

Modélisation de phénomènes aléatoires :
introduction aux chaînes de Markov et aux martingales

Thierry Bodineau

École Polytechnique
Département de Mathématiques Appliquées

1^{er} novembre 2020

Table des matières

I	Chaînes de Markov	9
1	De la marche aléatoire aux jeux de cartes	11
2	Propriété de Markov	17
2.1	Propriété de Markov et matrice de transition	17
2.2	Exemples de chaînes de Markov	19
2.3	Loi d'une chaîne de Markov	21
2.3.1	Équation de Chapman-Kolmogorov	21
2.3.2	Processus décalé en temps	23
2.4	Temps d'arrêt et propriété de Markov forte	25
2.5	Applications	27
2.5.1	Équation de la chaleur $\star$	27
2.5.2	Ruine du joueur	30
2.5.3	Méthode de Monte-Carlo pour un problème de Dirichlet $\star$	32
3	Mesures invariantes	37
3.1	Mesures de probabilité invariantes	37
3.1.1	Définition des mesures invariantes	37
3.1.2	Exemples de mesures invariantes	38
3.1.3	Premières propriétés des mesures invariantes	40
3.2	Irréductibilité et unicité des mesures invariantes	41
3.2.1	Irréductibilité	42
3.2.2	Unicité des mesures de probabilité invariantes	42
3.2.3	Construction de la mesure de probabilité invariante	44
3.3	Réversibilité et Théorème H	47
3.3.1	Réversibilité	47
3.3.2	Théorème H pour les chaînes de Markov $\star$	48
3.3.3	Application : modèle d'Ehrenfest	49
4	Espaces d'états dénombrables	53
4.1	Chaînes de Markov récurrentes et transitoires	53
4.2	Application : marches aléatoires	57
4.2.1	Marches aléatoires symétriques sur $\mathbb{Z}^d$	57
4.2.2	Un critère analytique	59
4.3	Mesures invariantes	60
4.4	Application : branchement et graphes aléatoires	64

4.4.1	Arbres aléatoires de Galton-Watson	64
4.4.2	Graphes aléatoires d'Erdős-Rényi	68
5	Ergodicité et convergence	73
5.1	Ergodicité	73
5.1.1	Théorème ergodique	73
5.1.2	Application : algorithme PageRank de Google	77
5.2	Convergence	78
5.2.1	Apériodicité et convergence	79
5.2.2	Distance en variation et couplage	83
5.2.3	Vitesses de convergence	90
6	Application aux algorithmes stochastiques	97
6.1	Optimisation	97
6.2	Algorithmes stochastiques	99
6.2.1	Algorithme de Metropolis-Hastings	99
6.2.2	Modèle d'Ising	101
6.3	Simulation parfaite : algorithme de Propp-Wilson ★	103
6.4	Algorithme de recuit simulé ★	106
6.4.1	Problème du voyageur de commerce	107
6.4.2	Traitement d'images	108
7	Application : la percolation ★	111
7.1	Description du modèle	111
7.2	Transition de phase	112
7.2.1	Absence de percolation pour p petit	113
7.2.2	Percolation pour p proche de 1	114
7.2.3	Point critique	115
7.2.4	Dimension 2	116
II	Martingales	119
8	De la marche aléatoire aux stratégies optimales	121
9	Espérance conditionnelle	125
9.1	Espérance conditionnelle sur un espace d'états discret	125
9.2	Définition de l'espérance conditionnelle	127
9.3	Propriétés de l'espérance conditionnelle	130
9.4	Processus aléatoire	132
10	Martingales et stratégies	135
10.1	Martingales	135
10.2	Stratégies et théorème d'arrêt	138
10.3	Stratégies de gestion et évaluation des options ★	143
10.4	Inégalités de martingales	151

11 Convergence des martingales	153
11.1 Convergence des martingales dans $\mathbb{L}^2$	153
11.2 Application : loi des grands nombres	155
11.3 Convergence des sous-martingales	157
11.4 Application : le modèle de Wright-Fisher	161
11.5 Martingales fermées	163
11.6 Théorème central limite	167
11.6.1 Théorème central limite pour les martingales	167
11.6.2 Inégalité de Hoeffding	170
11.6.3 Théorème central limite pour les chaînes de Markov	173
12 Applications des martingales	175
12.1 Cascades de Mandelbrot	175
12.2 Mécanismes de renforcement	180
12.2.1 Urne de Pólya	180
12.2.2 Graphes aléatoires de Barabási-Albert	181
12.3 Descente de gradient stochastique	184
12.3.1 Algorithme du gradient stochastique	184
12.3.2 Réseaux de neurones artificiels	189
12.4 L'algorithme de Robbins-Monro	191
12.5 Processus de Galton-Watson	195
13 Arrêt optimal et contrôle stochastique*	197
13.1 Arrêt optimal	197
13.1.1 Enveloppe de Snell	197
13.1.2 Le problème du parking	200
13.1.3 Problème des secrétaires	202
13.2 Contrôle stochastique	204
13.2.1 Équation de la programmation dynamique	205
13.2.2 Contrôle des chaînes de Markov	206
A Théorie de la mesure	209
A.1 Espaces mesurables et mesures	209
A.1.1 σ -algèbres	209
A.1.2 Mesures	210
A.2 Intégration	212
A.2.1 Fonctions mesurables	212
A.2.2 Intégration des fonctions positives	214
A.2.3 Fonctions intégrables	216
A.2.4 Théorème de convergence dominée	218
A.3 Espaces produits	220
A.3.1 Mesurabilité sur les espaces produits	220
A.3.2 Intégration sur les espaces produits	221

B	Théorie des probabilités	223
B.1	Variables aléatoires	223
B.2	Intégration de variables aléatoires et espaces $\mathbb{L}^p$	225
B.2.1	Théorèmes de convergence	225
B.2.2	Inégalités classiques	225
B.2.3	Espaces $\mathbb{L}^p$	228
B.3	Convergences	229
B.3.1	Convergence en probabilité	230
B.3.2	Convergence en loi	233
B.4	Indépendance	234
B.4.1	Variables aléatoires indépendantes	234
B.4.2	Comportement asymptotique	235
B.5	Espérance conditionnelle	237
B.5.1	Construction de l'espérance conditionnelle	237
B.5.2	Propriétés de l'espérance conditionnelle	239

L'aléa joue un rôle déterminant dans des contextes variés et il est souvent nécessaire de le prendre en compte dans de multiples aspects des sciences de l'ingénieur, citons notamment les télécommunications, la reconnaissance de formes ou l'administration des réseaux. Plus généralement, l'aléa intervient aussi en économie (gestion du risque), en médecine (propagation d'une épidémie), en biologie (évolution d'une population) ou en physique statistique (théorie des transitions de phases). Dans les applications, les données observées au cours du temps sont souvent modélisées par des variables aléatoires corrélées dont on aimerait prédire le comportement. L'objet de ce cours est de formaliser la notion de corrélation en étudiant deux types de *processus aléatoires* fondamentaux en théorie des probabilités : les chaînes de Markov et les martingales.

En 1913, A. Markov posait les fondements d'une théorie qui a permis d'étendre les lois des probabilités des variables aléatoires indépendantes à un cadre plus général susceptible de prendre en compte des corrélations. La première partie de ce cours décrit la théorie des chaînes de Markov et certaines de leurs applications. Le parti pris de ce cours est de considérer le cadre mathématique le plus simple possible en se focalisant sur des espaces d'états finis, voire dénombrables, pour éviter le recours à la théorie de la mesure.

Le comportement asymptotique des chaînes de Markov peut être classifié et prédit. Nous verrons que la structure de ces processus aléatoires corrélés est encodée dans une mesure invariante qui permet de rendre compte des propriétés ergodiques, généralisant ainsi les suites de variables aléatoires indépendantes. La convergence des chaînes de Markov vers leurs mesures invariantes constitue un aspect fondamental de la théorie des probabilités, mais elle joue aussi un rôle clef dans les applications.

Plusieurs exemples seront décrits pour illustrer le rôle majeur des chaînes de Markov dans différents domaines de l'ingénierie comme les problèmes numériques (méthodes de Dirichlet, optimisation), la reconnaissance de formes ou l'algorithme PageRank de Google. Des exemples issus de la physique statistique (irréversibilité en théorie cinétique des gaz, transitions de phases) ou de la dynamique des populations (arbres de Galton Watson) permettront aussi d'éclairer certains aspects des chaînes de Markov. D'autres applications des chaînes de Markov sont présentées dans le livre de M. Benaïm et N. El Karoui [2] et dans celui de J.F. Delmas et B. Jourdain [8]. Les ouvrages de C. Graham [13], P. Brémaud [5] et J. Norris [23] constituent aussi de très bonnes références sur la théorie des chaînes de Markov.

La seconde partie de ce cours porte sur la théorie des martingales qui permet d'étudier d'autres structures de dépendance que celles définies par les chaînes de Markov. Les martingales sont communément associées aux jeux de hasard et nous verrons comment des stratégies optimales peuvent être définies à l'aide de martingales et de temps d'arrêt. Les martingales forment une classe de processus aléatoires aux propriétés très riches.

En particulier, les fluctuations de ces processus peuvent être contrôlées et leur convergence facilement analysée. Les martingales permettent aussi d'étudier des mécanismes de renforcement pour mieux comprendre des comportements collectifs, des algorithmes stochastiques ou des phénomènes issus de l'écologie comme la dérive génétique.

D'autres aspects de la théorie des martingales figurent dans les livres de J. Neveu [22] et D. Williams [28]. Les ouvrages de B. Bercu, D. Chafaï [3] et M. Duflo [9] proposent de nombreux développements sur les applications des martingales aux algorithmes stochastiques.

Des compléments sur la théorie de la mesure et des probabilités figurent en annexe. Ils permettront d'approfondir certaines notions fondamentales de la théorie des probabilités et pourront servir de référence. Les parties du cours marquées par une étoile $\star$ servent d'illustration et de complément, elles peuvent être omises en première lecture.

Des développements plus complets figurent dans l'excellent cours de J.F. Le Gall [18] qui traite à la fois de la théorie de la mesure et des processus stochastiques. Les ouvrages anglophones de R. Durrett [10] et G. Grimmett, D. Stirzaker [14] proposent aussi de nombreux approfondissements sur la théorie des probabilités.

Je souhaite exprimer toute ma reconnaissance à Nizar Touzi pour m'avoir permis de reprendre des éléments de son cours [25]. Je tiens aussi à remercier chaleureusement Stéphanie Allasonnière, Anne de Bouard, Djailil Chafai, Jean-René Chazottes, Jean-François Delmas, Laurent Denis, Lucas Gerin, Carl Graham, Arnaud Guillin, Arnaud Guyader, Marc Hoffmann, Clément Mantoux et Sylvie Roelly qui m'ont aidé dans la rédaction de ce cours par leurs précieux conseils et leur relecture attentive.

Première partie

Chaînes de Markov

Chapitre 1

De la marche aléatoire aux jeux de cartes

La loi des grands nombres et le théorème central limite sont deux théorèmes clef de la théorie des probabilités. Ils montrent que la limite d'une somme de variables aléatoires indépendantes obéit à des lois simples qui permettent de prédire le comportement asymptotique. Prenons l'exemple classique d'une marche aléatoire symétrique (cf. figure 1.1)

$$X_0 = 0 \quad \text{et pour} \quad n \geq 1, \quad X_n = \sum_{i=1}^n \zeta_i \quad (1.1)$$

où les $\{\zeta_i\}_{i \geq 1}$ sont des variables indépendantes et identiquement distribuées

$$\forall i \geq 1, \quad \mathbb{P}(\zeta_i = 1) = \mathbb{P}(\zeta_i = -1) = 1/2.$$

La loi des grands nombres implique la convergence presque sûre

$$\lim_{n \rightarrow \infty} \frac{1}{n} X_n = \mathbb{E}(\zeta_1) = 0 \quad p.s. \quad (1.2)$$

et le théorème central limite assure la convergence en loi vers une gaussienne de moyenne nulle et de variance 1 que l'on notera γ

$$\frac{1}{\sqrt{n}} X_n \xrightarrow[n \rightarrow \infty]{(loi)} \gamma. \quad (1.3)$$

Pour de nombreuses applications, il est nécessaire d'ajouter des corrélations entre ces variables et d'enrichir ce formalisme au cas de processus aléatoires qui ne sont pas une simple somme de variables indépendantes. Par exemple, on voudrait décrire un mobile soumis à une force aléatoire et à une force de rappel qui le maintient près de l'origine (comme un atome qui vibre autour de sa position d'équilibre dans un cristal ou le prix d'une matière première soumise à la loi de l'offre et de la demande) et représenter sa position au cours du temps par Y_n . Une façon simple de prendre en compte une force de rappel est de construire récursivement une suite aléatoire

$$Y_0 = 0 \quad \text{et} \quad n \geq 1, \quad Y_n = Y_{n-1} + \text{signe}(Y_{n-1}) \zeta_n, \quad (1.4)$$

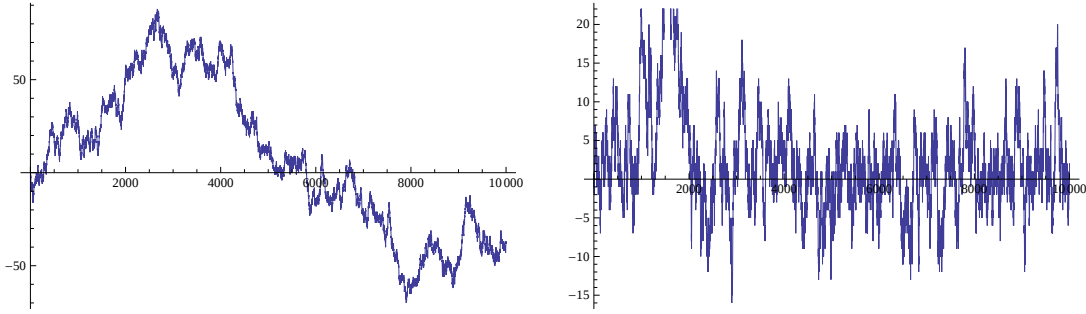


FIGURE 1.1 – À gauche, une réalisation de la marche aléatoire symétrique $n \rightarrow X_n$ représentée après 10000 pas. À droite, une réalisation de la trajectoire $n \rightarrow Y_n$ représentée après 10000 pas pour $p = 0.45$.

où les $\{\zeta_i\}_{i \geq 1}$ sont maintenant des variables indépendantes de Bernoulli de paramètre $p \in [0, 1/2]$

$$\forall i \geq 1, \quad \mathbb{P}(\zeta_i = 1) = 1 - \mathbb{P}(\zeta_i = -1) = p.$$

Dans (1.4), le signe de y est noté $\text{signe}(y) \in \{\pm 1\}$ et pour éviter toute ambiguïté, le signe de 0 sera pris égal à 1. Si $p = 1/2$, Y_n est égal en loi à X_n . Par contre si $p < 1/2$, le biais aléatoire dépend du signe de Y_n et il a tendance à ramener Y_n vers 0 car $\mathbb{E}(\zeta_1) < 0$. On voit sur la figure 1.1 que l'amplitude et la structure des processus $\{X_n\}$ et $\{Y_n\}$ sont très différentes. Même pour un biais petit $p = 0.45$, la trajectoire $n \rightarrow Y_n$ reste localisée autour de 0. Un des enjeux de ce cours sera de décrire le comportement asymptotique de Y_n . Pour le moment, essayons de deviner ce comportement limite. La figure 1.2 montre une trajectoire de Y_n dix fois plus longue que celle de la figure 1.1 et on constate que l'amplitude de cette trajectoire reste sensiblement inchangée. Ceci contraste avec l'amplitude de la marche aléatoire qui croît comme $\sqrt{n}$ par le théorème central limite. L'histogramme de la figure 1.2 représente la fréquence des passages de la trajectoire en un site, i.e. la mesure π_n définie par

$$\forall y \in \mathbb{Z}, \quad \pi_n(y) = \frac{1}{n} \sum_{k=1}^n 1_{\{Y_k=y\}}.$$

On montrera que π_n converge vers une mesure limite π qui encode le comportement asymptotique de Y_n

$$\forall y \in \mathbb{Z}, \quad \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{k=1}^n 1_{\{Y_k=y\}} = \pi(y) \quad p.s.$$

Cette convergence peut s'interpréter comme un analogue de la loi des grands nombres. Ceci pose plusieurs questions auxquelles nous essayerons de répondre dans ce cours

- Peut-on décrire la mesure π ?
- Quel est le temps nécessaire pour que π_n soit proche de π ?

Le processus $\{Y_n\}_{n \geq 0}$ en (1.4) a été obtenu par une *réurrence aléatoire* $Y_{n+1} = f(Y_n, \zeta_{n+1})$ pour une fonction f bien choisie. Sous cette forme, la structure additive de la marche aléatoire X_n a disparu et on peut ainsi envisager de construire des processus à valeurs dans un espace général. Par exemple, on peut définir une marche aléatoire sur le graphe $\mathcal{G}$ de

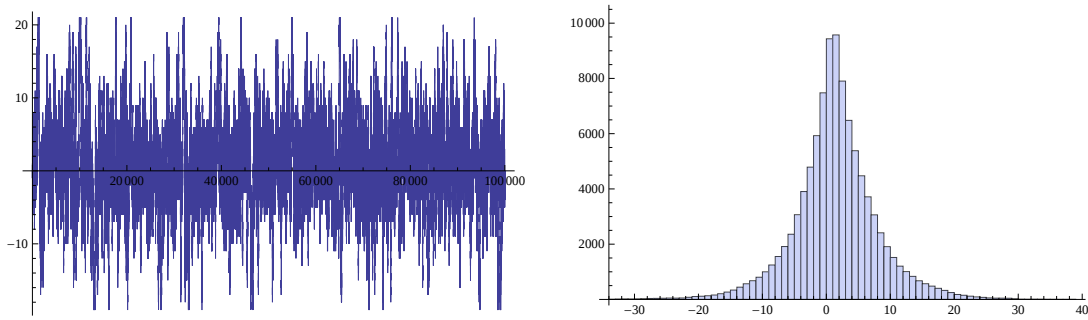


FIGURE 1.2 – À gauche, une réalisation de la trajectoire $n \rightarrow Y_n$ est représentée après 100000 pas. À droite, l'histogramme correspondant au nombre de passages par chaque site pour la trajectoire.

la figure 1.3 : le marcheur part d'un site donné et évolue à chaque pas de temps en sautant uniformément sur un des voisins du site occupé.

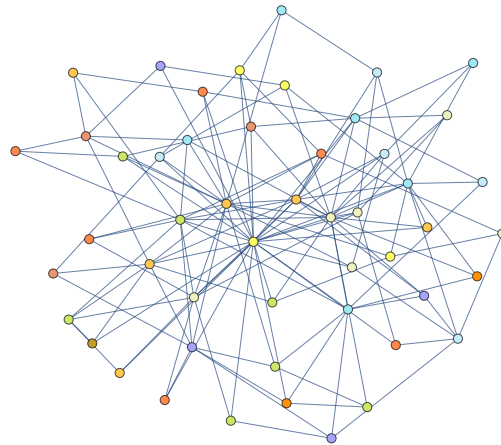


FIGURE 1.3 – Un graphe aléatoire (de Barabási-Albert) avec 50 sites.

Pour construire la récurrence aléatoire correspondante, on note $\mathcal{V}(x)$ l'ensemble des voisins d'un site x de $\mathcal{G}$, i.e. les sites reliés à x par une arête. On décrit maintenant une procédure pour choisir uniformément un des voisins à l'aide d'une variable aléatoire uniforme sur $[0, 1]$. On note $\deg(x)$ le degré de x , i.e. le cardinal de l'ensemble $\mathcal{V}(x)$ (qui peut varier en fonction du site x). On numérote (une fois pour toute) par un indice entre 0 et $\deg(x) - 1$ chaque arête entre x et ses voisins. Pour tout $x \in \mathcal{G}$ et $u \in [0, 1]$, on définit $f(x, u)$ comme le voisin de x dont l'arête a le numéro $\lfloor \deg(x)u \rfloor$, où $\lfloor \cdot \rfloor$ représente la partie entière. On a ainsi une construction explicite d'une marche aléatoire $\{Y_n\}_{n \geq 0}$ sur le graphe $\mathcal{G}$ en générant l'aléa à partir d'une suite $\{\zeta_n\}_{n \geq 1}$ de variables aléatoires indépendantes et uniformément distribuées sur $[0, 1]$

$$Y_0 = 0 \quad \text{et} \quad n \geq 1, \quad Y_{n+1} = f(Y_n, \zeta_{n+1}).$$

La marche étant construite, on voudrait comprendre son comportement asymptotique : après un temps très long, quelle est la probabilité que la marche soit sur un site donné ? On verra entre autres que cette probabilité est proportionnelle au degré de chaque site.

Le modèle peut encore être enrichi en orientant les arêtes du graphe (cf. figure 1.4) et en autorisant seulement les transitions selon les arêtes orientées. La probabilité de chaque saut peut aussi être pondérée selon les voisins, par exemple sur le graphe de la figure 1.4 : la marche peut passer du site 1 au site 2 avec la probabilité $P(1, 2) = 1/2$, au site 3 avec la probabilité $P(1, 3) = 1/4$ et rester sur place avec la probabilité $P(1, 1) = 1/4$. La seule contrainte étant d'ajuster la somme des probabilités à 1.

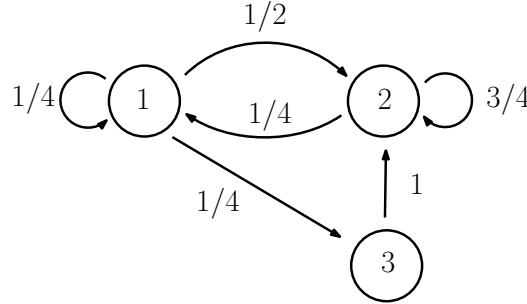


FIGURE 1.4 – Un graphe orienté avec 3 sites.

L'essentiel des exemples concrets que nous allons rencontrer dans ce cours peuvent se formaliser comme une marche aléatoire sur un graphe orienté avec des probabilités de transition associées à chaque lien. Parmi les exemples de marche aléatoire sur un graphe traités dans ce cours, nous évoquerons les robots d'indexation qui parcourent le World Wide Web pour collecter les données et indexer des pages Web. Certains graphes peuvent être compliqués et il est important de développer une théorie générale pour appréhender cette complexité.

Terminons ce tour d'horizon sur les chaînes de Markov par le mélange de cartes. On représente un jeu de 52 cartes en numérotant leurs positions dans le paquet de 1 à $K = 52$. Mélanger les cartes revient à appliquer des permutations successives sur leurs positions. Mathématiquement, cette procédure n'est rien d'autre qu'une marche aléatoire sur le groupe symétrique $\mathcal{S}_K$ des permutations sur $\{1, 2, \dots, K\}$. Initialement les cartes sont rangées dans l'ordre et l'état de départ est la permutation identité $\text{Id} = \{1, 2, \dots, K\}$. Étant donné une mesure μ de référence, on choisit au hasard une permutation σ_1 sous μ et le jeu de cartes est réordonné en $\sigma_1 = \sigma_1 \circ \text{Id}$. Pour battre les cartes, on itère plusieurs fois cette opération en tirant au hasard des permutations $\sigma_1, \sigma_2, \dots, \sigma_n$ et en les composant $\sigma_n \circ \dots \circ \sigma_2 \circ \sigma_1$. On peut imaginer différentes règles pour mélanger les cartes et choisir les permutation σ_k . Par exemple, permuter à chaque fois deux cartes choisies aléatoirement ou pour un modèle plus réaliste (mais mathématiquement plus compliqué) couper le jeu en 2 paquets et insérer l'un dans l'autre (riffle shuffle). On a ainsi construit une récurrence aléatoire sur le groupe symétrique $\mathcal{S}_K$.

Si on mélange suffisamment longtemps le paquet, on s'attend à ce que les positions des cartes soient réparties uniformément parmi les $52!$ choix possibles. Pour le joueur de poker, le comportement asymptotique ne sert à rien. La véritable question est de savoir combien de fois il faut mélanger le jeu de cartes ? En termes mathématiques, on veut évaluer la vitesse de convergence vers un état d'équilibre. Cette question est déterminante pour de nombreuses applications. Si on souhaite réaliser une simulation numérique par un algorithme stochastique, il faut pouvoir prédire à quel moment la simulation peut

être arrêtée. Dans ce cours, des critères théoriques sur les vitesses de convergence seront présentés.

Les modèles décrits précédemment sont tous des *chaînes de Markov* et peuvent être traités dans un formalisme unifié qui sera décrit dans les chapitres suivants.

Chapitre 2

Propriété de Markov

Dans ce chapitre, nous allons définir les chaînes de Markov et présenter leurs premières propriétés.

2.1 Propriété de Markov et matrice de transition

Une suite de variables aléatoires $\{X_n\}_{n \geq 0}$ prenant ses valeurs dans un ensemble E est appelée un *processus aléatoire discret* avec *espace d'états* E . Dans ce cours, on ne considérera que des espaces d'états E finis ou dénombrables.

Le point commun entre les exemples de processus aléatoires discrets présentés au chapitre 1 est la propriété de Markov : la dépendance du processus au temps $n + 1$ par rapport à son passé se résume à la connaissance de l'état X_n . On peut le formaliser ainsi

Définition 2.1 (Propriété de Markov). *Soit $\{X_n\}_{n \geq 0}$ un processus aléatoire discret sur un espace d'états dénombrable E . Le processus satisfait la propriété de Markov si pour toute collection d'états $\{x_0, x_1, \dots, x_n, y\}$ de E*

$$\mathbb{P}(X_{n+1} = y \mid X_0 = x_0, X_1 = x_1, \dots, X_n = x_n) = \mathbb{P}(X_{n+1} = y \mid X_n = x_n) \quad (2.1)$$

dès que les deux probabilités conditionnelles ci-dessus sont bien définies. Le processus $\{X_n\}_{n \geq 0}$ sera alors appelé une chaîne de Markov. Si le membre de droite de (2.1) ne dépend pas de n , on dira que la chaîne de Markov est homogène.

Les conditionnements dans (2.1) s'interprètent comme des probabilités conditionnelles

$$\mathbb{P}(A \mid B) = \frac{\mathbb{P}(A \cap B)}{\mathbb{P}(B)} \quad \text{si } \mathbb{P}(B) \neq 0$$

avec les événements $A = \{X_{n+1} = y\}$ et $B = \{X_0 = x_0, X_1 = x_1, \dots, X_n = x_n\}$. Dans ce cours, on ne considérera que des chaînes de Markov homogènes. La distribution d'une chaîne de Markov homogène $\{X_n\}_{n \geq 0}$ peut donc être codée simplement par une *matrice de transition* $P = \{P(x, y)\}_{x, y \in E}$. La matrice de transition décrit la probabilité de passer de x à y

$$\forall x, y \in E, \quad P(x, y) = \mathbb{P}(X_{n+1} = y \mid X_n = x) \quad (2.2)$$

et elle satisfait

$$\forall x, y \in E, \quad P(x, y) \geq 0 \quad \text{et} \quad \forall x \in E, \quad \sum_{y \in E} P(x, y) = 1.$$

Comme la chaîne est homogène les transitions ne dépendent pas du temps et la relation (2.2) est valable pour tout n .

Montrons maintenant que les processus définis au chapitre 1 vérifient la propriété de Markov.

Théorème 2.2 (Récurrence aléatoire). *Soit $\{\xi_n\}_{n \geq 1}$ une suite de variables aléatoires indépendantes et identiquement distribuées sur un espace F . Soit E un espace d'états dénombrable et f une fonction de $E \times F$ dans E . On considère aussi X_0 une variable aléatoire à valeurs dans E indépendante de la suite $\{\xi_n\}_{n \geq 1}$.*

La récurrence aléatoire $\{X_n\}_{n \geq 0}$ définie par la relation

$$n \geq 0, \quad X_{n+1} = f(X_n, \xi_{n+1}) \tag{2.3}$$

est une chaîne de Markov.

Démonstration. Nous allons établir la propriété de Markov (2.1)

$$\begin{aligned} \mathbb{P}(X_{n+1} = y \mid X_0 = x_0, \dots, X_n = x_n) &= \frac{\mathbb{P}(f(X_n, \xi_{n+1}) = y, X_0 = x_0, \dots, X_n = x_n)}{\mathbb{P}(X_0 = x_0, \dots, X_n = x_n)} \\ &= \frac{\mathbb{P}(f(x_n, \xi_{n+1}) = y, X_0 = x_0, X_1 = x_1, \dots, X_n = x_n)}{\mathbb{P}(X_0 = x_0, X_1 = x_1, \dots, X_n = x_n)}. \end{aligned}$$

L'évènement $\{X_0 = x_0, X_1 = x_1, \dots, X_n = x_n\}$ ne dépend que des variables $\{X_0, \xi_1, \dots, \xi_n\}$ qui sont indépendantes de l'évènement $\{f(x_n, \xi_{n+1}) = y\}$. Par conséquent, le numérateur est le produit de la probabilité de ces deux évènements indépendants et on en déduit

$$\mathbb{P}(X_{n+1} = y \mid X_0 = x_0, \dots, X_n = x_n) = \mathbb{P}(f(x_n, \xi_{n+1}) = y) = \mathbb{P}(X_{n+1} = y \mid X_n = x_n).$$

La propriété de Markov est donc vérifiée et la formule précédente permet de déterminer la matrice de transition

$$\forall x, y \in E, \quad P(x, y) = \mathbb{P}(f(x, \xi_1) = y).$$

□

Inversement à toute matrice de transition P indexée par $\mathbb{N} = \{0, 1, 2, \dots\}$ (l'ensemble des entiers naturels), on peut associer une chaîne de Markov en construisant une récurrence aléatoire. Étant donné un état $X_n = x$ de $\mathbb{N}$, on choisit au hasard une variable aléatoire ξ_{n+1} uniforme sur $[0, 1]$ et on attribue à X_{n+1} la valeur

$$X_{n+1} = y \quad \text{si} \quad \xi_{n+1} \in \left[\sum_{k=0}^{y-1} P(x, k), \sum_{k=0}^y P(x, k) \right]$$

avec la convention $\sum_{k=0}^{-1} = 0$. Cette relation définit la fonction $f : \mathbb{N} \times [0, 1] \mapsto \mathbb{N}$ de la récurrence aléatoire (2.3). La même procédure s'applique pour une matrice de transition sur un espace E dénombrable.

2.2 Exemples de chaînes de Markov

Variables indépendantes.

Une suite de variables aléatoires indépendantes et identiquement distribuées constitue un exemple élémentaire de chaîne de Markov. Dans ce cas, la représentation sous forme de récurrence aléatoire est évidente $X_{n+1} = f(\xi_{n+1})$ et la matrice de transition (2.2) ne dépend plus du site de départ

$$\forall x, y \in E, \quad P(x, y) = \mathbb{P}(X_{n+1} = y \mid X_n = x) = \mu(y),$$

où μ est la loi des variables $\{X_n\}_{n \geq 0}$ sur E .

Chaîne de Markov à 3 états.

On considère une chaîne de Markov à 3 états, notés $E = \{1, 2, 3\}$, dont le graphe de transition est représenté figure 2.1. Sa matrice de transition $P = \{P(i, j)\}$ est donnée par

$$P = \begin{pmatrix} 1/4 & 1/2 & 1/4 \\ 1/4 & 3/4 & 0 \\ 0 & 1 & 0 \end{pmatrix}. \quad (2.4)$$

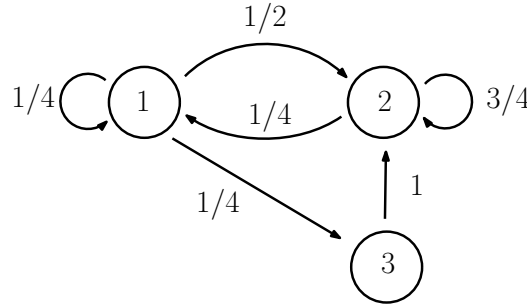


FIGURE 2.1 – Un graphe de transitions à 3 états.

Marche aléatoire.

Reprenons l'exemple (1.1) de la marche aléatoire sur $\mathbb{Z}$

$$X_0 = 0 \quad \text{et pour} \quad n \geq 0, \quad X_{n+1} = \sum_{i=1}^{n+1} \zeta_i = X_n + \zeta_{n+1}, \quad (2.5)$$

où les $\{\zeta_i\}_{i \geq 1}$ sont des variables indépendantes et identiquement distribuées

$$\forall i \geq 1, \quad \mathbb{P}(\zeta_i = 1) = p \quad \text{et} \quad \mathbb{P}(\zeta_i = -1) = q = 1 - p.$$

Cette récurrence aléatoire a pour matrice de transition

$$\forall x, y \in \mathbb{Z}, \quad P(x, y) = \begin{cases} p, & \text{si } y = x + 1 \\ q, & \text{si } y = x - 1 \\ 0, & \text{sinon} \end{cases}$$

La matrice de transition est cette fois indexée par $\mathbb{Z} \times \mathbb{Z}$ (mais la plupart de ses coefficients sont nuls).

On peut aussi considérer la marche aléatoire dans un domaine fini $\{1, \dots, L\}$ par exemple en supposant que le domaine est périodique. Dans ce cas si la marche aléatoire est en L , elle sautera en 1 avec probabilité p et réciproquement elle sautera de 1 à L avec probabilité q . La matrice de transition P sera une matrice $L \times L$

$$P = \begin{pmatrix} 0 & p & 0 & \dots & 0 & 0 & q \\ q & 0 & p & \dots & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & \dots & q & 0 & p \\ p & 0 & 0 & \dots & 0 & q & 0 \end{pmatrix}. \quad (2.6)$$

File d'attente.

