabilité de ne pas trouver de place ?
3. Supposons que vous n'avez pas trouvé de place et que vous essayez de nouveau après T secondes. Le parking affiche toujours complet. Comment les probabilités calculées en (a) et (b) vont-elles changer ?

CHAPITRE 3

Chaînes de Markov

3.1 Processus Stochastiques

En général les systèmes que nous considérons sont des systèmes dynamiques (dont l'état évolue au cours du temps) faisant intervenir de l'aléatoire. Par exemple, la longueur d'une file d'attente évolue au cours du temps en fonction de l'arrivée aléatoire des clients et du temps aléatoire qu'il faut pour les servir.

Un processus stochastique $(Y_t)_{t \geq 0}$ est une séquence de variables aléatoires indexées par le temps t . Le temps t peut être continu (resp. discret, i.e. $t = 1, 2, \dots$) et l'on parle alors de processus à temps continu (resp. à temps discret). Pour t fixé, Y_t est simplement une v.a. prenant ses valeurs dans un certain espace d'états suivant une distribution qui peut dépendre de t . L'ensemble des états possibles pour le processus peut lui-même être discret ($Y_t = 0, 1, \dots$) ou continu.

Pour caractériser un processus stochastique, nous ne considérerons ici que les statistiques de premier ordre, à savoir la distribution $F_{Y_t}(y) = P[Y_t \leq y]$ de l'état Y_t au cours du temps, et sa moyenne $E[Y_t]$. Une caractérisation au second ordre impliquerait d'établir la covariance du processus entre deux instants s et t . Cette covariance, donnée par,

$$R_{s,t} = E[(X_t - E[X_t])(X_s - E[X_s])]$$

permet d'étudier la relation entre les états du système à deux instants différents.

Si on tire aléatoirement une valeur y_t de Y_t pour chaque valeur de t , la séquence de valeurs obtenues est appelée une trajectoire du système. Cela représente

une succession au cours du temps d'états que le système peut prendre.

On peut s'intéresser au comportement transitoire du processus, c'est à dire à l'évolution des fonctions de répartition $F_{Y_t}(x)$ au cours du temps t . En général, on s'intéresse plutôt au comportement stationnaire (à l'équilibre) qui est obtenu si $F_{Y_t}(x)$ converge vers une distribution unique $F_Y(x)$ indépendante du temps quand $t \rightarrow \infty$. La figure 3.1 illustre la convergence des densités de probabilités d'un processus stochastique $(Y_i)_{i=1,2,\dots}$ vers une densité stationnaire. Attention : le régime permanent peut n'exister que sous certaines conditions et il faut soigneusement établir ces conditions quand on étudie un processus stochastique.

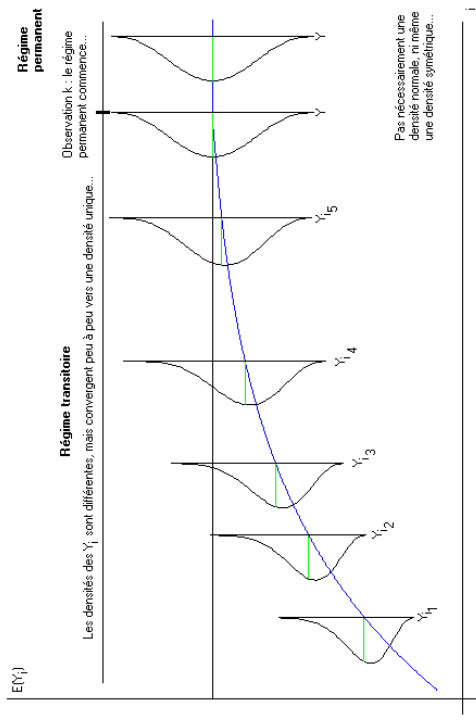


FIGURE 3.1 – Convergence de l'espérance et des densités de probabilités vers leurs valeurs stationnaires.

3.2 Processus Markovien

Soit un processus stochastique $(X_t)_{t \geq 0}$. Considérons des instants $t_1 < \dots < t_n$ situés dans le passé. On suppose connaître l'état x_i du processus à l'instant t_i , $i = 1, \dots, n$. L'état x_n est l'état courant et on s'intéresse à son état futur x_{n+1} à l'instant t_{n+1} . Le processus (X_t) a la propriété Markovienne si,

$$P[X_{t_{n+1}} \leq x_{n+1} \mid X_{t_n} = x_n, \dots, X_{t_1} = x_1] = P[X_{t_{n+1}} \leq x_{n+1} \mid X_{t_n} = x_n]$$

et ce quelque soit n et les instants considérés. Concrètement, cette équation signifie que le processus est Markovien si son état dans le futur ne dépend que de l'état présent, et pas de ses états passés. Autrement dit, l'état courant contient toute l'information permettant de caractériser l'évolution future du processus.

Un processus Markovien à espace d'états discret est appelée une chaîne de Markov.

3.3 Chaîne de Markov à temps discret

Une chaîne de Markov à temps discret est un processus Markovien $(X_n)_{n=1,2,\dots}$ à temps discret dont l'espace d'états est discret, i.e. $X_n \in \{0, 1, 2, \dots\}$ (le nombre d'états peut être fini).

3.3.1 Probabilités de transition

Une chaîne de Markov peut être caractérisée par ses probabilités de transition en une étape,

$$p_{i,j} = P[X_{n+1} = j \mid X_n = i]$$

Ainsi, $p_{i,j}$ est la probabilité de passer de i à j en un pas de temps. On notera que ces probabilités ne dépendent pas de l'instant n considéré : on dit dans ce cas que la chaîne de Markov est homogène.

Une chaîne de Markov à temps discret peut être représenté par un graphe tel que celui de la figure 3.2 (où p est un paramètre, e.g. $p = 1/3$). Les noeuds du graphe représentent les états que peut prendre le processus et les arcs représentent les transitions entre états. Le poids de l'arc (i, j) correspond à $p_{i,j}$.

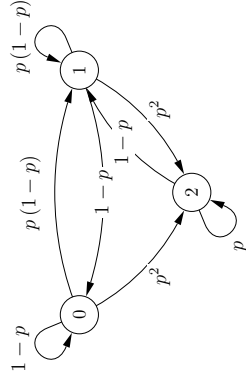


FIGURE 3.2 – Graphe de transitions d'une chaîne de Markov.

On appelle matrice des probabilités de transition la matrice $P = [p_{i,j}]$,

$$P = \begin{pmatrix} p_{0,0} & p_{0,1} & p_{0,2} & \dots \\ p_{1,0} & p_{1,1} & p_{1,2} & \dots \\ p_{2,0} & p_{2,1} & p_{2,2} & \dots \\ \vdots & \vdots & \vdots & \ddots \end{pmatrix}$$

La ligne i de cette matrice donne donc la probabilité de passer de l'état i vers tout autre état j (y compris l'état i lui même) en un pas de temps. Evidemment, comme une transition est effectuée à chaque pas de temps, la somme des éléments d'une ligne est égal à 1. On parle de matrice stochastique.

Si le système passe de l'état i à l'état j en deux étapes, c'est qu'il est passé par un état k (qui peut éventuellement être i ou j). On a donc :

$$p_{i,j}^{(2)} = \sum_k p_{i,k} p_{k,j}$$

Clairement $p_{i,j}^{(2)}$ est l'élément (i, j) de la matrice $P^2 = P P$. On peut ainsi établir que les probabilités de transition $p_{i,j}^{(n)}$ en n étapes s'obtiennent avec P^n , i.e. en mettant la matrice P à la puissance n .

Puisque on a toujours $P^n = P^m P^{n-m}$, il vient :

$$p_{i,j}^{(n)} = \sum_k p_{i,k}^{(m)} p_{k,j}^{(n-m)} \quad 0 \leq m \leq n$$

Cette équation est connue sous le nom d'équation de Chapman-Kolmogorov. Elle exprime que pour passer de i à j en n étapes, il faut passer par un état k à l'étape m .

3.3.2 Probabilités d'état

Notons $\pi_i^{(n)} = P[X_n = i]$ la probabilité que la chaîne de Markov soit dans l'état i à l'instant n . Pour manipuler ces probabilités d'états, on utilise en général le vecteur $\pi^{(n)} = [\pi_0^{(n)}, \pi_1^{(n)}, \dots]$. Puisque la matrice P donne les probabilités de transition en une étape, on a,

$$\pi^{(n)} = \pi^{(n-1)} P$$

d'où l'on déduit récursivement que,

$$\pi^{(n)} = \pi^{(0)} P^n$$

Reprenons l'exemple de la figure 3.2 avec $p = 1/3$. Dans ce cas, la matrice P est donnée par (aux arrondis près),

$$P = \begin{pmatrix} 0.6666 & 0.2222 & 0.1111 \\ 0.6666 & 0.2222 & 0.1111 \\ 0 & 0.6666 & 0.3333 \end{pmatrix}$$

la matrice de transitions à 2 pas est donnée par,

$$P^2 = \begin{pmatrix} 0.5926 & 0.2716 & 0.1358 \\ 0.5926 & 0.2716 & 0.1358 \\ 0.4444 & 0.3704 & 0.1852 \end{pmatrix}$$

Partant de l'état i à l'instant initial, la distribution de l'état au bout de 2 transitions est donnée par la ligne i de P^2 . Ainsi, si à l'instant initial la chaîne de Markov est dans l'état 0 de façon certaine, i.e. $\pi^{(0)} = [1, 0, 0]$, alors on sait que la distribution de probabilités de l'état à l'instant 2 est donnée par $\pi^{(2)} = [0.5926, 0.2716, 0.1358]$.

Si on continue le calcul jusqu'à l'ordre 8 (8 transitions), on a :

$$P^8 = \begin{pmatrix} 0.5714 & 0.2857 & 0.1429 \\ 0.5714 & 0.2857 & 0.1429 \\ 0.5714 & 0.2857 & 0.1429 \end{pmatrix}$$

Puisque toutes les lignes sont identiques, on voit que quel que soit l'état de départ, la distribution de probabilités atteinte est la même. On peut le confirmer par le calcul suivant,

$$\begin{aligned} \pi^{(8)} &= \begin{bmatrix} \pi_0^{(0)} & \pi_1^{(0)} & \pi_2^{(0)} \end{bmatrix} P^8 \\ &= \begin{bmatrix} 0.5714 \times (\pi_0^{(0)} + \pi_1^{(0)} + \pi_2^{(0)}) & \\ & 0.2857 \times (\pi_0^{(0)} + \pi_1^{(0)} + \pi_2^{(0)}) & \\ & 0.1429 \times (\pi_0^{(0)} + \pi_1^{(0)} + \pi_2^{(0)}) \end{bmatrix} \\ &= [0.5714, 0.2857, 0.1429] \end{aligned}$$

La chaîne de Markov converge donc vers un état d'équilibre indépendant de l'état initial. Le processus a oublié son état initial. Dans la suite, nous allons nous intéresser aux conditions sous lesquelles cet état d'équilibre existe et à son calcul.

3.3.3 Classification des états

On dit que l'état i conduit à l'état j si la chaîne de Markov peut passer de i à j avec une probabilité non nulle, i.e. $\exists n, p_{i,j}^{(n)} > 0$. Si la relation est réciproque, on dit que les états i et j communiquent. Une chaîne de Markov dont tous les états communiquent entre eux est dite irréductible.

Un état est absorbant si le processus reste bloqué dans cet état, i.e. $p_{i,i} = 1$. On distingue de plus les états transitoires des états récurrents. Un état est transitoire si, partant de cet état, la probabilité d'y revenir est strictement inférieure à 1. Si cette probabilité est égale à 1, l'état est récurrent.

Soit i un état récurrent. On note $T_{i,i}$ l'espérance du temps nécessaire pour revenir pour la première fois dans cet état quand $X_0 = i$. L'état i est dit positif récurrent si $T_{i,i} < \infty$, et nul récurrent si $T_{i,i} = \infty$.

Si, partant d'un état i , le temps nécessaire pour revenir dans cet état pour la première fois est nécessairement un multiple d'un entier d , cet état est dit périodique. Sinon, l'état est dit apériodique.

Un état positif récurrent et apériodique est dit ergodique. Si tous les états sont ergodiques, la chaîne de Markov est ergodique.

Toutes ces définitions sont illustrées sur la figure 3.3.

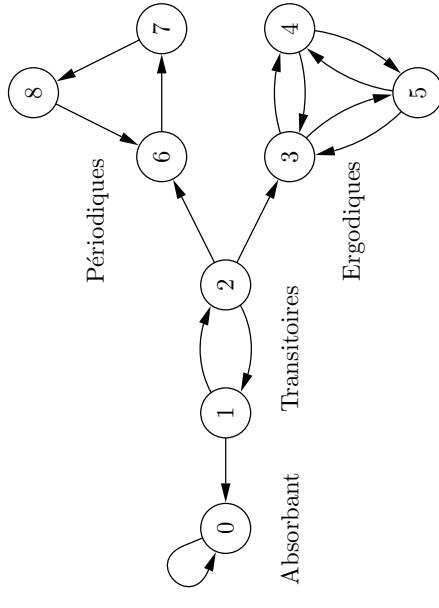


FIGURE 3.3 – Classification des états d'une chaîne de Markov.