Les files d'attente interviennent dans des contextes variés : dans les magasins, pour gérer des avions au décollage, pour le stockage de requêtes informatiques avant leur traitement, etc. Le modèle le plus simple consiste à supposer que ξ_n clients arrivent dans la file au temps n . On choisit les variables ξ_n indépendantes et identiquement distribuées à valeurs dans $\mathbb{N}$. Le serveur sert exactement 1 client à chaque pas de temps si la file n'est pas vide. Le nombre de clients X_n dans la file au temps n vérifie donc

$$X_n = (X_{n-1} - 1)^+ + \xi_n, \quad (2.7)$$

où la partie positive est notée $x^+ = \sup\{0, x\}$. Le processus $\{X_n\}_{n \geq 0}$ est une récurrence aléatoire sur $\mathbb{N}$ et donc une chaîne de Markov. Sa matrice de transition est donnée pour tous x, y dans $\mathbb{N}$ par

$$P(x, y) = \mathbb{P}(\xi_1 = y - (x - 1)^+).$$

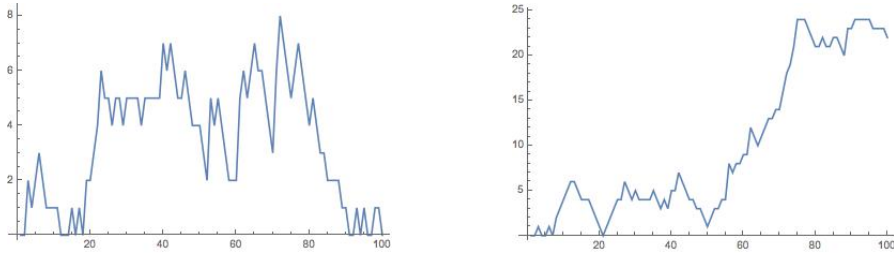


FIGURE 2.2 – Les deux graphiques représentent le nombre de clients au cours du temps dans une file d'attente du type (2.7) où les arrivées sont distribuées selon une loi de Poisson de paramètre λ . Dans le graphique de gauche, λ vaut 0.8 et le taux de service est suffisant pour résorber les demandes. Sur la figure de droite, le nombre de clients en attente diverge avec le temps car $\lambda = 1.2$ et les arrivées sont trop importantes pour la capacité de service.

La récurrence aléatoire (2.7) a une structure similaire à celle de la marche aléatoire (2.5), cependant leurs comportements asymptotiques peuvent être très différents. Dans la pratique, il s'avère essentiel de déterminer la stabilité d'une file d'attente en fonction

de la distribution des arrivées ξ_n , i.e. de déterminer si X_n diverge ou non quand le temps n tend vers l'infini. Cette question sera formalisée dans la suite du cours (cf. par exemple le théorème 4.8), mais la figure 2.2 permet déjà d'entrevoir différents comportements asymptotiques selon le taux des arrivées.

2.3 Loi d'une chaîne de Markov

Dans cette section, on considère une chaîne de Markov homogène $\{X_n\}_{n \geq 0}$ sur E dont la donnée initiale X_0 est choisie aléatoirement sur E selon la mesure de probabilité μ_0 qui attribue les probabilités $\{\mu_0(x)\}_{x \in E}$ aux éléments de E . On notera

$$\forall x \in E, \quad \mathbb{P}_{\mu_0}(X_0 = x) = \mu_0(x).$$

L'enjeu des paragraphes suivants est de déterminer la distribution de la chaîne de Markov au cours du temps.

2.3.1 Équation de Chapman-Kolmogorov

Après n pas de temps, X_n sera distribuée selon une loi que l'on notera μ_n

$$\forall x \in E, \quad \mu_n(x) := \mathbb{P}_{\mu_0}(X_n = x).$$

Avant d'énoncer l'équation de Chapman-Kolmogorov qui décrit l'évolution de la distribution μ_n , commençons par introduire quelques notations qui seront utilisées dans toute la suite du cours. Soit h une fonction de E dans $\mathbb{R}$. On définit

$$\forall x \in E, \quad Ph(x) = \sum_{y \in E} P(x, y)h(y). \quad (2.8)$$

Si E est fini, il s'agit simplement du produit à droite $P \cdot h$ entre la matrice P et $h = \{h(x)\}_{x \in E}$ vu comme un vecteur dont les coordonnées sont dans $\mathbb{R}$. Soit $\mu = \{\mu(x)\}_{x \in E}$ une mesure de probabilité sur E , on définit le produit

$$\forall y \in E, \quad \mu P(y) = \sum_{x \in E} \mu(x)P(x, y). \quad (2.9)$$

Si E est fini, il s'agit du produit à gauche $\mu^\dagger \cdot P$ entre un vecteur transposé et une matrice. Par convention, on omet le symbole transposé $\dagger$ dans (2.9). Pour $n \geq 1$, le produit matriciel P^n se définit par récurrence

$$\forall x, y \in E, \quad P^{n+1}(x, y) = P \times P^n(x, y) = \sum_{z \in E} P(x, z)P^n(z, y) = \sum_{z \in E} P^n(x, z)P(z, y) \quad (2.10)$$

avec la convention $P^1 = P$.

Théorème 2.3 (Chapman-Kolmogorov). *Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov sur E de matrice de transition P dont la donnée initiale X_0 est distribuée selon la loi μ_0 . Alors, la probabilité d'observer la trajectoire $\{x_0, x_1, \dots, x_n\}$ est donnée par*

$$\mathbb{P}_{\mu_0}(X_0 = x_0, X_1 = x_1, \dots, X_n = x_n) = \mu_0(x_0) P(x_0, x_1) P(x_1, x_2) \dots P(x_{n-1}, x_n). \quad (2.11)$$

La loi μ_n de X_n est déterminée par l'équation de Chapman-Kolmogorov

$$\forall y \in E, \quad \mu_n(y) = \mu_{n-1}P(y) = \mu_0 P^n(y). \quad (2.12)$$

Si initialement la chaîne de Markov part de $X_0 = x$, alors la distribution initiale est donnée par $\mu_0(y) = \mathbf{1}_{\{y=x\}}$ et (2.12) s'écrit

$$\forall y \in E, \quad \mathbb{P}(X_n = y \mid X_0 = x) = P^n(x, y). \quad (2.13)$$

Soit h une fonction bornée de E dans $\mathbb{R}$. Si initialement $X_0 = x$, l'espérance de $h(X_n)$ s'écrit

$$\mathbb{E}(h(X_n) \mid X_0 = x) = P^n h(x). \quad (2.14)$$

L'espérance $\mathbb{E}(h(X_n) \mid X_0 = x)$ sachant la donnée initiale est, à ce stade, posée comme une notation, mais on verra au chapitre 9 que celle-ci coïncide avec l'espérance conditionnelle.

On interprète l'équation de Chapman-Kolmogorov en disant que la probabilité d'observer X_n en y est la somme des probabilités de toutes les trajectoires possibles de la chaîne de Markov partant de x_0 et arrivant en y au temps n

$$\forall y \in E, \quad \mu_n(y) = \sum_{\{x_0, x_1, \dots, x_{n-1}\} \in E^n} \mu_0(x_0) P(x_0, x_1) P(x_1, x_2) \dots P(x_{n-1}, y).$$

Réciproquement, on remarquera que tout processus $\{X_n\}_{n \geq 0}$ dont la distribution satisfait (2.11) pour tout $n \geq 0$ est une chaîne de Markov.

Démonstration du théorème 2.3. Pour prouver (2.11), on procède par des conditionnements successifs et on applique la propriété de Markov (2.1) à chaque étape pour retrouver les probabilités de transition (2.2)

$$\begin{aligned} \mathbb{P}_{\mu_0}(X_0 = x_0, \dots, X_n = x_n) \\ &= \mathbb{P}_{\mu_0}(X_0 = x_0, \dots, X_{n-1} = x_{n-1}) \mathbb{P}_{\mu_0}(X_n = x_n \mid X_0 = x_0, \dots, X_{n-1} = x_{n-1}) \\ &= \mathbb{P}_{\mu_0}(X_0 = x_0, \dots, X_{n-1} = x_{n-1}) P(x_{n-1}, x_n) = \mu_0(x_0) P(x_0, x_1) P(x_1, x_2) \dots P(x_{n-1}, x_n). \end{aligned}$$

À la dernière étape, nous avons utilisé que X_0 est distribuée selon μ_0 .

Pour obtenir la loi de X_n , on écrit que X_{n-1} peut prendre toutes les valeurs possibles

$$\begin{aligned} \mu_n(y) &= \mathbb{P}_{\mu_0}(X_n = y) = \sum_{x \in E} \mathbb{P}_{\mu_0}(X_{n-1} = x, X_n = y) \\ &= \sum_{x \in E} \mathbb{P}_{\mu_0}(X_{n-1} = x) \mathbb{P}_{\mu_0}(X_n = y \mid X_{n-1} = x) \\ &= \sum_{x \in E} \mu_{n-1}(x) P(x, y) = \mu_{n-1} P(y) = \mu_0 P^n(y) \end{aligned}$$

où la dernière relation s'obtient par récurrence. L'identité (2.13) n'est qu'un cas particulier de l'équation de Chapman-Kolmogorov (2.12) pour une mesure initiale concentrée en x .

Si initialement $X_0 = x$, l'espérance de $h(X_n)$ peut se décomposer à l'aide de (2.13)

$$\mathbb{E}(h(X_n) \mid X_0 = x) = \sum_y h(y) \mathbb{P}(X_n = y \mid X_0 = x) = \sum_y P^n(x, y) h(y) = P^n h(x).$$

□

Reprenons l'exemple de la chaîne à 3 états dont la matrice de transition est donnée par (2.4). Les probabilités de transition après 2 pas de temps sont obtenues par produit matriciel $\mathbb{P}(X_2 = y | X_0 = x) = P^2(x, y)$

$$P^2 = \begin{pmatrix} 3/16 & 3/4 & 1/16 \\ 1/4 & 11/16 & 1/16 \\ 1/4 & 3/4 & 0 \end{pmatrix}.$$

Notations.

Soit A un évènement portant sur la trajectoire de la chaîne de Markov $\{X_n\}_{n \geq 0}$. Si la chaîne de Markov part d'un site x de E , la probabilité de A sera notée

$$\mathbb{P}_x(A) := \mathbb{P}(A | X_0 = x).$$

Si la donnée initiale X_0 est distribuée sous μ_0 , on notera

$$\mathbb{P}_{\mu_0}(A) := \sum_{x \in E} \mu_0(x) \mathbb{P}(A | X_0 = x). \quad (2.15)$$

On utilisera l'abréviation suivante pour décrire l'espérance au temps n d'une chaîne de Markov partant d'un site x de E

$$\mathbb{E}_x(h(X_n)) := \mathbb{E}(h(X_n) | X_0 = x) = \sum_{y \in E} h(y) \mathbb{P}(X_n = y | X_0 = x). \quad (2.16)$$

Si la variable X_0 est initialement distribuée sous une mesure μ_0 , on notera

$$\mathbb{E}_{\mu_0}(h(X_n)) := \mu_0 P^n h = \sum_{\{x_0, x_1, \dots, x_n\} \in E^{n+1}} \mu_0(x_0) P(x_0, x_1) \dots P(x_{n-1}, x_n) h(x_n). \quad (2.17)$$

2.3.2 Processus décalé en temps

Le théorème ci-dessous est une généralisation simple, mais très utile, de la propriété de Markov (2.1).

Théorème 2.4. *Pour toute collection d'états $\{x_0, \dots, x_n, y_1, \dots, y_K\}$ de E*

$$\mathbb{P}(X_{n+1} = y_1, \dots, X_{n+K} = y_K | X_0 = x_0, \dots, X_n = x_n) = \mathbb{P}(X_1 = y_1, \dots, X_K = y_K | X_0 = x_n). \quad (2.18)$$

Soit A un évènement de E^K dépendant uniquement des variables $\{X_{n+1}, \dots, X_{n+K}\}$ et B un évènement de E^n dépendant de $\{X_0, \dots, X_{n-1}\}$. Si $\{X_n = x_n\}$ est fixé, alors le conditionnement jusqu'au temps n n'est pas affecté par l'évènement B

$$\mathbb{P}(A | \{X_n = x_n\} \cap B) = \mathbb{P}(A | X_n = x_n). \quad (2.19)$$

De plus, conditionnellement à $\{X_n = x_n\}$, les évènements A et B sont indépendants

$$\mathbb{P}(A \cap B | X_n = x_n) = \mathbb{P}(A | X_n = x_n) \mathbb{P}(B | X_n = x_n). \quad (2.20)$$

L'identité (2.18) s'interprète en disant que conditionnellement à $\{X_n = x_n\}$, le processus décalé en temps $\{X_{n+k}\}_{k \geq 0}$ est une chaîne de Markov de matrice de transition P partant de x_n au temps 0 et indépendante du passé (cf. figure 2.3).

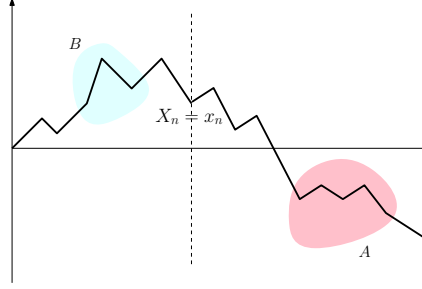


FIGURE 2.3 – Pour illustrer le théorème 2.4, on considère une marche aléatoire et deux événements. L'évènement A correspond au passage de la marche à travers la zone rose après le temps n et B au passage à travers la zone bleue. Si $X_n = x_n$ est fixé le comportement de la marche aléatoire après n est indépendante du fait que B soit réalisé avant le temps n .

Démonstration. On prouve le résultat à l'aide du théorème 2.3

$$\begin{aligned}
 & \mathbb{P}(X_{n+1} = y_1, \dots, X_{n+K} = y_K \mid X_0 = x_0, \dots, X_n = x_n) \\
 &= \frac{\mathbb{P}(X_0 = x_0, \dots, X_n = x_n, X_{n+1} = y_1, \dots, X_{n+K} = y_K)}{\mathbb{P}(X_0 = x_0, \dots, X_n = x_n)} \\
 &= P(x_n, y_1)P(y_1, y_2) \dots P(y_{K-1}, y_K) \\
 &= \mathbb{P}(X_1 = y_1, \dots, X_K = y_K \mid X_0 = x_n),
 \end{aligned}$$

où la seconde égalité s'obtient en utilisant la relation (2.11).

Pour montrer l'identité (2.19), on écrit

$$\mathbb{P}(A \mid \{X_n = x_n\} \cap B) = \frac{\mathbb{P}(A \cap \{X_n = x_n\} \cap B)}{\mathbb{P}(\{X_n = x_n\} \cap B)}. \quad (2.21)$$

La probabilité au numérateur se décompose en fonction des trajectoires

$$\begin{aligned}
 \mathbb{P}(A \cap \{X_n = x_n\} \cap B) &= \sum_{\substack{\{x_0, \dots, x_{n-1}\} \in B \\ \{y_1, \dots, y_K\} \in A}} \mathbb{P}(X_0 = x_0, \dots, X_n = x_n, X_{n+1} = y_1, \dots, X_{n+K} = y_K) \\
 &= \sum_{\substack{\{x_0, \dots, x_{n-1}\} \in B \\ \{y_1, \dots, y_K\} \in A}} \mathbb{P}(X_0 = x_0, \dots, X_n = x_n) \frac{\mathbb{P}(X_0 = x_0, \dots, X_n = x_n, X_{n+1} = y_1, \dots, X_{n+K} = y_K)}{\mathbb{P}(X_0 = x_0, \dots, X_n = x_n)} \\
 &= \sum_{\substack{\{x_0, \dots, x_{n-1}\} \in B \\ \{y_1, \dots, y_K\} \in A}} \mathbb{P}(X_0 = x_0, \dots, X_n = x_n) \mathbb{P}(X_{n+1} = y_1, \dots, X_{n+K} = y_K \mid X_0 = x_0, \dots, X_n = x_n).
 \end{aligned}$$

L'évènement $\{X_n = x_n\}$ étant fixé, la relation (2.18) permet de simplifier cette expression

$$\begin{aligned}
& \mathbb{P}(A \cap \{X_n = x_n\} \cap B) \\
&= \sum_{\substack{\{x_0, \dots, x_{n-1}\} \in B \\ \{y_1, \dots, y_K\} \in A}} \mathbb{P}(X_0 = x_0, \dots, X_n = x_n) \mathbb{P}(X_{n+1} = y_1, \dots, X_{n+K} = y_K \mid X_n = x_n) \\
&= \mathbb{P}(A \mid X_n = x_n) \sum_{\{x_0, \dots, x_{n-1}\} \in B} \mathbb{P}(X_0 = x_0, \dots, X_n = x_n) \\
&= \mathbb{P}(A \mid X_n = x_n) \mathbb{P}(B \cap \{X_n = x_n\}).
\end{aligned} \tag{2.22}$$

Le quotient (2.21) se simplifie et l'identité (2.19) est ainsi prouvée.

La preuve de (2.20) est similaire, il suffit d'appliquer une nouvelle fois (2.22)

$$\begin{aligned}
\mathbb{P}(A \cap B \mid X_n = x_n) &= \frac{\mathbb{P}(A \cap \{X_n = x_n\} \cap B)}{\mathbb{P}(X_n = x_n)} = \frac{\mathbb{P}(A \mid X_n = x_n) \mathbb{P}(B \cap \{X_n = x_n\})}{\mathbb{P}(X_n = x_n)} \\
&= \mathbb{P}(A \mid X_n = x_n) \mathbb{P}(B \mid X_n = x_n).
\end{aligned}$$

□

La propriété de Markov forte, qui sera établie dans la section suivante, généralise le théorème 2.4 en conditionnant en fonction de temps aléatoires.

2.4 Temps d'arrêt et propriété de Markov forte

Un *temps d'arrêt* T associé à un processus aléatoire discret $\{X_n\}_{n \geq 0}$ est une variable aléatoire à valeurs dans $\mathbb{N} \cup \{+\infty\}$ telle que pour tout $n \geq 0$, l'évènement $\{T = n\}$ est entièrement déterminé par les variables $\{X_0, \dots, X_n\}$, c'est-à-dire que pour tout n , il existe une fonction $\varphi_n : E^{n+1} \rightarrow \mathbb{R}$ telle que

$$1_{\{T=n\}} = \varphi_n(X_0, \dots, X_n).$$

Un temps d'arrêt très souvent utilisé est le premier temps d'atteinte d'un sous-ensemble $A \subset E$ par le processus $\{X_n\}_{n \geq 0}$

$$T_A = \inf \{n \geq 0; \quad X_n \in A\},$$

avec éventuellement $T_A = +\infty$ si le sous-ensemble A n'est jamais atteint. On a alors

$$1_{\{T_A=n\}} = 1_{\{X_0 \notin A, \dots, X_{n-1} \notin A, X_n \in A\}}.$$

Un temps d'arrêt T permet de stopper un processus $\{X_n\}_{n \geq 0}$ à un temps aléatoire dépendant uniquement du passé et du présent : l'évènement $\{T = n\}$ ne doit pas contenir d'information sur ce qui se passe au-delà du temps n . Par exemple, on peut chercher le meilleur moment pour convertir une devise sur le marché des changes. Le moment optimal sera choisi par rapport à la connaissance du passé et du présent, mais, à moins de délit d'initié, la décision ne pourra pas être influencée par le futur. Les temps d'arrêt jouent un rôle privilégié dans la théorie des processus aléatoires et nous les retrouverons

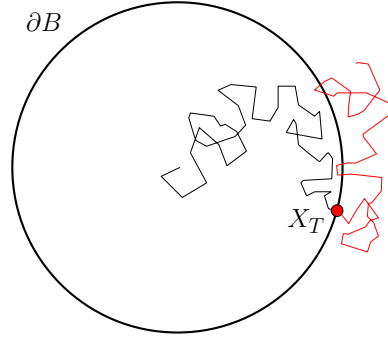


FIGURE 2.4 – La trajectoire d’une marche aléatoire partant de l’origine est représentée sur ce schéma. La marche atteint la frontière de la boule ∂B pour la première fois au temps d’arrêt T . Conditionnellement au point d’impact X_T , la seconde partie de la trajectoire est indépendante de la première.

tout au long de ce cours. En particulier, nous en donnerons une définition plus formelle au chapitre 9 (section 9.4).

Une conséquence importante de la propriété de Markov est que le processus décalé en temps $\{X_{n+k}\}_{k \geq 0}$ demeure, conditionnellement à $\{X_n = x\}$, une chaîne de Markov de matrice de transition P partant de x au temps 0 (cf. section 2.3.2). Cette propriété reste valable pour des décalages en temps par des temps d’arrêt.

Théorème 2.5 (Propriété de Markov forte). *Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov de matrice de transition P et de loi initiale μ_0 . On considère T un temps d’arrêt pour cette chaîne de Markov. Conditionnellement à $\{T < \infty\}$ et $X_T = x$, le processus décalé en temps $\{X_{T+k}\}_{k \geq 0}$ est une chaîne de Markov de matrice de transition P partant initialement de x .*

De plus, toujours conditionnellement à $\{T < \infty\}$ et $X_T = x$, la chaîne de Markov $\{X_{T+k}\}_{k \geq 0}$ est indépendante de $\{X_0, X_1, \dots, X_{T-1}\}$.

Démonstration. On rappelle que la notation $\mathbb{P}_{\mu_0}$ est définie en (2.15). Soit B un évènement dépendant uniquement de $\{X_0, X_1, \dots, X_{T-1}\}$. Alors pour tout entier ℓ , l’évènement $B \cap \{T = \ell\}$ est déterminé par $\{X_0, X_1, \dots, X_\ell\}$. On peut donc écrire pour tout $k \geq 1$

$$\begin{aligned} \mathbb{P}_{\mu_0}(\{X_{T+1} = x_1, \dots, X_{T+k} = x_k\} \cap B \cap \{T = \ell\} \cap \{X_T = x\}) \\ = \mathbb{P}_{\mu_0}(\{X_{\ell+1} = x_1, \dots, X_{\ell+k} = x_k\} \cap B \cap \{T = \ell\} \cap \{X_\ell = x\}) \\ = \mathbb{P}(X_1 = x_1, \dots, X_k = x_k \mid X_0 = x) \mathbb{P}_{\mu_0}(B \cap \{T = \ell\} \cap \{X_\ell = x\}) \end{aligned}$$

où la dernière égalité est une conséquence de la propriété de Markov (2.18) au temps ℓ . Il suffit de sommer ces équations pour toutes les valeurs de ℓ et on obtient

$$\begin{aligned} \mathbb{P}_{\mu_0}(\{X_{T+1} = x_1, \dots, X_{T+k} = x_k\} \cap B \cap \{X_T = x\} \cap \{T < \infty\}) \\ = \mathbb{P}(X_1 = x_1, \dots, X_k = x_k \mid X_0 = x) \mathbb{P}_{\mu_0}(B \cap \{X_T = x\} \cap \{T < \infty\}). \end{aligned}$$

Pour conclure le théorème, on reconstruit les probabilités conditionnelles en divisant les deux membres de l’équation par $\mathbb{P}_{\mu_0}(\{X_T = x\} \cap \{T < \infty\})$

$$\begin{aligned} \mathbb{P}_{\mu_0}(\{X_{T+1} = x_1, \dots, X_{T+k} = x_k\} \cap B \mid \{X_T = x\} \cap \{T < \infty\}) \\ = \mathbb{P}(X_1 = x_1, \dots, X_k = x_k \mid X_0 = x) \mathbb{P}_{\mu_0}(B \mid \{X_T = x\} \cap \{T < \infty\}). \end{aligned}$$

□

La propriété de Markov forte sera utilisée à plusieurs reprises dans la suite du cours. L'exemple du temps d'atteinte, illustré figure 2.4, montre l'intérêt de cette propriété.

2.5 Applications

Cette section regroupe trois applications de la propriété de Markov. La première met en évidence les liens entre la loi d'une marche aléatoire et l'équation de la chaleur en utilisant simplement l'équation de Chapman-Kolmogorov. L'importance des temps d'arrêt est ensuite illustrée dans les deux applications suivantes.

2.5.1 Équation de la chaleur ★

On considère la marche aléatoire symétrique $\{X_n\}_{n \geq 0}$ sur l'intervalle $\{1, \dots, L\}$ avec conditions périodiques dont la matrice de transition a été définie en (2.6) (en choisissant $p = q = 1/2$). L'équation de Chapman-Kolmogorov de la distribution au temps $n + 1$ s'écrit

$$\forall x \in \{1, \dots, L\}, \quad (2.23)$$

$$\mu_{n+1}(x) = \mu_n(x+1)P(x+1, x) + \mu_n(x-1)P(x-1, x) = \frac{1}{2}\mu_n(x+1) + \frac{1}{2}\mu_n(x-1)$$

où on identifie $L + 1$ avec le site 1 et 0 avec le site L . Cette équation dit simplement que pour être en x au temps $n + 1$, la marche devait se trouver en $x - 1$ ou $x + 1$ au temps précédent. On peut ainsi déterminer complètement μ_n au cours du temps. Pour ceci, nous allons étudier des lois initiales de la forme

$$\forall x \in \{1, \dots, L\}, \quad \mu_0(x) = \frac{1}{L} \left(1 + \sum_{k=1}^K a_k \cos\left(\frac{2\pi}{L}kx\right) + b_k \sin\left(\frac{2\pi}{L}kx\right) \right) \quad (2.24)$$

avec $K \in \{1, \dots, L\}$ et où les coefficients a_k, b_k vérifient

$$\inf_{r \in [0,1]} \left\{ \sum_{k=1}^K a_k \cos(2\pi kr) + b_k \sin(2\pi kr) \right\} \geq -1 \quad (2.25)$$

(cette dernière condition sert juste à assurer que la probabilité μ_0 ne prend pas des valeurs négatives). Si tous les a_k, b_k sont nuls, la donnée initiale est uniformément distribuée sur $\{1, \dots, L\}$. Les coefficients a_k, b_k sont simplement les coefficients de la transformée de Fourier (discrète) de μ_0 et par conséquent toute mesure initiale peut être décomposée sous la forme (2.24).

En utilisant les identités

$$\cos(a+b) + \cos(a-b) = 2\cos(b)\cos(a) \quad \text{et} \quad \sin(a+b) + \sin(a-b) = 2\cos(b)\sin(a)$$

on voit que

$$\mu_n(x) = \frac{1}{L} \left(1 + \sum_{k=1}^K \cos\left(\frac{2\pi}{L}k\right)^n \left[a_k \cos\left(\frac{2\pi}{L}kx\right) + b_k \sin\left(\frac{2\pi}{L}kx\right) \right] \right) \quad (2.26)$$

est bien la solution de la récurrence (2.23) avec donnée initiale μ_0 .

Nous allons appliquer ce résultat pour étudier le comportement d'une particule marquée dans un liquide. Le marquage de particules est souvent utilisé pour suivre les déplacements de matières dans une réaction chimique ou dans le sang. Pour modéliser le déplacement d'un marqueur placé dans une solution, nous allons supposer que le marqueur est distribué initialement dans une boîte $[0, 1]^d$ (de 1cm de côté) selon une loi de densité $\rho_0(r)$. Subdivisons la boîte $[0, 1]^d$ en L^d petits cubes de côté $1/L$ avec L très grand (disons que $1/L$ est de l'ordre de 10^{-6} cm soit quelques centaines d'ångströms). Le comportement microscopique précis du marqueur est très compliqué à décrire et nous nous contenterons d'une approximation en suivant simplement le cube où le marqueur se trouve. La solution étant à l'équilibre, le marqueur se déplace uniformément au gré des collisions microscopiques et on suppose qu'il peut passer d'un cube à un de ses voisins avec probabilité $\frac{1}{2d}$ (pour simplifier, on exclut les cubes voisins qui n'ont pas de face commune). Pour décrire l'évolution temporelle de ce marqueur, nous allons seulement considérer le cas de la dimension $d = 1$ qui correspond à un déplacement sur l'intervalle $[0, 1]$ (le cas général se traiterait de la même façon) et nous identifierons cet intervalle à un domaine périodique. Dans ce cadre simplifié, le marqueur a un comportement statistique proche de celui d'une marche aléatoire dans le domaine $\{1, \dots, L\}$ pour L extrêmement grand. On va donc analyser l'asymptotique de (2.26) quand L tend vers l'infini.

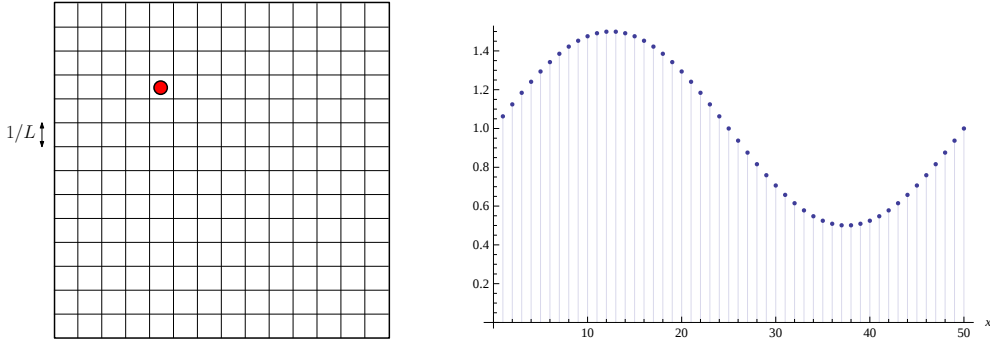


FIGURE 2.5 – Sur le schéma de gauche, la boîte $[0, 1]^2$ a été subdivisée en carrés de côté $1/L$ et la position du marqueur est identifiée par le carré dans lequel il se trouve. Le graphe de droite représente la discrétisation de la densité $\rho_0(r) = 1 + \sin(2\pi r)/2$ en subdivisant l'intervalle $[0, 1]$ avec $L = 50$.

La première étape est de décrire la donnée initiale. À l'échelle macroscopique, le marqueur est distribué dans la boîte $[0, 1]$ selon la densité $\rho_0(r)$ qui sera choisie de la forme

$$\rho_0(r) = 1 + \sum_{k=1}^K a_k \cos(2\pi kr) + b_k \sin(2\pi kr)$$

où K est un entier fixé et les a_k, b_k vérifient (2.25). On définit aussi

$$\forall t > 0, \quad \rho(t, r) = 1 + \sum_{k=1}^K \exp(-2\pi^2 k^2 t) [a_k \cos(2\pi kr) + b_k \sin(2\pi kr)]. \quad (2.27)$$

La répartition microscopique initiale s'obtient par discrétisation (cf. figure 2.5)

$$\forall x \in \{1, \dots, L\}, \quad \mu_0(x) = \frac{1}{L} \rho_0\left(\frac{x}{L}\right)$$

et satisfait donc (2.24). L'évolution de la mesure μ_n a été déterminée en (2.26). Soit $r \in [0, 1]$ une position macroscopique et $t > 0$ un temps macroscopique, on leur associe les suites d'entiers

$$x_L = \lfloor rL \rfloor \in \{1, \dots, L\}, \quad n_L = \lfloor tL^2 \rfloor \in \mathbb{N}$$

où $\lfloor \cdot \rfloor$ est la partie entière. Il faut interpréter x_L comme la position microscopique correspondant à r et n_L comme un temps microscopique

$$\lim_{L \rightarrow \infty} \frac{x_L}{L} = r, \quad \lim_{L \rightarrow \infty} \frac{n_L}{L^2} = t.$$

Pour tout $k \leq K$, on a

$$\begin{aligned} \lim_{L \rightarrow \infty} \cos\left(\frac{2\pi}{L}k\right)^{n_L} \left[a_k \cos\left(\frac{2\pi}{L}kx_L\right) + b_k \sin\left(\frac{2\pi}{L}kx_L\right) \right] \\ = \exp(-2\pi^2 k^2 t) [a_k \cos(2\pi kr) + b_k \sin(2\pi kr)] \end{aligned}$$

ceci s'obtient en utilisant le développement limité du cosinus à l'origine

$$\begin{aligned} \cos\left(\frac{2\pi}{L}k\right)^{n_L} &= \left(1 - \frac{1}{2}\left(\frac{2\pi}{L}k\right)^2 + o\left(\frac{1}{L^3}\right)\right)^{n_L} = \exp\left(-\frac{2\pi^2 k^2}{L^2}n_L + o\left(\frac{1}{L}\right)\right) \\ &\approx \exp(-2\pi^2 k^2 t). \end{aligned}$$

On en déduit que la densité microscopique μ_{n_L} converge vers la densité macroscopique $\rho(t, r)$ définie en (2.27)

$$\lim_{L \rightarrow \infty} L \mu_{n_L}(x_L) = \rho(t, r).$$

On vérifie facilement que $\rho(t, r)$ satisfait l'équation de la chaleur

$$\partial_t \rho(t, r) = \frac{1}{2} \partial_r^2 \rho(t, r).$$

Dans la pratique comme L est très grand, le comportement macroscopique du marqueur est bien décrit par l'équation de la chaleur. Le passage de modèles microscopiques où le comportement est aléatoire à des descriptions macroscopiques plus régulières (comme ici l'équation de la chaleur) est un problème très étudié en physique et en mathématiques. De nombreuses théories ont été développées pour comprendre comment des structures régulières peuvent émerger dans la limite macroscopique, mais des problèmes ouverts demeurent et il s'agit d'un domaine de recherche actuellement très actif en mathématiques.

On remarquera que le passage des coordonnées microscopiques x_L, n_L aux coordonnées macroscopiques r, t s'est fait en changeant l'échelle spatiale d'un facteur L et le temps d'un facteur L^2 . Ce changement d'échelle est lié au théorème central limite, en effet le marqueur effectue une marche aléatoire et il ne peut explorer que des distances de l'ordre $\sqrt{n}$ en un temps n . Pour que le marcheur puisse se déplacer d'une distance de l'ordre de L , il faut donc attendre des temps proportionnels à L^2 . Cette analogie entre la limite gaussienne de la marche aléatoire et l'équation de la chaleur est au cœur de la théorie du mouvement brownien.