3.3.4 Distribution stationnaire

Théorème de Kolmogorov

Pour une chaîne de Markov irréductible et apériodique, la distribution stationnaire,

$$\pi_j = \lim_{n \rightarrow \infty} \pi_j^{(n)}$$

existe et est indépendante de l'état initial. π_j représente la proportion de

temps que le processus passe dans l'état j . De plus, soit :

- Les états sont tous transitoires ou tous nuls récurrents, et dans les deux cas $\pi_j = 0, \forall j$,
- Les états sont tous positifs récurrents et la distribution stationnaire π (vecteur) est la solution des équations suivantes,

$$\pi = \pi P \quad \text{et} \quad \pi e^T = 1$$

où e est un vecteur ligne dont tous les composantes valent 1 et e^T est le vecteur colonne transposé. Ces équations peuvent également s'écrire :

$$\pi_j = \sum_i \pi_i p_{i,j} \quad \text{et} \quad \sum_j \pi_j = 1$$

Equations d'équilibre global

L'équation $\pi = \pi P$, ou $\pi_j = \sum_i \pi_i p_{i,j}$ pour $j = 0, 1, \dots$, est souvent appelé la condition d'équilibre global. En effet, on a $\sum_i p_{j,i} = 1$ puisqu'il y a forcément une transition depuis l'état j , soit vers j lui-même, soit vers un autre état. L'équation $\pi = \pi P$ implique donc,

$$\pi_j = \pi_j \sum_i p_{j,i} = \sum_i \pi_i p_{i,j} \quad j = 0, 1, \dots$$

Le terme $p_{j,j} \pi_j$ apparaissant des deux côtés, on peut le supprimer pour obtenir :

$$\pi_j \sum_{i \neq j} p_{j,i} = \sum_{i \neq j} \pi_i p_{i,j} \quad j = 0, 1, \dots$$

Cette équation exprime que la probabilité de sortir de l'état j vers un autre état est égale à la probabilité d'entrer dans l'état j en venant d'un autre état. C'est bien la notion d'équilibre (cf. figure 3.4).

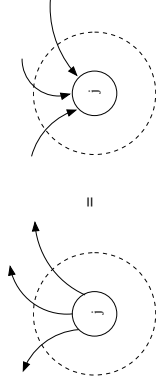


FIGURE 3.4 – Interprétation de l'équilibre global.

Exemple

Reprenons l'exemple de la figure 3.2 (avec p général). La condition d'équilibre de l'état 0 donne,

$$\pi_0 = (1-p)\pi_0 + (1-p)\pi_1 \Rightarrow \pi_0 = \frac{1-p}{p}\pi_1$$

La condition d'équilibre de l'état 1 donne,

$$\pi_1 = p(1-p)\pi_0 + p(1-p)\pi_1 + (1-p)\pi_2 \Rightarrow \pi_1 = \frac{1-p}{p}\pi_2, \pi_0 = \left(\frac{1-p}{p}\right)^2 \pi_2$$

La condition de normalisation $\pi_0 + \pi_1 + \pi_2 = 1$ permet de déduire la valeur de π_2 , et donc la solution stationnaire :

$$\pi = \left[\frac{(1-p)^2}{1-p(1-p)}, \frac{p(1-p)}{1-p(1-p)}, \frac{(p^2)}{1-p(1-p)} \right]$$

qui, pour $p = 1/3$, donne bien $\pi = [0.5714, 0.2857, 0.1429]$.

3.4 Chaînes de Markov à temps continu

Par rapport au cas discret, les transitions peuvent survenir à tout instant. De plus, il n'y a plus de transition d'un état vers lui-même.

Pour que le processus soit markovien, il faut que le temps de séjour dans un état soit sans mémoire : le temps restant ne dépend pas du temps déjà passé dans cet état. On a vu que la distribution exponentielle possède cette propriété, et c'est en fait la seule distribution continue ayant ce caractère "sans mémoire". Pour une chaîne de Markov continue, le temps de séjour dans un état est donc exponentiellement distribué.

Définissons le taux de transition de l'état i vers l'état j :

$$q_{i,j} = \lim_{\Delta t \rightarrow 0} \frac{P[X_{t+\Delta t} = j \mid X_t = i]}{\Delta t} \quad i \neq j$$

Quand le processus est dans l'état i , $q_{i,j} \Delta t$ représente ainsi la probabilité qu'une transition de i vers j survienne pendant la durée Δt . Le taux de sortie de l'état i est donc donnée par,

$$q_i = \sum_{j \neq i} q_{i,j}$$

La durée de séjour dans l'état i est donc une v.a. qui suit une distribution exponentielle de paramètre q_i .

La matrice des taux de transitions est définie par,

$$Q = \begin{pmatrix} -q_0 & q_{0,1} & q_{0,2} & \dots \\ q_{1,0} & -q_1 & q_{1,2} & \dots \\ q_{2,0} & q_{2,1} & -q_2 & \dots \\ \vdots & \vdots & \vdots & \ddots \end{pmatrix}$$

Soit $\pi(t) = [\pi_0(t), \pi_1(t), \dots]$ la distribution de probabilités de l'état à l'instant t . Examinons l'évolution de $\pi_i(t)$ sur Δt . Au 1^{er} ordre,

$$\pi_i(t + \Delta t) = (1 - q_i \Delta t) \pi_i(t) + \sum_{j \neq i} q_{j,i} \Delta t \pi_j(t) \quad \forall i$$

On en déduit que,

$$\frac{\pi_i(t + \Delta t) - \pi_i(t)}{\Delta t} = -q_i \pi_i(t) + \sum_{j \neq i} q_{j,i} \pi_j(t) \quad \forall i$$

Quand $\Delta t \rightarrow 0$, cela donne :

$$\frac{d\pi_i(t)}{dt} = -q_i \pi_i(t) + \sum_{j \neq i} q_{j,i} \pi_j(t) \quad \forall i$$

Soit sous forme matricielle :

$$\frac{d\pi(t)}{dt} = \pi(t) Q$$

La solution de cette équation différentielle matricielle est donnée par,

$$\pi(t) = \pi(0) e^{Qt} \quad \text{où} \quad e^{Qt} = I + Qt + \frac{1}{2!} Q^2 t^2 + \dots$$

Cette équation nous donne (au moins numériquement) la solution transitoire de la chaîne de Markov. Si elle existe, la solution stationnaire $\pi = \lim_{t \rightarrow \infty} \pi$ est indépendante du temps, i.e. $\frac{d\pi(t)}{dt} = 0$. La solution stationnaire vérifie donc :

$$\pi Q = 0$$

Cette équation nous donne la condition d'équilibre global. En effet, elle est équivalente à :

$$\pi_j \sum_{i \neq j} q_{j,i} = \sum_{i \neq j} \pi_i q_{i,j} \quad j = 0, 1, \dots$$

Autrement dit, à l'équilibre, la somme des taux de transition sortant d'un état j est égale à la somme des taux de transition qui y entrent.

Pour déterminer la solution stationnaire :

- utiliser la condition $\pi Q = 0$ pour exprimer tous les π_j , $j \neq 0$, en fonction de π_0 ,

- puis utiliser la normalisation des probabilités $\sum_i \pi_i = 1$, pour déterminer π_0 et donc les π_j , $j \neq 0$.

3.5 Processus de naissances et de morts

Les processus de naissances et de morts sont des cas particuliers de chaînes de Markov à temps continu très utilisés dans la théorie des files d'attente. Dans un processus de naissances et de morts, les transitions se font entre états adjacents : les transitions $i \rightarrow i + 1$ sont les naissances, et les transitions $i \rightarrow i - 1$ sont les morts. Dans la théorie des files d'attente, les naissances correspondent à des arrivées de clients (des appels ou des paquets par exemple) et les morts à leurs départs.

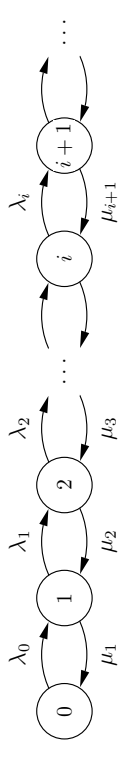


FIGURE 3.5 – Processus de naissances et de morts.

On notera λ_i le taux de naissance dans l'état i et μ_i le taux de mort, comme illustré sur la figure 3.5. Quand le processus est dans l'état i , la probabilité d'une naissance durant Δt est donc $\lambda_i \Delta t$ et celle d'une mort est $\mu_i \Delta t$ (temps de séjour exponentiellement distribué). Les taux de transitions $q_{i,j}$ sont donc donnés par :

$$q_{i,j} = \begin{cases} \lambda_i & \text{si } j = i + 1 \\ \mu_i & \text{si } j = i - 1 \\ 0 & \text{sinon} \end{cases}$$

Le générateur infinitésimal d'un processus de naissances et de morts est donc de la forme suivante :

$$Q = \begin{pmatrix} -\lambda_0 & \lambda_0 & 0 & \dots & \dots & \dots \\ \mu_1 & -(\lambda_1 + \mu_1) & \lambda_1 & 0 & \dots & \dots \\ 0 & \mu_2 & -(\lambda_2 + \mu_2) & \lambda_2 & 0 & \dots \\ \vdots & \vdots & \ddots & \ddots & \ddots & \ddots \end{pmatrix}$$

3.5.1 Distribution stationnaire

Le vecteur de probabilités stationnaires π est solution de $\pi Q = 0$. On obtient ainsi,

$$\begin{aligned} \mu_1 \pi_1 - \lambda_0 \pi_0 &= 0 \\ \lambda_{i-1} \pi_{i-1} - (\lambda_i + \mu_i) \pi_i + \mu_{i+1} \pi_{i+1} &= 0 \quad i \geq 1 \end{aligned}$$

Posons $v_i = \lambda_i \pi_i - \mu_{i+1} \pi_{i+1}$. Si on examine les deux équations ci-dessus, on voit que la première donne $v_0 = 0$ et la seconde $v_{i+1} - v_i = 0$. On en déduit que $v_i = 0$, $\forall i$, et donc que,

$$\pi_i = \frac{\lambda_{i-1}}{\mu_i} \pi_{i-1} = \frac{\lambda_{i-1} \lambda_{i-2}}{\mu_i \mu_{i-1}} \pi_{i-2} = \dots = \prod_{k=0}^{i-1} \frac{\lambda_k}{\mu_{k+1}} \pi_0 \quad i \geq 1$$

On a donc tous les π_i en fonction de π_0 . En utilisant la normalisation des probabilités, on obtient :

$$\sum_{i=0}^{\infty} \pi_i = 1 \Rightarrow \pi_0 = \frac{1}{1 + \sum_{i=1}^{\infty} \prod_{k=0}^{i-1} \frac{\lambda_k}{\mu_{k+1}}}$$

3.5.2 Distribution transitoire

On suppose connaître la distribution $\pi(0)$ de l'état à l'instant 0, par exemple $\pi_k(0) = 1$ et $\pi_j(0) = 0$ pour $j \neq k$. On veut déterminer son évolution transitoire $\pi(t)$, et plus seulement sa valeur limite $\pi = \lim_{t \rightarrow \infty} \pi(t)$. On a vu que la distribution transitoire est solution de,

$$\frac{d\pi(t)}{dt} = \pi(t) Q$$

Pour un processus de naissances et de morts, cela donne :

$$\begin{aligned} \frac{d\pi_0(t)}{dt} &= \mu_1 \pi_1(t) - \lambda_0 \pi_0(t) \\ \frac{d\pi_i(t)}{dt} &= \lambda_{i-1} \pi_{i-1}(t) - (\lambda_i + \mu_i) \pi_i(t) + \mu_{i+1} \pi_{i+1}(t) \quad i \geq 1 \end{aligned}$$

Ces équations permettent de déterminer $\pi(t)$ par intégration numérique. Soit Δt le pas d'intégration. Si on utilise la méthode d'Euler, on va ainsi déterminer les valeurs de $\pi(t)$ aux instants $t = \Delta t, t = 2\Delta t, \dots$, à partir de $\pi(0)$. Partant de $t = 0$, il suffit d'itérer sur les équations suivantes :

$$\begin{aligned} A/ \quad \pi_i(t + \Delta t) &= \pi_i(t) + \Delta t \frac{d\pi_i(t)}{dt}, \quad i \geq 0 \\ B/ \quad t &= t + \Delta t \end{aligned}$$

où les valeurs des dérivées ont été données ci-dessus.

3.6 Processus de Poisson

Le processus de Poisson peut être défini comme un processus de naissances pur tel que $\lambda_i = \lambda$ et $\mu_i = 0$, $\forall i$ (cf. figure 3.6). Le paramètre λ , qui est indépendant de l'état du processus, correspond au taux de naissance.