2.5.2 Ruine du joueur

Imaginons 2 joueurs A et B qui décident de miser leurs fortunes respectives a et b au jeu. À la fin de chaque partie, la fortune du gagnant augmente de 1 et celle du perdant diminue de 1. Le jeu s'arrête quand l'un des deux joueurs est ruiné. Retraduit en termes probabilistes, on se donne $p \in [0, 1]$ et une suite de variables aléatoires indépendantes, identiquement distribuées

$$\mathbb{P}(\xi_i = 1) = p \quad \text{et} \quad \mathbb{P}(\xi_i = -1) = 1 - p.$$

On notera $X_n = X_0 + \sum_{i=1}^n \xi_i$ la fortune de A au temps n et $X_0 = a$ sa fortune initiale. La fortune de B est alors donnée par $a + b - X_n$. Par construction, $\{X_n\}_{n \geq 0}$ est une chaîne de Markov partant de a au temps 0.

Si $p = 1/2$ le jeu est *équilibré*, sinon il est *biaisé* et un des joueurs a plus de chances de gagner que l'autre. On cherche à déterminer la probabilité que le joueur A soit ruiné avant B c'est-à-dire la probabilité avec laquelle le processus $\mathcal{X} = \{X_n\}_{n \geq 0}$, qui décrit la fortune de A , va atteindre 0 avant $a + b$ (cf. figure 2.6)

$$u(a) = \mathbb{P}_a(\mathcal{X} \text{ atteint } 0 \text{ avant } a + b).$$

On peut réécrire cette probabilité à l'aide des temps d'arrêt

$$T_0 = \inf\{n \geq 0; \quad X_n = 0\} \quad \text{et} \quad T_{a+b} = \inf\{n \geq 0; \quad X_n = a + b\}.$$

On admet que pour presque toute trajectoire la chaîne de Markov $\{X_n\}_{n \geq 0}$ atteindra 0 ou $a + b$ au bout d'un temps fini presque sûrement (ce résultat sera prouvé au lemme 3.8). Par conséquent $\inf\{T_0, T_{a+b}\}$ est fini presque sûrement et on peut réécrire

$$u(a) = \mathbb{P}_a(T_0 < T_{a+b}).$$

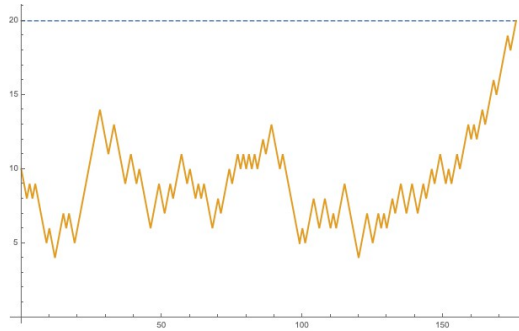


FIGURE 2.6 – Premier temps d'atteinte de $20 = a + b$ ou de 0 pour la marche aléatoire symétrique X_n partant de $a = 10$. Dans cette réalisation $T_{20} = 176 < T_0$.

Au lieu de chercher uniquement à calculer $u(a)$, nous allons généraliser la question et supposer que la chaîne de Markov puisse prendre toutes les valeurs initiales $i \in \{1, \dots, a + b - 1\}$

$$u(i) = \mathbb{P}_i(T_0 < T_{a+b}).$$

On pose aussi

$$u(0) = 1 \quad \text{et} \quad u(a+b) = 0.$$

Pour calculer u , nous allons effectuer une *analyse du premier pas*. L'idée consiste à conditionner la trajectoire en fonction des valeurs prises par X_1 pour démontrer que u satisfait les équations suivantes

$$i \in \{1, \dots, a+b-1\}, \quad u(i) = pu(i+1) + (1-p)u(i-1). \quad (2.28)$$

On définit $\tilde{\mathcal{X}}$ la chaîne de Markov décalée en temps $\tilde{X}_n = X_{n+1}$ pour $n \geq 0$. Si $X_0 \neq 0, a+b$ alors les événements $\{\mathcal{X} \text{ atteint } 0 \text{ avant } a+b\}$ et $\{\tilde{\mathcal{X}} \text{ atteint } 0 \text{ avant } a+b\}$ sont identiques. Par l'identité (2.18), X_0 et $\{\tilde{X}_n\}_{n \geq 0}$ sont indépendants sachant X_1 . On obtient donc pour $i \in \{1, \dots, a+b-1\}$

$$\begin{aligned} \mathbb{P}(\{T_0 < T_{a+b}\}, X_0 = i, X_1 = i+1) &= \mathbb{P}(\{\mathcal{X} \text{ atteint } 0 \text{ avant } a+b\}, X_0 = i, X_1 = i+1) \\ &= \mathbb{P}(\{\tilde{\mathcal{X}} \text{ atteint } 0 \text{ avant } a+b\}, X_0 = i, X_1 = i+1) \\ &= \mathbb{P}(X_0 = i, X_1 = i+1) \mathbb{P}(\{\tilde{\mathcal{X}} \text{ atteint } 0 \text{ avant } a+b\} \mid \tilde{X}_0 = i+1). \end{aligned}$$

La chaîne de Markov décalée en temps $\tilde{\mathcal{X}}$ ayant la même loi que $\{X_n\}_{n \geq 0}$, l'identité se réécrit donc

$$\mathbb{P}(\{T_0 < T_{a+b}\}, X_0 = i, X_1 = i+1) = \mathbb{P}(X_0 = i) p u(i+1). \quad (2.29)$$

De la même façon, si $X_1 = i-1$, on obtient

$$\mathbb{P}(\{T_0 < T_{a+b}\}, X_0 = i, X_1 = i-1) = \mathbb{P}(X_0 = i) (1-p) u(i-1). \quad (2.30)$$

Étant donné $X_0 = i$, le pas suivant sera $X_1 = i \pm 1$

$$u(i) = \mathbb{P}_i(\{T_0 < T_{a+b}\}, X_1 = i+1) + \mathbb{P}_i(\{T_0 < T_{a+b}\}, X_1 = i-1).$$

En utilisant (2.29) et (2.30), on obtient la relation annoncée en (2.28)

$$i \in \{1, \dots, a+b-1\}, \quad u(i) = pu(i+1) + (1-p)u(i-1)$$

avec $u(0) = 1$ et $u(a+b) = 0$. On peut ensuite déterminer explicitement les solutions de cette récurrence linéaire en remarquant que les racines du polynôme caractéristique associé $py^2 - y + (1-p)$ sont 1 et $(1-p)/p$. On distingue deux cas.

Jeu biaisé.

Quand $p \neq 1/2$, les racines du polynôme sont distinctes et la solution de (2.28) s'écrit sous la forme $u(i) = c_1 + c_2 \left(\frac{1-p}{p}\right)^i$. Il suffit ensuite d'ajuster les constantes c_1 et c_2 en fonction des conditions aux bords pour conclure

$$u(i) = \frac{\left(\frac{1-p}{p}\right)^{(a+b)} - \left(\frac{1-p}{p}\right)^i}{\left(\frac{1-p}{p}\right)^{(a+b)} - 1}.$$

Jeu équilibré.

Quand $p = 1/2$, les deux racines du polynôme valent 1 et on trouve

$$u(i) = 1 - \frac{i}{a+b}.$$

Ces résultats permettent de retrouver la valeur $u(a)$ cherchée qui vaut

$$u(a) = \frac{\left(\frac{1-p}{p}\right)^{(a+b)} - \left(\frac{1-p}{p}\right)^a}{\left(\frac{1-p}{p}\right)^{(a+b)} - 1}.$$

pour un jeu biaisé et $u(a) = \frac{b}{a+b}$ dans le cas d'un jeu équilibré.

On définit $T = \inf\{T_0, T_{a+b}\}$ le temps où le jeu s'arrête. Une méthode analogue permet de calculer l'espérance $\mathbb{E}_i(T)$. En utilisant la chaîne de Markov décalée en temps $\tilde{X}_n = X_{n+1}$, on obtient pour tout i dans $\{1, \dots, a+b-1\}$

$$\mathbb{E}_i(T) = 1 + p \mathbb{E}_{i+1}(T) + (1-p) \mathbb{E}_{i-1}(T).$$

Il suffit donc de résoudre le système linéaire satisfait par $v(i) = \mathbb{E}_i(T)$

$$v(i) = 1 + p v(i+1) + (1-p) v(i-1)$$

avec les conditions aux bords $v(0) = v(a+b) = 0$. Dans le cas d'un jeu équilibré ($p = 1/2$), on trouve pour tout i de $\{0, \dots, a+b\}$

$$\mathbb{E}_i(T) = i(a+b-i). \quad (2.31)$$

Le problème de la ruine du joueur constitue un cas d'école, mais des questions similaires se posent aux compagnies d'assurance qui cherchent à estimer la probabilité d'incidents aléatoires pouvant les conduire à la faillite.

2.5.3 Méthode de Monte-Carlo pour un problème de Dirichlet ★

L'équation de Laplace intervient dans de nombreux domaines de la physique (mécanique des fluides, électromagnétisme...) et elle joue un rôle clef en analyse. Le problème de Dirichlet associé se formule de la façon suivante. On considère un domaine D borné, connexe et régulier de $\mathbb{R}^2$. Les résultats suivants se généralisent facilement à toute dimension $d \geq 1$. Le bord de D sera noté ∂D . Étant donnée une fonction régulière φ définie sur ∂D , on cherche à déterminer $f : D \rightarrow \mathbb{R}$ telle que

$$\Delta f(r) = \sum_{k=1}^2 \partial_k^2 f(r) = 0 \quad \text{et} \quad \forall r \in \partial D, \quad f(r) = \varphi(r) \quad (2.32)$$

où ∂_k^2 est la dérivée seconde par rapport à la $k^{\text{ième}}$ coordonnée. Cette équation modélise par exemple la variation de température dans une plaque de métal en contact avec différentes sources de chaleur sur son bord. La plaque de métal est représentée par le domaine

D , la température au point $r \in D$ par $f(r)$ et les températures au bord de la plaque sont égales à φ .

Il existe différentes méthodes pour résoudre numériquement l'équation (2.32). Nous allons décrire une approche probabiliste dite méthode de *Monte-Carlo*. La première étape consiste à discrétiser le domaine D avec un maillage de taille $1/L$, on notera D_L le réseau discret correspondant

$$D_L = \left\{ \left(\frac{i}{L}, \frac{j}{L} \right) \in D \quad \text{avec} \quad i, j \in \mathbb{Z}^2 \right\} = D \cap \frac{1}{L} \mathbb{Z}^2.$$

Le bord discret ∂D_L est constitué des sites de $D^c \cap \frac{1}{L} \mathbb{Z}^2$ à distance inférieure à $1/L$ de D (cf. figure 2.7).

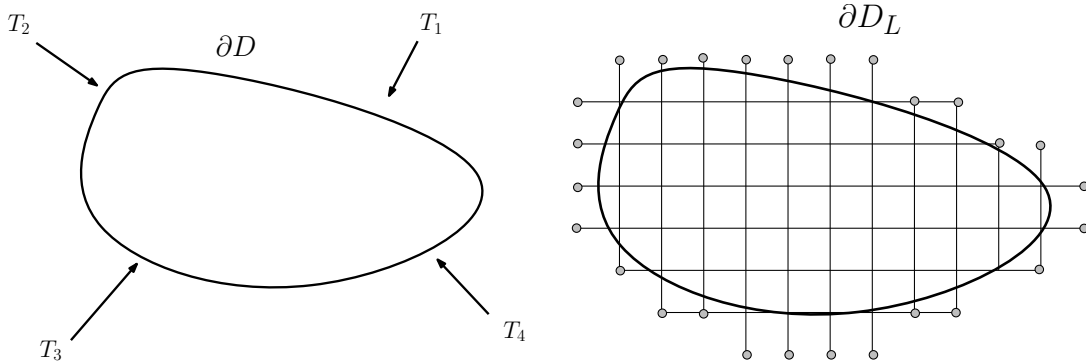


FIGURE 2.7 – Un domaine borné $D \subset \mathbb{R}^2$ avec différentes températures imposées sur son bord ∂D . Le maillage de D induit une frontière discrète ∂D_L représentée par les sites gris.

Considérons f une fonction C^3 sur D . La formule de Taylor implique que

$$\begin{cases} f\left(\left(\frac{i+1}{L}, \frac{j}{L}\right)\right) - f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) = \frac{1}{L} \partial_1 f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) + \frac{1}{2L^2} \partial_1^2 f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) + O(1/L^3) \\ f\left(\left(\frac{i-1}{L}, \frac{j}{L}\right)\right) - f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) = -\frac{1}{L} \partial_1 f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) + \frac{1}{2L^2} \partial_1^2 f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) + O(1/L^3) \end{cases}$$

En sommant ces deux équations, on obtient une approximation de la dérivée seconde ∂_1^2 quand le pas du maillage tend vers 0

$$\partial_1^2 f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) = L^2 \left[f\left(\left(\frac{i+1}{L}, \frac{j}{L}\right)\right) + f\left(\left(\frac{i-1}{L}, \frac{j}{L}\right)\right) - 2f\left(\left(\frac{i}{L}, \frac{j}{L}\right)\right) \right] + O(1/L).$$

Pour toute fonction F de $D_L \cup \partial D_L$ dans $\mathbb{R}$, on définit le *Laplacien discret* en tout point x de D_L par

$$\bar{\Delta} F(x) = \sum_{\substack{y \in D_L \cup \partial D_L \\ y \sim x}} (F(y) - F(x))$$

où la notation $y \sim x$ signifie qu'on somme sur les voisins y de x , c'est-à-dire les sites de $D_L \cup \partial D_L$ à distance $1/L$ de x . En particulier si x est proche du bord, les valeurs de F sur la frontière ∂D_L interviennent. Les calculs précédents justifient cette définition car pour des fonctions f régulières, le Laplacien discret est une bonne approximation du Laplacien

$\Delta f(x) = L^2 \bar{\Delta} f(x) + O(1/L)$. Le problème de Dirichlet continu (2.32) peut être approché par le problème de Dirichlet discret

$$\forall x \in D_L, \quad \bar{\Delta} F(x) = 0 \quad \text{et} \quad \forall y \in \partial D_L, \quad F(y) = \varphi_L(y) \quad (2.33)$$

où la contrainte de Dirichlet φ sur ∂D a été discrétisée en une fonction φ_L sur ∂D_L .

La solution du problème de Dirichlet discret peut s'écrire à l'aide d'une marche aléatoire $\{X_n\}_{n \geq 0}$ sur $\frac{1}{L}\mathbb{Z}^2$ qui saute uniformément d'un site à chacun de ses voisins avec probabilité $\frac{1}{4}$. On note $T_{\partial D_L}$ le premier temps d'atteinte du bord ∂D_L par cette marche et $X_{T_{\partial D_L}}$ le site de ∂D_L où la marche est sortie. On admet (provisoirement) que $T_{\partial D_L}$ est fini presque sûrement, c'est-à-dire qu'une marche aléatoire finit toujours par sortir du domaine (ceci sera vérifié au lemme 3.8). On définit

$$\forall x \in D_L, \quad F(x) = \mathbb{E}_x \left(\varphi_L(X_{T_{\partial D_L}}) \right). \quad (2.34)$$

Nous vérifierons dans la suite que la fonction F est l'unique solution du problème de Dirichlet discret (2.33). Avant cela, nous allons utiliser cette représentation probabiliste pour déterminer numériquement la solution de (2.33).

Pour chaque site x de D_L , la valeur de $F(x)$ se calcule en évaluant l'espérance de la fonction φ_L au point où la marche aléatoire partant de $x \in D_L$ a touché ∂D_L pour la première fois. Concrètement pour calculer $F(x)$, il suffit de construire K réalisations de la marche aléatoire $\{X_n^{(i)}\}_{i \leq K}$ partant de $x \in D_L$ et de prendre la moyenne sur les différents points de sortie $X_{T_{\partial D_L}}^{(i)}$ du domaine D_L . Les K marches étant indépendantes, les variables $\varphi_L(X_{T_{\partial D_L}}^{(i)})$ sont indépendantes, identiquement distribuées et la loi des grands nombres fournit une approximation de F quand K tend vers l'infini

$$\lim_{K \rightarrow \infty} \frac{1}{K} \sum_{i=1}^K \varphi_L(X_{T_{\partial D_L}}^{(i)}) = \mathbb{E}_x \left(\varphi_L(X_{T_{\partial D_L}}) \right) \quad \text{presque sûrement.} \quad (2.35)$$

Il faudra choisir K de façon optimale pour que l'approximation (2.35) donne des résultats pertinents sans nécessiter des temps de calcul trop longs. Le temps nécessaire pour réaliser ces simulations sera proportionnel à $K \mathbb{E}_x(T_{\partial D_L})$.

Il reste à vérifier que F est bien solution du problème de Dirichlet discret (2.33). On remarque que le cas de la dimension 1 a déjà été traité avec la ruine du joueur en (2.28) : on avait alors $L = a + b$ et $\varphi_L(0) = 1, \varphi_L(L) = 1$. On procède de la même façon en décomposant la trajectoire de la marche aléatoire en fonction du premier pas. Étant donné $X_0 = x$ dans D_L , le pas suivant sera $X_1 = y$ pour $y \sim x$

$$F(x) = \sum_{y \sim x} \mathbb{E}_x(\varphi_L(X_{T_{\partial D_L}}) \mathbf{1}_{X_1=y}).$$

En considérant, comme dans la ruine du joueur, la marche décalée en temps $\tilde{X}_n = X_{n+1}$, on voit que F est la solution de l'équation de Laplace discrète

$$F(x) = \frac{1}{4} \sum_{y \sim x} F(y) \quad \Rightarrow \quad \bar{\Delta} F(x) = 0.$$

De plus si x appartient au bord ∂D_L alors $X_{T_{\partial D_L}} = x$. La fonction F satisfait bien la contrainte de Dirichlet sur le bord. Par conséquent (2.34) fournit une représentation explicite de la solution du problème de Dirichlet discret (2.33).

On peut facilement vérifier que la solution de (2.33) est unique. En effet, si F_1 et F_2 sont deux solutions alors $\psi = F_2 - F_1$ satisfait $\Delta \psi = 0$ et vaut 0 sur le bord ∂D_L . Supposons que ψ atteigne son maximum en $x_0 \in D_L$. Comme

$$\psi(x_0) = \frac{1}{4} \sum_{y \sim x_0} \psi(y) \quad \text{et} \quad \psi(y) \leq \psi(x_0)$$

alors $\psi(y) = \psi(x_0)$ pour tous les voisins de y de x_0 . En itérant cette procédure, on peut trouver un chemin de sites $x_0, x_1, x_2, \dots, x_\ell$ avec $x_\ell \in \partial D_L$ tels que $x_i \sim x_{i+1}$ pour tout $i \geq 0$ et $\psi(x_i) = \psi(x_0)$. Comme $x_\ell \in \partial D_L$, on en déduit que $0 = \psi(x_0)$. Le maximum de ψ étant pris en x_0 , ceci implique que $\psi \leq 0$. Par symétrie $\psi \geq 0$ et on a donc prouvé l'unicité de la solution du problème de Dirichlet discret (2.33).

La méthode de Monte-Carlo permet d'évaluer la formulation probabiliste (2.33) du problème de Dirichlet discret. Dans la pratique plusieurs questions se posent pour mettre en œuvre cette méthode de Monte-Carlo. Quel pas de maillage $1/L$ doit-on prendre pour que le problème de Dirichlet discret (2.33) approche correctement l'équation limite. La valeur de L étant choisie, combien de marches indépendantes en chaque site doit-on lancer pour que la moyenne empirique (2.35) décrive correctement la solution du problème de Dirichlet discret (2.33).

L'intérêt de la méthode de Monte-Carlo est d'être facile à implémenter et d'être performante quand la dimension d devient grande. On peut aussi généraliser cette méthode pour résoudre des équations aux dérivées partielles de la forme

$$\nabla(\sigma(r) \cdot \nabla f) + b(r) \cdot \nabla f = 0 \quad \text{et} \quad \forall r \in \partial D, \quad f(r) = \varphi(r)$$

où σ et b sont des champs de vecteurs sur D .

Chapitre 3

Mesures invariantes

Dans le chapitre précédent, nous avons vu au théorème 2.3 que la distribution d'une chaîne de Markov de matrice de transition P évolue à chaque pas de temps selon les équations (2.12) de Chapman-Kolmogorov $\mu_{n+1} = \mu_n P$. Nous allons, dans ce chapitre, étudier les mesures de probabilité μ invariantes par ces équations c'est-à-dire les mesures satisfaisant $\mu = \mu P$. Ces mesures joueront par la suite un rôle clef dans le comportement asymptotique des chaînes de Markov.

Dans ce chapitre nous ne considérerons que des chaînes de Markov sur des espaces d'états E finis. On notera $|E|$ le cardinal de E . Le cas des espaces d'états dénombrables sera abordé au chapitre 4.

3.1 Mesures de probabilité invariantes

3.1.1 Définition des mesures invariantes

Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov sur un espace E fini de matrice de transition P .

Définition 3.1. La mesure π sur E est une mesure invariante pour la chaîne de Markov $\{X_n\}_{n \geq 0}$ si $\pi = \pi P$, c'est-à-dire

$$\forall y \in E, \quad \pi(y) = \sum_{x \in E} \pi(x) P(x, y).$$

Dans ce chapitre, nous ne considérerons que des mesures de probabilité invariantes, i.e. des mesures invariantes satisfaisant la normalisation $\sum_{x \in E} \pi(x) = 1$.

Au lemme 3.2, nous verrons que si la chaîne de Markov est distribuée initialement selon une mesure invariante π (on note $\mu_0 = \pi$) alors la distribution à tout temps n reste $\mu_n = \pi$. Une mesure invariante π décrit donc un système dans un état stationnaire. En théorie des probabilités, on utilise de façon équivalente la terminologie *mesure stationnaire* ou mesure invariante.

Comme exemple concret de mesure invariante, on peut imaginer un gaz à l'équilibre dans une pièce (confinée) : les positions des atomes sont aléatoires, les atomes se déplacent au cours du temps, mais à chaque instant, ils restent uniformément répartis dans la pièce. Leur distribution est donc invariante. Par contre si on ouvre un flacon de parfum au centre de cette pièce, le parfum se répand et la distribution des molécules n'est

pas stationnaire au cours du temps. En anticipant un peu sur les prochains chapitres, on sait cependant qu'au bout d'un temps très long le parfum se sera entièrement répandu et que ses molécules seront distribuées uniformément dans toute la pièce, le système aura donc convergé vers la mesure invariante. Nous reviendrons sur l'interprétation d'une mesure invariante en utilisant l'analogie avec un gaz dans la section 3.3.3.

3.1.2 Exemples de mesures invariantes

Considérons un graphe $\mathcal{G} = (\mathcal{S}, \mathcal{E})$ fini et sans boucles (un site n'est jamais relié à lui-même). On notera $\mathcal{S}$ l'ensemble des sites du graphe et $\mathcal{E}$ l'ensemble des arêtes entre les sites. On définit une marche aléatoire $\{X_n\}_{n \geq 0}$ sur $\mathcal{S}$ dont les probabilités de transition d'un site vers ses voisins sont uniformes

$$\forall x, y \in \mathcal{S}, \quad P(x, y) = \mathbf{1}_{\{x \sim y\}} \frac{1}{\deg(x)} \quad (3.1)$$

où la notation $x \sim y$ signifie que x et y sont reliés par une arête du graphe ($(x, y) \in \mathcal{E}$) et $\deg(x)$ est le nombre de voisins de x , i.e. le degré de x .

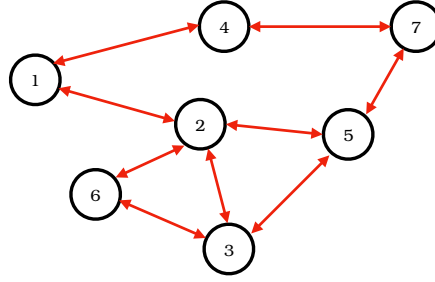


FIGURE 3.1 – Le graphe $\mathcal{G}$ représenté ci-dessus a pour sommets $\mathcal{S} = \{1, \dots, 7\}$. La matrice de transition P définie en (3.1) est indexée en fonction des arêtes de $\mathcal{E}$ dessinées en rouge. Sur cet exemple $\deg(1) = 2$ et $\deg(2) = 4$. Ainsi la marche aléatoire peut aller de 1 vers 2 avec probabilité $P(1, 2) = 1/2$ et de 2 vers 1 avec probabilité $P(2, 1) = 1/4$.

On définit la probabilité π sur $\mathcal{S}$ par

$$\forall x \in \mathcal{S}, \quad \pi(x) = \frac{\deg(x)}{2|\mathcal{E}|} \quad (3.2)$$

où $|\mathcal{E}|$ est le cardinal du nombre d'arêtes. Le graphe ne contenant pas de boucles, on voit facilement que π est bien normalisée car

$$\sum_{x \in \mathcal{S}} \deg(x) = 2|\mathcal{E}|.$$

En effet, chaque arête du graphe est comptée deux fois dans la somme. La mesure de probabilité π est une mesure invariante car pour tout y dans $\mathcal{S}$

$$\sum_{x \in \mathcal{S}} \pi(x) P(x, y) = \sum_{x \in \mathcal{S}} \frac{\deg(x)}{2|\mathcal{E}|} \frac{1}{\deg(x)} \mathbf{1}_{\{x \sim y\}} = \frac{1}{2|\mathcal{E}|} \sum_{x \sim y} \mathbf{1}_{\{x \sim y\}} = \pi(y).$$

Le résultat précédent implique que la marche aléatoire symétrique sur $\{1, \dots, L\}$, définie par la matrice de transition (2.6) pour $p = q = 1/2$, avec conditions périodiques (cf. figure 3.2) a pour mesure de probabilité invariante la mesure uniforme

$$\forall x \in \{1, \dots, L\}, \quad \pi(x) = \frac{1}{L}.$$

Le choix de la matrice de transition (3.1) ne s'applique pas aux probabilités de sauts

$$P(x, x+1) = p, \quad P(x, x-1) = 1-p,$$

si $p \neq 1/2$. Cependant dans le cas particulier de la marche sur $\{1, \dots, L\}$ avec conditions périodiques, on peut facilement vérifier que la mesure uniforme $\pi(x) = \frac{1}{L}$ est encore invariante pour tout $p \in [0, 1]$.

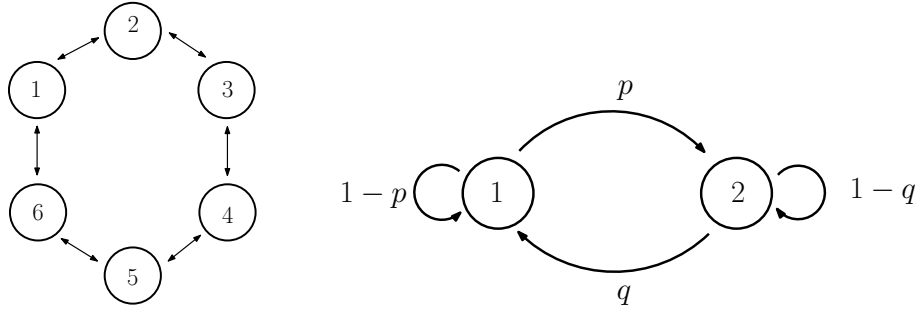


FIGURE 3.2 – À gauche, le graphe des transitions associé à la marche aléatoire symétrique sur le domaine périodique $\{1, \dots, 6\}$. Le graphe des transitions pour la chaîne à deux états est représenté à droite.

En général, si la matrice de transition n'est pas de la forme (3.1), la mesure invariante peut être très difficile à calculer. Considérons l'exemple simple, d'une chaîne de Markov à deux états $\{1, 2\}$ (cf. figure 3.2) dont la matrice de transition est donnée par

$$P = \begin{pmatrix} 1-p & p \\ q & 1-q \end{pmatrix} \quad (3.3)$$

avec $p, q \in]0, 1[$. La mesure de probabilité

$$\pi(1) = \frac{q}{p+q}, \quad \pi(2) = \frac{p}{p+q}$$

est invariante. Pour le montrer, il suffit de calculer

$$\pi P(1) = \pi(1)P(1,1) + \pi(2)P(2,1) = \frac{q}{p+q}(1-p) + \frac{p}{p+q}q = \pi(1).$$

On a de même $\pi P(2) = \pi(2)$. Si p ou q est différent de 0, on peut vérifier que π est l'unique mesure de probabilité invariante. La distribution π reflète le comportement de la chaîne de Markov. En effet, si p est proche de 0 et q proche de 1, alors la chaîne de Markov va avoir tendance à sauter de 2 vers 1 et à rester sur place une fois qu'elle est au site 1. On remarque que $\pi(1)$ est effectivement proche de 1 et $\pi(2)$ proche de 0.

3.1.3 Premières propriétés des mesures invariantes

Lemme 3.2. *Si π est une mesure invariante alors $\pi = \pi P^n$ pour tout $n \geq 1$. De plus, si la chaîne de Markov est initialement distribuée sous la mesure invariante π , alors*

$$\forall n, m \geq 0, \quad (X_0, X_1, \dots, X_n) \stackrel{(loi)}{=} (X_m, X_{m+1}, \dots, X_{m+n}). \quad (3.4)$$

Ce lemme traduit le fait qu'une mesure invariante est préservée pour tout temps $\pi(x) = \mathbb{P}_\pi(X_n = x)$. Ceci explique le rôle clef joué par les mesures invariantes dans le comportement asymptotique des chaînes de Markov (cf. chapitre 5).

Démonstration. Comme π est invariante, on a $\pi = \pi P$ et en composant à droite par la matrice de transition, on en déduit $\pi P = \pi P^2$, ce qui donne $\pi = \pi P^2$. Il suffit ensuite d'itérer cette relation pour conclure le Lemme 3.2. On peut décomposer le calcul matriciel précédant en écrivant

$$\pi(y) = \sum_{z \in E} \pi(z) P(z, y).$$

On en déduit en remplaçant $\pi(z)$ par $\pi P(z)$ que

$$\pi(y) = \sum_{z \in E} \sum_{x \in E} \pi(x) P(x, z) P(z, y) = \sum_{x \in E} \pi(x) \left[\sum_{z \in E} P(x, z) P(z, y) \right] = \sum_{x \in E} \pi(x) P^2(x, y).$$

On retrouve bien la relation $\pi = \pi P^2$.

Pour démontrer l'identité en loi (3.4), commençons par appliquer la propriété de Markov au temps m pour un événement A dans E^{n+1}

$$\begin{aligned} \mathbb{P}_\pi((X_m, \dots, X_{m+n}) \in A) &= \sum_{x \in E} \mathbb{P}((X_m, \dots, X_{m+n}) \in A \mid X_m = x) \mathbb{P}_\pi(X_m = x) \\ &= \sum_{x \in E} \mathbb{P}((X_0, \dots, X_n) \in A \mid X_0 = x) \mathbb{P}_\pi(X_m = x). \end{aligned} \quad (3.5)$$

Comme la mesure π est invariante, on sait par la première partie du lemme que

$$\mathbb{P}_\pi(X_m = x) = \pi P^m(x) = \pi(x).$$

Ceci conclut la preuve de (3.4) car l'égalité (3.5) se réécrit pour tout événement A

$$\mathbb{P}_\pi((X_m, \dots, X_{m+n}) \in A) = \mathbb{P}_\pi((X_0, \dots, X_n) \in A).$$

□

Dans le cas d'un espace E fini, la mesure de probabilité invariante π peut être interprétée comme un vecteur à valeurs dans $[0, 1]^{|E|}$ qui est un espace compact. Un argument de compacité va donc nous permettre de justifier l'existence d'au moins une mesure de probabilité invariante π .

Théorème 3.3. *Pour toute chaîne de Markov sur un espace d'états fini E , il existe une mesure de probabilité invariante.*

Notons que ce théorème ne dit rien sur l'unicité de la mesure de probabilité invariante.

Démonstration. À une mesure de probabilité ν donnée sur E , on associe la suite de mesures sur E

$$\nu_n = \frac{1}{n} \sum_{k=1}^n \nu P^k.$$

Les vecteurs $(\nu_n(x))_{x \in E}$ prennent leurs valeurs dans l'ensemble compact $[0, 1]^{|E|}$. Il existe donc une suite extraite ν_{n_k} qui converge vers une mesure π dans E

$$\forall x \in E, \quad \lim_{k \rightarrow \infty} \nu_{n_k}(x) = \pi(x).$$

Par construction π est bien une mesure de probabilité car E est fini et pour tout n , on a

$$\sum_{x \in E} \nu_n(x) = 1 \quad \Rightarrow \quad \sum_{x \in E} \pi(x) = 1.$$

Nous allons vérifier que π est une mesure invariante. Par construction

$$\nu_n P = \frac{1}{n} \sum_{k=1}^n \nu P^{k+1} = \nu_n + \frac{1}{n} (\nu P^{n+1} - \nu P).$$

Pour chaque x de E , la suite $\{\nu_n P(x) - \nu_n(x)\}_{n \geq 1}$ converge donc vers 0. En passant à la limite, on en déduit que π est invariante car

$$\forall x \in E, \quad \pi P(x) - \pi(x) = \lim_{k \rightarrow \infty} \nu_{n_k} P(x) - \nu_{n_k}(x) = 0.$$

□

3.2 Irréductibilité et unicité des mesures de probabilité invariantes

La structure des mesures invariantes dépend de la matrice de transition P . Reprenons l'exemple de la chaîne sur deux sites dont la matrice de transition est définie en (3.3) dans le cas particulier où il n'existe aucune transition entre les états 1 et 2, i.e. $p = q = 0$. Les deux mesures

$$\pi_1(x) = \mathbf{1}_{\{x=1\}} \quad \text{et} \quad \pi_2(x) = \mathbf{1}_{\{x=2\}}$$

sont invariantes car si la chaîne part d'un état, elle y restera tout le temps. On peut vérifier que toute combinaison linéaire $\lambda \pi_1 + (1 - \lambda) \pi_2$ (avec $\lambda \in [0, 1]$) est une mesure invariante. Cet exemple très simple montre qu'il peut exister une infinité de mesures invariantes. Considérons maintenant le cas où il n'existe pas de transition de 1 vers 2, i.e. $p = 0$ et $q \neq 0$. Alors l'unique mesure invariante est $\pi_1(x) = \mathbf{1}_{\{x=1\}}$ et elle n'attribue aucun poids au site 2 (ce qui n'est pas le cas si $p > 0$). Par conséquent, le support des mesures invariantes dépend de la structure de P .

Avant d'analyser l'unicité des mesures invariantes, nous allons introduire la notion d'*irréductibilité* qui est équivalente à la connexité du graphe des transitions associé à la matrice P .

3.2.1 Irréductibilité

Soit $X = \{X_n\}_{n \geq 0}$ une chaîne de Markov sur E de matrice de transition P .

Définition 3.4. Soient x, y deux états de E . On dit que

- (i) x communique avec y , on note $x \rightarrow y$, s'il existe un entier $n \geq 0$ et des états $x_0 = x, x_1, \dots, x_n = y \in E$ tels que $P(x_0, x_1) \cdots P(x_{n-1}, x_n) > 0$, i.e. si

$$\mathbb{P}_x(X_n = y) = \mathbb{P}(X_n = y | X_0 = x) > 0. \quad (3.6)$$

- (ii) x et y communiquent, on note $x \leftrightarrow y$, si x communique avec y et y communique avec x .

Définition 3.5.

- (i) Une classe $E_0 \subset E$ est dite fermée si pour tous $x, y \in E$

$$x \in E_0 \quad \text{et} \quad x \rightarrow y \quad \text{alors} \quad y \in E_0.$$

- (ii) Une classe $E_0 \subset E$ est dite irréductible si $x \leftrightarrow y$ pour tous $x, y \in E_0$ et E_0 est fermée.

- (iii) La chaîne de Markov X est dite irréductible si l'espace d'états E est irréductible.

Les définitions (i) et (ii) sont illustrées figure 3.3. La définition (iii) est la plus importante en pratique car les chaînes irréductibles vont constituer notre principal cadre d'étude.

La restriction de la chaîne de Markov à une classe fermée E_0 est ainsi une chaîne de Markov d'espace d'états E_0 . Enfin si $E_0 = \{x_0\}$ est fermée, alors l'état x_0 est dit *absorbant* car une fois que la chaîne de Markov l'a atteint, elle reste bloquée dans cet état pour toujours.

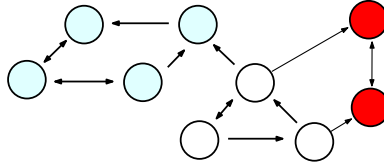


FIGURE 3.3 – Dans ce graphe de transition, on distingue 2 classes irréductibles fermées en rouge et en bleu clair. Aucun site de ces deux classes ne communique avec les sites blancs qui ne forment pas une classe fermée. La chaîne de Markov peut être restreinte à chacune des classes irréductibles.

3.2.2 Unicité des mesures de probabilité invariantes

Pour E fini, l'existence d'une mesure invariante a été prouvée au théorème 3.3. L'hypothèse d'irréductibilité de la chaîne permet de renforcer ce résultat.

Théorème 3.6. Pour toute chaîne de Markov irréductible sur un espace d'états fini E , il existe une unique mesure de probabilité invariante π . De plus $\pi(x) > 0$ pour tout $x \in E$.

Démonstration. Soit π une mesure invariante (son existence est assurée par le théorème 3.3). Nous allons d'abord vérifier que $\pi(x) > 0$ pour tout $x \in E$. Comme $\sum_{y \in E} \pi(y) = 1$, il existe x_0 de E tel que $\pi(x_0) > 0$. La chaîne étant irréductible, x_0 communique avec tout y de E et pour chaque y , il existe un entier n tel que $P^n(x_0, y) > 0$. La mesure π est invariante $\pi = \pi P^n$ (cf. lemme 3.2) et on en déduit

$$\pi(y) = \sum_{z \in E} \pi(z) P^n(z, y) \geq \pi(x_0) P^n(x_0, y) > 0.$$

Pour montrer l'unicité, nous allons d'abord établir un résultat préliminaire et prouver que toute fonction h de E dans $\mathbb{R}$ vérifiant

$$\forall x \in E, \quad h(x) = \sum_{y \in E} P(x, y) h(y) \quad (3.7)$$

est nécessairement constante. Une fonction h satisfaisant (3.7) est dite *harmonique*. Nous avons déjà rencontré des fonctions harmoniques section 2.5.3. Comme E est fini, il existe un état x_0 où la fonction atteint son minimum $h(x_0) = \min_{y \in E} h(y)$. S'il existait z de E tel que $P(x_0, z) > 0$ et $h(z) > h(x_0)$, on obtiendrait une contradiction en utilisant le fait que $\sum_{y \in E} P(x, y) = 1$

$$h(x_0) = \sum_{y \in E} P(x_0, y) h(y) > P(x_0, z) h(x_0) + \sum_{y \neq z} P(x_0, y) h(y) \geq h(x_0).$$

Ceci étant impossible, la fonction h est donc égale à $h(x_0)$ pour les états y connectés à x_0 , i.e. tels que $P(x_0, y) > 0$. Comme la chaîne est irréductible, on déduit en itérant cette procédure que h est constante sur E .

On remarque qu'une fonction harmonique h est un vecteur propre à droite pour P car $h = Ph$ tandis qu'une mesure invariante π est un vecteur propre à gauche car $\pi = \pi P$. Le résultat précédent sur les fonctions harmoniques implique que la matrice $P - \text{Id}$ a un noyau de dimension 1 (les vecteurs de la forme $\lambda(1, \dots, 1)$). La valeur propre 0 étant de multiplicité 1, elle est aussi valeur propre de multiplicité 1 pour la transposée $P^\dagger - \text{Id}$. Par conséquent s'il existe 2 mesures invariantes π_1, π_2 (que l'on peut interpréter comme des vecteurs) telles que $\pi_1 = \pi_1 P$ et $\pi_2 = \pi_2 P$ alors les deux vecteurs π_1, π_2 sont dans le noyau de $P^\dagger - \text{Id}$. Le noyau étant de dimension 1, il existe une constante c telle que $\pi_1 = c \pi_2$. Comme les deux mesures sont normalisées par 1, on en déduit que $\pi_1 = \pi_2$. $\square$

Exercice 3.7. On propose une preuve alternative de l'unicité des mesures invariantes du théorème 3.6. Supposons que π_1, π_2 soient deux mesures invariantes strictement positives sur E . Montrer que

$$\forall x, y \in E, \quad Q(x, y) = P(y, x) \frac{\pi_1(y)}{\pi_1(x)}$$

est une matrice de transition irréductible. Vérifier que $f(x) = \frac{\pi_2(x)}{\pi_1(x)}$ est harmonique pour Q , i.e. $f = Qf$. En utilisant le résultat sur l'unicité des fonctions harmoniques, en déduire que $\pi_1 = \pi_2$.

3.2.3 Construction de la mesure de probabilité invariante

Le théorème 3.3 a permis d'obtenir l'existence d'une mesure invariante de façon implicite. Dans cette section, nous allons construire la mesure invariante et en donner une expression explicite qui pourra se généraliser facilement aux espaces d'états dénombrables.

On rappelle la définition du premier temps d'atteinte T_x d'un élément x de E

$$T_x = \inf \{n \geq 0; \quad X_n = x\}.$$

On définit aussi

$$T_x^+ = \inf \{n \geq 1; \quad X_n = x\}.$$

Ces deux temps d'arrêt coïncident sauf si le site initial est x car dans ce cas $T_x = 0$ et T_x^+ est le *premier temps de retour* en x .

Une propriété importante des temps de retour dans le cas des espaces finis est la suivante :

Lemme 3.8. *Pour une chaîne de Markov irréductible sur un espace d'états E fini*

$$\forall x, y \in E, \quad \mathbb{E}_x(T_y^+) < \infty.$$

Démonstration. La chaîne étant irréductible et E fini, il existe $\varepsilon > 0$ et un entier n tels que pour tous x, y dans E

$$\exists j \leq n, \quad P^j(x, y) \geq \varepsilon.$$

La valeur j peut varier selon les couples x, y mais reste bornée par n . La probabilité d'atteindre y avant le temps n en partant de n'importe quel point est au moins ε . L'inégalité suivante est donc vraie uniformément en x, y

$$\mathbb{P}_x(T_y^+ > n) \leq 1 - \varepsilon.$$

Nous allons itérer ce résultat en conditionnant le processus par le passé jusqu'au temps $(k-1)n$

$$\begin{aligned} \mathbb{P}_x(T_y^+ > kn) &= \sum_{\substack{z \in E \\ z \neq y}} \mathbb{P}_x(\{T_y^+ > (k-1)n\} \cap \{X_{(k-1)n} = z\} \bigcap_{i=1}^n \{X_{(k-1)n+i} \neq y\}) \\ &= \sum_{\substack{z \in E \\ z \neq y}} \mathbb{P}_x(\{T_y^+ > (k-1)n\} \cap \{X_{(k-1)n} = z\}) \\ &\quad \mathbb{P}\left(\bigcap_{i=1}^n \{X_{(k-1)n+i} \neq y\} \mid \{T_y^+ > (k-1)n\} \cap \{X_{(k-1)n} = z\}\right) \end{aligned}$$

où z représente toutes les valeurs possibles pouvant être prises par $X_{(k-1)n}$. Par la propriété de Markov appliquée au temps $(k-1)n$ (cf. théorème 2.4), on en déduit que le conditionnement ne dépend que de la valeur de $X_{(k-1)n}$

$$\begin{aligned} \mathbb{P}_x(T_y^+ > kn) &= \sum_{\substack{z \in E \\ z \neq y}} \mathbb{P}_x\left(\{T_y^+ > (k-1)n\} \cap \{X_{(k-1)n} = z\}\right) \mathbb{P}\left(\bigcap_{i=1}^n \{X_{(k-1)n+i} \neq y\} \mid X_{(k-1)n} = z\right) \end{aligned}$$

Le dernier terme peut s'exprimer par la propriété de Markov comme l'évènement $\{T_y^+ > n\}$ pour la chaîne décalée en temps

$$\mathbb{P}\left(\bigcap_{i=1}^n \{X_{(k-1)n+i} \neq y\} \mid X_{(k-1)n} = z\right) = \mathbb{P}\left(\bigcap_{i=1}^n \{X_i \neq y\} \mid X_0 = z\right) = \mathbb{P}_z(T_y^+ > n).$$

Conditionnellement à $\{X_{(k-1)n} = z\}$, on peut le borner uniformément en z (avec $z \neq y$) par $1 - \varepsilon$ pour obtenir

$$\mathbb{P}_x(T_y^+ > kn) \leq (1 - \varepsilon)\mathbb{P}_x(T_y^+ > (k-1)n).$$

En itérant on obtient

$$\mathbb{P}_x(T_y^+ > kn) \leq (1 - \varepsilon)^k.$$

Pour toute variable aléatoire Z à valeurs dans $\mathbb{N}$, on rappelle l'identité classique

$$\mathbb{E}(Z) = \sum_{\ell \geq 1} \mathbb{P}(Z \geq \ell).$$

On obtient donc

$$\mathbb{E}_x(T_y^+) = \sum_{\ell \geq 1} \mathbb{P}_x(T_y^+ \geq \ell) \leq n \sum_{k \geq 0} \mathbb{P}_x(T_y^+ > kn) \leq n \sum_{k \geq 0} (1 - \varepsilon)^k < \infty.$$

Ce qui permet de conclure le lemme. $\square$

Le théorème suivant permet d'exprimer la mesure invariante en fonction de la fréquence à laquelle la chaîne de Markov visite les sites de E .

Théorème 3.9. *Pour une chaîne de Markov irréductible $\{X_n\}_{n \geq 0}$ sur un espace d'états E fini, l'unique mesure de probabilité invariante est donnée par*

$$\forall x \in E, \quad \pi(x) = \frac{1}{\mathbb{E}_x(T_x^+)}.$$

Démonstration. Étant donné x un élément de E , nous allons définir la mesure $\tilde{\pi}$ comme la moyenne du temps passé en chaque site par la chaîne de Markov entre deux passages par x

$$\forall y \in E, \quad \tilde{\pi}(y) = \mathbb{E}_x(\text{nombre de visites en } y \text{ avant de retourner en } x).$$

Commençons par réécrire la mesure $\tilde{\pi}$ sous une forme plus maniable

$$\tilde{\pi}(y) = \mathbb{E}_x\left(\sum_{n=0}^{T_x^+-1} 1_{\{X_n=y\}}\right) = \mathbb{E}_x\left(\sum_{n \geq 0} 1_{\{X_n=y\}} 1_{\{T_x^+ > n\}}\right) = \sum_{n \geq 0} \mathbb{P}_x(X_n = y, T_x^+ > n),$$

où nous avons utilisé le théorème de Fubini dans la dernière égalité. A priori, $\tilde{\pi}$ dépend du choix de x , mais nous allons montrer que ce n'est pas le cas.

Le lemme 3.8 implique que $\tilde{\pi}(y)$ est bien définie pour tout y car

$$\tilde{\pi}(y) \leq \sum_{n \geq 0} \mathbb{P}_x(T_x^+ > n) = \mathbb{E}_x(T_x^+) < \infty.$$

Par contre $\tilde{\pi}$ n'est pas une mesure de probabilité car elle n'est pas normalisée.

Pour montrer que $\tilde{\pi}$ est stationnaire, nous calculons

$$\sum_{z \in E} \tilde{\pi}(z)P(z, y) = \sum_{z \in E} \sum_{n \geq 0} \mathbb{P}_x(X_n = z, T_x^+ > n)P(z, y).$$

Le point clef de la preuve est de constater que l'évènement $\{T_x^+ > n\} = \{T_x^+ \geq n+1\}$ ne dépend que de $\{X_0, \dots, X_n\}$ (on sait juste que la marche n'est pas revenue en x avant le temps n). Par conséquent, on peut appliquer la propriété de Markov et écrire

$$\begin{aligned} \mathbb{P}_x(X_n = z, T_x^+ \geq n+1, X_{n+1} = y) &= \mathbb{P}_x(X_n = z, T_x^+ > n) \mathbb{P}(X_{n+1} = y \mid X_n = z, T_x^+ > n) \\ &= \mathbb{P}_x(X_n = z, T_x^+ > n)P(z, y). \end{aligned}$$

On déduit de cette relation que

$$\begin{aligned} \sum_{z \in E} \tilde{\pi}(z)P(z, y) &= \sum_{z \in E} \sum_{n \geq 0} \mathbb{P}_x(X_n = z, T_x^+ \geq n+1, X_{n+1} = y) \\ &= \sum_{n \geq 0} \mathbb{P}_x(T_x^+ \geq n+1, X_{n+1} = y) = \sum_{n \geq 1} \mathbb{P}_x(T_x^+ \geq n, X_n = y). \end{aligned}$$

Cette expression est très proche de la définition de $\tilde{\pi}$

$$\begin{aligned} \sum_{n \geq 1} \mathbb{P}_x(T_x^+ \geq n, X_n = y) &= \sum_{n \geq 0} \mathbb{P}_x(X_n = y, T_x^+ > n) - \mathbb{P}_x(T_x^+ > 0, X_0 = y) + \sum_{n \geq 1} \mathbb{P}_x(T_x^+ = n, X_n = y) \\ &= \tilde{\pi}(y) - \mathbb{P}_x(X_0 = y) + \mathbb{P}_x(X_{T_x^+} = y). \end{aligned}$$

Il suffit maintenant de considérer les deux cas :

- Si $y = x$ alors $\mathbb{P}_x(X_0 = y) = \mathbb{P}_x(X_{T_x^+} = y) = 1$.
- Si $y \neq x$ alors $\mathbb{P}_x(X_0 = y) = \mathbb{P}_x(X_{T_x^+} = y) = 0$.

On a donc prouvé que la mesure $\tilde{\pi}$ est invariante, i.e. $\sum_{z \in E} \tilde{\pi}(z)P(z, y) = \tilde{\pi}(y)$. Pour obtenir une mesure de probabilité, il suffit de la normaliser. Comme $\sum_{z \in E} \tilde{\pi}(z) = \mathbb{E}_x(T_x^+)$, la mesure de probabilité

$$\forall y \in E, \quad \pi(y) = \frac{\tilde{\pi}(y)}{\mathbb{E}_x(T_x^+)}$$

est l'unique mesure de probabilité invariante (cf. théorème 3.6). En particulier, elle vérifie

$$\pi(x) = \frac{\tilde{\pi}(x)}{\mathbb{E}_x(T_x^+)} = \frac{1}{\mathbb{E}_x(T_x^+)}.$$

Le site x qui servait de site de référence pour indexer les excursions de la chaîne a été choisi arbitrairement. L'identité ci-dessus est donc vérifiée pour tout x car la mesure de probabilité invariante est unique pour une chaîne de Markov irréductible. $\square$

3.3 Réversibilité et Théorème H

3.3.1 Réversibilité

Les mesures invariantes sont caractérisées par le système d'équations linéaires

$$\forall y \in E, \quad \pi(y) = \sum_{x \in E} \pi(x)P(x, y)$$

dont le nombre d'inconnues est égal au cardinal de E . Dans la pratique, ce cardinal peut être très grand (parfois infini) et il est souvent difficile de déterminer π analytiquement. La *réversibilité* est une condition suffisante et facile à vérifier qui assure l'existence d'une mesure invariante.

Définition 3.10. Une chaîne de Markov de matrice de transition P sur E est dite *réversible par rapport à la mesure π* si elle satisfait

$$\forall x, y \in E, \quad \pi(x)P(x, y) = \pi(y)P(y, x).$$

La réversibilité caractérise les systèmes à l'équilibre. Sous la mesure initiale π , si on observe une trajectoire $x_0, x_1, \dots, x_n$ pour une chaîne de Markov réversible alors la trajectoire inverse aura la même probabilité

$$\mathbb{P}_\pi(X_0 = x_0, X_1 = x_1, \dots, X_n = x_n) = \mathbb{P}_\pi(X_0 = x_n, X_1 = x_{n-1}, \dots, X_n = x_0).$$

Pour prouver cette relation, on écrit la propriété de Markov

$$\mathbb{P}_\pi(X_0 = x_0, X_1 = x_1, \dots, X_n = x_n) = \pi(x_0)P(x_0, x_1)P(x_1, x_2) \dots P(x_{n-1}, x_n).$$

En itérant de proche en proche la relation de réversibilité, on obtient

$$\begin{aligned} \mathbb{P}_\pi(X_0 = x_0, X_1 = x_1, \dots, X_n = x_n) &= P(x_1, x_0)\pi(x_1)P(x_1, x_2) \dots P(x_{n-1}, x_n) \\ &= P(x_1, x_0)P(x_2, x_1)\pi(x_2) \dots P(x_{n-1}, x_n) \\ &= \pi(x_n)P(x_n, x_{n-1})P(x_{n-1}, x_{n-2}) \dots P(x_1, x_0). \end{aligned}$$

Exemples.

- D'après la section 3.1.2, la marche aléatoire sur $\{1, \dots, L\}$ avec conditions périodiques et probabilités de sauts

$$P(x, x+1) = p, \quad P(x, x-1) = 1-p$$

a pour mesure invariante la mesure uniforme $\pi(x) = \frac{1}{L}$ pour tout $p \in [0, 1]$. Cette chaîne de Markov n'est réversible que pour $p = \frac{1}{2}$ car pour $p \neq \frac{1}{2}$

$$\pi(x)P(x, x+1) \neq \pi(x+1)P(x+1, x).$$

En effet, si $p \neq \frac{1}{2}$ la marche va tourner de façon privilégiée dans un sens et la probabilité de la voir tourner en sens inverse sera très faible. Si $p = \frac{1}{2}$, la marche est symétrique et toute fluctuation dans un sens sera aussi probable qu'une fluctuation en sens inverse.

- La chaîne de Markov à deux états de matrice de transition P donnée par (3.3) est aussi réversible pour la mesure $\pi(1) = \frac{q}{p+q}$, $\pi(2) = \frac{p}{p+q}$ car

$$\pi(1) P(1, 2) = \frac{qp}{p+q} = \pi(2) P(2, 1).$$

Dans la pratique, la réversibilité permet de vérifier facilement qu'une mesure est invariante.

Théorème 3.11. *Si une chaîne de Markov de matrice de transition P est réversible par rapport à la mesure π , alors π est une mesure invariante.*

Démonstration. Pour tout y dans E , on obtient en utilisant la réversibilité et $\sum_{x \in E} P(y, x) = 1$

$$\sum_{x \in E} \pi(x) P(x, y) = \sum_{x \in E} \pi(y) P(y, x) = \pi(y).$$

Par conséquent, π est bien une mesure invariante. $\square$

3.3.2 Théorème H pour les chaînes de Markov $\star$

Soient μ et ν deux mesures de probabilité sur un espace E . On suppose que $\mu(x) > 0$ pour tout x de E et on note $\mathbb{H}$ l'entropie relative de ν par rapport à μ

$$\mathbb{H}(\nu|\mu) = \sum_{x \in E} \left(\frac{\nu(x)}{\mu(x)} \log \frac{\nu(x)}{\mu(x)} \right) \mu(x).$$

L'entropie relative permet de mesurer la distance entre μ et ν car $\mathbb{H}$ est positive et ne s'annule que si $\nu = \mu$. Pour le vérifier, il suffit de remarquer que $\phi(u) = u \log(u)$ est strictement convexe et par l'inégalité de Jensen

$$\mathbb{H}(\nu|\mu) = \sum_{x \in E} \phi \left(\frac{\nu(x)}{\mu(x)} \right) \mu(x) \geq \phi \left(\sum_{x \in E} \frac{\nu(x)}{\mu(x)} \mu(x) \right) = 0.$$

L'inégalité est stricte dès qu'il existe un x pour lequel $\frac{\nu(x)}{\mu(x)} \neq 1$.

Théorème 3.12. *On considère une chaîne de Markov irréductible de matrice de transition P et de mesure invariante π . Alors pour toute probabilité μ , on a*

$$\mathbb{H}(\mu P|\pi) \leq \mathbb{H}(\mu|\pi).$$

Par conséquent, l'entropie relative $n \rightarrow \mathbb{H}(\mu P^n|\pi)$ décroît avec le temps.

Ce résultat fait écho au Théorème H pour l'équation de Boltzmann, néanmoins les physiciens utilisent la convention opposée pour l'entropie et considèrent plutôt la croissance de $-\mathbb{H}$. Nous verrons au chapitre 5 que les mesures μP^n relaxent vers la mesure d'équilibre π . La décroissance de l'entropie relative permet déjà d'entrevoir cette relaxation car la distance entre μP^n et π se réduit au cours du temps.

Démonstration. En utilisant la fonction $\phi(u) = u \log(u)$, on a

$$\mathbb{H}(\mu P | \pi) = \sum_{x \in E} \phi \left(\frac{1}{\pi(x)} \sum_{y \in E} \mu(y) P(y, x) \right) \pi(x) = \sum_{x \in E} \phi \left(\sum_{y \in E} \frac{\mu(y)}{\pi(y)} \frac{\pi(y) P(y, x)}{\pi(x)} \right) \pi(x).$$

La mesure π étant invariante, on peut vérifier que $y \rightarrow \frac{\pi(y) P(y, x)}{\pi(x)}$ est une mesure de probabilité sur E

$$\sum_{y \in E} \frac{\pi(y) P(y, x)}{\pi(x)} = \frac{1}{\pi(x)} \sum_{y \in E} \pi(y) P(y, x) = \frac{\pi(x)}{\pi(x)} = 1$$

on obtient par l'inégalité de Jensen

$$\mathbb{H}(\mu P | \pi) \leq \sum_{x \in E} \pi(x) \sum_{y \in E} \frac{\pi(y) P(y, x)}{\pi(x)} \phi \left(\frac{\mu(y)}{\pi(y)} \right) = \sum_{y \in E} \sum_{x \in E} \pi(y) P(y, x) \phi \left(\frac{\mu(y)}{\pi(y)} \right).$$

Comme $\sum_{x \in E} P(y, x) = 1$, on conclut que

$$\mathbb{H}(\mu P | \pi) \leq \mathbb{H}(\mu | \pi).$$

□

3.3.3 Application : modèle d'Ehrenfest

Nous allons illustrer les concepts de réversibilité et de décroissance de l'entropie avec une chaîne de Markov proposée par Paul et Tatiana Ehrenfest en 1907. Ce modèle a joué un rôle important pour établir les fondements de la mécanique statistique et comprendre le paradoxe de l'irréversibilité en théorie cinétique des gaz.

Commençons par un rappel historique pour expliquer les motivations qui ont conduit à introduire ce modèle. En 1872, Boltzmann propose une équation pour décrire l'évolution d'un gaz peu dense hors équilibre. Cette équation deviendra l'élément fondateur de la théorie cinétique des gaz. Le point de départ est de représenter les molécules dans un gaz comme un ensemble de sphères dures qui avancent en ligne droite à des vitesses différentes et qui rebondissent à la manière de boules de billard quand elles se touchent. Pour décrire le comportement d'un tel gaz, il n'est pas possible (ni souhaitable) de rendre compte de l'évolution de tous les atomes (l'ordre de grandeur de leur nombre étant 10^{23}), par contre une quantité plus globale comme la densité $f(t, x, v)$ de particules au temps t avec la position x et la vitesse v suffit à décrire le transport de matière. L'équation de Boltzmann régit l'évolution de la densité de particules

$$\begin{aligned} \partial_t f(t, x, v) + v \cdot \nabla_x f(t, x, v) &= Q(f, f) \\ Q(f, f)(v) &= \iint_{\mathbf{S}^2 \times \mathbb{R}^3} [f(v') f(v'_1) - f(v) f(v_1)] ((v - v_1) \cdot \nu)_+ dv_1 d\nu \end{aligned}$$

où ν est un vecteur intégré sur la sphère unité $\mathbf{S}^2$ et les vitesses après collisions s'écrivent

$$v' = v + \nu \cdot (v_1 - v) \nu, \quad v'_1 = v_1 - \nu \cdot (v_1 - v) \nu.$$

Une propriété fondamentale de cette équation connue comme le Théorème H, est la croissance de l'entropie au cours du temps

$$H(t) = - \iint dx dv f(t, x, v) \log f(t, x, v) \quad \text{et} \quad \partial_t H(t) \geq 0$$

(La convention en physique est opposée à celle des mathématiciens qui considèrent plutôt la décroissance de $-H(t)$, cf. section 3.3.2). Cette croissance traduit l'irréversibilité du système : si on perce un ballon rempli de gaz, le gaz s'échappe et le ballon se dégonfle. Ce mécanisme est irréversible, en effet il est très rare d'observer un ballon se regonflant spontanément.

L'équation de Boltzmann a pourtant suscité de nombreuses controverses car l'irréversibilité des solutions de cette équation semble incompatible avec la réversibilité de la dynamique microscopique. La dynamique microscopique est un immense billard avec 10^{23} boules rebondissant les unes sur les autres. Si on observe l'évolution de cette dynamique jusqu'au temps t et qu'à l'instant t toutes les vitesses sont renversées ($v \rightarrow -v$) alors le système microscopique revient en arrière en suivant exactement l'évolution inverse. En 1876, Loschmidt objectait que l'équation de Boltzmann ne pouvait pas rendre compte du système microscopique qui lui était réversible. Un second paradoxe est signalé par Zermelo en 1896 car le théorème de Poincaré assure que cette dynamique microscopique va repasser au cours du temps arbitrairement près de sa condition initiale et ce pour presque toutes les conditions initiales. Ceci pose encore la question de l'irréversibilité de l'équation de Boltzmann. Le modèle des époux Ehrenfest a permis de comprendre ces deux paradoxes.

On considère un récipient isolé coupé en deux par une paroi, la partie gauche est remplie d'un gaz et celle de droite est vide (cf. figure 3.4). À l'instant initial, un trou minuscule est percé dans la paroi pour permettre au gaz de passer d'un compartiment à l'autre. Pour simplifier le modèle, on imagine qu'à chaque pas de temps, un atome est choisi au hasard et transféré d'un compartiment à l'autre. On note X_n le nombre d'atomes dans la partie gauche au temps n et on suppose qu'initialement le récipient contient K atomes, i.e. $X_0 = K$. Cette chaîne de Markov a pour espace d'états $\{0, \dots, K\}$ et les probabilités de transition sont données par

$$\mathbb{P}(X_{n+1} = \ell - 1 | X_n = \ell) = \frac{\ell}{K}, \quad \mathbb{P}(X_{n+1} = \ell + 1 | X_n = \ell) = \frac{K - \ell}{K}.$$

Quand le système est à l'équilibre, les molécules sont réparties uniformément et la mesure invariante devrait intuitivement être une loi binomiale $\pi(\ell) = \frac{1}{2^K} \binom{K}{\ell}$ (on choisit ℓ molécules parmi K et on les place dans la partie gauche, les $K - \ell$ autres seront alors dans la partie droite). Pour le vérifier, il suffit de remarquer que cette chaîne de Markov est réversible pour la mesure invariante π

$$\pi(\ell)P(\ell, \ell + 1) = \frac{1}{2^K} \frac{K!}{\ell!(K - \ell)!} \frac{K - \ell}{K} = \pi(\ell + 1)P(\ell + 1, \ell).$$

La distribution de π est représentée figure 3.4. La réversibilité stochastique peut être vue comme l'analogie de la réversibilité des équations du mouvement pour la dynamique microscopique du gaz de sphères dures.

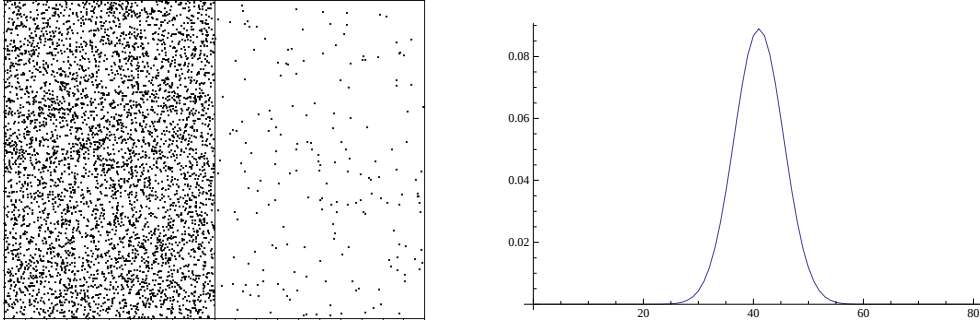


FIGURE 3.4 – Sur le schéma de gauche, le modèle d’Ehrenfest est représenté : le compartiment de gauche est rempli de molécules, celui de droite est presque vide. Un passage est ouvert entre les deux compartiments pour permettre le transfert des molécules. Le graphe de droite représente la distribution de π pour $K = 80$. Cette distribution est symétrique autour de sa moyenne et elle décrit l’état d’équilibre du gaz.

L’irréductibilité de la chaîne de Markov implique par le lemme 3.8 que la chaîne va revenir en chacun des points presque sûrement. Si initialement le compartiment de gauche est rempli de gaz et celui de droite est vide, les molécules vont d’abord se répartir assez rapidement dans tout le récipient mais si on attend assez longtemps toutes les molécules finiront par retourner dans le compartiment de gauche. Cette propriété est l’analogie du théorème de récurrence de Poincaré pour les systèmes dynamiques que Zermelo opposait à Boltzmann. Dans le cas de la chaîne d’Ehrenfest, le théorème 3.9 permet de calculer l’espérance du temps de retour

$$\mathbb{E}_\ell(T_\ell^+) = \frac{1}{\pi(\ell)} = 2^K \frac{\ell!(K-\ell)!}{K!}.$$

En utilisant la formule de Stirling $n! \simeq \sqrt{2\pi n}(n/e)^n$, on remarque que pour K très grand, par exemple de l’ordre du nombre d’Avogadro 10^{23}

$$\mathbb{E}_K(T_K^+) = 2^K \quad \text{et} \quad \mathbb{E}_{K/2}(T_{K/2}^+) \simeq \sqrt{\frac{\pi K}{2}}.$$

Par conséquent, le temps de retour en $K/2$ (qui est la valeur d’équilibre) sera infiniment moins long que le temps de retour en K . Ce dernier est tellement grand qu’il peut être supérieur à la durée de vie de l’univers. Il faudra donc s’armer de patience avant de voir un ballon percé se regonfler spontanément. Les temps de récurrence d’événements rares étant extrêmement longs, il n’y a donc pas de contradiction avec la validité de l’équation de Boltzmann sur des échelles de temps plus courtes.

Chapitre 4

Espaces d'états dénombrables

Ce chapitre est consacré aux chaînes de Markov sur des espaces d'états infinis (mais dénombrables) pour lesquelles des phénomènes nouveaux apparaissent concernant la fréquence des visites d'un état. On verra aussi que pour des espaces d'états infinis, l'irréductibilité ne suffit pas à garantir l'existence d'une unique mesure de probabilité invariante.