FIGURE 3.6 – Processus de Poisson.

Le processus de Poisson est un des modèles les plus utilisés dans la théorie des files d'attente. En effet, supposons que chaque naissance corresponde à l'arrivée d'un client (un appel ou un paquet par exemple) dans la file d'attente. Notons $N(t)$ l'état du processus de Poisson en supposant que $N(0) = 0$. On voit bien que $N(t) = i$ signifie qu'il y a eu i arrivées de clients entre 0 et t : le processus de Poisson permet de compter le nombre d'arrivées (cf. figure 3.7).

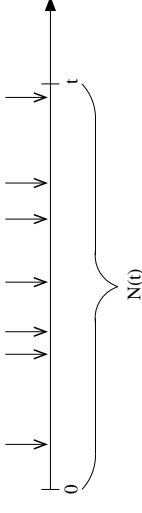


FIGURE 3.7 – Comptage du nombre d'arrivées.

Avec cette analogie, on voit bien que le temps de séjour dans un état correspond à la durée entre deux arrivées. Par conséquent :

- Le temps inter-arrivée est exponentiellement distribué de paramètre λ , i.e.,

$$P[\text{durée inter-arrivée} > t] = e^{-\lambda t}$$

- La durée moyenne entre deux arrivées est $1/\lambda$ et sa variance est $1/\lambda^2$.
- Pour dt "suffisamment petit", la probabilité d'avoir une arrivée durant dt est λdt et la probabilité de plus d'une arrivée est négligeable.

Toujours en supposant que $N(0) = 0$, intéressons nous à la probabilité d'avoir i arrivées entre 0 et t , c'est à dire à la probabilité $\pi_i(t)$ que $N(t) = i$. On peut pour cela utiliser les résultats sur la distribution transitoire d'un processus de naissances et de morts établis dans la section précédente. Avec $\lambda_i = \lambda$ et $\mu_i = 0$, $\forall i$, ces résultats deviennent :

$$\begin{aligned}\frac{d\pi_0(t)}{dt} &= -\lambda \pi_0(t) \\ \frac{d\pi_i(t)}{dt} &= \lambda \pi_{i-1}(t) - \lambda \pi_i(t) \quad i \geq 1\end{aligned}$$

La première équation a pour solution évidente $\pi_0(t) = e^{-\lambda t}$. C'est la probabilité qu'il n'y ait aucune arrivée entre 0 et t . Pour $i = 1$, si on multiplie la seconde équation par $e^{\lambda t}$, on a :

$$\frac{d\pi_1(t)}{dt} e^{\lambda t} + \lambda e^{\lambda t} \pi_1(t) = \lambda e^{\lambda t} \pi_0(t) = \lambda e^{\lambda t} e^{-\lambda t} = \lambda$$

On reconnaît dans le membre de gauche la dérivée d'un produit :

$$\frac{d}{dt} [\pi_1(t) e^{\lambda t}] = \lambda$$

La solution de cette équation est évidemment :

$$\pi_1(t) e^{\lambda t} = \lambda t \Rightarrow \pi_1(t) = \lambda t e^{-\lambda t}$$

C'est la probabilité d'avoir une seule arrivée entre 0 et t . En procédant par récurrence, on montre ainsi que,

$$\pi_i(t) = \frac{(\lambda t)^i}{i!} e^{-\lambda t}$$

On voit ainsi que le nombre d'arrivée entre 0 et t d'un processus de Poisson est déterminé par la distribution de Poisson.

L'hypothèse de durées inter-arrivée exponentiellement distribuées peut sembler a priori restrictives. Bien qu'elle ne soit pas toujours vérifiée, le processus de Poisson est souvent un très bon modèle pour beaucoup de phénomènes physiques. L'explication tient dans un résultat très fort démontré par Khintchine (théorème de Palm-Khintchine). En gros, ce résultat stipule que si on observe un processus d'arrivées qui, au niveau microscopique est constitué d'un très grand nombre de micro-processus générant un événement d'arrivée très rarement, alors le processus d'arrivée correspond à un processus de Poisson. C'est ce qui explique par exemple que l'on considère que les arrivées d'appels dans un réseau téléphoniques se font suivant un processus de Poisson (grand nombre d'abonnés).

Enfin, citons quelques propriétés importantes du processus de Poisson :

- **Stabilité par agrégation** : la superposition de deux flots d'arrivées Poissonniennes de taux λ_1 et λ_2 donne un processus de Poisson de taux $\lambda_1 + \lambda_2$.

- **Stabilité par désagrégation** : Si un flot d'arrivées Poissonniennes de taux λ est divisé en n sous-flots, chaque arrivée étant affectée au sous-flot i avec la probabilité p_i , alors les n sous-flots sont Poissonniens de paramètres $p_i \lambda$, $i = 1, \dots, n$.

- **Propriété PASTA** : Poisson Arrival See Time Averages (les arrivées poissonniennes "voient" les moyennes temporelles). Concrètement, cela signifie que si un système à la probabilité π_s d'être dans l'état s à l'équilibre, des arrivées suivant un processus de Poisson vont voir en moyenne le système dans l'état s .

Exercices

Exercice 1 Soit R une v.a. positive discrète. Considérons,

$$\begin{aligned}P(R = 1) + P(R = 2) + P(R = 3) + \dots \\ + P(R = 2) + P(R = 3) + \dots \\ + P(R = 3) + \dots \\ + \dots\end{aligned}$$

- Qu'obtient on en sommant sur une colonne ?
- Qu'obtient on en sommant sur une ligne ?
- Conclure que $E(R) = \sum_{i=0}^{\infty} P(R > i)$.

Exercice 2 On considère une chaîne de Markov ayant les probabilités de transition suivantes : $p_{12} = 1, p_{21} = 1/2, p_{23} = 1/2, p_{32} = 1/2$ et $p_{33} = 1/2$. Déterminer la distribution stationnaire de cette chaîne de Markov.

Exercice 3 Une base de données est l'objet de deux types d'accès : des accès en lecture et des accès en écriture. On suppose que ces accès se font suivant deux processus de Poisson indépendants, de taux λ_R pour les lectures et de taux λ_W pour les écritures.

- Quelle est la probabilité que le temps entre deux accès consécutifs en lecture soit supérieur à t ?
- Quelle est la probabilité que le prochain accès soit un accès en lecture ?
- Quelle est la probabilité que pendant l'intervalle de temps $[0, t]$ au plus trois accès en écriture se produisent ?
- Quelle est la probabilité que sur le même intervalle de temps il y ait au moins deux accès à la base ?

Exercice 4 La page web du célèbre professeur T.E. Acher est régulièrement consultée par des étudiants du monde entier. De nouvelles visites se produisent suivant un processus de Poisson avec en moyenne 10 visites par heure. Mr. Acher

est également très respecté par ses collègues. Le nombre moyen de collègues qui visitent sa page web en une heure est 2 (également suivant un processus de Poisson).

- (a) Quelle est la probabilité que la page web soit visitée au moins deux fois en une heure ?
- (b) Quelle est la probabilité que la page web ne soit pas visitée du tout en 15 minutes ?

Un étudiant visitant la page web “clique” sur le lien présentant les activités de recherche du professeur Acher avec une probabilité $2/5$.

- (c) Déterminer la probabilité que durant une journée de travail (8 heures) exactement un étudiant consulte la page décrivant les activités de recherche du professeur Acher ?

Exercice 5 Un processeur est examiné chaque semaine pour déterminer son état. L'état peut être parfait, bon, raisonnable ou mauvais. Un nouveau processeur est toujours parfait après une semaine avec la probabilité 0.7, avec la probabilité 0.2 son état est bon, et avec la probabilité 0.1 son état est raisonnable. Un processeur dans un état bon est toujours bon après une semaine avec la probabilité 0.6, dans un état raisonnable avec la probabilité 0.2 et dans un état mauvais avec la probabilité 0.2. Un processeur dans un état raisonnable a 50% de chances de l'être encore une semaine après, et 50% de chances d'être dans un état mauvais. Un processeur dans un état mauvais doit être réparé, ce qui prend une semaine, après quoi il se retrouve dans un état parfait.

- (a) Dessiner le diagramme de transition de la chaîne de Markov décrivant l'évolution de l'état du processeur de semaine en semaine.
- (b) Déterminer la distribution stationnaire de l'état du processeur.

Exercice 6 Considérons le processus markovien à temps continu décrivant le nombre d'utilisateurs dans un système (l'état i veut donc dire qu'il y a i utilisateurs dans le système). Le système est tel qu'il a les taux de transition suivant : $q_{01} = \lambda$, $q_{10} = \mu$, $q_{12} = \lambda$ et $q_{20} = \mu$.

- a) Ecrire le diagramme de transition.
- b) Ecrire la matrice des taux de transition Q .
- c) Déterminer la distribution stationnaire et le nombre moyen d'utilisateurs dans le système.

Exercice 7 Une chaîne de Markov dont les états sont $1, \dots, 4$ a la matrice de probabilités de transition ci-dessous :

$$P = \begin{pmatrix} 1-p & 0 & p & 0 \\ q & 0 & q & 0 \\ 0 & 0 & 1-p & p \\ 1 & 0 & 0 & 0 \end{pmatrix}$$

- a) Quelle doit être la valeur de q ?
- b) Ecrire le diagramme de transition et préciser le type de chaque état
- c) Quelle est la probabilité que le système soit dans l'état 4 à l'instant 4 en supposant qu'il est dans l'état 2 à l'instant 2 ?

Exercice 8 Soit un processus de naissances et de morts ayant pour espace d'états $i = 0, 1, 2, \dots$, et pour taux de transition $q_{i,i+1} = \lambda$, $q_{i+1,i} = (i+1)\mu$.

- a. Déterminer la distribution à l'équilibre..
- b. Reconnaissez vous cette distribution ?

CHAPITRE 4

Théorie des files d'attente

4.1 Introduction

La théorie des files d'attente est l'un des formalismes les plus utilisés pour la modélisation des réseaux de communication. De manière générale, la théorie des files d'attente s'intéresse à des systèmes dans lesquels des "clients" partagent une ou plusieurs ressources. Les dates d'arrivée des clients et la quantité de service qu'ils demandent sont imprévisibles. Ils sont en concurrence pour l'accès aux ressources. L'objectif est de déterminer les performances du système et notamment la qualité de service offerte aux clients en fonction de la capacité des ressources et de la politique de partage mise en oeuvre.

Dans les cas les plus simples, une file d'attente peut être représentée comme sur la figure 4.1. Le système est constitué d'un serveur muni d'une file d'attente. Des clients arrivent et demandent une certaine durée de service. Quand un client ne peut être servi parce que le serveur est déjà occupé, ce client est mis en attente dans la file d'attente.

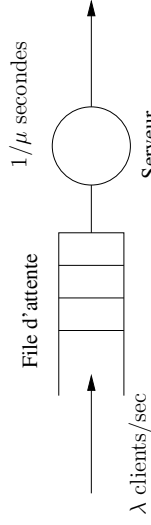


FIGURE 4.1 – File d'attente.

L'étude de cette file d'attente correspond à celle du processus stochastique

gouvernant le nombre de clients dans le système. L'aléatoire réside ici dans la variabilité des durées services et dans l'imprévisibilité des arrivées de clients. Les instants d'arrivées constituent un processus stochastique ; plus précisément, on considère que le temps qui sépare les instants d'arrivée de deux clients consécutifs sont des variables aléatoires (v.a.) indépendantes et identiquement distribuées (i.i.d.). De même, les temps de service requis par les clients constituent des v.a. i.i.d. L'évolution du nombre de clients est donc aléatoire, et dépend du choix de la discipline de service, c'est à dire, de l'ordre dans lequel les client seront servis.

Dans ce chapitre, nous commençons par introduire les concepts de base de la théorie des files d'attente. Nous étudions ensuite le cas simple de quelques files d'attente markoviennes. Puis, nous présentons quelques résultats sur des systèmes non markoviens. Enfin, nous énonçons les principaux théorèmes de la théorie des réseaux de files d'attente.

4.2 Les concepts de base des files d'attente

Nous introduisons ci-dessous les principaux concepts utilisés dans l'analyse des files d'attente (cf. Figure 4.2) :

- Arrivées, services et départs,
- Discipline de service,
- Notation de Kendall,
- Nombre de clients,
- Temps de réponse,
- Courbes de charge,,
- Périodes d'activité et d'inactivité,

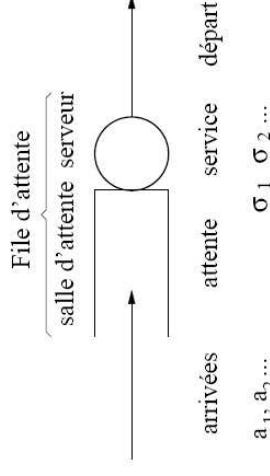


FIGURE 4.2 – Représentation graphique d'une file d'attente

Arrivées, services et départs

Les processus d'arrivée est décrit par la suite des dates successives d'arrivée :

$$a_1, a_2, \dots, a_n, a_{n+1}.$$

où a_n est la date d'arrivée du n ème client. La différence entre deux dates d'arrivée consécutives est appelée inter-arrivée. On note par τ_n la n -ième inter-arrivée : $\tau_n = a_{n+1} - a_n$. De façon équivalente, on peut donc décrire le processus d'arrivée par la suite des inter-arrivées,

$$\tau_1, \tau_2, \dots, \tau_n, \tau_{n+1}, \dots$$

Définition 1 Taux d'arrivée Le taux d'arrivée des clients (également appelé intensité des arrivées), est le nombre moyen de clients qui arrivent dans le système par unité de temps.