4.1 Chaînes de Markov récurrentes et transitoires

Un exemple simple de chaîne de Markov sur un espace d'états infini est la marche aléatoire symétrique sur $\mathbb{Z}^d$ avec $d \geq 1$ et de matrice de transition

$$\forall x, y \in \mathbb{Z}^d, \quad P(x, y) = \frac{1}{2d} \mathbf{1}_{\{\|x-y\|_2=1\}}.$$

La marche saute de façon équiprobable d'un site vers ses voisins. Elle est irréductible car elle peut rejoindre n'importe quel point de $\mathbb{Z}^d$. Dans le cas des espaces d'états finis, le lemme 3.8 assure que, pour toute chaîne de Markov irréductible, le temps d'atteinte d'un état y de E

$$T_y^+ = \inf \{n \geq 1; \quad X_n = y\}$$

en partant d'un état x est toujours intégrable $\mathbb{E}_x(T_y^+) < \infty$. En particulier, toute trajectoire de la chaîne de Markov issue de x finira par toucher presque sûrement n'importe quel état y de E . Pour des chaînes de Markov sur des espaces d'états infinis, cette propriété n'est plus vraie en général et il convient donc de distinguer plusieurs cas.

Définition 4.1. Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov sur un espace d'états E dénombrable. Un état x de E est dit

- transitoire si $\mathbb{P}_x(T_x^+ < \infty) < 1$.
- récurrent si $\mathbb{P}_x(T_x^+ < \infty) = 1$.

Les états récurrents peuvent être de deux types :

- Les états récurrents nuls si $\mathbb{E}_x(T_x^+) = \infty$.
- Les états récurrents positifs si $\mathbb{E}_x(T_x^+) < \infty$.

Nous allons montrer qu'une chaîne de Markov repasse presque sûrement une infinité de fois par un état récurrent alors qu'elle ne passera qu'un nombre fini de fois par un état transitoire. Si l'espace d'états est fini, le lemme 3.8 implique que tous les états sont récurrents positifs pour une chaîne de Markov irréductible.

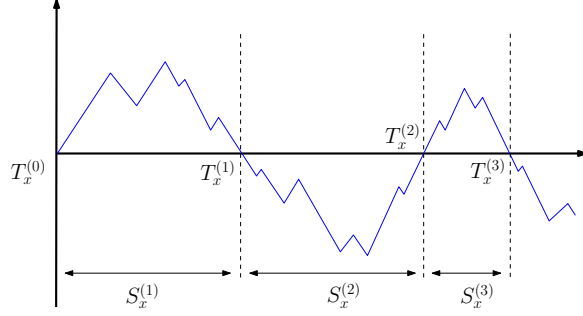


FIGURE 4.1 – Les excursions d'une marche aléatoire dans $\mathbb{Z}$ sont représentées en utilisant comme état de référence $x = 0$.

Pour tout entier k , on définit le $k^{\text{ième}}$ temps de retour en x par

$$T_x^{(k+1)} = \inf \{n \geq T_x^{(k)} + 1; \quad X_n = x\} \in \mathbb{N} \cup \{\infty\}$$

où $T_x^{(0)} = 0$ et $T_x^{(1)}$ coïncide avec T_x^+ . Si la chaîne ne passe que k fois en x alors tous les temps $T_x^{(\ell)}$ avec $\ell \geq k + 1$ seront infinis. La chaîne de Markov fait des *excursions* entre deux passages en x (cf. figure 4.1) dont les longueurs sont données par

$$S_x^{(k)} = \begin{cases} T_x^{(k)} - T_x^{(k-1)}, & \text{si } T_x^{(k-1)} < \infty \\ 0, & \text{sinon} \end{cases}$$

Le nombre de visites d'un état x par la chaîne de Markov $\{X_n\}_{n \geq 0}$ est donné par

$$\mathcal{N}_x = \sum_{n \geq 0} \mathbf{1}_{\{X_n = x\}}.$$

Théorème 4.2. *On distingue deux comportements différents :*

— Si un état x est récurrent, alors une chaîne de Markov issue de x repasse infiniment souvent en x

$$\mathbb{P}_x(\mathcal{N}_x = \infty) = 1.$$

— Si un état x est transitoire, le nombre de visites $\mathcal{N}_x$ d'une chaîne de Markov issue de x suit une loi géométrique sur $\mathbb{N}$ de paramètre $\mathbb{P}_x(T_x^+ = \infty) > 0$. En particulier

$$\mathbb{E}_x(\mathcal{N}_x) = \frac{1}{\mathbb{P}_x(T_x^+ = \infty)} \quad \Rightarrow \quad \mathbb{P}_x(\mathcal{N}_x < \infty) = 1.$$

Démonstration. Le point crucial de la preuve consiste à remarquer que pour $k \geq 2$, si on conditionne par l'évènement $\{T_x^{(k-1)} < \infty\}$, alors la longueur de l'excursion $S_x^{(k)}$ est indépendante de la trajectoire de la chaîne avant l'instant $T_x^{(k-1)}$, c'est-à-dire de $\{X_n; \quad n \leq T_x^{(k-1)}\}$ et

$$\forall \ell \in \mathbb{N}, \quad \mathbb{P}(S_x^{(k)} = \ell \mid T_x^{(k-1)} < \infty) = \mathbb{P}_x(T_x^{(1)} = \ell). \quad (4.1)$$

Intuitivement le résultat est évident : quand la chaîne de Markov repart de x , il n'y a plus besoin de connaître son passé pour déterminer son futur. En particulier, les excursions représentées figure 4.1 ont toutes la même loi et sont indépendantes. Pour le démontrer, il suffit d'appliquer la propriété de Markov forte établie au théorème 2.5. Au temps d'arrêt $T_x^{(k-1)}$ la chaîne de Markov est dans l'état x et la chaîne de Markov décalée en temps $\{X_{T_x^{(k-1)}+n}\}_{n \geq 0}$ a la même loi que la chaîne $\{X_n\}_{n \geq 0}$ partant de x . La longueur de l'excursion peut s'écrire comme un temps de retour

$$S_x^{(k)} = \inf \left\{ n \geq 1, \quad X_{T_x^{(k-1)}+n} = x \right\}.$$

Conditionnellement à l'évènement $\{T_x^{(k-1)} < \infty\}$, la longueur $S_x^{(k)}$ a la même loi que $T_x^{(1)}$. On en déduit donc (4.1) par la propriété de Markov forte.

Comme $X_0 = x$, on remarque que $\mathcal{N}_x \geq 1$. Pour tout $k \geq 1$, on obtient

$$\begin{aligned} \mathbb{P}_x(\mathcal{N}_x \geq k+1) &= \mathbb{P}_x(T_x^{(k)} < \infty) = \mathbb{P}_x(T_x^{(k-1)} < \infty \text{ et } S_x^{(k)} < \infty) \\ &= \mathbb{P}_x(T_x^{(k-1)} < \infty) \mathbb{P}_x(S_x^{(k)} < \infty \mid T_x^{(k-1)} < \infty). \end{aligned}$$

En utilisant (4.1) et en itérant, on conclut que

$$\mathbb{P}_x(\mathcal{N}_x \geq k+1) = \mathbb{P}_x(T_x^{(k-1)} < \infty) \mathbb{P}_x(T_x^{(1)} < \infty) = \mathbb{P}_x(T_x^+ < \infty)^k.$$

Dans la dernière égalité, on a identifié $T_x^{(1)} = T_x^+$. Si x est récurrent, alors $\mathbb{P}_x(\mathcal{N}_x \geq k) = 1$ pour tout k . Comme l'évènement $\{\mathcal{N}_x = \infty\}$ est la limite (décroissante) des évènements $\{\mathcal{N}_x \geq k\}$, on en déduit que $\mathbb{P}_x(\mathcal{N}_x = \infty) = 1$.

Un calcul similaire au précédent permet de montrer que pour $k \geq 1$

$$\mathbb{P}_x(\mathcal{N}_x = k) = \mathbb{P}_x(T_x^+ < \infty)^{k-1} \left(1 - \mathbb{P}_x(T_x^+ < \infty) \right).$$

Si x est transitoire alors $\mathbb{P}_x(T_x^+ < \infty) < 1$. Le nombre de visites $\mathcal{N}_x$ d'une chaîne de Markov issue de x suit une loi géométrique sur $\mathbb{N}$ de paramètre $\mathbb{P}_x(T_x^+ < \infty)$ et l'espérance du nombre de retours est finie.

□

La propriété de Chapman-Kolmogorov permet d'écrire pour tout $n \geq 1$

$$\forall x, y \in E, \quad P^n(x, y) = \mathbb{P}_x(X_n = y)$$

où P^n est la puissance $n^{\text{ième}}$ de la matrice P . Le théorème 4.2 peut être reformulé à l'aide d'un critère plus simple à utiliser.

Théorème 4.3. *Pour tout état x de E , il n'existe que deux possibilités :*

- *x est transitoire si et seulement si $\sum_{n \geq 0} P^n(x, x) < \infty$.*
- *x est récurrent si et seulement si $\sum_{n \geq 0} P^n(x, x) = \infty$.*

Si x communique avec y , i.e. $x \rightarrow y$, et x est un état récurrent alors y est un état récurrent. Par conséquent, si la chaîne est irréductible alors les états sont tous transitoires ou tous récurrents. De plus, dans ce dernier cas, on repasse infiniment souvent par n'importe quel point

$$\forall x, y \in E, \quad \mathbb{P}_y(T_x^+ < \infty) = 1 \quad \text{et} \quad \mathbb{P}_y(\mathcal{N}_x = \infty) = 1. \quad (4.2)$$

Démonstration. L'espérance du nombre de visites peut se réécrire en utilisant le théorème de Fubini (pour permuter l'espérance et la somme)

$$\mathbb{E}_x(\mathcal{N}_x) = \mathbb{E}_x \left(\sum_{n \geq 0} \mathbf{1}_{X_n=x} \right) = \sum_{n \geq 0} \mathbb{E}_x(\mathbf{1}_{X_n=x}) = \sum_{n \geq 0} P^n(x, x).$$

Le théorème 4.2 suffit donc à prouver l'alternative entre les deux possibilités.

Soit x un état récurrent communiquant avec y . Il existe donc un entier ℓ tel que $P^\ell(x, y) > 0$. Supposons que y ne communique pas avec x , c'est-à-dire que $P^k(y, x) = 0$ pour tout k alors

$$\mathbb{P}_y(T_x^+ < \infty) \leq \sum_{k \geq 1} \mathbb{P}_y(X_k = x) = \sum_{k \geq 1} P^k(y, x) = 0.$$

Dans ce cas x ne pourrait pas être récurrent car la chaîne a une probabilité non nulle de passer dans l'état y et ainsi de ne plus revenir en x

$$\mathbb{P}_x(T_x^+ = \infty) \geq P^\ell(x, y) > 0.$$

Par conséquent si $x \rightarrow y$ et x est récurrent alors $y \rightarrow x$. Il existe donc deux entiers $\ell, k \geq 1$ tels que $P^\ell(x, y) > 0$ et $P^k(y, x) > 0$. On peut décomposer les trajectoires partant de y

$$\sum_{n \geq 0} P^n(y, y) \geq \sum_{n \geq 0} P^k(y, x) P^n(x, x) P^\ell(x, y) = c \sum_{n \geq 0} P^n(x, x) = \infty$$

où $c > 0$ est une constante. La dernière égalité provient du critère de récurrence appliqué à x . On en déduit que y doit aussi être récurrent.

Si la chaîne de Markov est irréductible, tous les états communiquent et il suffit que l'un soit récurrent pour que les autres le soient.

Supposons que la chaîne soit irréductible et que tous les états soient récurrents. Cette hypothèse assure, en particulier, que $\mathbb{P}_x(T_x^+ = \infty) = 0$ pour tout x . Nous allons montrer la première assertion de (4.2). Étant donné y dans E , l'irréductibilité implique l'existence d'un chemin, de probabilité positive, de longueur n_0 allant de x à y et ne repassant pas par x . Ceci permet d'affirmer que

$$\mathbb{P}_x(\{X_{n_0} = y\} \cap \{T_x^+ > n_0\}) > 0.$$

On en déduit, en utilisant la propriété de Markov au temps n_0 , que

$$\mathbb{P}_x(T_x^+ = \infty) \geq \mathbb{P}_x(\{X_{n_0} = y\} \cap \{T_x^+ = \infty\}) = \mathbb{P}_x(\{X_{n_0} = y\} \cap \{T_x^+ > n_0\}) \mathbb{P}_y(T_x^+ = \infty),$$

ce qui implique le résultat souhaité $\mathbb{P}_y(T_x^+ = \infty) = 0$.

La seconde assertion de (4.2), s'obtient en remarquant que

$$\mathbb{P}_y(\mathcal{N}_x = \infty) = \mathbb{P}_y(\{\mathcal{N}_x = \infty\} \cap \{T_x^+ < \infty\}) = \mathbb{P}_y(T_x^+ < \infty) \mathbb{P}_x(\mathcal{N}_x = \infty),$$

où la dernière égalité est une conséquence de la propriété de Markov forte appliquée au temps T_x^+ . Comme x est récurrent $\mathbb{P}_x(\mathcal{N}_x = \infty) = 1$ et $\mathbb{P}_y(T_x^+ < \infty) = 1$, on en déduit que $\mathbb{P}_y(\mathcal{N}_x = \infty) = 1$. $\square$

4.2 Application : marches aléatoires

Les états récurrents et transitoires peuvent être illustrés dans le cadre des marches aléatoires sur $\mathbb{Z}^d$.

4.2.1 Marches aléatoires symétriques sur $\mathbb{Z}^d$

Le comportement de la marche aléatoire symétrique sur $\mathbb{Z}^d$ de matrice de transition

$$\forall x, y \in \mathbb{Z}^d, \quad P(x, y) = \frac{1}{2d} \mathbf{1}_{\{\|x-y\|_2=1\}}$$

dépend de la dimension d . Le théorème suivant a été prouvé par Polya en 1921.

Théorème 4.4. *La marche aléatoire symétrique sur $\mathbb{Z}$ ou $\mathbb{Z}^2$ est récurrente. Pour $d \geq 3$, la marche symétrique sur $\mathbb{Z}^d$ est transitoire.*

Démonstration. La chaîne étant irréductible tous les états sont de la même nature. D'après le théorème 4.3, il suffit donc de déterminer si la série $\sum_{n \geq 0} P^n(0, 0)$ est divergente ou convergente. L'étude se fait pour chaque dimension.

d = 1.

Une marche aléatoire ne peut revenir en 0 qu'après un nombre pair de pas. Pour revenir en 0 au temps $2n$, il faut qu'il y ait eu exactement n accroissements égaux à 1 et n accroissements égaux à -1 . On a donc

$$P^{2n}(0, 0) = \frac{1}{2^{2n}} \binom{2n}{n} \simeq \frac{1}{\sqrt{\pi n}}, \quad P^{2n+1}(0, 0) = 0$$

où l'asymptotique a été obtenue en utilisant la formule de Stirling $n! \simeq \sqrt{2\pi n}(n/e)^n$. Par conséquent, la série $\sum_{n \geq 0} P^n(0, 0)$ est divergente et 0 est un état récurrent.

d = 2.

En inclinant la tête de 45 degrés (cf. figure 4.2), on voit qu'une marche aléatoire X_n sur $\mathbb{Z}^2$ se réécrit $X_n = (\frac{X_n^+ + X_n^-}{2}, \frac{X_n^+ - X_n^-}{2})$ en fonction de X_n^+, X_n^- , deux marches aléatoires indépendantes sur $\mathbb{Z}$ partant initialement de 0.

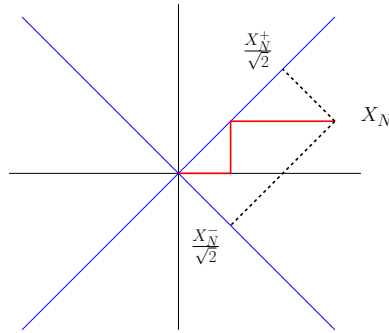


FIGURE 4.2 – Les marches aléatoires $\frac{X_n^+}{\sqrt{2}}$ et $\frac{X_n^-}{\sqrt{2}}$ sont les projections de X_n sur les axes du réseau à 45 degrés.

Par l'indépendance des marches X_n^+ et X_n^- , on déduit du cas unidimensionnel que

$$P^{2n}(0,0) = \mathbb{P}_0(X_{2n} = 0) = \mathbb{P}_0(X_{2n}^+ = 0)\mathbb{P}_0(X_{2n}^- = 0) \simeq \frac{1}{\pi n}$$

où 0 représente (par abus de notation) l'origine de $\mathbb{Z}$ et de $\mathbb{Z}^2$. La série $\sum_{n \geq 0} P^n(0,0)$ de terme principal $1/n$ est divergente et 0 est un état récurrent.



FIGURE 4.3 – Deux réalisations de la marche aléatoire dans $\mathbb{Z}^2$ pour 10^4 et 10^5 pas.

d = 3.

Pour calculer $P^{2n}(0,0)$, on décompose les trajectoires de longueur $2n$ en $2i, 2j, 2k$ sauts selon chacun des axes

$$P^{2n}(0,0) = \sum_{\substack{i,j,k \geq 0 \\ i+j+k=n}} \frac{(2n)!}{(i!j!k!)^2} \left(\frac{1}{6}\right)^{2n} = \binom{2n}{n} \left(\frac{1}{2}\right)^{2n} \sum_{\substack{i,j,k \geq 0 \\ i+j+k=n}} \binom{n}{i,j,k}^2 \left(\frac{1}{3}\right)^{2n}$$

où $\binom{n}{i,j,k} = \frac{n!}{i!j!k!}$ est le nombre de façons de ranger n boules dans 3 boîtes en mettant i boules dans la première, j dans la seconde et k dans la troisième. On rappelle

$$\sum_{\substack{i,j,k \geq 0 \\ i+j+k=n}} \binom{n}{i,j,k} \left(\frac{1}{3}\right)^n = 1 \quad \text{et} \quad \binom{3n}{i,j,k} \leq \binom{3n}{n,n,n}.$$

Commençons par considérer le cas $n = 3\ell$. Par les identités précédentes, on a

$$\begin{aligned} P^{2n}(0,0) &\leq \binom{2n}{n} \left(\frac{1}{2}\right)^{2n} \binom{3\ell}{\ell,\ell,\ell} \left(\frac{1}{3}\right)^n \sum_{\substack{i,j,k \geq 0 \\ i+j+k=n}} \binom{n}{i,j,k} \left(\frac{1}{3}\right)^n \\ &= \binom{2n}{n} \left(\frac{1}{12}\right)^n \binom{n}{\ell,\ell,\ell} \simeq \frac{1}{2(2\pi)^{3/2}} \left(\frac{6}{n}\right)^{3/2} \end{aligned}$$

où l'expression asymptotique pour n grand est une conséquence de la formule de Stirling. Les autres valeurs de n peuvent être bornées par comparaison

$$P^{2(3\ell)}(0,0) \geq \left(\frac{1}{6}\right)^2 P^{2(3\ell-1)}(0,0) \quad \text{et} \quad P^{2(3\ell)}(0,0) \geq \left(\frac{1}{6}\right)^4 P^{2(3\ell-2)}(0,0).$$

La série $\sum_n P^{2n}(0,0)$ converge car un majorant de son terme principal tend vers 0 comme $\frac{1}{n^{3/2}}$. La marche aléatoire en dimension 3 est donc transitoire.

Une approche combinatoire similaire pour les dimensions $d \geq 4$ montre que $P^{2n}(0,0)$ est asymptotiquement de l'ordre de $\frac{1}{n^{d/2}}$. Par conséquent la marche aléatoire symétrique est transitoire dès que $d \geq 3$. $\square$

On remarquera que la mesure $\pi(x) = 1$ pour tout x de $\mathbb{Z}^d$ est une mesure invariante pour la marche aléatoire. Par contre, il n'existe pas de mesure de probabilité invariante, i.e. de mesure normalisée par 1. Ce comportement spécifique des espaces d'états infinis sera expliqué section 4.3.

4.2.2 Un critère analytique

Il est parfois plus facile d'étudier la série génératrice des temps de retours d'une chaîne de Markov. Pour x, y dans E et s dans $[0, 1]$, on pose

$$U(x, y, s) = \mathbb{E}_x(s^{T_y^+} 1_{\{T_y^+ < \infty\}}) = \sum_{n \geq 1} s^n \mathbb{P}_x(T_y^+ = n).$$

On remarque que $U(x, x, 1) = \mathbb{P}_x(T_x^+ < \infty)$.

Nous allons illustrer l'utilisation de U en étudiant la marche aléatoire dans $\mathbb{Z}$ de probabilité de transition

$$P(x, x+1) = p, \quad P(x, x-1) = q$$

où $p + q = 1$.

Théorème 4.5.

Si $p \neq 1/2$, la marche aléatoire est transitoire car $\mathbb{P}_0(T_0^+ < \infty) = 1 - |1 - 2p| < 1$.

Si $p = 1/2$, la marche aléatoire est récurrente nulle car $\mathbb{E}_0(T_0^+) = \infty$.

Démonstration. En utilisant la propriété de Markov après un pas de temps, on obtient

$$U(1, 0, s) = s(pU(2, 0, s) + q) \quad \text{et} \quad U(-1, 0, s) = s(p + qU(-2, 0, s)). \quad (4.3)$$

Pour aller de 2 à 0, la marche doit d'abord passer par 1. Le temps nécessaire pour atteindre 0 se décompose donc sous la forme $T_0^+ = T_1^+ + T_{1 \rightarrow 0}$ où $T_{1 \rightarrow 0}$ est le premier temps d'atteinte de 0 après avoir touché 1. La propriété de Markov forte appliquée après le temps d'arrêt T_1^+ permet d'écrire

$$\begin{aligned} U(2, 0, s) &= \mathbb{E}_2(s^{T_1^+} 1_{\{T_1^+ < \infty\}} s^{T_{1 \rightarrow 0}} 1_{\{T_{1 \rightarrow 0} < \infty\}}) = \mathbb{E}_1(s^{T_{1 \rightarrow 0}} 1_{\{T_{1 \rightarrow 0} < \infty\}}) U(2, 1, s) \\ &= U(1, 0, s) U(2, 1, s) = U(1, 0, s)^2. \end{aligned}$$

Par symétrie $U(-2, 0, s) = U(-1, 0, s)^2$. Ces relations permettent de réécrire les équations (4.3) et d'obtenir, en remarquant que $U(1, 0, 0) = U(-1, 0, 0) = 0$, une solution explicite

$$\begin{aligned} U(1, 0, s) &= s(pU(1, 0, s)^2 + q) \quad \Rightarrow \quad U(1, 0, s) = \frac{1 - \sqrt{1 - 4pqs^2}}{2ps} \\ U(-1, 0, s) &= s(p + qU(1, 0, s)^2) \quad \Rightarrow \quad U(-1, 0, s) = \frac{1 - \sqrt{1 - 4pqs^2}}{2qs}, \end{aligned}$$

En appliquant une nouvelle fois la propriété de Markov, on en déduit

$$U(0,0,s) = s(pU(1,0,s) + qU(-1,0,s)) = 1 - \sqrt{1 - 4pqs^2}.$$

Si $p \neq q$, $U(0,0,s)$ admet une limite quand s tend vers 1

$$\mathbb{P}_0(T_0^+ < \infty) = U(0,0,1) = 1 - \sqrt{1 - 4pq} = 1 - |1 - 2p|.$$

Si $p = q = 1/2$, la limite vaut 1 et la marche aléatoire sur $\mathbb{Z}$ est bien récurrente, comme nous l'avons déjà montré par le théorème de Polya 4.4. Pour tout $s < 1$, on peut dériver U

$$\partial_s U(0,0,s) = \mathbb{E}_0(T_0^+ s^{T_0^+ - 1} 1_{T_0^+ < \infty}) = \frac{s}{\sqrt{1 - s^2}}.$$

Comme la limite diverge quand s tend vers 1, on en déduit que $\mathbb{E}_0(T_0^+ 1_{T_0^+ < \infty}) = \infty$. La marche est donc récurrente nulle : elle revient infiniment souvent en 0 mais l'espérance du temps de retour est infinie. $\square$

4.3 Mesures invariantes

Le théorème 3.9 sur les mesures invariantes dans le cas des espaces d'états finis se généralise aux chaînes récurrentes positives sur des espaces d'états dénombrables.

Théorème 4.6. *Pour toute chaîne de Markov irréductible sur un espace d'états E dénombrable, les deux assertions suivantes sont équivalentes :*

- (i) *La chaîne est récurrente positive.*
- (ii) *Il existe une mesure de probabilité invariante.*

De plus s'il existe une mesure de probabilité invariante alors elle est unique et est donnée par

$$\forall x \in E, \quad \pi(x) = \frac{1}{\mathbb{E}_x(T_x^+)} > 0.$$

Pour une chaîne de Markov irréductible et récurrente nulle, on peut montrer qu'il existe toujours une mesure invariante (voir par exemple [13]), cependant celle-ci ne peut pas être normalisée en une mesure de probabilité. Ainsi, la mesure $\pi(x) = 1$ est invariante pour la marche aléatoire dans $\mathbb{Z}^d$ (que la chaîne soit récurrente nulle ou transitoire). Certaines chaînes de Markov irréductibles et transitoires peuvent ne pas admettre de mesures invariantes.

*Démonstration ** L'implication (i) donne (ii) est une conséquence directe de la preuve du théorème 3.9 dans le cas des espaces d'états finis.

Considérons maintenant l'implication inverse et supposons que (ii) soit vérifiée. La preuve se décompose en trois étapes.

Étape 1. Montrons que tous les états sont récurrents.

La chaîne étant irréductible, les états sont tous transitoires ou tous récurrents. Supposons qu'ils soient tous transitoires alors pour tous x et y de E

$$\lim_{n \rightarrow \infty} P^n(x, y) = 0$$

car le nombre de visites en y est fini

$$\mathbb{E}_x(\mathcal{N}_y) = \sum_{n \geq 0} P^n(x, y) < \infty.$$

S'il existe une mesure de probabilité invariante π , elle vérifie pour tout temps n les relations

$$\forall y \in E, \quad \pi(y) = \sum_{x \in E} \pi(x) P^n(x, y).$$

Comme $\sum_{x \in E} \pi(x) = 1$, on en déduit (par exemple en utilisant le théorème de convergence dominée) que

$$\forall y \in E, \quad \pi(y) = \lim_{n \rightarrow \infty} \sum_{x \in E} \pi(x) P^n(x, y) = \sum_{x \in E} \pi(x) \lim_{n \rightarrow \infty} P^n(x, y) = 0.$$

Ce qui contredit l'existence de la mesure π . Par conséquent, la chaîne de Markov doit être récurrente.

Étape 2. Montrons que la chaîne est récurrente positive.

Soit x un état de référence fixé. Comme dans le théorème 3.9 pour les espaces d'états finis, nous allons montrer que la mesure définie par les excursions issues de x

$$\forall y \in E, \quad \tilde{\pi}(y) = \sum_{n \geq 0} \mathbb{P}_x(X_n = y, T_x^+ > n) \quad (4.4)$$

est invariante. Nous n'avons pas encore établi que la chaîne est récurrente positive, par conséquent il faut vérifier que la mesure $\tilde{\pi}$ est bien définie. Par l'hypothèse (ii), il existe π une mesure invariante qui vérifie

$$\pi(y) = \sum_{z_1 \in E} \pi(z_1) P(z_1, y) = \pi(x) P(x, y) + \sum_{z_1 \neq x} \pi(z_1) P(z_1, y).$$

En appliquant la propriété d'invariance aux états $z_1 \neq x$

$$\begin{aligned} \pi(y) &= \pi(x) P(x, y) + \sum_{z_1 \neq x} \sum_{z_2 \in E} \pi(z_2) P(z_2, z_1) P(z_1, y) \\ &= \pi(x) P(x, y) + \sum_{z_1 \neq x} \pi(x) P(x, z_1) P(z_1, y) + \sum_{z_1 \neq x} \sum_{z_2 \neq x} \pi(z_2) P(z_2, z_1) P(z_1, y). \end{aligned}$$

On itère ℓ fois cette procédure

$$\begin{aligned} \pi(y) &= \pi(x) \left[P(x, y) + \sum_{k=1}^{\ell-1} \sum_{\substack{z_1 \neq x, z_2 \neq x, \\ \dots, z_k \neq x}} P(x, z_k) \dots P(z_1, y) \right] \\ &\quad + \sum_{\substack{z_1 \neq x, z_2 \neq x, \\ \dots, z_\ell \neq x}} \pi(z_\ell) P(z_\ell, z_{\ell-1}) \dots P(z_1, y) \\ &\geq \pi(x) \sum_{n=1}^{\ell-1} \mathbb{P}_x(X_n = y, T_x^+ > n) \end{aligned}$$

où la dernière inégalité a été obtenue en identifiant le terme entre crochets par l'espérance du nombre de passages en y avant de revenir en x et en négligeant le second terme qui est positif. Quand ℓ tend vers l'infini, ceci implique que pour tout y de E

$$\pi(y) \geq \pi(x) \tilde{\pi}(y). \quad (4.5)$$

La chaîne étant irréductible, la mesure π est strictement positive pour tout x (cf. théorème 3.6). La mesure $\tilde{\pi}$ est donc bien définie et

$$\mathbb{E}_x(T_x^+) = \sum_{y \in E} \tilde{\pi}(y) \leq \frac{1}{\pi(x)} \sum_{y \in E} \pi(y) < \infty.$$

L'état x est donc récurrent positif. En répétant la preuve pour d'autres états de référence, on en déduit que tous les états de E sont récurrents positifs, i.e. que la chaîne de Markov est récurrente positive.

Étape 3. Représentation de la mesure invariante.

Un calcul identique à celui du théorème 3.9 montre que $\tilde{\pi}$ (définie en (4.4)) est une mesure invariante. On rappelle que x est l'état de référence pour construire $\tilde{\pi}$ et que $\tilde{\pi}(x) = 1$. Supposons que $\pi(x) = 1$ (quitte à multiplier tous les coefficients par le facteur $1/\pi(x)$) alors l'argument utilisé au théorème 3.6 permet de conclure à l'unicité $\pi = \tilde{\pi}$. En effet la mesure $\nu = \{\pi(y) - \tilde{\pi}(y)\}_{y \in E}$ a tous ses termes positifs ou nuls d'après l'inégalité (4.5). De plus ν est invariante car π et $\tilde{\pi}$ le sont. Comme ν atteint son minimum en x car $\nu(x) = \pi(x) - \tilde{\pi}(x) = 0$, il suffit de suivre la preuve du théorème 3.6 pour montrer $\pi = \tilde{\pi}$. On a donc identifié l'unique mesure invariante telle que $\pi(x) = 1$. Pour obtenir une mesure de probabilité, il suffit maintenant de la normaliser et on retrouve

$$\forall x \in E, \quad \pi(x) = \frac{1}{\mathbb{E}_x(T_x^+)}.$$

□

Corollaire 4.7. *Les états d'une chaîne de Markov irréductible et récurrente sont tous récurrents positifs ou tous récurrents nuls.*

Démonstration. S'il existe un état récurrent positif, on peut construire une mesure de probabilité invariante (4.4) et on en déduit par le théorème 4.6 que tous les états sont récurrents positifs. □

Le processus de naissance et de mort est un exemple classique de chaîne de Markov à valeurs dans $\mathbb{N}$. Au temps n , on note X_n le nombre d'individus dans une population ou de clients dans une file d'attente et on suppose que ce processus évolue comme une chaîne de Markov de matrice de transition

$$P(x, x+1) = p, \quad P(0, 0) = 1 - p \quad \text{et} \quad P(x, x-1) = 1 - p \quad \text{si } x \geq 1.$$

Le processus de naissance et de mort s'interprète comme une marche aléatoire sur $\mathbb{N}$ réfléchie quand elle touche 0. Si $p < 1/2$ la chaîne de Markov aura tendance à revenir vers 0. Quel que soit l'état initial, on s'attend donc à ce que la chaîne atteigne un régime stationnaire décrit par une mesure invariante localisée autour de 0. Si $p > 1/2$ la chaîne de Markov va croître en moyenne et elle va diverger vers l'infini.

Théorème 4.8. *On distingue deux comportements :*

- Si $p \leq 1/2$, la chaîne est récurrente.
- Si $p > 1/2$, la chaîne est transitoire.

Les mesures invariantes de cette chaîne de Markov sont de la forme

$$\forall x \geq 1, \quad \pi(x) = \pi(0) \left(\frac{p}{1-p} \right)^x$$

et elle est récurrente positive si et seulement si $p < 1/2$.

Démonstration. Tant que la chaîne n'a pas touché 0, elle se comporte comme une marche aléatoire sur $\mathbb{Z}$ avec probabilités de transition $(p, 1-p)$. On peut donc reprendre les notations et l'argument de la preuve du théorème 4.5 pour obtenir

$$U(1, 0, s) = \mathbb{E}_1(s^{T_0^+} \mathbf{1}_{\{T_0^+ < \infty\}}) = \frac{1 - \sqrt{1 - 4p(1-p)s^2}}{2ps}.$$

Au point 0, les nouvelles probabilités de transition s'appliquent

$$U(0, 0, s) = s((1-p) + pU(1, 0, s)) = s(1-p) + \frac{1 - \sqrt{1 - 4p(1-p)s^2}}{2}.$$

Finalement, on obtient

$$\mathbb{P}_0(T_0^+ < \infty) = U(0, 0, 1) = (1-p) + \frac{1 - \sqrt{1 - 4p(1-p)}}{2} = (1-p) + \frac{1}{2} - \left| \frac{1}{2} - p \right|.$$

Ceci conclut la première partie du théorème

$$\mathbb{P}_0(T_0^+ < \infty) = \begin{cases} 1, & \text{si } p \leq 1/2 \\ 2-2p, & \text{si } p > 1/2 \end{cases}$$

Une mesure invariante π satisfait les relations

$$\forall x \geq 1, \quad \pi(x) = p\pi(x-1) + (1-p)\pi(x+1) \quad \text{et} \quad \pi(0) = (1-p)\pi(0) + (1-p)\pi(1)$$

qui se réécrivent $\pi(1) = \left(\frac{p}{1-p} \right) \pi(0)$ et

$$\forall i \geq 1, \quad \pi(x+1) - \pi(x) = \left(\frac{p}{1-p} \right) (\pi(x) - \pi(x-1)) = \left(\frac{p}{1-p} \right)^x (\pi(1) - \pi(0)).$$

La famille des mesures invariantes est donc indexée par un paramètre $\pi(0)$

$$\forall x \geq 1, \quad \pi(x) = \pi(0) \left(\frac{p}{1-p} \right)^x.$$

Ces mesures ne peuvent être normalisées que pour $p < 1/2$ et dans ce cas la chaîne de Markov est récurrente positive. On remarque que la chaîne de Markov est réversible pour ces mesures invariantes

$$P(x, x+1)\pi(x) = P(x+1, x)\pi(x+1).$$

□

Exercice 4.9. Montrer qu'un processus de naissance et de mort de matrice de transition générale

$$P(x, x+1) = p_x > 0, \quad P(0, 0) = 1 - p_0 \quad \text{et} \quad P(x, x-1) = 1 - p_x \quad \text{si} \quad x \geq 1$$

admet une mesure de probabilité invariante si et seulement si

$$\sum_{n \geq 1} \prod_{x=1}^n \frac{p_{x-1}}{1 - p_x} < \infty.$$

4.4 Application : processus de branchement et graphes aléatoires

Les processus de branchement ont de nombreuses applications allant de la démographie aux arbres phylogénétiques en passant par la fission nucléaire. Nous allons décrire un exemple de processus de branchement, les arbres aléatoires de Galton-Watson, et montrer comment ces arbres permettent d'étudier des graphes aléatoires.