On écrit :

$$\lambda = \lim_{t \rightarrow \infty} \frac{A(t)}{t},$$

où :

$$A(t) = \max_n \{n | a_n \leq t\}.$$

A chaque client entrant dans la file d'attente est associé une demande de service. Le service demandé par le n -ième client sera noté σ_n .

Après son arrivée dans le système à la date a_n , le n -ème client passe éventuellement un certain temps dans la file d'attente en attendant son tour. Puis ce client entre en service pour une durée σ_n . Quand il a reçu la quantité de service demandée, ce client quitte le système. On notera d_n sa date de départ. Le processus aléatoire des départs est alors défini par la suite,

$$d_1, d_2, \dots, d_n, d_{n+1}, \dots$$

Discipline de service

La discipline de service précise la politique de partage de la capacité du (des) serveur(s) entre les clients. Deux grandes politiques peuvent être utilisées : soit la capacité d'un serveur est complètement dédiée à un seul client à la fois, soit elle est partagée entre tous les clients présents.

Dans le premier cas, la discipline de service décrit l'ordre dans lequel les clients sont servis. Les disciplines les plus courantes sont alors :

- FIFO (First-In-First-Out) ou, en français, FAPS (Premier Arrivé Premier Servi) : service dans l'ordre défini par les arrivées. On rencontre également l'acronyme FCFS (First-Come- First-Served).

- LIFO (Last-In-First-Out) ou, en français, DAPS (Dernier Arrivé Premier Servi),
- Priorités : Il y a plusieurs classes de clients, chacune avec un niveau de priorité. C'est le client qui a la priorité la plus haute parmi ceux présents qui est servi en premier. S'il y en a plusieurs, c'est le premier arrivé (discipline FIFO par classe).

Concernant les politiques partageant la capacité d'un serveur entre les clients présents, nous n'étudierons ici que la discipline PS (Processor Sharing) ou Politique à temps partagé. Avec cette discipline, la capacité du serveur est partagée équitablement entre tous les clients présents.

La notation de Kendall

Il est clair qu'en fonction des hypothèses faites sur le processus des arrivées, sur la loi de service et sur la discipline de service, il existe une multitude de modèles de files d'attente. La notation de Kendall constitue un moyen commode de préciser le modèle étudié. Cette notation a la forme suivante :

$$A/S/P/K/D$$

où :

A et S désignent respectivement la loi des inter-arrivées et la loi de service. A et S peuvent prendre les valeurs suivantes (liste non exhaustive) :

- M pour la distribution exponentielle (M comme Markov) ; Si $A=M$, le processus d'arrivée est un processus de Poisson.
- D pour la distribution déterministe ; Si $A=D$, le processus d'arrivée est périodique, tandis que si $S=D$, les durées de service demandées correspondent à une constante.
- G (ou GI) pour une distribution générale (arbitraire). Un résultat établi dans le cas $A=G$ sera donc valable quel que soit le processus des arrivées. De même, un résultat établi dans le cas $S=G$ sera valable pour toute loi de service.

P désigne le nombre de serveurs. P peut prendre n'importe quelle valeur entière, y compris l'infini. Par défaut, $P=1$.

K désigne la capacité du système, c'est-à-dire le nombre maximum de clients qui peuvent être présents simultanément, en comptant le(s) client(s) en service. K peut prendre n'importe quelle valeur supérieure à P , y compris l'infini. Par défaut, $K = \infty$.

D désigne la discipline de service. Dans ce cours, nous étudierons surtout les cas où $D=FIFO$ et $D=PS$. Par défaut, $D=FIFO$.

Pour donner un exemple, la file $M/M/1$ correspond à la file $M/M/1/\infty/FIFO$, c'est à dire une file d'attente avec un serveur, des arrivées suivant un processus de Poisson, des demandes de service exponentiellement distribuées, une file d'attente de capacité infinie et une discipline FIFO. Dans la suite de ce chapitre, nous allons étudier attentivement quelques exemples de tels systèmes.

Client	Date d'arrivée	durée de service
1	0	4
2	2	8
3	6	4
4	8	4
5	15	2

TABLE 4.1 – Arrivées et services

Nombre de clients

L'un des principales fonctions que l'on utilise pour étudier le comportement d'une file d'attente est $N(t)$, le nombre de clients présents dans la file d'attente à l'instant $t \geq 0$. $N(t)$ est une fonction à valeurs dans l'ensemble $\{0, 1, \dots, K\}$, où K est la capacité de la file d'attente, donc $N(t)$ est une courbe en escalier.

Illustrons ce concept sur un exemple. Considérons une file d'attente avec un serveur, de capacité d'attente illimitée et avec la discipline de service FIFO. Supposons que les dates d'arrivée et les durées de service des 5 premiers clients soient celles données dans le Tableau 4.1.

L'évolution du nombre de clients correspond alors à la courbe représentée sur la Figure 4.3. On a représenté en dessous les instants de début et de fin de service des cinq clients.

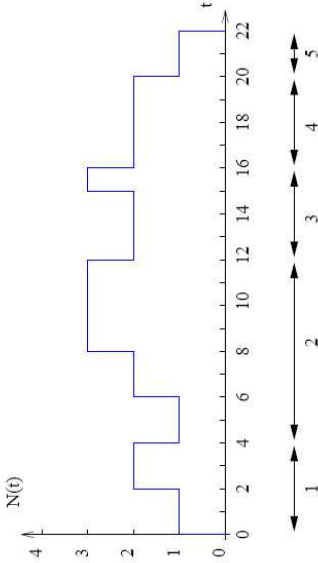


FIGURE 4.3 – Courbe du nombre de clients

Temps de réponse

Le temps de réponse (ou temps de séjour) d'un client dans une file d'attente est la différence entre sa date de départ et sa date d'arrivée. Avec les notations introduites précédemment, le temps de réponse du client n vaut

$$T_n = d_n - a_n.$$

Généralement, le temps de réponse se décompose en d'une part le temps d'attente (temps passé dans la file d'attente) et la durée de service (temps passé dans un des serveurs).

4.2.1 La formule de Little

La formule de Little est l'un des résultats les plus utilisés dans la théorie des files d'attente. Cette formule permet d'établir un lien entre le nombre moyen de clients et le temps de réponse moyen des clients.

Proposition 1 Soit λ le taux d'arrivée dans une file d'attente à l'équilibre. Soit $\bar{N} = E[X]$ le nombre moyen de clients dans le système et soit $\bar{T}$ le temps de réponse moyen d'un client. Alors

$$\bar{N} = \lambda \bar{T}.$$

Cette formule peut s'appliquer quand le système est stable. Elle est très utile car souvent on connaît $\bar{N}$ et λ . Elle dit que le nombre moyen de clients dans le système est égal au taux moyen d'arrivée multiplié par le temps moyen passé dans le système. Notons que ce résultat est extrêmement général puisqu'il ne fait aucune hypothèse sur le processus des arrivées, sur la loi de service, sur le nombre de serveurs ou encore sur la discipline de service.

Une explication intuitive est la suivante. Le rapport $\bar{N}/\bar{T}$ est égal au taux de départ, qui doit être égal au taux d'arrivée λ car le système se trouve en équilibre.

Charge et courbe de charge

La charge (workload) du système à l'instant t , $W(t)$, est la quantité de travail qui reste à effectuer à cet instant. Une courbe de charge est un graphe de la fonction $W(t)$. Quand on la construit, on a le choix de l'unité de travail. Celle-ci dépend en fait de l'application. Par exemple dans les réseaux de communication la quantité de travail pourra être des octets ou des bits. En général, on se ramène à des secondes en utilisant la vitesse du serveur, mesurée en nombre d'unités de travail par unité de temps. Par exemple, dans un réseau de communication :

$$\frac{N_{octets}}{C_{octets/sec}} = \frac{N}{C} sec.$$

La courbe de charge correspondant au Tableau 4.1 est représenté dans la Figure 4.4.

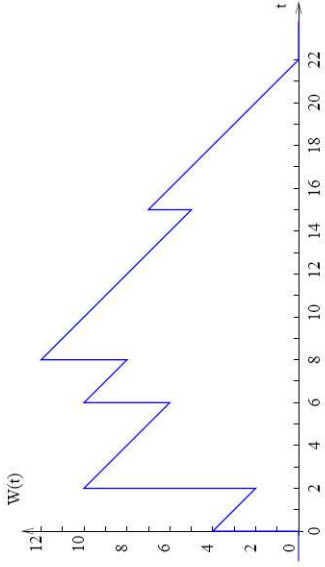


FIGURE 4.4 – Courbe de la charge $W(t)$

On voit sur cette figure les propriétés principales de cette fonction :

- C'est une courbe en dents de scie.
- Quand $W(t) > 0$, $W(t)$ décroît à vitesse constante sauf quand il y a une arrivée. La pente de la courbe est égale à la vitesse du serveur. Si on a choisi la même unité en abscisse et en ordonnée, la pente vaut exactement -1 .
- Elle croît par sauts aux instants d'arrivée. L'amplitude du saut est égal au temps de service requis par le client qui vient d'arriver.
- Quand $W(t)$ atteint 0, elle reste constamment nulle jusqu'à la prochaine arrivée.

La courbe de charge $W(t)$ et le nombre de clients $N(t)$ sont deux manières complémentaires de décrire l'évolution d'une file d'attente. Bien évidemment, ces deux courbes ont des relations étroites. D'abord nous allons observer comment on peut relier ces quantités entre elles graphiquement. Sur la Figure 4.5, on a représenté les courbes $W(t)$ et $N(t)$ pour l'exemple considéré précédemment. On voit comment on peut "lire" les dates de départ d_n sur la courbe de $W(t)$ en prolongeant les parties décroissantes de la courbe jusqu'à l'axe horizontal. Ceci s'explique de la manière suivante. Quand un client arrive à la date a_n , il trouve dans la file d'attente une quantité de travail $W(a_n)$. Comme la politique de service est FIFO, ce client entrera en service quand cette charge sera servie, soit donc à la date $a_n + W(a_n)$. Il finira son propre service à la date

$$d_n = a_n + W(a_n) + \sigma_n.$$

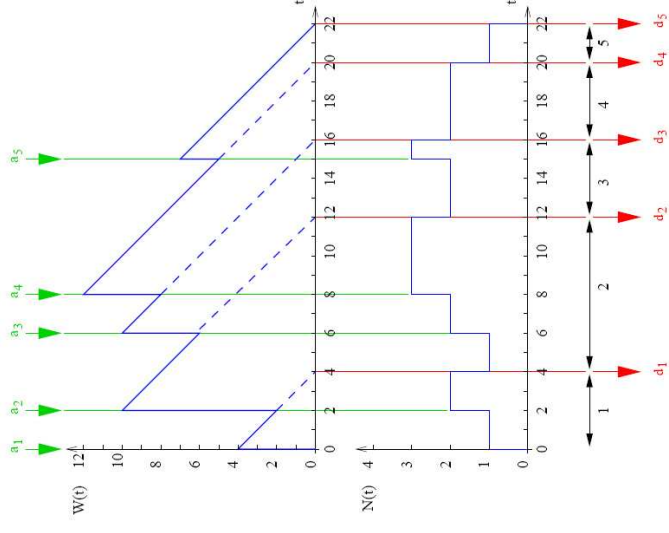


FIGURE 4.5 – Relation entre les courbes $W(t)$ et $N(t)$

Cette date se lit donc directement sur la courbe de $W(t)$. On note que le raisonnement qui précède ne vaut que si la discipline de service est FIFO.

Périodes d'activité et d'inactivité

Une période d'activité (busy period) du serveur (ou des serveurs) est un intervalle de temps pendant lequel des clients sont servis continuellement. De façon similaire, une période d'inactivité (idle period) est un intervalle de temps pendant lequel aucun service n'est délivré. Il est clair qu'une file d'attente suit une succession de périodes d'activité et d'inactivité.

Dans le cas des disciplines que nous considérons le serveur travaille quand, et seulement quand il y a des clients. Donc les périodes d'activité sont les intervalles où $W(t) > 0$ et $N(t) > 0$. On note que la durée des périodes d'activité et d'inactivité ne dépend pas de la discipline. On dit que la discipline est "work

conserving”.

4.3 Files d'attente Markoviennes

Dans cette partie, nous utilisons les résultats obtenus dans le chapitre précédent sur les processus de naissances et de morts pour résoudre quelques modèles simples mais importants de files d'attente.

4.3.1 La file $M/M/1$

Dans cette file les clients arrivent selon un processus de Poisson de taux λ (les inter-arrivées sont donc exponentiellement distribuées de moyenne $1/\lambda$ secondes). Chaque client requiert une quantité de service aléatoire suivant une loi exponentielle de paramètre μ , c'est à dire que la durée moyenne d'un service est $1/\mu$ secondes. Il y a un seul serveur et les clients sont servis un par un dans l'ordre de leur arrivée (FIFO). La capacité de la file d'attente est infinie.

Soit $N(t)$ le nombre de clients à l'instant t . Le processus $\{N(t), t \geq 0\}$ prend ses valeurs dans $\{0, 1, 2, \dots\}$. Les inter-arrivées et les durées de service étant exponentiellement distribuées, le temps de séjour dans un état $N(t) = i$ est lui aussi exponentiellement distribué. De plus, les transitions se font vers les états $N(t) = i - 1$ (départ d'un client) et $N(t) = i + 1$ (arrivée d'un client). Il s'en suit que le processus $\{N(t), t \geq 0\}$ est un processus de naissances et de morts dont le diagramme de transition est représenté sur la figure 4.6.

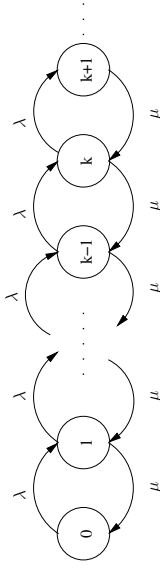


FIGURE 4.6 – Diagramme de transition de la file $M/M/1$.

Soit π_k , $k \geq 0$ la probabilité stationnaire qu'il y ait k clients dans le système, si elle existe. D'après les résultats obtenus sur les processus de naissances et de morts, on sait que si le régime d'équilibre existe, on a :

$$\pi_k = \prod_{i=0}^{k-1} \frac{\lambda_i}{\mu_{i+1}} \pi_0 = \left(\frac{\lambda}{\mu}\right)^k \pi_0 \quad k \geq 1$$

où,

$$R = (1 - \pi_0)\mu = \lambda.$$

$$\pi_0 = \frac{1}{1 + \sum_{k=1}^{\infty} (\lambda/\mu)^k} = 1 - \lambda/\mu$$

Définissons $\rho = \lambda/\mu$ l'intensité du trafic, c'est à dire la quantité moyenne de travail qui arrive au système par unité de temps. On a alors le résultat suivant.

Proposition 2 Soit une file d'attente $M/M/1$ avec un taux d'arrivée λ et durée moyenne de service $1/\mu$. Si $\rho < 1$

$$\pi_k = (1 - \rho) \rho^k \quad (4.1)$$

pour tout $k \geq 0$.

La condition $\rho < 1$ est dite condition de stabilité de la file $M/M/1$. Elle s'interprète par la nécessité qu'en moyenne, la quantité de clients qui arrive au système dans une unité de temps doit être inférieure à la quantité moyenne de clients que le serveur peut servir dans ce même temps.

Une conclusion importante que l'on peut déduire du résultat précédent est que la distribution stationnaire ne dépend de λ et μ qu'à travers de leur rapport, la charge ρ .

Sous la condition de stabilité, le nombre moyen de clients à l'équilibre est donné par,

$$E[N] = \sum_{k=0}^{\infty} k \pi_k = (1 - \rho) \sum_{k=0}^{\infty} k \rho^k = \frac{\rho}{1 - \rho}.$$

Observons que $E[N] \rightarrow \infty$ si $\rho \rightarrow 1$, donc dans la pratique, si le système n'est pas stable le nombre moyen de clients divergera.

On peut aussi s'intéresser à la probabilité qu'il y ait plus de K clients dans la file d'attente :

$$P(N \geq K) = \sum_{k=K}^{\infty} \pi_k = (1 - \rho) \sum_{k=K}^{\infty} \rho^k = \rho^K.$$

En utilisant la loi de Little, on obtient directement le temps de réponse moyen des clients,

$$T = E[N]/\lambda = \frac{1}{\mu - \lambda}.$$

Quel est le taux de sortie des clients (débit R) d'une file d'attente $M/M/1$ à l'équilibre ? Puisque on est à l'équilibre, tout ce qui entre devrait sortir, et par conséquent la réponse devrait être $R = \lambda$. Vérifions-le :

4.3.2 La file $M/M/1/K$

Dans la pratique, les files d'attente sont toujours de capacité finie, et le dimensionnement des buffers est un problème important. Ainsi, un nouveau client sera perdu (rejeté) lorsqu'il trouvera le système plein en arrivant.

Une file $M/M/1/K$ est représentée sur la figure 4.7. Le fonctionnement de cette file d'attente est similaire à celui de la file $M/M/1$, sauf en ce qui concerne le rejet des clients quand le nombre de clients dans le système atteint K (y compris celui en service). S'il y a K clients dans le système et qu'un nouveau client se présente, ce dernier sera rejeté.

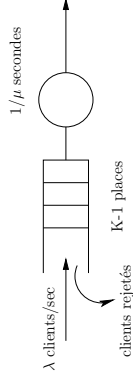


FIGURE 4.7 – File $M/M/1/K$.

Observons que pour cette file d'attente il n'y a pas de condition de stabilité. Le nombre de clients présents dans le système étant au maximum K , il ne peut pas diverger vers l'infini. Si le débit d'arrivée des clients λ est supérieur à la capacité de traitement μ du serveur ($\rho > 1$), alors la distribution stationnaire aura tendance à charger les états proches de K .

De façon évidente, le nombre de clients dans une file $M/M/1/K$ correspond à un processus de naissances et de morts, dont le diagramme de transition est représenté sur la figure 4.8.

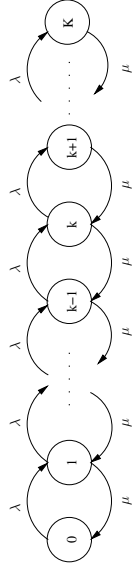


FIGURE 4.8 – Diagramme de transition de la file $M/M/1/K$.

En utilisant là encore les résultats obtenus sur les processus de naissances et de morts, on obtient,

$$\pi_k = \prod_{i=0}^{k-1} \frac{\lambda_i}{\mu_{i+1}} \pi_0 = \rho^k \pi_0 \quad k \geq 1$$

où,

$$\pi_0 = \frac{1}{\frac{K}{\rho} + \sum_{k=1}^K \rho^k} = \begin{cases} \frac{1-\rho}{1-\rho^{K+1}} & \rho \neq 1 \\ \frac{1}{K+1} & \rho = 1 \end{cases}$$

Proposition 3 Soit une file d'attente $M/M/1/K$ avec un taux d'arrivée λ et durée moyenne de service $1/\mu$. Si $\rho \neq 1$

$$\pi_k = \frac{(1-\rho)\rho^k}{1-\rho^{K+1}} \quad k = 0, 1, \dots, K \quad (4.2)$$

Si $\rho = 1$

$$\pi_k = \frac{1}{K+1} \quad k = 0, 1, \dots, K \quad (4.3)$$

Enfin, notons que la probabilité de rejet est π_K (propriété PASTA).

4.3.3 La file $M/M/C$

Dans cette file d'attente, il y a un nombre C de serveurs et la file d'attente à une capacité infinie (cf figure 4.9). Si lorsque un client arrive il y a au moins un serveur disponible, le client entre en service immédiatement. Dans le cas contraire, il est placé dans la file d'attente.

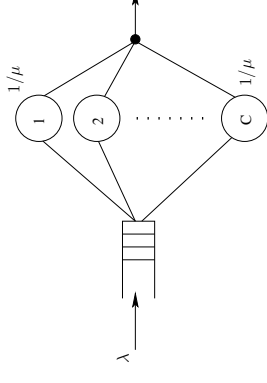


FIGURE 4.9 – File $M/M/C$.

Soit à nouveau $\{N(t), t \geq 0\}$ le nombre de clients dans le système. $N(t)$ est là encore un processus de naissances et de morts prenant ses valeurs dans $\{0, 1, 2, \dots\}$. Soit i un état arbitraire. Le taux de naissance est $\lambda_i = \lambda, \forall i \geq 0$. Le taux de mort est $\mu_i = \min(i\mu, C\mu)$. Avec ces valeurs de λ_i et μ_i , on obtient immédiatement le résultat suivant avec les mêmes techniques que précédemment.

Proposition 4 Soit une file d'attente $M/M/C$ avec un taux d'arrivée λ et une durée moyenne de service $1/\mu$. Si $\rho = \lambda/\mu < C$

$$\pi_i = \begin{cases} \pi_0 \frac{\rho^i}{i!} & \text{pour } i = 0, 1, 2, \dots, C \\ \pi_0 \frac{\rho^i}{C!} & \text{pour } i \geq C \end{cases}$$

où

$$\pi_0 = \left(\sum_{i=0}^{C-1} \frac{\rho^i}{i!} + \frac{\rho^C}{C!(1-\rho/C)} \right)^{-1}. \quad (4.4)$$

La probabilité qu'un client soit placé dans la file d'attente en arrivant est,

$$\begin{aligned} P(\text{attente}) &= \sum_{i=C}^{\infty} \pi(i) = \sum_{i=C}^{\infty} \pi(0) \frac{\rho^{i-C-i}}{C!} \\ &= \frac{\left(\frac{\rho^C}{C!}\right) \left(\frac{1}{1-\rho/C}\right)}{\left(\sum_{i=0}^{C-1} \frac{\rho^i}{i!} + \frac{\rho^C}{C!(1-\rho/C)}\right)} \end{aligned}$$

Cette probabilité est largement utilisé dans le dimensionnement des centres d'appels (hotlines) par exemple. Elle donne la probabilité que il n'y ait aucun serveur disponible pour un appel qui arrive. Elle est connue sous le nom de formule d'Erlang- C .

Le cas $C = \infty$

Dans ce cas, il y a un nombre infini de serveurs. Un client entre donc toujours immédiatement en service à son arrivée. La résolution du processus de naissances et de morts associé donne :

$$\pi_i = \pi_0 \frac{\rho^i}{i!}.$$

Avec $1 = \sum_{i=0}^{\infty} \pi_i$, on obtient $\pi_0 = e^{-\rho}$. Par conséquent la distribution stationnaire de la file $M/M/\infty$ est

$$\pi_i = \frac{\rho^i}{i!} e^{-\rho}, \quad i \geq 0.$$

Cette distribution est identique à une distribution de Poisson de paramètre ρ . Cette distribution existe toujours, et donc il n'y a pas de condition de stabilité pour la file $M/M/\infty$. Il se trouve que cette distribution de Poisson de paramètre ρ est aussi la distribution stationnaire de la file $M/GI/1$. Donc le résultat ne dépend pas de la distribution du temps de service, mais uniquement de sa moyenne. On parle d'insensibilité.

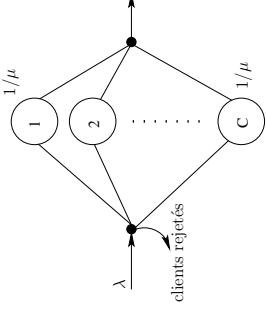


FIGURE 4.10 – File $M/M/C$.

Le nombre de clients dans le système peut être modélisé comme un processus de naissance et de morts. Le taux d'arrivée dans un état $i < C$ est $\lambda_i = \lambda$ tandis que $\lambda_C = 0$. Le taux de mort est $\mu_i = i\mu$ dans l'état $i = 0, 1, \dots, C$.

Proposition 5 Soit une file d'attente $M/M/C$ avec un taux d'arrivée λ et une durée moyenne de service $1/\mu$. La distribution stationnaire est :

$$\pi_i = \pi_0 \frac{\rho^i}{i!} \quad i = 0, 1, \dots, C \quad (4.5)$$

où,

$$\pi_0 = \frac{1}{\sum_{i=0}^C \rho^i / i!}.$$

Comme on le verra dans la seconde partie du cours, ce système est d'un grand intérêt en téléphonie. En particulier, π_C donne la probabilité que tous les serveurs soient occupés, c'est à dire la probabilité qu'une arrivée soit rejetée (propriété PASTA). Elle est donnée par :

$$\pi_C = \frac{\rho^C / C!}{\sum_{j=0}^C \rho^j / j!}$$

Cette formule est connue sous le nom de formule d'Erlang-B. Signalons de plus que de façon similaire à la file $M/M/\infty$, la Proposition 5 est valable pour n'importe quelle distribution du temps de service (insensibilité).