4.4.1 Arbres aléatoires de Galton-Watson

Les processus de branchement que nous allons considérer modélisent le nombre d'individus au cours du temps d'une population en fonction de règles de reproduction. Le modèle a été proposé par Sir Francis Galton en 1873 pour décrire l'évolution des noms de famille en Angleterre. À l'époque les noms de famille étant transmis exclusivement par les hommes, il suffisait de suivre le nombre de descendants masculins dans chaque famille. Cette hypothèse permet de considérer un seul type d'individus et de supposer qu'à chaque génération les individus se reproduisent selon la même loi de probabilité. On s'intéressera à la taille de la population à chaque génération. Ce modèle peut aussi décrire la fission des neutrons dans une réaction nucléaire. Si cette fission s'opère trop rapidement cela peut conduire à une explosion. La mutation de gènes dans une population peut être modélisée par ces processus de branchement. D'autres applications des processus de branchement en écologie et dans les modèles d'évolution sont détaillées dans le cours *Modèles aléatoires en écologie et évolution* [21].

Les *arbres aléatoires de Galton-Watson* sont des processus de branchement définis par récurrence à l'aide d'une suite $\{\zeta_i^t\}_{\substack{i \geq 1, \\ t \geq 1}}$ (à 2 indices) de variables aléatoires indépendantes et identiquement distribuées sur les entiers selon la loi

$$\forall k \geq 0, \quad \mathbb{P}(\zeta_i^t = k) = p_k. \quad (4.6)$$

L'arbre de Galton-Watson décrit la croissance d'une population dont le nombre d'individus au temps t sera noté par Z_t . Initialement pour $t = 0$, on suppose que cette population n'est formée que par un seul individu et on pose $Z_0 = 1$. Au temps $t = 1$, cet individu a $Z_1 = \zeta_1^1$ enfants. Chaque enfant a ensuite un nombre aléatoire de descendants selon la loi de reproduction (4.6). Par récurrence, le nombre Z_{t+1} d'individus dans la population au temps $t + 1$ évolue en fonction des Z_t individus présents au temps t :

$$Z_{t+1} = \begin{cases} \zeta_1^{t+1} + \dots + \zeta_{Z_t}^{t+1}, & \text{si } Z_t > 0 \\ 0, & \text{si } Z_t = 0 \end{cases} \quad (4.7)$$

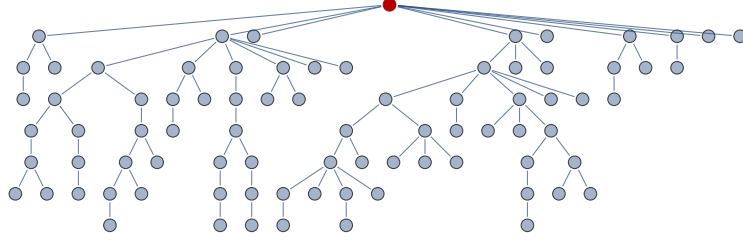


FIGURE 4.4 – Un exemple d'arbre aléatoire après 7 générations.

S'il n'y a plus de descendants à partir d'un temps t , i.e. $Z_t = 0$, alors la population restera éteinte à jamais.

La connaissance de la population au temps $t + 1$ ne dépend que de Z_t et des variables $\{\zeta_i^{t+1}\}_{i \geq 1}$ qui sont indépendantes du passé. Le processus $\{Z_t\}_{t \geq 0}$ est donc une chaîne de Markov. Pour calculer sa matrice de transition, on suppose qu'il existe n individus au temps initial

$$\begin{aligned} P(n, k) &= \mathbb{P}(Z_1 = k \mid Z_0 = n) = \mathbb{P}\left(\sum_{\ell=1}^{Z_0} \zeta_\ell^1 = k \mid Z_0 = n\right) \\ &= \mathbb{P}\left(\sum_{\ell=1}^n \zeta_\ell^1 = k\right) = \sum_{\substack{i_1, \dots, i_n \in \mathbb{N} \\ i_1 + \dots + i_n = k}} p_{i_1} p_{i_2} \dots p_{i_n}. \end{aligned}$$

On exclut les lois de reproduction pathologiques (4.6) telles que

- $p_1 = 1$: la taille de population reste constante
- $p_0 + p_1 = 1$ et $p_0 > 0$: il y a au maximum un seul individu qui finit forcément par mourir.
- $p_0 = 0$ et $p_1 < 1$: la population ne fait que croître.

Dans la suite, nous étudierons donc le comportement des arbres de Galton-Watson quand t tend vers l'infini pour des lois de reproduction satisfaisant

$$0 < p_0 \quad \text{et} \quad p_0 + p_1 < 1. \quad (4.8)$$

Si la population disparaît au temps t alors $Z_s = 0$ pour tous les temps suivants $s \geq t$. L'état 0 est donc absorbant pour cette chaîne de Markov. On remarque qu'aucun autre état ne peut être récurrent car tous les états communiquent avec 0. En effet, la population peut disparaître en un seul pas de temps

$$\forall n \geq 1, \quad \mathbb{P}(Z_1 = 0 \mid Z_0 = n) = p_0^n > 0.$$

Par conséquent, il n'existe que deux comportements possibles : la population disparaît ou sa taille tend vers l'infini.

Le comportement asymptotique du processus de branchement est déterminé par le nombre moyen d'enfants par individu $\mu = \mathbb{E}(\zeta_1^1)$. Pour s'en convaincre, il suffit d'analyser $\mathbb{E}(Z_t)$ la taille moyenne de la population au temps t . En utilisant la propriété de

Markov et en conditionnant par la génération précédente, on a

$$\begin{aligned}\mathbb{E}(Z_{t+1}) &= \sum_{n=0}^{\infty} \mathbb{E}(Z_{t+1} | Z_t = n) \mathbb{P}(Z_t = n) = \sum_{n=0}^{\infty} \mathbb{E}(\zeta_1^{t+1} + \cdots + \zeta_n^{t+1}) \mathbb{P}(Z_t = n) \\ &= \mathbb{E}(\zeta_1^{t+1}) \sum_{n=0}^{\infty} n \mathbb{P}(Z_t = n) = \mu \mathbb{E}(Z_t) = \mu^{t+1} \mathbb{E}(Z_0) = \mu^{t+1}\end{aligned}\quad (4.9)$$

où on a utilisé que la loi de reproduction est identique pour tous les individus. On distingue donc trois régimes

- *Le régime sous-critique* $\mu < 1$: la taille moyenne de la population tend vers 0 exponentiellement vite et la population va disparaître presque sûrement. Pour le démontrer, il suffit de remarquer que

$$\mathbb{P}(Z_t \geq 1) \leq \mathbb{E}(Z_t) \leq \mu^t$$

et d'utiliser le théorème de Borel-Cantelli.

- *Le régime sur-critique* $\mu > 1$: la taille moyenne de la population tend vers l'infini exponentiellement vite et nous allons montrer que la taille de la population diverge avec une probabilité positive.
- *Le régime critique* $\mu = 1$: la taille moyenne de la population reste constante, mais la population va s'éteindre presque sûrement.

Pour préciser ces comportements, nous allons étudier le premier temps d'extinction, i.e. le temps d'atteinte de 0 par la chaîne de Markov

$$T_0 = \inf \{t \geq 1; \quad Z_t = 0\}.$$

On définit la fonction génératrice de la loi de reproduction

$$\forall u \in [0, 1], \quad \varphi(u) = \mathbb{E}(u^{\zeta_1^t}) = \sum_{n \geq 0} u^n p_n.$$

Le théorème suivant confirme l'heuristique établie par le calcul de la taille moyenne de la population

Théorème 4.10. *Si $\mu \leq 1$, la population s'éteint presque sûrement*

$$\mathbb{P}(T_0 < \infty) = 1.$$

Si $\mu > 1$, la population s'éteint avec probabilité $\rho \in]0, 1[$

$$\mathbb{P}(T_0 < \infty) = \rho$$

où ρ est l'unique point fixe dans $]0, 1[$ de $\varphi(\rho) = \rho$. Par conséquent la taille de la population diverge avec probabilité $1 - \rho > 0$.

Démonstration. Commençons par calculer la fonction génératrice de Z_t

$$\forall u \in [0, 1], \quad \Phi_t(u) = \mathbb{E}(u^{Z_t}) = \sum_{n \geq 0} u^n \mathbb{P}(Z_t = n).$$

La propriété de Markov et l'indépendance des variables $\{\zeta_i^{t+1}\}_{i \geq 0}$ permettent d'écrire

$$\begin{aligned}\Phi_{t+1}(u) &= \sum_{n=0}^{\infty} \mathbb{E}(u^{Z_{t+1}} | Z_t = n) \mathbb{P}(Z_t = n) = \sum_{n=0}^{\infty} \mathbb{E}(u^{\zeta_1^{t+1} + \dots + \zeta_n^{t+1}}) \mathbb{P}(Z_t = n) \\ &= \sum_{n=0}^{\infty} \prod_{\ell=1}^n \mathbb{E}(u^{\zeta_{\ell}^{t+1}}) \mathbb{P}(Z_t = n).\end{aligned}$$

avec la convention $\prod_{\ell=1}^0 \mathbb{E}(u^{\zeta_{\ell}^{t+1}}) = 1$ pour $n = 0$. En identifiant la fonction génératrice de la loi de reproduction, on obtient la relation de récurrence

$$\Phi_{t+1}(u) = \sum_{n=0}^{\infty} \varphi(u)^n \mathbb{P}(Z_t = n) = \Phi_t(\varphi(u)).$$

On en déduit

$$\Phi_{t+1}(u) = \underbrace{\varphi \circ \varphi \circ \dots \circ \varphi}_{t+1 \text{ fois}}(u).$$

Ceci peut aussi s'écrire

$$\Phi_{t+1}(u) = \varphi(\Phi_t(u)).$$

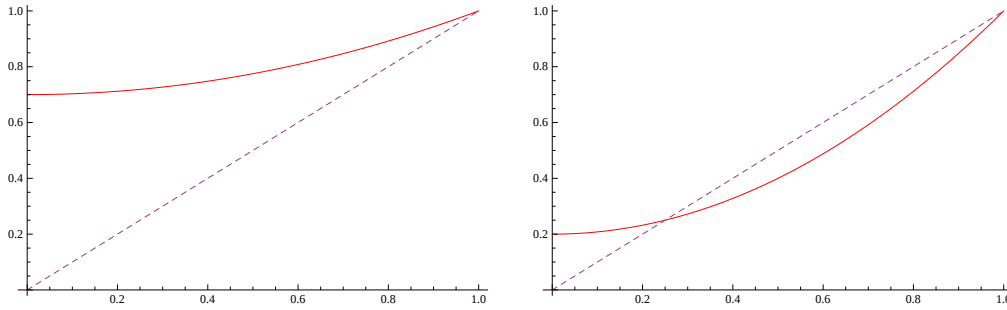


FIGURE 4.5 – Graphes de la fonction génératrice φ et de $x \rightarrow x$. Dans le cas sous-critique $\mu < 1$, représenté à gauche, l'unique point fixe $\varphi(x) = x$ est $x = 1$. Dans le cas sur-critique $\mu > 1$, représenté à droite, il existe un point fixe $\rho < 1$.

On note $x_t = \mathbb{P}(Z_t = 0)$ la probabilité que la population ait disparu au temps t . On remarque que $x_t = \Phi_t(0)$, par conséquent cette suite satisfait la récurrence

$$x_{t+1} = \varphi(x_t) \quad \text{avec} \quad x_0 = 0.$$

L'asymptotique de x_t quand t tend vers l'infini dépend des points fixes de $\varphi(x) = x$. On vérifie que pour $u \in [0, 1[$

$$\varphi(1) = 1, \quad \varphi'(u) = \sum_{n \geq 1} n u^{n-1} p_n \quad \text{et} \quad \varphi''(u) = \sum_{n \geq 2} n(n-1) u^{n-2} p_n. \quad (4.10)$$

Par l'hypothèse (4.8), on sait que $p_0 + p_1 < 1$ et donc $\varphi''(u) > 0$ sur $[0, 1[$. Ainsi, la fonction φ est strictement convexe et elle va intersecter la droite $x \rightarrow x$ uniquement en 1 si $\varphi'(1) \leq 1$ et en un autre point $\rho < 1$ si $\varphi'(1) > 1$ (cf. figure 4.5). La stricte convexité de

φ implique l'unicité d'un tel point ρ . Comme $\mu < \infty$, l'identité (4.10) permet d'identifier la dérivée

$$\varphi'(1) = \sum_{n \geq 0} n p_n = \mu.$$

Il n'existe donc que deux comportements possibles

$$\lim_{t \rightarrow \infty} x_t = \begin{cases} 1, & \text{si } \mu \leq 1 \\ \rho, & \text{si } \mu > 1 \end{cases}$$

On sait que

$$\{T_0 < \infty\} = \bigcup_{t=0}^{\infty} \{Z_t = 0\}$$

qui est une réunion croissante d'évènements car $\{Z_t = 0\} \subset \{Z_{t+1} = 0\}$. On en déduit que

$$\mathbb{P}(T_0 < \infty) = \lim_{t \rightarrow \infty} \mathbb{P}(Z_t = 0) = \begin{cases} 1, & \text{si } \mu \leq 1 \\ \rho, & \text{si } \mu > 1 \end{cases}$$

Si $\mu > 1$, la probabilité que la population ne s'éteigne jamais est $\mathbb{P}(T_0 = \infty) = 1 - \rho$ et comme les états non nuls sont transitoires la taille de la population diverge avec probabilité $1 - \rho$. $\square$

Le comportement asymptotique des arbres peut être étudié plus précisément. Nous reviendrons sur le comportement asymptotique du cas sur-critique au chapitre 12.

4.4.2 Graphes aléatoires d'Erdős-Rényi

Les graphes aléatoires interviennent dans des contextes variés pour modéliser par exemple des réseaux sociaux, le réseau internet ou des réseaux neuronaux. Selon les applications, les détails de chaque graphe aléatoire diffèrent mais il est intéressant de classer ces graphes en fonction de structures invariantes et de propriétés communes. Dans cette section, nous allons considérer un modèle spécifique de graphes aléatoires qui a été inventé par Erdős et Rényi en 1959 [11] et montrer que les propriétés de ces graphes peuvent s'analyser à l'aide d'un processus de branchement.

Pour N un entier donné, on considère $\mathcal{S} = \{1, \dots, N\}$ un ensemble de sites reliés aléatoirement selon la procédure suivante. Soit $\lambda > 0$, on définit $N(N-1)/2$ variables aléatoires de Bernoulli indépendantes indexées par les couples $(x, y) \in \mathcal{S}^2$ avec $x \neq y$

$$\mathbb{P}(\eta_{x,y} = 1) = 1 - \mathbb{P}(\eta_{x,y} = 0) = \frac{\lambda}{N}. \quad (4.11)$$

On supposera que N est très grand et donc que $\lambda < N$. On ne distingue pas l'orientation des arêtes (x, y) et on pose $\eta_{x,y} = \eta_{y,x}$. Étant donnée une réalisation $\{\eta_{x,y}\}_{(x,y) \in \mathcal{S}^2}$, on construit un graphe $\mathcal{G} = (\mathcal{S}, \mathcal{E})$ dont les arêtes relient uniquement les sites x et y de $\mathcal{S}$ tels que $\eta_{x,y} = 1$. Ce procédé permet de générer un graphe aléatoire d'Erdős-Rényi à N sites (cf. figure 4.6).

Deux sites x et y sont connectés dans $\mathcal{G}$ s'il existe une suite d'arêtes allant de x à y , c'est-à-dire k sites de $\mathcal{S}$ avec $k \leq N$ tels que $\eta_{x,x_1} = \eta_{x_1,x_2} = \dots = \eta_{x_{k-1},x_k} = \eta_{x_k,y} = 1$.

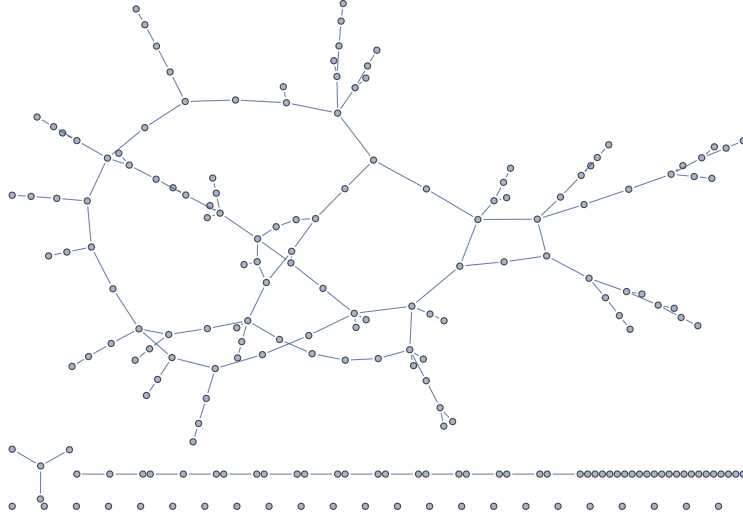


FIGURE 4.6 – Toutes les parties connexes d’un graphe aléatoire d’Erdős-Rényi sont représentées : au-dessus la partie connexe principale et en dessous les plus petites composantes connexes (certaines sont réduites à un seul site).

Pour tout x dans $\mathcal{S}$, on définit $\mathcal{C}(x)$ la composante connexe de x comme l’ensemble des sites y connectés à x .

Nous allons étudier la structure des connections dans les graphes d’Erdős-Rényi en fonction de la valeur du paramètre λ . Étant donnée une réalisation du graphe, on note C^* le cardinal de la plus grande de ses composantes connexes (il se peut qu’il y ait plusieurs composantes connexes de taille C^*). La figure 4.7 représente une simulation numérique de la densité de la composante connexe maximale $\lambda \rightarrow \mathbb{E}(C^*)/N$ pour différentes valeurs de N . Pour obtenir une approximation de l’espérance (λ et N étant fixés), on simule un grand nombre K de réalisations de graphes et on prend la moyenne des tailles $\{C_i^*\}_{i \leq K}$ des composantes maximales de chacun de ces graphes

$$\mathbb{E}(C^*) \simeq \frac{1}{K} \sum_{i=1}^K C_i^*.$$

La loi des grands nombres permet d’affirmer que l’approximation est correcte quand K tend vers l’infini. Les simulations de la figure 4.7 ont été faites pour $K = 1000$ et des fluctuations persistent.

On remarque que pour $\lambda < 1$ cette densité semble tendre vers 0 quand N augmente. Ceci veut dire qu’aucune composante connexe ne recouvre une fraction d’ordre N des sites. Nous allons montrer que pour $\lambda < 1$, les composantes connexes typiques d’un graphe d’Erdős-Rényi sont de taille finie même quand N tend vers l’infini. Dans ce cas, le graphe n’est qu’une collection de petits sous-graphes disjoints voire même de sites isolés.

Pour $\lambda > 1$, le comportement change radicalement et la plus grande composante contient une densité positive de sites. On dit qu’il y a une *transition de phase* au point critique $\lambda_c = 1$. La figure 4.6 représente une réalisation d’un graphe pour $\lambda > 1$. On remarque qu’il existe une composante connexe principale qui relie une grande partie des

sites et que les autres composantes connexes sont beaucoup plus petites. Il existe un lien entre les composantes connexes et les arbres de Galton-Watson, en particulier on peut observer que la composante principale du graphe ressemble à un arbre au voisinage de chaque site, même si à une plus grande échelle des boucles se forment.

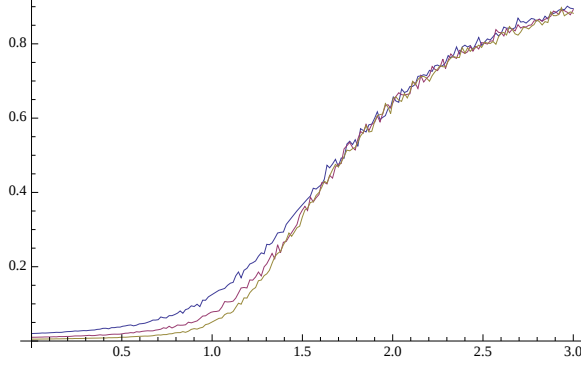


FIGURE 4.7 – Densité moyenne de la composante connexe maximale d'un graphe aléatoire d'Erdős-Rényi pour λ variant entre 0 et 3 avec $N = 50$ pour la courbe bleue, $N = 100$ pour la courbe violette, $N = 200$ pour la courbe verte.

Théorème 4.11. *Pour $\lambda < 1$, la composante connexe associée au site 1 a une taille moyenne bornée uniformément en N*

$$\mathbb{E}(|\mathcal{C}(1)|) \leq \frac{1}{1-\lambda}$$

où $|\mathcal{C}(1)|$ est le cardinal de $\mathcal{C}(1)$.

Démonstration. L'idée de la preuve consiste à remarquer qu'un site a en moyenne $\lambda \frac{N-1}{N}$ voisins car

$$\sum_{y \in V \setminus \{x\}} \mathbb{E}(\eta_{x,y}) = (N-1) \frac{\lambda}{N}.$$

Ses voisins eux-mêmes seront reliés à environ λ voisins et ainsi de suite. Cette structure ressemble à celle d'un arbre de Galton-Watson et elle permet de prédire l'existence de comportements différents selon que λ est plus grand ou plus petit que 1. Cependant la topologie des graphes est plus complexe que celle des arbres car il peut exister des boucles et il faut une preuve spécifique pour préciser cette analogie.

Nous allons explorer la composante connexe $\mathcal{C}(1)$ à la manière d'un arbre. Au temps initial $t = 0$, on pose $\mathcal{A}_0 = \{1\}$, $\mathcal{I}_0 = \{2, 3, \dots, N\}$ et $\mathcal{R}_0 = \emptyset$. Les trois ensembles vont évoluer au cours du temps selon la règle suivante (cf. figure 4.8)

$$\begin{cases} \mathcal{R}_{t+1} &= \mathcal{R}_t \cup \mathcal{A}_t \\ \mathcal{A}_{t+1} &= \bigcup_{x \in \mathcal{A}_t} \{y \in \mathcal{I}_t; \quad \eta_{x,y} = 1\} \\ \mathcal{I}_{t+1} &= \mathcal{I}_t \setminus \mathcal{A}_{t+1} \end{cases} \quad (4.12)$$

L'ensemble $\mathcal{A}_t$ représente les sites actifs au temps t , ceux-ci vont s'apparier avec les sites inactifs de $\mathcal{I}_t$ qui sont liés à $\mathcal{A}_t$ dans le graphe d'Erdős-Rényi. Ces nouveaux sites deviennent actifs au temps $t + 1$ et les sites de $\mathcal{A}_t$ viennent grossir l'ensemble $\mathcal{R}_{t+1}$. Ainsi

la composante connexe $\mathcal{C}(1)$ est explorée entièrement au cours de ce processus qui se termine au temps $\tau \leq N$ quand $\mathcal{A}_\tau = \emptyset$ et $\mathcal{C}(1) = \mathcal{R}_\tau$.

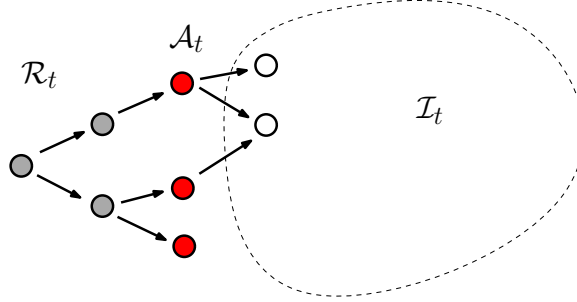


FIGURE 4.8 – La composante connexe $\mathcal{C}(1)$ est explorée en découvrant les sites voisins à chaque étape. Les sites actifs $\mathcal{A}_t$ sont représentés en rouge à distance 2 et ils sont connectés aux sites de $\mathcal{I}_t$ marqués en blanc. On remarque que cette exploration peut conduire à découvrir le même site de $\mathcal{I}_t$ s'il est relié à plusieurs sites de $\mathcal{A}_t$. Contrairement aux arbres, les composantes connexes du graphe peuvent donc avoir des boucles.

Le processus d'exploration $\mathcal{R}_t$ ressemble à un arbre de Galton-Watson. La différence étant que la distribution des descendants des sites actifs dépend de t car les descendants sont choisis dans l'ensemble $\mathcal{I}_t$ qui se réduit au cours du temps. Nous allons montrer que l'arbre de Galton-Watson permet de contrôler la croissance du processus d'exploration.

On se donne une collection de variables aléatoires indépendantes $\{\zeta_{x,y}^t\}$ avec $t \geq 1$, $x \geq 1$ et $y \in \{1, \dots, N\}$. Ces variables sont identiquement distribuées selon une loi de Bernoulli

$$\mathbb{P}(\zeta_{x,y}^t = 1) = 1 - \mathbb{P}(\zeta_{x,y}^t = 0) = \frac{\lambda}{N}.$$

On définit $U_0 = 1$ et on construit l'arbre de Galton-Watson dont la population au temps $t + 1$ est donnée par

$$U_{t+1} = \sum_{x \in \mathcal{A}_t, y \in \mathcal{I}_t} \eta_{x,y} + \sum_{x \in \mathcal{A}_t, y \in \mathcal{I}_t^c} \zeta_{x,y}^t + \sum_{x=N+1}^{N+U_t-|\mathcal{A}_t|} \sum_{y=1}^N \zeta_{x,y}^t \quad (4.13)$$

où $|\mathcal{A}_t|$ désigne le cardinal de $\mathcal{A}_t$. Le second terme ajoute des descendants fictifs dans $\{1, \dots, N\} \setminus \mathcal{I}_t$ pour compenser la réduction du cardinal de $\mathcal{I}_t$ à chaque pas. Ces descendants fictifs ont ensuite eux-même une descendance qui est prise en compte dans le troisième terme de (4.13). Par conséquent $\{U_t\}$ est un processus de branchement dont la loi de reproduction est une loi binomiale de paramètres $(N, \frac{\lambda}{N})$, i.e. que la distribution des enfants de chaque site a la même loi que $\sum_{i=1}^N \omega_i$ où les ω_i sont des variables de Bernoulli indépendantes de paramètre λ/N . Il est important de remarquer que quand N tend vers l'infini la loi de reproduction converge vers une loi de Poisson de paramètre λ .

$$\mathbb{P}\left(\sum_{i=1}^N \omega_i = k\right) = \binom{N}{k} \left(\frac{\lambda}{N}\right)^k \left(1 - \frac{\lambda}{N}\right)^{N-k} \xrightarrow{N \rightarrow \infty} \exp(-\lambda) \frac{\lambda^k}{k!}.$$

On s'attend donc à ce que de très grands systèmes ($N \rightarrow \infty$) convergent vers une structure limite et soient bien décrits par des arbres de Galton-Watson de loi de reproduction

donnée par cette loi de Poisson. La moyenne de la loi de reproduction est λ . Si $\lambda < 1$, le théorème 4.10 implique que les arbres seront finis presque sûrement.

Comme le processus $\{U_t\}$ est construit en ajoutant des sites fictifs (4.13) par rapport à ceux existant dans la composante $\mathcal{C}(1)$ du graphe, son cardinal domine toujours le cardinal de $\mathcal{C}(1)$. On en déduit que

$$\mathbb{E}(|\mathcal{C}(1)|) = \sum_{t=0}^{\infty} \mathbb{E}(|\mathcal{A}_t|) \leq \sum_{t=0}^{\infty} \mathbb{E}(U_t) = \sum_{t=0}^{\infty} \lambda^t = \frac{1}{1-\lambda} \quad (4.14)$$

car l'identité (4.9) implique que $\mathbb{E}(U_t) = \lambda^t$. Ceci conclut le théorème. $\square$

Pour $\lambda > 1$, la comparaison avec un arbre permet de montrer qu'il existe une composante connexe contenant une proportion de sites proportionnelle à N . Le nombre moyen d'enfants étant égal à λ , la population de l'arbre a une probabilité positive de diverger ce qui indique que la composante connexe $\mathcal{C}(1)$ doit être très grande. La preuve est délicate car la comparaison entre un arbre et le processus d'exploration de $\mathcal{C}(1)$ (décrit dans la preuve du théorème 4.11) n'est plus valable quand la composante connexe explorée $\mathcal{R}_t$ est trop grande : les boucles ne peuvent plus être négligées et l'ajout des sites fictifs devient trop important. La preuve complète est faite dans le livre de R. Durrett [11].

La structure des graphes d'Erdős-Rényi est très bien comprise mathématiquement. On peut par exemple montrer que pour $\lambda > 1$, deux sites appartenant à la composante connexe principale sont typiquement à distance $\log N$ (bien que cette composante contienne un nombre de sites proportionnel à N).

Une application possible est d'interpréter un graphe aléatoire comme un ensemble d'agents en interaction (réseau informatique, système financier) et d'étudier la résistance de ce graphe à une perturbation (virus informatique, défaut de paiement). Un exemple simple consiste à retirer aléatoirement des liens avec une probabilité p et à les garder avec probabilité $(1-p)$. On souhaite déterminer s'il existera toujours une composante connexe d'ordre N après cette modification du réseau. Dans le cas particulier des graphes d'Erdős-Rényi, le réseau modifié reste équivalent à un graphe d'Erdős-Rényi de paramètre $(1-p)\lambda$. Si $(1-p)\lambda$ est plus grand que 1 alors le graphe restera fortement connecté, sinon la composante connexe principale sera décomposée en une multitude de composantes disjointes. Un autre type de graphes aléatoires sera construit au chapitre 12 et de nombreux autres modèles de graphes aléatoires figurent dans le livre [11].

Chapitre 5

Ergodicité et convergence des chaînes de Markov

L'étude des comportements asymptotiques de variables aléatoires constitue un aspect essentiel des probabilités. La loi des grands nombres et le théorème central limite en sont deux exemples très importants. Ce chapitre décrit les comportements asymptotiques des chaînes de Markov pour lesquels les mesures invariantes jouent un rôle clef.

5.1 Ergodicité

Soient $\{Y_n\}_{n \geq 0}$ des variables aléatoires indépendantes et identiquement distribuées à valeurs dans $\mathbb{R}$ telles que l'espérance $\mathbb{E}(|f(Y_0)|)$ soit finie pour une fonction f donnée. Le théorème de la loi des grands nombres implique la convergence presque sûre

$$\frac{1}{n} \sum_{i=0}^{n-1} f(Y_i) \xrightarrow{n \rightarrow \infty} \mathbb{E}(f(Y_0)). \quad (5.1)$$

Ce théorème se généralise aux chaînes de Markov récurrentes positives et on parle alors de *théorème ergodique*.

5.1.1 Théorème ergodique

Théorème 5.1. Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov irréductible, récurrente positive sur un espace d'états E dénombrable. On notera π son unique mesure de probabilité invariante. Soit F une fonction de E dans $\mathbb{R}$ dont l'espérance sous π est finie $\mathbb{E}_\pi(|F|) = \sum_{x \in E} |F(x)|\pi(x) < \infty$.

On suppose que la donnée initiale X_0 est distribuée selon une mesure de probabilité μ sur E . Les moyennes le long des trajectoires convergent presque sûrement

$$\frac{1}{n} \sum_{i=0}^{n-1} F(X_i) \xrightarrow{n \rightarrow \infty} \mathbb{E}_\pi(F). \quad (5.2)$$

Si $F(x, y)$ est une fonction de $E \times E$ dans $\mathbb{R}$ telle que $\sum_{x, y \in E} \pi(x)P(x, y)|F(x, y)|$ est finie alors

$$\frac{1}{n} \sum_{i=1}^n F(X_{i-1}, X_i) \xrightarrow{n \rightarrow \infty} \mathbb{E}_\pi(F(X_0, X_1)) = \sum_{x, y \in E} \pi(x)P(x, y)F(x, y). \quad (5.3)$$

Démonstration. Supposons que la chaîne parte initialement d'un état x de E fixé, i.e. que $\mu = \delta_x$. Cet état x va servir d'état de référence pour représenter l'unique mesure invariante π associée à cette chaîne de Markov récurrente positive

$$\forall y \in E, \quad \pi(y) = \frac{\sum_{n \geq 0} \mathbb{P}_x(X_n = y, T_x^+ > n)}{\mathbb{E}_x(T_x^+)} = \frac{\mathbb{E}_x\left(\sum_{n=0}^{T_x^+-1} \mathbf{1}_{X_n=y}\right)}{\mathbb{E}_x(T_x^+)}. \quad (5.4)$$

Contrairement à l'expression (4.4), le facteur $1/\mathbb{E}_x(T_x^+)$ sert à normaliser la mesure. La probabilité $\pi(y)$ décrit la statistique des passages dans l'état y pendant une excursion de la chaîne de Markov. Par ailleurs, les différentes excursions (cf. figure 5.1) sont indépendantes par la propriété de Markov forte. Nous allons donc décomposer les trajectoires de la chaîne de Markov en excursions pour établir la correspondance avec π .