4.4 Exemples d'application

4.4.1 Comparaison de systèmes

Une entreprise fournissant des services web désire renouveler son parc de serveurs informatiques. Elle a le choix entre trois options illustrées sur la figure 4.11 :

- L'option A consiste à utiliser un seul serveur de capacité 2μ ,
- L'option B consiste à utiliser deux serveurs de capacité μ en parallèle, chacun avec sa propre file d'attente, et à transmettre 50% des requêtes vers chaque serveur,
- L'option C consiste à utiliser deux serveurs de capacité μ en parallèle, mais partageant cette fois-ci la même file d'attente. Dans ce cas, s'il y a une requête en attente de traitement, elle est transmise vers le premier serveur qui se libère.

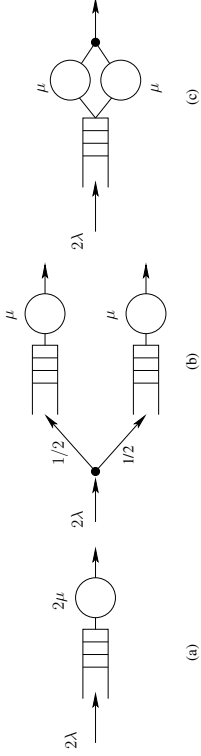


FIGURE 4.11 – Les trois options envisagées.

Dans chacune des options envisagées, les files d'attente sont de capacités infinies et sont gérées suivant la discipline FIFO. On suppose que les requêtes arrivent suivant un processus de Poisson de taux 2λ . Une fois prise en charge par un serveur, la durée de traitement d'une requête est aléatoire et suit une distribution exponentielle de moyenne $\frac{1}{\mu}$ dans l'option A, et de moyenne $\frac{1}{\mu}$ dans les options B et C.

On désire déterminer la configuration qui permet de minimiser le temps de réponse moyen aux requêtes. On remarquera que cette comparaison est équitable dans la mesure où le taux d'utilisation d'un serveur, dans les trois configurations envisagées, est toujours $\nu = \lambda/\mu$. On supposera dans la suite que $\nu < 1$.

Etude de l'option A

Cette option correspond à une file d'attente $M/M/1$ avec un taux d'arrivée 2λ , un taux de service 2μ , et donc un taux d'utilisation $\nu = 2\lambda/(2\mu)$. Le nombre moyen N_A de requêtes présentes dans le système (en attente ou en traitement) et le temps de séjour moyen T_A d'une requête sont donnés par,

$$N_A = \frac{\nu}{1 - \nu} \quad \text{et} \quad T_A = N_A / (2\lambda) = \frac{1}{2\mu(1 - \nu)}$$

Etude de l'option B

Le processus de Poisson étant stable par désagrégation, le processus d'arrivée des requêtes sur chacun des deux serveurs est un processus de Poisson de taux λ . Chacun des serveurs peut donc être modélisé par une file $M/M/1$, de taux d'arrivée λ , de taux de service μ et de taux d'utilisation $\nu = \lambda/\mu$. Le temps de séjour moyen T_B^i d'une requête sur le serveur i est donc,

$$T_B^i = \frac{1}{\mu - \lambda} = \frac{1}{\mu(1 - \nu)} \quad i = 1, 2$$

Le temps de séjour moyen T_B d'une requête dans la configuration B est,

$$T_B = \frac{T_B^1 + T_B^2}{2} = \frac{1}{\mu(1 - \nu)}$$

On voit ainsi que $T_A = T_B/2$. Autrement dit, le temps de réponse moyen de l'option A est deux fois plus court que celui de l'option B! Intuitivement, ceci s'explique par le fait que dans l'option B, quand une des deux files est vide, la moitié de la capacité n'est pas utilisée. Au contraire, dans l'option A, la totalité de la capacité est utilisée dès lors qu'il y a au moins une requête à servir.

Etude de l'option C

Cette configuration peut être modélisée par une file d'attente $M/M/2/\infty$. En utilisant les résultats obtenus sur la file $M/M/C$ avec $C = 2$ et $\rho = 2\lambda/\mu = 2\nu$, on obtient,

$$\pi_i = \begin{cases} \pi_0 \frac{\rho^i}{i!} & i = 0, 1 \\ \pi_0 \rho^i \frac{2^{i-2}}{2!} & i \geq 2 \end{cases}$$

Pour $i \geq 2$, on a ainsi :

$$\pi_i = \pi_0 (2\nu)^i \frac{2^{i-2}}{2} = \pi_0 \nu^i 2^{i+2-i-1} = 2\pi_0 \nu^i$$

Vu que $\pi_1 = 2\pi_0\nu$, on en déduit que $\pi_i = 2\pi_0\nu^i$ pour $i \geq 1$. Avec la condition de normalisation des probabilités, on a alors,

$$\pi_0 = \left(1 + 2 \sum_{i=1}^{\infty} \nu^i \right)^{-1} = \left(1 + 2 \left(\frac{1}{1 - \nu} - 1 \right) \right)^{-1} = \frac{1 - \nu}{1 + \nu}$$

Le nombre moyen de requêtes N_C présentes dans le système est alors,

$$N_C = \sum_{i=1}^{\infty} i\pi_i = 2\pi_0\nu \sum_{i=1}^{\infty} i\nu^{i-1} = 2 \frac{1-\nu}{1+\nu} \nu \frac{1}{(1-\nu)^2} = \frac{2\nu}{1-\nu^2}$$

Avec la loi de Little, on a le temps de séjour moyen T_C d'une requête,

$$T_C = N_C / (2\lambda) = \frac{1}{\mu(1-\nu^2)}$$

Comparons alors avec l'option A :

$$\frac{T_A}{T_C} = \frac{\mu(1-\nu^2)}{2\mu(1-\nu)} = \frac{1+\nu}{2}$$

Sachant que $\nu < 1$, on voit ainsi que $T_A < T_C$. L'explication intuitive est liée au fait que quand il y a un seul client dans le système, il est servi avec le débit μ dans l'option C, alors qu'il est servi avec le débit 2μ dans l'option A. On voit d'ailleurs avec l'équation ci-dessus que quand le système est proche de la saturation ($\nu \approx 1$), les deux options deviennent équivalentes.

4.4.2 Etude d'un mécanisme de token bucket

Un "token bucket" est un mécanisme de policing utilisé par les opérateurs de réseaux IP pour réguler le trafic en entrée de leur réseau et garantir qu'un client n'émet pas plus de trafic en moyenne que ce que prévoit son contrat. Le principe est illustré sur la figure 4.12.

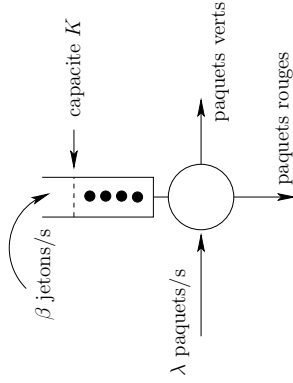


FIGURE 4.12 – Principe d'un token bucket.

Le fonctionnement est le suivant. Le token bucket correspond à une file d'attente FIFO dans laquelle des jetons sont placés régulièrement. On note β le taux d'arrivée des jetons. Le token bucket a d'autre part une capacité limitée : il ne peut contenir plus de K jetons. Lorsqu'un paquet IP arrive pour être transmis, il est marqué vert s'il reste des jetons dans le token bucket, et rouge sinon. Les paquets verts sont ensuite transmis de façon prioritaire, alors que les paquets

rouges sont soit déclassés (priorité moins importante), soit directement jetés. Ainsi, le token bucket permet de garantir que le débit λ_{vert} du trafic vert ne pourra pas dépasser β paquets/seconde.

Pour simplifier, on suppose ici que les paquets IP arrivent suivant un processus de Poisson. La durée entre l'arrivée de deux paquets IP consécutifs est donc exponentiellement distribuée de moyenne $1/\lambda$. De la même façon, on supposera que les jetons arrivent à l'entrée du token bucket suivant un processus de Poisson de taux β . Enfin, on supposera que $\beta \neq \lambda$.

Supposons que le token bucket ne soit pas vide, et considérons le jeton le plus ancien dans cette file. Ce jeton va être consommé quand un paquet arrivera. Sachant que les inter-arrivées de paquets sont exponentiellement distribuées de moyenne $1/\lambda$, on en déduit que le jeton va rester en tête du token bucket une durée exponentiellement distribuée de moyenne $1/\lambda$.

On peut alors modéliser l'évolution du nombre de jetons du token bucket par une file d'attente $M/M/1/K$. Les arrivées de jeton se faisant suivant un processus de Poisson de taux β et les durées de service (attente du prochain paquet) étant exponentiellement distribuées de moyenne $1/\lambda$, le taux d'utilisation ρ dans ce modèle est $\rho = \beta/\lambda \neq 1$. On en déduit la probabilité π_0 que le token bucket soit vide,

$$\pi_0 = \frac{1-\rho}{1-\rho^{K+1}}$$

En utilisant la propriété PASTA du processus de Poisson, on voit ainsi que la probabilité qu'un paquet soit marqué vert est $1 - \pi_0$. On en déduit le débit λ_{vert} du trafic vert en paquets/seconde :

$$\lambda_{vert} = (1 - \pi_0)\lambda = \frac{1 - \rho^K}{1 - \rho^{K+1}} \beta$$

Observons que l'on a bien $\lambda_{vert} < \beta$ puisque pour $\rho \neq 1$ on a toujours $1 - \rho^K < 1 - \rho^{K+1}$.

4.5 Files d'attente à loi de service générale

4.5.1 La file M/G/1/FCFS

Pour cette file, le processus d'arrivée est toujours un processus de Poisson de taux λ , mais les durées de service suivent maintenant une loi générale :

$$P(S \leq x) = B(x),$$

dont la moyenne est notée $E[S] = 1/\mu$. Ici, le processus $N(t)$ n'est plus un processus de Markov, et les techniques précédentes ne s'appliquent plus. En fait, la probabilité d'avoir une transition de l'état $\{N = n\}$ vers l'état $\{N = n - 1\}$

dépend maintenant de la quantité de service que le client en service a déjà reçu. Cette information n'est pas incluse dans la variable $N(t)$, sauf dans le cas d'une loi de service exponentielle (grâce à sa propriété sans-mémoire) qui est donc le seul cas pour lequel le système est markovien.

Néanmoins, le nombre moyen de clients, le temps moyen d'attente et le temps moyen de réponse dans le système peuvent être calculés. En fait, même la distribution de ces variables peut être calculée en étudiant la chaîne de Markov incluse (en n'observant le système qu'aux instants de départs).

Dans cette partie, nous montrons en suivant une autre approche comment calculer le temps d'attente moyen dans une file $M/G/1/FCFS$. La formule obtenue est connue sous le nom de formule de Pollaczek-Khinchin.

Soit $N_q(t)$ le nombre de clients qui attendent dans la file à l'instant $t \geq 0$. On a la relation $N_q(t) = \max\{N(t) - 1, 0\}$. De plus, notons $R(t)$ le temps de service résiduel du client en service avant de quitter le système (voir Figure 4.13). On note que si $N(t) = 0$ alors $R(t) = 0$. Enfin, notons $\rho = \lambda/\mu$ le taux d'utilisation du système. On suppose ci-dessous que $\rho < 1$ et on admettra que c'est une condition nécessaire et suffisante pour que le régime stationnaire existe.

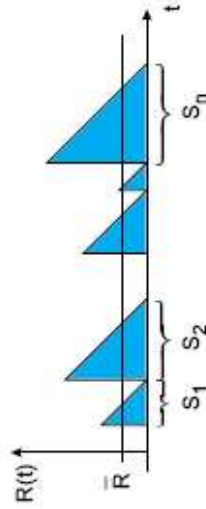


FIGURE 4.13 – Temps de service résiduel

Grâce à la propriété PASTA, nous avons la relation

$$E[W] = E[N_q] \frac{1}{\mu} + E[R]$$

L'idée clé est de remarquer que si l'on applique la loi de Little à la file d'attente (sans le serveur), on obtient :

$$E[W] = \lambda E[N_q],$$

et par conséquent,

$$E[W] = \frac{E[R]}{1 - \rho},$$

Ainsi, il ne nous reste qu'à déterminer la valeur de $E[R]$ pour avoir une expression de $E[W]$. Pour cela, considérons la Figure 4.13 sur un intervalle de temps t très long. La valeur moyenne de $R(t)$ peut être calculée de la façon suivante :

$$\begin{aligned} E[R] &= \lim_{t \rightarrow \infty} \frac{1}{t} \int_0^t R(s) ds \approx \lim_{t \rightarrow \infty} \frac{1}{t} \sum_{i=1}^n \frac{1}{2} S_i^2 \\ &= \lim_{t \rightarrow \infty} \frac{n}{t} \frac{1}{n} \sum_{i=1}^n \frac{1}{2} S_i^2 \\ &= \frac{1}{\lambda} E[S^2]. \end{aligned}$$

On obtient ainsi le résultat suivant.