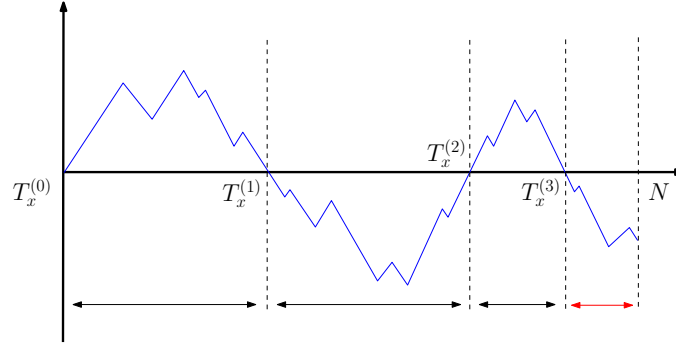


FIGURE 5.1 – Décomposition de la trajectoire d'une chaîne de Markov en trois excursions entre les passages en x . Le dernier segment jusqu'au temps N n'est pas une excursion complète.

Rappelons les notations sur les temps de retour. Pour tout entier $k \geq 1$, on définit le $k^{\text{ième}}$ temps de retour en x par

$$T_x^{(k)} = \inf \{n \geq T_x^{(k-1)} + 1; \quad X_n = x\} \in \mathbb{N} \cup \{\infty\}$$

où $T_x^{(0)} = 0$ et $T_x^{(1)}$ coïncide avec T_x^+ . Par hypothèse, la chaîne est récurrente et tous les temps d'arrêt sont finis presque sûrement. On définit les variables aléatoires $\{Y_k\}_{k \geq 0}$ associées à la contribution de chaque excursion

$$Y_k = \sum_{\ell=T_x^{(k)}}^{T_x^{(k+1)}-1} F(X_\ell).$$

Par la propriété de Markov forte démontrée au théorème 2.5, les variables $\{Y_k\}$ sont indépendantes et identiquement distribuées. Leur espérance s'écrit en fonction de la mesure invariante

$$\begin{aligned} \mathbb{E}(Y_1) &= \mathbb{E}_x \left(\sum_{\ell=0}^{T_x^+-1} F(X_\ell) \right) = \mathbb{E}_x \left(\sum_{\ell=0}^{T_x^+-1} \sum_{y \in E} F(y) \mathbf{1}_{X_\ell=y} \right) = \sum_{y \in E} F(y) \mathbb{E}_x \left(\sum_{\ell=0}^{T_x^+-1} \mathbf{1}_{X_\ell=y} \right) \\ &= \mathbb{E}_x(T_x^+) \sum_{y \in E} F(y) \pi(y) = \mathbb{E}_x(T_x^+) \mathbb{E}_\pi(F) \end{aligned}$$

où l'hypothèse $\mathbb{E}_\pi(|F|) < \infty$ nous a permis d'utiliser le théorème de Fubini pour permuter la somme et l'espérance. La loi des grands nombres (5.1) pour les variables indépendantes implique la convergence presque sûre

$$\frac{1}{k} \sum_{\ell=0}^{T_x^{(k)}-1} F(X_\ell) = \frac{1}{k} \sum_{i=0}^{k-1} Y_i \xrightarrow{k \rightarrow \infty} \mathbb{E}_x(T_x^+) \mathbb{E}_\pi(F).$$

Ce résultat appliqué à $F = 1$ permet d'écrire

$$\frac{T_x^{(k)}}{k} \xrightarrow{k \rightarrow \infty} \mathbb{E}_x(T_x^+). \quad (5.5)$$

Par conséquent en indexant la trajectoire par les temps de retours, on a prouvé la convergence presque sûre

$$\frac{1}{T_x^{(k)}} \sum_{\ell=0}^{T_x^{(k)}-1} F(X_\ell) \xrightarrow{k \rightarrow \infty} \mathbb{E}_\pi(F). \quad (5.6)$$

Pour obtenir le théorème ergodique, il suffit de contrôler la contribution de la trajectoire après le dernier passage en x et de montrer qu'elle ne joue aucun rôle à la limite (cf. figure 5.1). On note $\mathfrak{N}_n$ le nombre de passages en x avant le temps n , c'est-à-dire

$$\mathfrak{N}_n = \sum_{\ell=1}^n \mathbf{1}_{X_\ell=x} \quad \Rightarrow \quad T_x^{(\mathfrak{N}_n)} \leq n < T_x^{(\mathfrak{N}_n+1)}.$$

La chaîne étant récurrente $\mathfrak{N}_n$ diverge quand n tend vers l'infini. On peut décomposer F en une partie positive et négative $F = F^+ - F^-$ et traiter chaque terme séparément. Supposons donc que F soit positive, alors

$$\frac{T_x^{(\mathfrak{N}_n)}}{T_x^{(\mathfrak{N}_n+1)}} \frac{1}{T_x^{(\mathfrak{N}_n)}} \sum_{\ell=0}^{T_x^{(\mathfrak{N}_n)}-1} F(X_\ell) \leq \frac{1}{n} \sum_{\ell=0}^{n-1} F(X_\ell) \leq \frac{T_x^{(\mathfrak{N}_n+1)}}{T_x^{(\mathfrak{N}_n)}} \frac{1}{T_x^{(\mathfrak{N}_n+1)}} \sum_{\ell=0}^{T_x^{(\mathfrak{N}_n+1)}-1} F(X_\ell),$$

où on a utilisé

$$\frac{T_x^{(\mathfrak{N}_n)}}{T_x^{(\mathfrak{N}_n+1)}} \leq \frac{T_x^{(\mathfrak{N}_n)}}{n} \quad \text{et} \quad \frac{T_x^{(\mathfrak{N}_n+1)}}{n} \leq \frac{T_x^{(\mathfrak{N}_n+1)}}{T_x^{(\mathfrak{N}_n)}}.$$

La convergence de (5.5) implique que presque sûrement

$$\frac{T_x^{(\mathfrak{N}_n)}}{T_x^{(\mathfrak{N}_n+1)}} \xrightarrow{n \rightarrow \infty} 1.$$

Il suffit donc d'appliquer (5.6) pour conclure que pour une donnée initiale $X_0 = x$

$$\frac{1}{n} \sum_{\ell=0}^{n-1} F(X_\ell) \xrightarrow{n \rightarrow \infty} \mathbb{E}_\pi(F).$$

La limite ne dépend pas de l'état de référence x choisi. Par conséquent, si la donnée initiale X_0 est choisie selon une mesure μ , il suffit de sommer sur la probabilité de chaque état initial et d'appliquer la relation précédente.

Pour la seconde partie du théorème, il suffit de remarquer que $Z_n = (X_{n-1}, X_n)$ est une chaîne de Markov à valeurs dans un sous-ensemble Γ de $E \times E$ de matrice de transition

$$\mathbb{P}(Z_2 = (y_1, y_2) \mid Z_1 = (x_1, x_2)) = \mathbf{1}_{y_1=x_2} P(x_2, y_2)$$

et de mesure de probabilité invariante $\tilde{\pi}(x, y) = \pi(x)P(x, y)$. Insistons sur le fait que la chaîne de Markov $\{Z_n\}_{n \geq 1}$ n'est pas irréductible sur $E \times E$ mais seulement sur le sous-ensemble Γ où elle prend ses valeurs. La limite (5.3) se déduit du théorème ergodique (5.2) appliqué à la chaîne $\{Z_n\}_{n \geq 1}$. Le résultat se généralise facilement aux fonctions F à k variables. $\square$

Pour une chaîne de Markov récurrente positive, le théorème ergodique (5.2) appliqué à la fonction $F(y) = \mathbf{1}_{y=x}$ permet d'interpréter la mesure invariante π comme la fréquence de visites des états par la chaîne de Markov

$$\frac{1}{n} \sum_{\ell=0}^{n-1} \mathbf{1}_{X_\ell=x} \xrightarrow{n \rightarrow \infty} \pi(x) = \frac{1}{\mathbb{E}_x(T_x^+)} \quad \text{presque sûrement.}$$

Cette convergence reste vraie pour les chaînes de Markov récurrentes nulles et transitoires pour lesquelles $\mathbb{E}_x(T_x^+) = \infty$.

Théorème 5.2. Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov récurrente nulle ou transitoire à valeurs dans E dénombrable alors pour tout x de E

$$\frac{1}{n} \sum_{\ell=0}^{n-1} \mathbf{1}_{X_\ell=x} \xrightarrow{n \rightarrow \infty} 0 \quad \text{presque sûrement.} \quad (5.7)$$

Démonstration. Si la chaîne de Markov est transitoire alors elle ne repassera qu'un nombre fini de fois par un état par conséquent $\sum_{\ell=0}^{\infty} \mathbf{1}_{X_\ell=x}$ est fini presque sûrement et (5.7) est vérifié.

On rappelle que pour une fonction f positive, la loi des grands nombres (5.1) reste vraie même si $\mathbb{E}(f(Y_0)) = \infty$ (cf. corollaire B.22)

$$\frac{1}{n} \sum_{i=0}^n f(Y_i) \xrightarrow{n \rightarrow \infty} \infty.$$

Pour le démontrer, il suffit d'appliquer la loi des grands nombres à la fonction tronquée $\inf\{f, K\}$ puis de faire tendre K vers l'infini.

Si la chaîne de Markov est récurrente nulle, on peut appliquer la preuve du théorème 5.1 pour obtenir par (5.5) la convergence presque sûre

$$\frac{T_x^{(k+1)}}{k} \xrightarrow{k \rightarrow \infty} \infty.$$

En utilisant les notations du théorème 5.1 et le fait que $T_x^{(\mathfrak{N}_n)} \leq n$, on peut écrire

$$\frac{1}{n} \sum_{\ell=0}^n \mathbf{1}_{X_\ell=x} = \frac{\mathfrak{N}_n}{n} \leq \frac{\mathfrak{N}_n}{T_x^{(\mathfrak{N}_n)}} \xrightarrow{n \rightarrow \infty} 0$$

car la chaîne de Markov est récurrente et $\mathfrak{N}_n$ converge presque sûrement vers l'infini. $\square$

La décomposition des trajectoires en excursions indépendantes permet aussi de démontrer l'analogie du théorème central limite pour des chaînes de Markov récurrentes positives. Une preuve différente du théorème central limite sera faite au chapitre 11.

5.1.2 Application : algorithme PageRank de Google

Le principe du moteur de recherche Google consiste à indexer les pages du web par des robots (web crawler). Ces logiciels sondent et conservent automatiquement le contenu de chaque site web dans un immense catalogue qui indexe un nombre de pages web de taille gigantesque, de l'ordre de $3 \cdot 10^{13}$ (d'après Google). La difficulté consiste ensuite à attribuer un ordre de priorité dans cette masse d'informations pour satisfaire au mieux les différentes requêtes des internautes. Un élément clef utilisé par le moteur de recherche Google pour hiérarchiser l'importance des données est l'algorithme PageRank. Celui-ci a été développé à l'Université de Stanford dans le cadre d'un projet de recherche commencé en 1995 par Larry Page et rejoint plus tard par Sergey Brin.

L'information disponible sur le web n'est pas structurée sur le modèle des bases de données traditionnelles, elle s'est plutôt auto-organisée au fil du temps. De plus, la taille gigantesque du web ne permet pas d'utiliser les méthodes classiques de recherche documentaire. Le principe de PageRank consiste à laisser faire le hasard pour découvrir les pages web aléatoirement sans utiliser une approche déterministe comme on pourrait le faire dans un environnement bien structuré. Cet algorithme a complètement révolutionné le fonctionnement des moteurs de recherche.

Nous allons décrire le principe de fonctionnement de PageRank en gardant à l'esprit que de nombreuses améliorations ont été introduites depuis afin de répondre aux divers détournements des utilisateurs. L'objectif est d'associer à une page web i un indice de popularité $\pi(i)$. L'idée est de dire qu'une page web i est importante si de nombreux liens pointent sur cette page. En particulier si un site j très populaire pointe sur la page i , il va générer beaucoup de connections sur i et ainsi augmenter la popularité de i . Cette règle empirique conduit à la relation suivante

$$\pi(i) = \sum_{j \rightarrow i} \frac{1}{\deg(j)} \pi(j)$$

où on note $j \rightarrow i$ si la page j pointe sur la page i et $\deg(j)$ le nombre de liens partant de la page j . La page j transmet son indice de popularité proportionnellement entre les $\deg(j)$ pages web auxquelles elle renvoie. Cette relation définit la mesure invariante d'une marche aléatoire sur le graphe $\mathcal{G} = (\mathcal{S}, \mathcal{E})$ dont les sites $\mathcal{S}$ sont indexés par les pages web et les arêtes $\mathcal{E}$ par les liens entre les sites. Quand les liens du graphe ne sont pas orientés, une telle marche a été définie en (3.1) et sa probabilité invariante $\pi(x) = \frac{\deg(x)}{2|\mathcal{E}|}$ calculée en (3.2). Pour évaluer π , une possibilité est d'indexer de façon systématique tous les liens, mais le graphe du web est compliqué et il évolue sans cesse. L'option retenue par l'algorithme PageRank consiste à laisser faire le hasard en suivant des marches aléatoires qui évoluent de pages en pages selon les liens et en indexant le contenu à chaque fois. Le théorème ergodique 5.1 permet ensuite de retrouver la mesure invariante π , i.e. les indices de popularité, en moyennant sur les trajectoires des marcheurs.

La vitesse de convergence d'un algorithme est fondamentale dans les applications industrielles. Pour cette raison, Brin et Page ont introduit la matrice de transition de Google G de composantes

$$G(i, j) = \alpha \frac{1}{\deg(i)} \mathbf{1}_{i \rightarrow j} + (1 - \alpha) \frac{1}{N} \quad \text{pour tous } i, j \in S, \quad (5.8)$$

où N est le cardinal de S et α un paramètre dans $]0, 1[$ appelé facteur d'amortissement. Avec probabilité $1 - \alpha$, les marches aléatoires sont relancées sur un site choisi au hasard parmi les N sites de S . Cette modification de la matrice de transition rappelle le comportement de l'internaute qui suit quelques liens puis au bout d'un moment (avec une probabilité $1 - \alpha$) se dirige vers un de ses liens favoris que l'on suppose distribués uniformément sur l'ensemble des pages. Le choix de la valeur du paramètre α est délicat. Le souci de rapidité de la convergence de l'algorithme nous pousse à choisir α proche de 0, mais ceci conduirait à une mesure invariante qui ne refléterait plus la vraie structure du web (toutes les pages auraient la même probabilité $1/N$). La valeur exacte du paramètre α est un secret gardé de Google.

5.2 Convergence

Reprenons l'exemple de la chaîne de Markov irréductible à deux états $\{1, 2\}$ (cf. figure 3.2 et équation (3.3)) dont la matrice de transition et la mesure invariante π sont données par

$$P = \begin{pmatrix} 1-p & p \\ q & 1-q \end{pmatrix}, \quad \pi(1) = \frac{q}{p+q}, \quad \pi(2) = \frac{p}{p+q}$$

avec $p, q \in]0, 1[$. Cette matrice de transition se diagonalise facilement

$$P = \begin{pmatrix} 1 & -\pi(2) \\ 1 & \pi(1) \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & 1-p-q \end{pmatrix} \begin{pmatrix} \pi(1) & \pi(2) \\ -1 & 1 \end{pmatrix}$$

et peut être multipliée n fois

$$P^n = \begin{pmatrix} 1 & -\pi(2) \\ 1 & \pi(1) \end{pmatrix} \begin{pmatrix} 1 & 0 \\ 0 & (1-p-q)^n \end{pmatrix} \begin{pmatrix} \pi(1) & \pi(2) \\ -1 & 1 \end{pmatrix} \xrightarrow{n \rightarrow \infty} \begin{pmatrix} \pi(1) & \pi(2) \\ \pi(1) & \pi(2) \end{pmatrix}.$$

Quand le temps tend vers l'infini, les probabilités de transition convergent exponentiellement vite vers la mesure invariante

$$\forall x, y \in \{1, 2\}, \quad \lim_{n \rightarrow \infty} \mathbb{P}_x(X_n = y) = \pi(y).$$

L'état initial n'apparaît plus dans la limite.

L'enjeu de cette section est de démontrer que la convergence en temps long vers la mesure invariante est une propriété très générale des chaînes de Markov récurrentes positives et de quantifier la vitesse de convergence.

5.2.1 Apériodicité et convergence

Une conséquence du théorème ergodique 5.1 est la convergence pour tout état initial x de

$$\frac{1}{n} \sum_{\ell=0}^n \mathbb{P}_x(X_\ell = y) \xrightarrow{n \rightarrow \infty} \pi(y).$$

Ceci ne suffit pas à impliquer que $\lim_{n \rightarrow \infty} \mathbb{P}_x(X_n = y) = \pi(y)$. Pour s'en convaincre, considérons la chaîne de Markov irréductible à deux états de matrice de transition

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

dont le comportement est périodique $\mathbb{P}_1(X_n = 1) = 1 - \mathbb{P}_1(X_n = 2) = (1 + (-1)^n)/2$. Un autre exemple est la marche aléatoire sur le domaine périodique $\{1, \dots, 2L\}$ avec un nombre pair de sites (cf. figure 5.2). Si la marche part de 1 au temps 0, elle ne pourra atteindre un site pair qu'à des temps impairs.

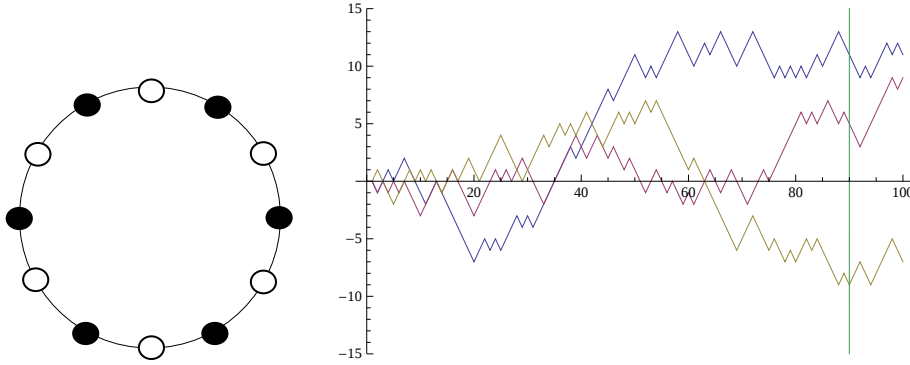


FIGURE 5.2 – La marche aléatoire sur l'intervalle périodique $\{1, \dots, 2L\}$ reste sur les sites noirs au temps pairs si elle est partie d'un site noir. Le schéma de droite représente 3 réalisations d'une même chaîne de Markov. Sous les hypothèses du théorème 5.5, la mesure invariante peut s'obtenir en moyennant les valeurs de plusieurs trajectoires à un temps donné.

Pour démontrer la convergence, il faut restreindre la classe des matrices de transition aux chaînes de Markov apériodiques.

Définition 5.3. Une chaîne de Markov irréductible sur E est apériodique si pour tous x, y de E il existe $n(x, y) \in \mathbb{N}$ tel que la probabilité $\mathbb{P}_x(X_n = y) = P^n(x, y)$ est strictement positive dès que $n \geq n(x, y)$.

Cette définition permet d'éviter les pathologies décrites précédemment car une chaîne de Markov apériodique a une probabilité positive de connecter 2 états dès que le temps est assez grand. On peut facilement se ramener à des chaînes de Markov apériodiques en transformant la matrice de transition. La matrice $Q = (I + P)/2$ est associée à la version "fainéante" de la chaîne de Markov : avec probabilité 1/2 la chaîne reste sur place et avec probabilité 1/2 elle fait un saut selon la matrice P . La matrice Q a la même probabilité invariante que P et donc un comportement asymptotique similaire.

Lemme 5.4. *Si la chaîne de Markov sur E est irréductible et si un état x est apériodique, c'est-à-dire qu'il vérifie*

$$\mathbb{P}_x(X_n = x) = P^n(x, x) > 0 \quad \text{dès que } n \text{ est suffisamment grand}$$

alors la chaîne est apériodique.

Démonstration. Soient y, z deux états de E . Comme la chaîne est irréductible, il existe deux entiers r et s tels que $P^r(y, x) > 0$ et $P^s(x, z) > 0$, on en déduit que pour tout n suffisamment grand

$$P^{r+n+s}(y, z) \geq P^r(y, x)P^n(x, x)P^s(x, z) > 0.$$

La chaîne est donc apériodique. □

La définition 5.3 n'est pas la seule caractérisation des chaînes apériodiques et nous reviendrons par la suite sur cette notion. Pour le moment, nous allons montrer une conséquence importante de l'apériodicité.

Théorème 5.5. *Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov irréductible et apériodique de mesure de probabilité invariante π sur un espace d'états E dénombrable. Pour toute distribution initiale μ_0 sur E , la distribution de $\{X_n\}_{n \geq 0}$ converge vers π quand n tend vers l'infini*

$$\lim_{n \rightarrow \infty} \mathbb{P}_{\mu_0}(X_n = x) = \pi(x).$$

Ce théorème peut s'interpréter de la façon suivante. Pour n très grand, la mesure invariante est bien approchée par $\mathbb{P}_\mu(X_n = x)$ et cette distribution au temps n peut elle même être obtenue par la loi des grands nombres en simulant plusieurs réalisations indépendantes de la chaîne de Markov

$$\pi(x) \simeq \mathbb{P}_\mu(X_n = x) = \lim_{K \rightarrow \infty} \frac{1}{K} \sum_{k=1}^K \mathbf{1}_{\{X_n^{(k)} = x\}}$$

où $\{X_n^{(k)}\}_{n \geq 0}$ sont des réalisations indépendantes de la chaîne de Markov. Il y a donc deux approches complémentaires pour estimer $\pi(x)$: on moyenne la fréquence de passage en x le long d'une trajectoire (c'est le théorème ergodique 5.1) ou on fixe un temps n et on construit un histogramme à partir de plusieurs simulations indépendantes (cf. figure 5.2).

Démonstration. La preuve repose sur une méthode de couplage et nécessite plusieurs étapes. On considère $\{X_n\}_{n \geq 0}$ et $\{Y_n\}_{n \geq 0}$ deux réalisations indépendantes de la chaîne de Markov dont les états initiaux diffèrent : $X_0 = x$ et $Y_0 = y$ pour distribution initiale la mesure invariante π .

Étape 1 : *La chaîne de Markov jointe $W_n = (X_n, Y_n)$ est irréductible et récurrente positive.*

On voit facilement que processus $W_n = (X_n, Y_n)$ est une chaîne de Markov dans $E \times E$ de matrice de transition

$$\hat{P}\left((x_1, y_1), (x_2, y_2)\right) = P(x_1, x_2)P(y_1, y_2).$$

L'irréductibilité est une conséquence de l'apériodicité car pour tous les (x_1, y_1) et (x_2, y_2) dans $E \times E$

$$\hat{P}^n((x_1, y_1), (x_2, y_2)) = P^n(x_1, x_2)P^n(y_1, y_2) > 0$$

dès que n est assez grand. L'apériodicité de la chaîne est essentielle pour prouver l'irréductibilité, en effet si $\{X_n\}_{n \geq 0}$ et $\{Y_n\}_{n \geq 0}$ étaient deux réalisations d'une marche aléatoire sur $\{1, \dots, 2L\}$ (cf. figure 5.2) l'une partant d'un nombre pair et l'autre d'un nombre impair, alors le couple formé par W_n ne pourra jamais atteindre tous les sites $\{1, \dots, 2L\}^2$ et en particulier les trajectoires $\{X_n\}_{n \geq 0}$ et $\{Y_n\}_{n \geq 0}$ ne se rencontreront jamais.

Comme X_n et Y_n ont pour mesure invariante π , il est facile de vérifier que $\{W_n\}_{n \geq 0}$ a pour mesure invariante la mesure produit

$$\forall (x, y) \in E \times E, \quad \pi \otimes \pi(x, y) = \pi(x)\pi(y).$$

Par le théorème 4.6, la chaîne $\{W_n\}_{n \geq 0}$ est donc récurrente positive.

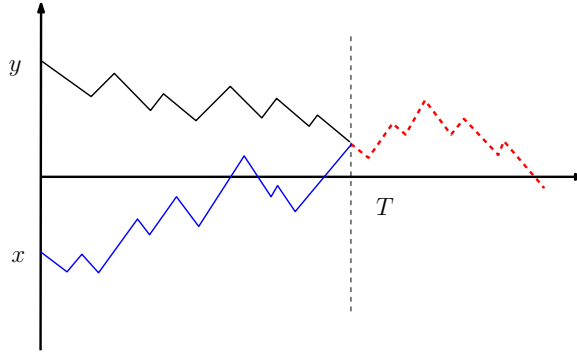


FIGURE 5.3 – Le schéma représente un couplage entre 2 trajectoires issues des états x et y . Leur partie commune après le temps T est dessinée en pointillés.

Étape 2 : Construction d'un couplage.

On définit le temps d'arrêt T comme le premier temps où les chaînes X_n, Y_n se touchent (cf. figure 5.3)

$$T = \inf \left\{ n \geq 0; \quad X_n = Y_n \right\} = \inf \left\{ n \geq 0; \quad W_n \in A \right\}$$

avec $A = \{(x, x); \quad x \in E\} \subset E \times E$. Ainsi T peut s'interpréter comme un temps d'atteinte pour la chaîne W_n . Comme W_n est irréductible et récurrente, T est fini presque sûrement. On définit le processus

$$Z_n = \begin{cases} X_n, & \text{si } n < T \\ Y_n, & \text{si } n \geq T \end{cases}$$

Nous allons vérifier que $\{Z_n\}_{n \geq 0}$ est une chaîne de Markov et a la même distribution que $\{X_n\}_{n \geq 0}$. Par la propriété de Markov forte démontrée au théorème 2.5, la chaîne de Markov décalée en temps $\{W_{T+n}\}_{n \geq 0}$ est indépendante de $\{(X_0, Y_0), \dots, (X_T, Y_T)\}$ conditionnellement à (X_T, Y_T) . Comme leurs données initiales coïncident les chaînes de

Markov $\{X_{T+n}\}_{n \geq 0}$ et $\{Y_{T+n}\}_{n \geq 0}$ ont la même loi et il est donc équivalent de suivre la trajectoire associée à Y_{T+n} plutôt que celle de X_{T+n} . Par conséquent $\{Z_n\}_{n \geq 0}$ est bien une chaîne de Markov de même loi que $\{X_n\}_{n \geq 0}$.

Étape 3 : Convergence.

Les données initiales sont $X_0 = x$ et la mesure invariante π pour Y_0 , on peut donc écrire pour tout état y de E

$$\mathbb{P}_x(X_n = y) - \pi(y) = \mathbb{P}_x(X_n = y) - \mathbb{P}_\pi(Y_n = y).$$

Par l'étape 2, la chaîne $\{Z_n\}_{n \geq 0}$ a la même loi que $\{X_n\}_{n \geq 0}$

$$\mathbb{P}_x(X_n = y) = \mathbb{P}_x(Z_n = y) = \hat{\mathbb{P}}(X_n = y, T > n) + \hat{\mathbb{P}}(Y_n = y, T \leq n),$$

où $\hat{\mathbb{P}}$ fait référence à la mesure jointe des trajectoires $\{X_n\}$ et $\{Y_n\}$. On en déduit que

$$\left| \mathbb{P}_x(X_n = y) - \pi(y) \right| = \left| \hat{\mathbb{P}}(X_n = y, T > n) - \hat{\mathbb{P}}(Y_n = y, T > n) \right| \leq \hat{\mathbb{P}}(T > n).$$

D'après la première étape, la chaîne $\{W_n\}_{n \geq 0}$ est récurrente positive. Par conséquent le temps d'arrêt T est fini presque sûrement et $\mathbb{P}(T > n)$ tend vers 0 quand n tend vers l'infini. Ceci conclut la preuve du théorème. $\square$

La définition 5.3 de l'apériodicité est particulièrement bien adaptée pour les preuves. Il existe un autre point de vue complémentaire qui justifie le choix du mot *apériodique* et que nous présentons dans la suite. Cette partie peut être omise en première lecture.

Pour tout état $x \in E$, on définit

$$\mathbf{p}(x) = \text{PGCD}[I(x)] \quad \text{où} \quad I(x) = \{n \geq 1; \quad P^n(x, x) > 0\}$$

où PGCD désigne le *plus grand commun diviseur*.

Proposition 5.6. Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov irréductible sur l'espace d'états E . Alors la fonction $x \rightarrow \mathbf{p}(x)$ est constante sur E .

Démonstration. Soient $x, y \in E$ deux états. Comme la chaîne de Markov est irréductible, il existe deux entiers i, j tels que $P^i(x, y) > 0$ et $P^j(y, x) > 0$. La propriété de Markov implique que pour tout $r \in I(y)$

$$P^{i+j}(x, x) > 0 \quad \text{et} \quad P^{i+j+r}(x, x) > 0.$$

Ainsi $\mathbf{p}(x)$ divise $i + j$ et $i + j + r$ et donc $\mathbf{p}(x)$ divise la différence r de ces deux entiers. Comme r est arbitraire dans $I(y)$, on déduit que $\mathbf{p}(x)$ divise $\mathbf{p}(y)$. En inversant les rôles de x et y , on montre l'égalité $\mathbf{p}(x) = \mathbf{p}(y)$. $\square$

Pour la marche aléatoire de la figure 5.3, la période est égale à 2. Nous donnons maintenant une seconde définition de l'apériodicité équivalente à celle de la définition 5.3

Définition 5.7. Une chaîne de Markov irréductible $\{X_n\}_{n \geq 0}$ est dite apériodique si tout état x vérifie $\mathbf{p}(x) = 1$.

Le lemme qui suit permet de faire le lien entre les deux définitions.

Lemme 5.8. *Soit une chaîne de Markov irréductible de matrice de transition P sur E . Pour tout x dans E , les deux assertions suivantes sont équivalentes :*

- (i) $\mathbf{p}(x) = 1$,
- (ii) *il existe $\mathbf{n}(x) \geq 1$ tel que $P^n(x, x) > 0$ pour tout $n \geq \mathbf{n}(x)$.*

Démonstration. L'implication (ii) $\Rightarrow$ (i) est immédiate. Pour établir la réciproque, on fixe x tel que $\mathbf{p}(x) = 1$. Considérons des entiers $n_1, \dots, n_k$ dans $I(x)$ tels que $\text{PGCD}[n_1, \dots, n_k] = 1$. Le théorème de Bézout assure l'existence de $q_1, \dots, q_k \in \mathbb{Z}$ vérifiant

$$\sum_{i=1}^k q_i n_i = 1. \quad (5.9)$$

En utilisant la notation $q_i^+ = \max(q_i, 0)$ et $q_i^- = -\min(q_i, 0)$, on définit deux entiers

$$a = \sum_{i=1}^k q_i^+ n_i \quad \text{et} \quad b = \sum_{i=1}^k q_i^- n_i$$

qui satisfont $a - b = 1$ d'après (5.9).

Remarquons que si $b = 1$ alors l'assertion (ii) est vérifiée car un des entiers $n_1, \dots, n_k$ doit être égal à 1 et donc $P(x, x) > 0$. De même (ii) est vraie si $b = 0$ car alors $a = 1$. Par conséquent, il suffit de considérer le cas $b > 1$. Posons

$$\mathbf{n}(x) = b^2 - 1 > 0.$$

Alors pour tout $n \geq \mathbf{n}(x)$, la division euclidienne de n par $b(x)$ s'écrit

$$n = db + r \quad \text{avec} \quad d \geq r \quad \text{et} \quad 0 \leq r \leq b(x) - 1$$

Par la relation $a - b = 1$, on peut réécrire n comme combinaison linéaire de $n_1, \dots, n_k$ à coefficients dans $\mathbb{N}$

$$n = (d - r)b + ra = (d - r) \sum_{i=1}^k q_i^- n_i + r \sum_{i=1}^k q_i^+ n_i.$$

Comme $P^{n_i}(x, x) > 0$ pour tout $i \leq k$, on en déduit que $P^n(x, x) > 0$. □

5.2.2 Distance en variation et couplage

Dans les applications, il est important de quantifier la vitesse de relaxation de la mesure $\{\mathbb{P}_\mu(X_n = y)\}_{y \in E}$ vers la mesure invariante π . Il faut donc préciser la convergence du théorème 5.5. Pour cela nous commencerons par définir une distance entre les mesures sur E .

Définition 5.9. *Soient μ et ν deux mesures de probabilité sur un espace dénombrable E . On définit la distance en variation totale entre ces deux mesures par*

$$\|\mu - \nu\|_{VT} = \frac{1}{2} \sum_{x \in E} |\mu(x) - \nu(x)|.$$

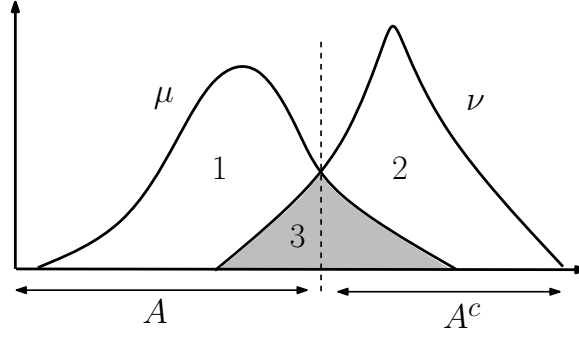


FIGURE 5.4 – Les densités des mesures μ et ν sont représentées. Leur partie commune est dessinée en gris et la distance en variation est proportionnelle à l'aire des zones blanches 1 et 2. Si les 2 mesures étaient identiques les zones 1 et 2 n'existeraient pas et leur distance en variation serait nulle. Inversement, si les supports des mesures sont disjoints leur distance est maximale.