Proposition 6 Si $\rho = \lambda/\mu < 1$, le temps d'attente moyen dans une file $M/G/1/FCFS$ est donné par l'expression suivante (formule P-K) :

$$E[W] = \frac{\lambda E[S^2]}{2(1 - \rho)}. \quad (4.6)$$

Comme conséquence directe, on obtient le temps de réponse moyen du système,

$$E[T] = \frac{1}{\mu} + E[W] = \frac{1}{\mu} + \frac{\lambda E[S^2]}{2(1 - \rho)}$$

En utilisant la formule de Little, on en déduit le nombre moyen de clients en attente ainsi que le nombre moyen de clients dans le système :

$$E[N_q] = \frac{E[S^2]}{2(1 - \rho)}, \quad \text{et} \quad E[N] = \rho + \frac{\lambda^2 E[S^2]}{2(1 - \rho)}$$

On notera en particulier que le nombre moyen de clients, ainsi que le temps de séjour augmentent avec la "variabilité" de la distribution du temps de service.

Exemples

A titre d'exemple, considérons tout d'abord le cas d'une file $M/M/1$ dont nous connaissons déjà la solution. La loi de service est exponentielle, ce qui implique que $E[S^2] = 2(\frac{1}{\mu})^2$. Par conséquent,

$$\begin{aligned}
E[N] &= \rho + \frac{\lambda^2 E[S^2]}{2(1-\rho)} \\
&= \rho + \frac{\rho^2}{1-\rho} = \frac{\rho}{1-\rho},
\end{aligned}$$

et

$$\begin{aligned}
E[T] &= \frac{1}{\mu} + \frac{\lambda E[S^2]}{2(1-\rho)} \\
&= \frac{1}{\mu} \left(1 + \frac{\rho}{1-\rho} \right) = \frac{1}{\mu} \frac{1}{1-\rho}.
\end{aligned}$$

Considérons maintenant le cas d'une file $M/D/1$. La durée de service est une constante, et on a donc $E[S^2] = E[S]^2$. On en déduit,

$$E[N] = \rho + \frac{1}{2} \frac{\rho^2}{1-\rho},$$

et

$$E[T] = \frac{1}{\mu} \left(1 + \frac{1}{2} \frac{\rho}{1-\rho} \right).$$

Application au contrôle d'admission

Le contrôle d'admission joue un rôle central dans les réseaux IP multimédia. En effet, certains trafics, comme la vidéo, ayant un débit très variable dans le temps, l'arrivée d'une nouvelle session peut avoir un impact très néfaste sur la qualité de service des sessions déjà établies, particulièrement quand le réseau est proche de la congestion. Il faut donc pouvoir refuser une session dans certains cas pour pouvoir préserver la qualité de service des autres sessions. En même temps, l'opérateur n'a pas intérêt à refuser trop de sessions. Il faut donc un critère simple et fiable pour accepter ou refuser l'établissement d'une nouvelle session.

Considérons un seul lien de communication. On suppose qu'il y a K différents types de session. On suppose qu'il y a n_k sessions du type $k = 1, \dots, K$ établies. Une nouvelle session du type i arrive. La question est la suivante : faut-il accepter ou refuser cette session pour que le délai moyen introduit sur les paquets reste inférieur à une constante $\alpha > 0$, c'est à dire pour que $E[W] < \alpha$.

Supposons qu'une session de type $k = 1, \dots, K$ génère des paquets suivant un processus de Poisson de taux λ_k . De tels paquets sont dits de classe k . Le processus d'arrivée des paquets est alors un processus de Poisson de taux $\lambda = \sum_{k=1}^K n_k \lambda_k$.

Considérons maintenant la distribution $B(x)$ du temps de service des paquets. Notons $B_k(x) = P(S_k \leq x)$ la distribution aléatoire régissant le temps de transmission d'un paquet de classe k , $1/\mu_k$ sa moyenne et $E[S_k^2]$ son moment d'ordre 2. A la fin de la transmission d'un paquet, la probabilité que le prochain paquet à transmettre soit de classe k est $n_k \lambda_k / \lambda$. Notons Y la variable aléatoire représentant la classe du prochain paquet à transmettre. En appliquant la loi des probabilités totales, on a ainsi :

$$\begin{aligned}
B(x) &= P(S \leq x) \\
&= \sum_{k=1}^K P(S \leq x, Y = k) \\
&= \sum_{k=1}^K P(S \leq x \mid Y = k) \frac{n_k \lambda_k}{\lambda} \\
&= \sum_{k=1}^K \frac{n_k \lambda_k}{\lambda} B_k(x)
\end{aligned}$$

La durée moyenne de transmission d'un paquet est donnée par :

$$\frac{1}{\mu} = \int_0^\infty x dB(x) = \sum_{k=1}^K \frac{n_k \lambda_k}{\lambda} \int_0^\infty x dB_k(x) = \sum_{k=1}^K \frac{n_k \rho_k}{\lambda},$$

où $\rho_k = \lambda_k / \mu_k$ est le trafic offert d'une session du type k . On supposera ici que le trafic offert $\rho = \lambda / \mu < 1$. Le moment d'ordre 2 de la durée de transmission d'un paquet est donné par :

$$E[S^2] = \int_0^\infty x^2 dB(x) = \sum_{k=1}^K \frac{n_k \lambda_k}{\lambda} \int_0^\infty x^2 dB_k(x) = \sum_{k=1}^K \frac{n_k \lambda_k}{\lambda} E[S_k^2],$$

En analysant le système comme une file $M/G/1/FCFS$ et en utilisant la formule PK, on voit que la condition $E[W] < \alpha$ peut s'écrire $\lambda E[S^2] < 2\alpha(1 - \rho)$, c'est à dire :

$$\sum_{k=1}^K n_k \lambda_k E[S_k^2] < 2\alpha \left(1 - \sum_{k=1}^K n_k \rho_k \right)$$

En posant $\alpha_k = \rho_k + \frac{\lambda_k E[S_k^2]}{2\alpha}$, la condition ci-dessus s'écrit simplement :

$$\sum_{k=1}^K n_k \alpha_k < 1 \quad (4.7)$$

On voit ainsi qu'il faut accepter la nouvelle session si $\sum_{k=1}^K n_k \alpha_k + \alpha_i < 1$, et la refuser sinon. La quantité α_k est appelée bande-passante équivalente d'une

session de type k . On voit qu'elle est toujours supérieure ou égale au trafic offert ρ_k . Elle n'est égale à ρ_k que dans le cas d'une session à débit constant (CBR : Constant Bit Rate). La différence entre les deux augmente avec la variabilité du trafic, mesurée ici par le terme $E[S_k^2]$ (variabilité de la taille des paquets). Il faut aussi noter que la bande-passante équivalente dépend de la qualité de service que l'on veut garantir au travers de la borne sur le délai α . La bande-passante équivalente est d'autant plus grande que la borne α est petite. Au contraire, si on renonce à garantir un délai minimum, i.e. quand $\alpha \rightarrow \infty$, alors on obtient $\alpha_k = \rho_k$ et la condition 4.7 est en fait juste la condition de stabilité $\rho < 1$.

4.5.2 La file M/G/1/PS

On considère un système avec un seul serveur de capacité μ dans lequel les arrivées se produisent suivant un processus de Poisson de taux λ . A chaque instant, le serveur partage équitablement sa capacité entre les clients présents. Ainsi, s'il y a n clients, chacun est servi avec le taux de service μ/n . On a le résultat suivant.

Proposition 7 Si $\rho = \lambda/\mu < 1$, alors la distribution stationnaire de la file M/G/1/PS existe et elle est donnée par,

$$\pi_n = \pi_0 \rho^n, \quad n = 0, 1, \dots \quad \text{où} \quad \pi_0 = 1 - \rho$$

On peut faire deux remarques importantes. Tout d'abord, on remarque que cette distribution est identique à celle d'une file M/M/1/FIFS de mêmes paramètres. Ensuite, cette distribution ne dépend pas réellement de la distribution des durées de service, mais seulement de sa moyenne.

4.6 Réseaux de files d'attente

Dans cette section, on présente les principaux résultats connus sur les réseaux de files d'attente. Les modèles à file d'attente unique ne permettent pas de modéliser des systèmes complexes. Par exemple, le problème du passage aux guichets d'une poste ne se résout pas avec une seule file, car chaque serveur à sa propre file d'attente. L'idée est donc de représenter chaque attente par une file simple, et de connecter ces files entre elles dans un réseau. Les questions que l'on se pose sont les mêmes : étant donnés les paramètres du modèle, quelle est la distribution du nombre de clients dans un site donné (ou à défaut sa moyenne), que valent les temps d'attente, quelles sont les conditions de stabilité du système, etc...

Le problème est évidemment plus compliqué que dans le cas d'une file simple. On ne peut pas espérer obtenir de résultats sans faire d'hypothèses plus restrictives (mais le moins possible) sur les lois d'arrivée, de service, le type des files simples qui constituent le réseau, le routage (la façon dont les clients se déplacent

dans le réseau) ou la topologie de ce réseau. Dans certains cas que nous allons voir, les résultats obtenus dans la Section 4.3 peuvent être généralisés.

4.6.1 Les réseaux de Jackson

Les réseaux de Jackson ont pour particularité la spécification du routage des clients de manière aléatoire. Les paramètres d'un tel réseau sont :

- Le nombre N de stations
- Le vecteur $\vec{\lambda}^0 = (\lambda_1^0, \dots, \lambda_N^0)$ des taux d'arrivées extérieures dans chaque file,
- Le vecteur $(\mu_1, \dots, \mu_N)$ des taux de service,
- La matrice carrée $N \times N$ de routage interne $\mathbf{R}$, dont chaque composante $r_{i,j}$ est la probabilité qu'un client sortant du noeud i aille vers le noeud j .

La probabilité que le client sorte du système est alors $r_{i,0} = 1 - \sum_j r_{i,j}$. La matrice $\mathbf{R}$ est une matrice sous-stochastique, ce qui revient à dire qu'un client a toujours un ou plusieurs chemins qui mènent vers l'extérieur, et finit donc par en prendre un. On suppose donc qu'il existe au moins un $r_{i,0} > 0$.

Le processus $N(t) = (N_1(t), \dots, N_N(t))$ qui décrit le nombre de clients dans chaque file d'attente du réseau au temps t , est une chaîne de Markov en temps continu sur $\mathbb{N}^N$. En utilisant une notation vectorielle : $\mathbf{n} = (n_1, \dots, n_N)$ et $\mathbf{e}_i = (0, \dots, 1, \dots, 0)$ pour le i -ième vecteur unitaire, on décrit le générateur $\mathbf{Q}$ par le tableau de transitions :

$$\begin{array}{lll} \mathbf{n} & \rightarrow & \mathbf{n} + \mathbf{e}_i & \text{avec taux} & \lambda_i^0 \\ \mathbf{n} & \rightarrow & \mathbf{n} - \mathbf{e}_i & \text{avec taux} & \mu_i r_{i,0} \\ \mathbf{n} & \rightarrow & \mathbf{n} - \mathbf{e}_i + \mathbf{e}_j & \text{avec taux} & \mu_i r_{i,j} \end{array}$$

On a alors :

Théorème 4.1 de Jackson Soit le vecteur $\vec{\lambda} = (\lambda_1, \dots, \lambda_N)$ solution du système linéaire :

$$\vec{\lambda} = \vec{\lambda}^0 + \vec{\lambda} \mathbf{R}.$$

Si la condition de stabilité $\lambda_i < \mu_i, i = 1, \dots, N$, est satisfaite, alors le processus $\{N(t), t \geq 0\}$ admet une distribution stationnaire unique donnée par :

$$P(N_1 = n_1, \dots, N_N = n_N) = \prod_{i=1}^N \left(1 - \frac{\lambda_i}{\mu_i}\right) \left(\frac{\lambda_i}{\mu_i}\right)^{n_i} \quad (4.8)$$

La preuve du Théorème de Jackson et des théorèmes similaires peut se faire par vérification directe : il suffit juste de vérifier que la distribution proposée satisfait les équations d'équilibre. Dans le cas présent, on doit alors avoir pour

tout vecteur $\mathbf{n}$:

$$\pi(\mathbf{n}) \left(\sum_{i=0}^N \mu_i \mathbf{1}_{\{n_i > 0\}} + \lambda_i^0 \right) = \sum_{j=0}^N \sum_{i=0}^N \pi(\mathbf{n} + \mathbf{e}_j - \mathbf{e}_i) \mu_j r_{j,i} \mathbf{1}_{\{n_i > 0\}} + \sum_{j=1}^N \pi(\mathbf{n} + \mathbf{e}_j) \mu_j r_{j,0} + \sum_{j=1}^N \pi(\mathbf{n} - \mathbf{e}_j) \lambda_j^0 \mathbf{1}_{\{n_j > 0\}}$$

Ce qu'on peut vérifier avec un peu de patience en utilisant (4.8). L'unicité se démontre à partir du fait que la chaîne est irréductible.

D'après le Théorème 4.1, la distribution stationnaire de la file i est donnée par

$$P(n_i) = \left(1 - \frac{\lambda_i}{\mu_i}\right) \left(\frac{\lambda_i}{\mu_i}\right)^{n_i}.$$

Donc la distribution stationnaire de la file i est identique à la distribution dans une file $M/M/1$ avec un taux d'arrivée λ_i et un temps de service moyen μ_i . La probabilité d'un état du réseau s'obtient comme produit des probabilités marginales : on dit que cette distribution a une "forme produit". Ainsi, dans l'état stationnaire, l'état de chaque file d'attente est celui d'une file $M/M/1$, indépendante du reste du réseau.

Exemple

Considérons l'exécution de jobs sur un ordinateur (cf. figure 4.14). Les jobs sont soumis par des utilisateurs supposés en grand nombre suivant un processus de Poisson de taux λ . Chaque job alterne ensuite entre une phase d'exécution sur la CPU, de durée exponentiellement distribuée de moyenne $1/\mu_1$, et une phase d'accès disque, de durée exponentiellement distribuée de moyenne $1/\mu_2$. Le nombre de phases d'exécution sur la CPU est géométriquement distribué de paramètre q , et donc un job s'exécute en moyenne $1/q$ fois sur la CPU.

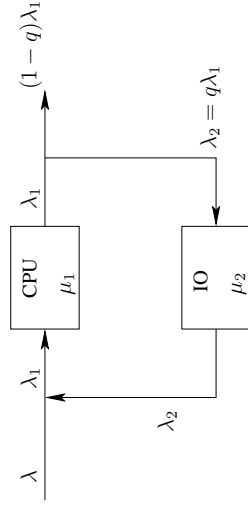


FIGURE 4.14 – Exécution de jobs sur un ordinateur.

Les équations de trafic s'écrivent :

$$\begin{aligned} \lambda_1 &= \lambda + \lambda_2 \\ \lambda_2 &= q \lambda_1 \end{aligned}$$

On en déduit $\lambda_1 = \lambda/(1-q)$ et $\lambda_2 = \lambda q/(1-q)$. En posant $\rho_1 = \lambda_1/\mu_1 < 1$ et $\rho_2 = \lambda_2/\mu_2 < 1$, on obtient :

$$P(N_1 = n_1, N_2 = n_2) = (1 - \rho_1)(1 - \rho_2) \rho_1^{n_1} \rho_2^{n_2} \quad n_1 \geq 0, n_2 \geq 0$$

On en déduit que $E[N_i] = \rho_i/(1 - \rho_i)$, $i = 1, 2$ et le temps moyen d'exécution d'un job est alors donné par $(E[N_1] + E[N_2])/\lambda$.

4.6.2 Réseaux fermés

Il existe aussi des systèmes dans lesquels le nombre de clients est fixé. Un exemple classique est le problème machines/réparateurs, où les machines d'un site (nombre fini) sont susceptibles de tomber en panne et sont réparées par des personnes compétentes (en nombre fini aussi). Les réparateurs sont alors représentés par des serveurs et les machines par des clients, qui alternativement passent par une période de bon fonctionnement, puis vont se faire réparer par un réparateur libre.

Les paramètres d'un tel réseau sont :

- Le nombre N de stations
- Le vecteur $(\mu_1, \dots, \mu_N)$ des taux de service,
- La matrice carrée $N \times N$ de routage interne $\mathbf{R}$.
- Le nombre K de clients présents.

Comme aucun client n'entre ni ne sort, le nombre de clients total dans le réseau est conservé. Comme précédemment, $\{N(t), t \geq 0\}$ est une CMTC, mais son espace d'états est (au plus) $\{\mathbf{n} \in \mathbb{N}^N \mid n_1 + \dots + n_N = K\}$. On suppose que $\mathbf{R}$ est irréductible, et par conséquent chaque client visite toutes les files d'attente du réseau.

On a le résultat suivant :

Théorème 4.2 Jackson pour les réseaux fermés. *Soit un réseau de Jackson fermé avec K clients et N serveurs. Soit $\theta = (\theta_1, \dots, \theta_N)$ une solution positive du système linéaire :*

$$\theta = \theta \mathbf{R}.$$

Le processus $\{N(t), t \geq 0\}$ admet une distribution de probabilité stationnaire donnée par :

$$P(N_1 = n_1, \dots, N_N = n_N) = G \prod_{i=1}^N \left(\frac{\theta_i}{\mu_i} \right)^{n_i} \mathbf{1}_{\{n_1 + \dots + n_N = K\}},$$

où

$$G = \frac{1}{\sum_{n=1}^{\infty} n! P(N_1 = n_1, \dots, N_N = n_N)}.$$

La solution pour θ n'est pas unique. Cependant, comme l'ensemble des solutions forme une droite vectorielle, on peut facilement vérifier que la distribution obtenue ne dépend pas du choix de la solution de cette équation.

Le calcul des probabilités et des autres mesures de performances (nombre moyen de clients, débits) passe par le calcul de la constante de normalisation G .

Il s'agit d'une somme avec un nombre de termes $\binom{N+K-1}{N}$, ce qui n'est pas toujours faisable. Il existe toutefois des algorithmes récursifs rapides pour ce calcul.

Pour finir cette section, on note qu'il est possible de généraliser ces résultats avec des hypothèses plus faibles sur les disciplines et les lois de service. Ainsi, le résultat de l'existence d'une forme produite, reste valable pour certaines valeurs plus générales de la discipline de service et de la loi de service. Il est également permis d'attribuer des classes aux clients, et des lois de service dépendant de la classe. Ces réseaux généralisant ceux de Jackson sont dits "BCMP" d'après les initiales de leurs découvreurs Basket, Chandy, Muntz et Palacios.

Exercices

Exercice 1 Un réparateur de télévision met un temps exponentiellement distribué pour en réparer une. En moyenne, il lui faut 30 minutes. Ses clients lui amènent des télévisions en panne suivant un processus de Poisson, avec en moyenne 10 télé par jour (en 24 heures). Cet homme là n'a pas besoin de dormir, et travaille donc 24 heures par jour.

- Suggérer un modèle de file d'attente pour ce problème ?
- Quel est le pourcentage du temps où le réparateur n'a rien à faire ?
- Combien y a-t'il en moyenne de télévisions dans son atelier ?
- Combien de temps une télé reste-t-elle dans son atelier en moyenne ?

Exercice 2 N camions font des aller-retours entre une plate-forme de chargement et une plate-forme de déchargement à l'intérieur d'un terminal portuaire. Il y a une seule grue pour charger les camions et une seule grue pour les décharger, ce qui fait que les camions doivent faire la queue pour charger et pour décharger. Les opérations de chargement et de déchargement ont des durées exponentiellement distribuées, de moyennes respectives $1/\lambda$ et $1/\mu$. On suppose que $\lambda \neq \mu$ et que le temps aller-retour peut être négligé par rapport aux durées de chargement et de déchargement.

- La dynamique de ce système est identique à celle d'une file d'attente vue en cours. Décrire cette file d'attente.

- Quelle est la probabilité que k des N camions soient à la plate-forme de chargement ?

Indication. Définir l'état X comme le nombre de camions à la plate-forme de déchargement et dessiner le diagramme de transition de X .

Exercice 3 Un magasin de photocopieuses emploie trois réparateurs. Les demandes de réparation arrivent suivant un processus de Poisson de taux $\lambda = 5$ réparations par jour. Les durées de réparation sont exponentiellement distribuées avec une durée moyenne de $1/\mu = 0.5$ jour. On suppose que le système est à l'équilibre.

- Déterminer le diagramme de transition entre les différents états et les taux de transition.
- Quel est le modèle de file d'attente correspondant ?
- Quelle est le nombre moyen de photocopieuses dans l'atelier ?
- Combien de temps une photocopieuse reste-t-elle dans l'atelier en moyenne ?
- Combien y a-t'il de réparateurs en activité en moyenne ?
- Quelle est la proportion de temps où un réparateur est occupé ?

Exercice 4 Deux types de clients arrivent à un bureau de poste ayant un seul guichet. Les clients de type 1 sont patients, et se mettent toujours dans la queue pour attendre leur tour. Les clients du type 2 n'ont pas cette patience et ne restent que s'il y a moins de K clients dans le bureau de poste. Les clients des types 1 et 2 arrivent suivant des processus de Poissons indépendants, de taux respectifs λ et γ . La durée de service d'un client ne dépend pas de son type et est exponentiellement distribuée de moyenne $1/\mu$.

- Ecrire les équations d'équilibre et déterminer les probabilités stationnaires.
- Donner une expression de la proportion de clients quittant le bureau de poste sans avoir été servi.

Exercice 5 Soit une file $M/G/1$ de taux d'arrivée λ , de durée moyenne de service $E[B]$ et dont le second moment de la loi de service est $E[B^2]$.

- Prenons $\lambda = 3/2$, $E[B] = 1/2$ et $E[B^2] = 1/2$. Calculer les valeurs moyennes du temps d'attente, du temps de séjour, du nombre de clients en attente et du nombre de clients dans le système.
- Supposons que des mesures nous apprennent que la durée moyenne d'attente est 5, que le taux d'utilisation est $2/3$ et que le nombre moyen de clients dans le système est 8. Déterminer le taux d'arrivée et les deux premiers moments de la loi de service. Quelle sont alors la longueur moyenne de la file d'attente et le temps de séjour moyen ?

Exercice 6 Un bureau de poste avec un seul guichet voit des clients arriver suivant un processus de Poisson de taux 60 clients par heure. La moitié des clients

ont une durée de service qui est la somme d'une durée fixe de 15 secondes et d'un temps exponentiellement distribué de moyenne 15 secondes. L'autre moitié ont une durée de service exponentiellement distribuée de moyenne 1 minute. Déterminer la durée moyenne d'attente et la longueur moyenne de la file d'attente.

Exercice 7 Soit une file $M/G/1$ de taux d'arrivée $\lambda = 1/3$ ayant pour loi de service la distribution hyper-exponentielle ci-dessous :

$$B(x) = \frac{1+a}{2} \left(1 - e^{-\frac{1+a}{2}x} \right) + \frac{1-a}{2} \left(1 - e^{-\frac{1-a}{2}x} \right), \quad x \geq 0,$$

avec $0 \leq a < 1$.

- Quelles sont la moyenne et la variance des durées de service ?
- Pour quelle valeur de a le temps d'attente moyen devient plus grand que la durée moyenne d'une période d'activité ?
- Vérifier que quand $a \rightarrow 1$, cette file d'attente a un comportement très différent que dans le cas $a = 1$.

Exercice 8 En moyenne 10 clients par heure arrive a un guichet. Le service prend en moyenne 6 minutes. Si le guichet est occupé, il y a une place pour attendre son tour. S'il y a plus de deux clients (un en service et l'autre en attente), les clients qui arrivent ne restent pas.

- Dessiner le diagramme de transition correspondant et résoudre les équations d'équilibre en supposant des arrivées suivant un processus de Poisson et une durée de service exponentiellement distribuée.
- Combien de clients sont servis au cours d'une heure en moyenne ?
- Supposons maintenant qu'il y a deux guichets, et pas de place pour attendre. Répondre aux questions a) et b) dans ce cas.

Exercice 9 Soit une file $M/G/1$ de taux d'arrivée λ et de durée moyenne de service β . En utilisant la loi de Little, montrer que,

$$P(N > 0) = \rho,$$

où $\rho = \lambda\beta$.

Exercice 10 Considérons un modèle simplifié d'un lien TCP. Supposons que les paquets arrivent suivant un processus de Poisson d'intensité $\lambda = 100$ paquets/sec sur un modem ADSL à 2 Mbps. La distribution de la taille des paquets et les durées de service correspondante sont les suivantes :

length	proportion	time/ms
40	0.1	0.16
576	0.3	2.3
1500	0.6	5.9

Déterminer le temps d'attente moyen d'un paquet dans le buffer quand la discipline de service est FCFS.

Exercice 11 Soit le réseau de communication décrit sur la Figure 4.15. Des paquets venant de l'extérieur arrivent aux noeuds 1, 2 et 5 suivant un processus de Poisson de taux $\lambda = 2$ paquets/s. Dans tous les noeuds, chaque lien a son propre buffer. Le flux de paquets entrants dans un noeud est routé aléatoirement suivant les probabilités décrites sur la figure. Le lien issu du noeud 4 a une capacité $\mu = 8$ paquets/s, tandis que la capacité des autres liens est $\mu = 3$ paquets/s.

- Quel est le délai moyen des paquets prenant les routes 1-2-3 et 1-5-4 ?
- Combien y a-t-il en moyenne de paquets dans ce réseau.
- Quelle est le temps de séjour moyen d'un paquet dans ce réseau ?

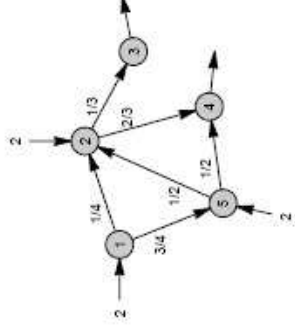


FIGURE 4.15 – Figure pour l'exercice 11