Cette distance s'interprète comme la moitié de l'aire des régions 1 et 2 sur la figure 5.4. Le point clef de la preuve de convergence du théorème 5.5 résidait dans la construction d'un couplage entre les trajectoires. Nous allons maintenant revenir sur la notion de couplage et montrer qu'elle est intimement liée à la distance en variation totale.

Un couplage entre les deux mesures de probabilité μ et ν est une paire de variables aléatoires (X, Y) telles que X ait pour distribution μ et Y pour distribution ν

$$\mathbb{P}(X = x) = \sum_{y \in E} \hat{\mathbb{P}}(X = x, Y = y) = \mu(x) \quad (5.10)$$

$$\mathbb{P}(Y = y) = \sum_{x \in E} \hat{\mathbb{P}}(X = x, Y = y) = \nu(y) \quad (5.11)$$

où $\hat{\mathbb{P}}$ est la probabilité jointe des deux variables X et Y . Il existe de multiples façons de coupler deux mesures. Supposons par exemple que $\mu = \nu = \frac{1}{2}(\delta_0 + \delta_1)$. Un couplage possible est de choisir X et Y indépendamment

$$\forall x, y \in \{0, 1\}, \quad \hat{\mathbb{P}}(X = x, Y = y) = \frac{1}{4} \quad \Rightarrow \quad \hat{\mathbb{P}}(X \neq Y) = \frac{1}{2}.$$

Un autre couplage consiste à corrélérer fortement les 2 variables en choisissant X selon une loi de Bernoulli de paramètre 1/2 puis en posant $Y = X$

$$\forall x, y \in \{0, 1\}, \quad \hat{\mathbb{P}}(X = x, Y = y) = \frac{1}{2} \mathbf{1}_{y=x} \quad \Rightarrow \quad \hat{\mathbb{P}}(X \neq Y) = 0.$$

Les deux couplages respectent la propriété des lois marginales (5.10), mais leurs lois jointes sont très différentes. Certains couplages sont plus intéressants que d'autres comme le montre le lemme qui suit.

Lemme 5.10. Soient μ et ν deux mesures de probabilité sur un espace dénombrable E , alors

$$\|\mu - \nu\|_{VT} = \max_{B \subseteq E} |\mu(B) - \nu(B)| \quad (5.12)$$

$$= \inf \left\{ \hat{\mathbb{P}}(X \neq Y); \quad (X, Y) \text{ est un couplage de } \mu \text{ et } \nu \right\} \quad (5.13)$$

où l'infimum est pris sur tous les couplages possibles de μ et ν . Les couplages qui réalisent l'égalité sont dits optimaux.

Dans l'exemple précédent des mesures $\mu = \nu = \frac{1}{2}(\delta_0 + \delta_1)$, le second couplage est optimal mais pas le premier. Dans la suite du cours, seule l'identité (5.13) sera utilisée.

Démonstration. Commençons par montrer l'identité (5.12). On définit le sous-ensemble A (cf. figure 5.4) comme

$$A = \{x \in E; \quad \mu(x) \geq \nu(x)\}.$$

On a

$$\begin{aligned} \|\mu - \nu\|_{VT} &= \frac{1}{2} \sum_{x \in E} |\mu(x) - \nu(x)| = \frac{1}{2} \sum_{x \in A} \mu(x) - \nu(x) - \frac{1}{2} \sum_{x \in A^c} \mu(x) - \nu(x) \\ &= \frac{1}{2} (\mu(A) - \nu(A) - \mu(A^c) + \nu(A^c)) = \mu(A) - \nu(A) = -\mu(A^c) + \nu(A^c) \end{aligned}$$

où on a utilisé $1 = \mu(A) + \mu(A^c) = \nu(A) + \nu(A^c)$. Par ailleurs pour tout $B \subset E$

$$\mu(B) - \nu(B) \leq \mu(B \cap A) - \nu(B \cap A) \leq \mu(A) - \nu(A)$$

et aussi

$$\nu(B) - \mu(B) \leq \nu(A^c) - \mu(A^c).$$

On en déduit que

$$\max_{B \subset E} |\mu(B) - \nu(B)| = \mu(A) - \nu(A) = \|\mu - \nu\|_{VT}. \quad (5.14)$$

Pour montrer (5.13), vérifions d'abord que

$$\|\mu - \nu\|_{VT} \leq \inf \left\{ \widehat{\mathbb{P}}(X \neq Y); \quad (X, Y) \text{ est un couplage de } \mu \text{ et } \nu \right\}. \quad (5.15)$$

Soit B un sous-ensemble de E

$$\begin{aligned} \mu(B) - \nu(B) &= \mathbb{P}(X \in B) - \mathbb{P}(Y \in B) \\ &= \widehat{\mathbb{P}}(X \in B, Y \notin B) + \widehat{\mathbb{P}}(X \in B, Y \in B) - \mathbb{P}(Y \in B) \\ &\leq \widehat{\mathbb{P}}(X \in B, Y \notin B) \leq \widehat{\mathbb{P}}(X \neq Y). \end{aligned}$$

Par symétrie

$$\nu(B) - \mu(B) \leq \widehat{\mathbb{P}}(X \neq Y).$$

Il suffit d'utiliser l'identité (5.12) pour en déduire l'inégalité (5.15).

Pour montrer la réciproque, il suffit de construire un couplage qui réalise l'égalité dans (5.13). On définit

$$p = \sum_{x \in E} \inf\{\mu(x), \nu(x)\}$$

qui s'interprète comme l'aire de la région 3 dans la figure 5.4. En utilisant la relation (5.14), on peut réécrire p sous la forme

$$\begin{aligned} p &= \sum_{\substack{x \in E \\ \mu(x) \leq \nu(x)}} \mu(x) + \sum_{\substack{x \in E \\ \mu(x) > \nu(x)}} \nu(x) = 1 + \sum_{\substack{x \in E \\ \mu(x) > \nu(x)}} \nu(x) - \mu(x) = 1 - (\mu(A) - \nu(A)) \\ &= 1 - \|\mu - \nu\|_{VT}. \end{aligned}$$

Le couplage consiste à choisir une variable B de loi de Bernoulli

$$\mathbb{P}(B = 0) = p, \quad \mathbb{P}(B = 1) = 1 - p$$

et à construire X, Y en fonction de la valeur de B .

(i) Si $B = 0$, alors on choisit une variable Z selon la probabilité sur E

$$m_3(x) = \frac{1}{p} \inf\{\mu(x), \nu(x)\}$$

qui est concentrée dans la région 3 de la figure 5.4. On pose ensuite $X = Y = Z$. Remarquons que si $p = 0$ alors B n'est jamais égal à 0.

(ii) Si $B = 1$, alors on choisit X selon la probabilité

$$m_1(x) = \begin{cases} \frac{\mu(x) - \nu(x)}{\|\mu - \nu\|_{VT}} & \text{si } \mu(x) > \nu(x) \\ 0 & \text{sinon} \end{cases}$$

et Y est choisi indépendamment selon la probabilité

$$m_2(x) = \begin{cases} \frac{\nu(x) - \mu(x)}{\|\mu - \nu\|_{VT}} & \text{si } \nu(x) > \mu(x) \\ 0 & \text{sinon} \end{cases}$$

La relation (5.14) permet de vérifier que les 2 mesures sont normalisées par 1.

Cette procédure construit bien un couplage car X et Y ont les bonnes lois marginales

$$pm_3(x) + (1 - p)m_1(x) = \mu(x) \quad \text{et} \quad pm_3(x) + (1 - p)m_2(x) = \nu(x).$$

Les mesures m_1 et m_2 ont des supports disjoints (associés aux régions 1 et 2 de la figure 5.4). Par conséquent $X \neq Y$ si et seulement si $B = 1$. L'égalité dans (5.13) est donc bien vérifiée car

$$\widehat{\mathbb{P}}(X \neq Y) = \mathbb{P}(B = 1) = 1 - p = \|\mu - \nu\|_{VT}.$$

□

Le Lemme 5.10 va nous permettre de renforcer le théorème 5.5 sous une hypothèse introduite par Doeblin.

Théorème 5.11. Soit $\{X_n\}_{n \geq 0}$ une chaîne de Markov irréductible, apériodique sur un espace E dénombrable. On suppose que sa matrice de transition vérifie la condition de Doeblin, i.e. qu'il existe $r \geq 1$, $\delta > 0$ et une mesure de probabilité ν sur E tels que

$$P^r(x, z) \geq \delta \nu(z) \quad \text{pour tous } x, z \text{ dans } E. \quad (5.16)$$

Alors la chaîne de Markov admet une mesure de probabilité invariante π vers laquelle la distribution de la chaîne de Markov converge exponentiellement vite (uniformément par rapport aux états initiaux)

$$\sup_{x \in E} \|P^n(x, \cdot) - \pi\|_{VT} = \frac{1}{2} \sup_{x \in E} \sum_{y \in E} |P^n(x, y) - \pi(y)| \leq (1 - \delta)^{\lfloor n/r \rfloor}$$

où $P^n(x, \cdot)$ est la distribution au temps n en partant de x et $\lfloor \cdot \rfloor$ représente la partie entière. Pour une chaîne de Markov distribuée initialement selon μ , on a aussi

$$\left\| \sum_{x \in E} \mu(x) P^n(x, \cdot) - \pi(\cdot) \right\|_{VT} = \frac{1}{2} \sum_{y \in E} |\mathbb{P}_\mu(X_n = y) - \pi(y)| \leq (1 - \delta)^{\lfloor n/r \rfloor}. \quad (5.17)$$

Dans une chaîne de Markov irréductible, tous les états communiquent mais la probabilité de passer d'un état vers un autre peut être arbitrairement petite (3.6). La condition de Doeblin (5.16) impose une borne inférieure uniforme sur les probabilités de transition entre 2 états. En particulier, si l'espace d'états E est fini, une chaîne de Markov irréductible et apériodique satisfait toujours la condition de Doeblin. En effet, il existe $r \geq 1$ tel que pour tout couple x, y de E , on ait $P^r(x, y) > 0$. Il suffit de choisir

$$\delta = \sum_{y \in E} \min_{x \in E} P^r(x, y) > 0 \quad \text{et} \quad \nu(y) = \frac{1}{\delta} \min_{x \in E} P^r(x, y). \quad (5.18)$$

Démonstration. La preuve se décompose en 3 temps.

Étape 1. Existence d'une mesure invariante.

Vérifions d'abord que la condition de Doeblin (5.16) implique l'existence d'une mesure invariante. Comme il existe au moins un état z tel que $\nu(z) > 0$, on peut majorer la probabilité de ne pas repasser en z avant un temps kr . En utilisant successivement la propriété de Markov au temps $(k-1)r$ et la condition de Doeblin (5.16), on obtient la décroissance exponentielle

$$\begin{aligned} \mathbb{P}_z(T_z^+ > kr) &= \sum_{x \in E} \mathbb{P}_z(T_z^+ > (k-1)r, X_{(k-1)r} = x) \mathbb{P}(T_z^+ > r \mid X_0 = x) \\ &\leq \sum_{x \in E} \mathbb{P}_z(T_z^+ > (k-1)r, X_{(k-1)r} = x) (1 - \delta \nu(z)) \\ &\leq \mathbb{P}_z(T_z^+ > (k-1)r) (1 - \delta \nu(z)) \leq (1 - \delta \nu(z))^k. \end{aligned}$$

Ceci implique que l'état z est récurrent positif car

$$\mathbb{E}_z(T_z^+) = \sum_{n \geq 1} \mathbb{P}_z(T_z^+ \geq n) \leq r \sum_{k \geq 0} \mathbb{P}_z(T_z^+ > kr) < \infty$$

et par le théorème 4.6 qu'il existe une mesure invariante.

Étape 2. *Convergence pour une condition de Doeblin simplifiée.*

Pour démontrer la convergence, nous supposons d'abord que (5.16) est vraie avec $r = 1$, c'est-à-dire

$$P(x, z) \geq \delta \nu(z) \quad \text{pour tous } x, z \text{ dans } E.$$

La borne inférieure $\delta \nu(z)$ peut être interprétée comme la région grise de la figure 5.4 : c'est la partie commune des mesures de probabilité $P(x, \cdot)$ pour différentes valeurs de x .

Nous allons construire un couplage qui s'inspire de la preuve du lemme 5.10. Soient $\{X_n\}_{n \geq 0}$ et $\{Y_n\}_{n \geq 0}$ deux réalisations de la chaîne de Markov, la première partant de x et l'autre de y . Supposons que les deux trajectoires soient construites jusqu'au temps n . Si $X_n = Y_n$, on préserve l'égalité au temps $n + 1$ (cf. figure 5.3) et on pose

$$\hat{\mathbb{P}}(X_{n+1} = x_{n+1}, Y_{n+1} = y_{n+1} | X_n = x_n, Y_n = x_n) = P(x_n, x_{n+1}) \mathbf{1}_{\{y_{n+1} = x_{n+1}\}}.$$

Si $X_n \neq Y_n$ alors au pas de temps $n + 1$, on tire au hasard une variable aléatoire B_{n+1} de Bernoulli de paramètre $1 - \delta$ (indépendante de tout le passé)

$$\mathbb{P}(B_{n+1} = 1) = 1 - \delta, \quad \mathbb{P}(B_{n+1} = 0) = \delta.$$

— Si $B_{n+1} = 0$, on choisit un site z de E selon la loi ν et on pose

$$X_{n+1} = Y_{n+1} = z.$$

Ce choix ne dépend pas des valeurs de X_n et Y_n .

— Si $B_{n+1} = 1$, X_{n+1} et Y_{n+1} sont choisis indépendamment en fonction de lois qui dépendent des valeurs X_n et Y_n

$$\hat{\mathbb{P}}(X_{n+1} = x_{n+1} | X_n = x_n, B_{n+1} = 1) = \frac{1}{1 - \delta} (P(x_n, x_{n+1}) - \delta \nu(x_{n+1}))$$

$$\hat{\mathbb{P}}(Y_{n+1} = y_{n+1} | Y_n = y_n, B_{n+1} = 1) = \frac{1}{1 - \delta} (P(y_n, y_{n+1}) - \delta \nu(y_{n+1}))$$

On remarque que la matrice modifiée est bien une matrice de transition car ses termes sont positifs et

$$\forall z_1 \in E, \quad \sum_{z_2 \in E} \frac{1}{1 - \delta} (P(z_1, z_2) - \delta \nu(z_2)) = 1.$$

De plus les processus $\{X_n\}_{n \geq 0}$ et $\{Y_n\}_{n \geq 0}$ sont chacun des chaînes de Markov de matrice de transition P car

$$\begin{aligned} \mathbb{P}(X_{n+1} = x_{n+1} | X_n = x_n) &= \delta \mathbb{P}(X_{n+1} = x_{n+1} | X_n = x_n, B_{n+1} = 0) \\ &\quad + (1 - \delta) \mathbb{P}(X_{n+1} = x_{n+1} | X_n = x_n, B_{n+1} = 1) \\ &= \delta \nu(x_{n+1}) + (P(x_n, x_{n+1}) - \delta \nu(x_{n+1})) = P(x_n, x_{n+1}). \end{aligned}$$

Le couplage étant construit, il ne reste plus qu'à estimer le temps d'arrêt T quand les marches se rejoignent (cf. figure 5.3)

$$T = \inf \left\{ n \geq 0; \quad X_n = Y_n \right\}.$$

La condition de Doeblin (5.16) assure que pour tous $x_0 \neq y_0$ dans E

$$\widehat{\mathbb{P}}(T = 1) = \widehat{\mathbb{P}}(X_1 = Y_1 \mid X_0 = x_0, Y_0 = y_0) \geq \mathbb{P}(B_1 = 0) \geq \delta.$$

En utilisant la propriété de Markov et cette borne inférieure, on obtient

$$\begin{aligned} \widehat{\mathbb{P}}(T > n) &= \sum_{\substack{z, z' \in E \\ z \neq z'}} \widehat{\mathbb{P}}(T > (n-1), X_{n-1} = z, Y_{n-1} = z') \widehat{\mathbb{P}}(T > 1 \mid X_0 = z, Y_0 = z') \\ &\leq \widehat{\mathbb{P}}(T > (n-1)) (1 - \delta) \leq (1 - \delta)^n. \end{aligned}$$

Il est important de remarquer que cette borne est uniforme pour tous les états de départ x et y de E . Pour estimer l'écart entre les distributions au temps n , il ne reste plus qu'à utiliser le lemme 5.10 en choisissant le couplage que nous venons de construire

$$\left\| P^n(x, \cdot) - P^n(y, \cdot) \right\|_{VT} \leq \widehat{\mathbb{P}}(X_n \neq Y_n) = \widehat{\mathbb{P}}(T > n) \leq (1 - \delta)^n. \quad (5.19)$$

Si la chaîne de Markov $\{Y_n\}_{n \geq 0}$ était issue de la mesure invariante π alors sa distribution à tout temps serait égale à π

$$\forall y \in E, \quad \pi(y) = \mathbb{P}_\pi(Y_n = y) = \sum_{z \in E} \pi(z) P^n(z, y).$$

Pour conclure le théorème, il suffit de considérer un couplage entre $\{X_n\}_{n \geq 0}$ partant de $X_0 = x$ et $\{Y_n\}_{n \geq 0}$ distribuée initialement sous π . Plus généralement si X_0 est distribuée sous la mesure μ , on obtient par le même argument de couplage l'inégalité (5.17).

Étape 3. Cas général.

Supposons maintenant que la condition de Doeblin (5.16) soit satisfaite avec un paramètre $r \geq 1$. On remarque que la chaîne de Markov $\{\tilde{X}_k = X_{kr}\}_{k \geq 0}$ a pour matrice de transition P^r et pour mesure invariante π . On peut décomposer tout entier n sous la forme $n = kr + \ell$ avec $\ell \in \{0, \dots, r-1\}$ et écrire la distribution au temps n d'une chaîne partant de x comme la distribution de $\tilde{X}_k$ partant initialement de la mesure $P^\ell(x, \cdot)$

$$\mathbb{P}_x(X_n = y) = \sum_{z \in E} P^\ell(x, z) \mathbb{P}_z(\tilde{X}_k = y).$$

Cette égalité est une simple réécriture de l'équation de Chapman-Kolmogorov $P^n = P^\ell P^{kr}$. La chaîne $\{\tilde{X}_k\}_{k \geq 0}$ vérifie le critère de Doeblin pour $r = 1$, il suffit donc d'appliquer à $\{\tilde{X}_k\}_{k \geq 0}$ le résultat de la seconde étape pour établir la convergence exponentielle de la distribution $P^n(x, \cdot)$. $\square$

5.2.3 Vitesses de convergence

Dans de nombreuses applications (cf. chapitre 6) la vitesse de convergence vers la mesure d'équilibre est fondamentale, car elle permet de déterminer le temps nécessaire pour qu'une simulation donne un résultat avec la précision voulue. La condition de Doeblin définie au théorème 5.11 est importante, car elle fournit un cadre théorique simple pour montrer une convergence exponentielle. Cependant, de meilleures vitesses de convergence peuvent souvent être prouvées par une étude spécifique de chaque modèle. Nous allons l'illustrer sur un exemple.

Marche aléatoire.

Considérons une marche aléatoire fainéante symétrique sur le domaine périodique $E = \{1, \dots, L\}$ de matrice de transition P donnée par

$$\forall i \in E, \quad P(i, i+1) = P(i, i-1) = 1/4, \quad P(i, i) = 1/2$$

où on identifie $L+1 \equiv 1$ et $0 \equiv L$. Comme la probabilité de rester sur place est non nulle, cette chaîne de Markov est apériodique.

On construit le couplage (Y_n^1, Y_n^2) sur $E \times E$ partant initialement de (x, y) . Au temps n , si $Y_n^1 = Y_n^2$ alors les 2 coordonnées évoluent de la même manière selon la matrice de transition P et on a $Y_{n+1}^1 = Y_{n+1}^2$ (cf. figure 5.3). Si $Y_n^1 \neq Y_n^2$, on choisit une variable $B_{n+1} \in \{1, 2\}$ avec probabilité $1/2$, puis seule la coordonnée B_{n+1} est mise à jour et saute à droite ou à gauche avec probabilité $1/2$: $Y_{n+1}^{B_{n+1}} = Y_n^{B_{n+1}} \pm 1$.

On note T le premier temps où les deux trajectoires se rencontrent. En appliquant le lemme 5.10, on peut donc contrôler la convergence en fonction de T

$$\|P^n(x, \cdot) - P^n(y, \cdot)\|_{VT} \leq \hat{\mathbb{P}}(T > n) \leq \frac{1}{n} \hat{\mathbb{E}}(T)$$

où $\hat{\mathbb{E}}$ correspond à l'espérance pour la mesure jointe du couplage (Y_n^1, Y_n^2) .

Supposons $x > y$ et analysons la différence $Z_n = Y_n^1 - Y_n^2$. On constate que $\{Z_n\}_{n \geq 0}$ est une marche aléatoire partant de $x - y$ et sautant à chaque pas de temps à gauche ou à droite avec probabilité $1/2$ tant qu'elle n'a pas atteint 0 ou L , c'est-à-dire tant que les marches Y_n^1, Y_n^2 ne se sont pas rejointes. Le temps d'arrêt T correspond donc au moment où le processus Z_n est absorbé en 0 ou en L , c'est l'analogie du temps défini dans la ruine du joueur dont l'espérance a été calculée section 2.5.2 en (2.31)

$$\hat{\mathbb{E}}(T) = (x - y)(L - (x - y)).$$

On obtient donc uniformément en x et y

$$\|P^n(x, \cdot) - P^n(y, \cdot)\|_{VT} \leq \frac{L^2}{4n}.$$

Ceci montre que pour une taille L assez grande, la chaîne de Markov sera proche de l'équilibre dès que le temps est de l'ordre de L^2 . Cet ordre de grandeur est optimal comme l'indique le théorème central limite : une marche aléatoire au temps n visite des régions de taille $\sqrt{n}$, par conséquent pour recouvrir le domaine $\{1, \dots, L\}$, il faudra au moins attendre des temps de l'ordre L^2 .

Comparons maintenant ce résultat avec celui donné par le théorème 5.11. La condition de Doeblin (5.16) suppose de trouver un paramètre r tel que tous les états puissent être connectés en r sauts. Il faut au minimum choisir $r \geq L/2$. Pour $r = L/2$, la constante δ est alors de l'ordre $\frac{1}{4^{L/2}}$. Pour ces valeurs, le théorème 5.11 implique

$$\sup_{x \in E} \|P^n(x, \cdot) - \pi\|_{VT} \leq \left(1 - \frac{1}{4^{L/2}}\right)^{\lfloor \frac{n}{L/2} \rfloor} \simeq \exp\left(-c \frac{n}{2^L L}\right)$$

où la dernière égalité est un équivalent pour L grand et c est une constante. Dans cet exemple la condition de Doeblin assure seulement la convergence pour des temps de l'ordre $2^L L$ et d'un point de vue pratique, elle n'est pas pertinente car elle ne prédit pas l'ordre L^2 .

Considérons maintenant une marche modifiée qui au lieu de rester sur place avec probabilité $1/2$ peut sauter uniformément sur tous les sites selon la probabilité de transition

$$\forall i, j \in E, \quad P(i, j) = \frac{1}{4} \mathbf{1}_{\{j=i \pm 1\}} + \frac{1}{2L}$$

où on identifie $L+1 \equiv 1$ et $0 \equiv L$. Dans ce cas la condition de Doeblin s'applique avec $r = 1$, $\delta = 1/2$ et $\nu(y) = 1/L$. Par le théorème 5.11, on obtient

$$\sup_{x \in E} \|P^n(x, \cdot) - \pi\|_{VT} \leq \frac{1}{2^n}.$$

Cette fois la convergence est beaucoup plus rapide, elle ne dépend plus de L et la condition de Doeblin fournit une information précise. Cette modification des probabilités de transition est en fait identique à celle introduite dans la matrice de transition de Google (5.8) afin d'accélérer la vitesse de convergence.

Mélange de cartes.

Pour mélanger un jeu de N cartes, on répète plusieurs fois la procédure suivante : une carte est choisie au hasard et est déplacée en haut du paquet (si la carte en haut du paquet est choisie, elle reste sur place). Les cartes sont numérotées de 1 à N et cette procédure définit une chaîne de Markov $\{\Sigma_k\}_{k \geq 0}$ sur l'ensemble des permutations $\mathcal{S}_N$ de N éléments. On notera P sa matrice de transition et $P^k(\text{Id}, \cdot)$ la distribution au temps k de la chaîne partant de la configuration $\text{Id} = \{1, \dots, N\}$ de $\mathcal{S}_N$. L'exemple ci-dessous représente un jeu de 5 cartes après deux pas de temps. Les cartes 4 et 1 ont été choisies successivement avec probabilité $1/5$.

$$\Sigma_0 = \begin{vmatrix} 1 \\ 2 \\ 3 \\ 4 \\ 5 \end{vmatrix} \xrightarrow[\text{déplacée}]{\text{carte 4}} \Sigma_1 = \begin{vmatrix} 4 \\ 1 \\ 2 \\ 3 \\ 5 \end{vmatrix} \xrightarrow[\text{déplacée}]{\text{carte 1}} \Sigma_2 = \begin{vmatrix} 1 \\ 4 \\ 2 \\ 3 \\ 5 \end{vmatrix}$$

La mesure invariante π de cette chaîne de Markov est uniforme sur les $N!$ permutations possibles. Pour le voir il suffit de remarquer qu'une permutation σ donnée a exactement N antécédents possibles et que chacun peut atteindre σ avec probabilité $1/N$

$$\pi(\sigma) = \sum_{\eta \in \mathcal{S}_N} \pi(\eta) P(\eta, \sigma) = \frac{1}{N!} \sum_{\eta \in \mathcal{S}_N} P(\eta, \sigma) = \frac{1}{N!}.$$

La proposition suivante permet d'estimer la vitesse de convergence pour des jeux de cartes de grande taille (proche de l'infini !)

Proposition 5.12. *Pour tout $\varepsilon > 0$, il suffit de mélanger $k_N \geq (1 + \varepsilon)N \log N$ fois un paquet de N cartes pour que sa distribution soit uniforme (quand N est très grand)*

$$\lim_{N \rightarrow \infty} \|P^{k_N}(\text{Id}, \cdot) - \pi\|_{VT} = 0. \quad (5.20)$$

Inversement, il ne faut pas mélanger moins de fois le paquet car pour toute suite $\ell_N \leq (1 - \varepsilon)N \log N$ la distribution au temps ℓ_N est encore très loin de π

$$\lim_{N \rightarrow \infty} \|P^{\ell_N}(\text{Id}, \cdot) - \pi\|_{VT} = 1. \quad (5.21)$$

Cette proposition illustre un phénomène de seuil (ou cutoff) que l'on retrouve dans de nombreuses chaînes de Markov : la convergence pour la norme en variation totale ne s'opère pas régulièrement au cours du temps. Dans l'exemple du jeu de cartes, la distribution de la chaîne de Markov reste très loin de la mesure invariante avant le temps $N \log N$, puis devient très proche après $N \log N$. La convergence se produit très rapidement autour de $N \log N$ dans un intervalle de temps plus court que $\varepsilon N \log N$ (cf. figure 5.5).

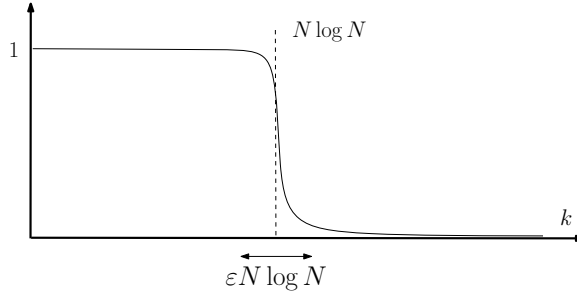


FIGURE 5.5 – Distance en variation en fonction du temps k entre $P^k(\text{Id}, \cdot)$ et la mesure invariante. La convergence est localisée autour de $N \log N$.

Démonstration. Commençons par prouver la convergence vers l'équilibre, i.e. la limite (5.20). Pour cela, construisons un couplage entre deux trajectoires de la chaîne $\{\Sigma_k^1\}$ et $\{\Sigma_k^2\}$. La première a pour donnée initiale $\Sigma_0^1 = \text{Id}$ et l'état initial Σ_0^2 de la seconde est choisi selon π . À chaque pas de temps, on tire un numéro de carte au hasard entre 1 et N et cette carte est posée en haut du paquet dans les deux jeux.

$$\Sigma_0^1 = \begin{vmatrix} 1 \\ 2 \\ 3 \\ 4 \\ 5 \end{vmatrix} \quad \text{et} \quad \Sigma_0^2 = \begin{vmatrix} 5 \\ 4 \\ 1 \\ 3 \\ 2 \end{vmatrix} \quad \xrightarrow[\text{de mélange}]{\text{une étape}} \quad \Sigma_1^1 = \begin{vmatrix} 3 \\ 1 \\ 2 \\ 4 \\ 5 \end{vmatrix} \quad \text{et} \quad \Sigma_1^2 = \begin{vmatrix} 3 \\ 5 \\ 4 \\ 1 \\ 2 \end{vmatrix}$$

Dans l'exemple ci-dessus la carte 3 a été déplacée dans les deux configurations. Chacun des jeux évolue selon la bonne probabilité de transition. Si une carte est choisie au temps

k alors elle occupera la même position dans les deux jeux à tous les temps suivants. En effet, elle sera mise au-dessus du paquet au temps k puis descendra dès que d'autres cartes seront ajoutées. Elle pourra se retrouver encore au-dessus du paquet si elle est choisie une seconde fois, mais tous ses déplacements seront les mêmes dans les deux jeux. Par conséquent si toutes les cartes ont été choisies au moins une fois, les deux jeux sont identiques. On définit le temps d'arrêt

$$\tau = \inf \{k \geq 0, \text{ toutes les cartes ont été choisies au moins une fois au temps } k\}$$

qui permet de borner le temps de couplage

$$T = \inf \{k \geq 0, \Sigma_k^1 = \Sigma_k^2\} \leq \tau.$$

Pour estimer la vitesse de convergence, il suffit donc de contrôler τ

$$\|P^k(\text{Id}, \cdot) - \pi\|_{VT} \leq \widehat{\mathbb{P}}(T > k) \leq \widehat{\mathbb{P}}(\tau > k). \quad (5.22)$$

Le temps d'arrêt τ se décompose facilement en

$$\tau = \tau_1 + \tau_2 + \cdots + \tau_N \quad (5.23)$$

où τ_i est le temps nécessaire pour obtenir la $i^{\text{ème}}$ nouvelle carte sachant qu'on a déjà $i - 1$ cartes différentes. La première étape est très simple $\tau_1 = 1$. Supposons que $i - 1$ cartes différentes ont été trouvées, la probabilité de découvrir une nouvelle carte au temps k suit une loi géométrique de paramètre $1 - (i - 1)/N$

$$\mathbb{P}(\tau_i = k) = \left(\frac{i-1}{N}\right)^{k-1} \left(1 - \frac{i-1}{N}\right)$$

car la probabilité de trouver une nouvelle carte est $1 - \frac{i-1}{N}$. On en déduit que

$$\mathbb{E}(\tau_{N-i}) = \frac{N}{i+1} \quad \text{et} \quad \text{Var}(\tau_{N-i}) = \mathbb{E}(\tau_{N-i}^2) - \mathbb{E}(\tau_{N-i})^2 = \frac{N(N-i-1)}{(i+1)^2} \leq \frac{N^2}{(i+1)^2}.$$

Ceci permet de calculer l'espérance asymptotique de τ quand N est grand

$$\mathbb{E}(\tau) = \sum_{i=1}^N \mathbb{E}(\tau_i) = \sum_{i=0}^{N-1} \frac{N}{i+1} \simeq N \log N$$

ainsi qu'une borne sur sa variance car les τ_i sont indépendants

$$\text{Var}(\tau) = \mathbb{E}((\tau - \mathbb{E}(\tau))^2) = \sum_{i=1}^N \text{Var}(\tau_i) \leq \sum_{i=0}^{N-1} \frac{N^2}{(i+1)^2} \leq cN^2$$

où c est une constante.

Supposons que $k > \mathbb{E}(\tau)$ alors l'inégalité de Tchebyshev permet d'écrire

$$\mathbb{P}(\tau \geq k) = \mathbb{P}(\tau - \mathbb{E}(\tau) \geq k - \mathbb{E}(\tau)) \leq \frac{\mathbb{E}((\tau - \mathbb{E}(\tau))^2)}{(k - \mathbb{E}(\tau))^2} \leq \frac{cN^2}{(k - \mathbb{E}(\tau))^2}.$$

En utilisant (5.22), ce résultat fournit une vitesse de convergence vers la mesure invariante pour N fixé

$$\|P^k(\text{Id}, \cdot) - \pi\|_{VT} \leq \frac{cN^2}{(k - \mathbb{E}(\tau))^2}.$$

Pour N grand, si $k_N \geq (1 + \varepsilon)N \log N > \mathbb{E}(\tau)$, on en déduit que

$$\|P^{k_N}(\text{Id}, \cdot) - \pi\|_{VT} \leq \frac{c(\varepsilon)N^2}{(N \log N)^2} \leq \frac{c(\varepsilon)}{(\log N)^2}$$

où $c(\varepsilon)$ est une constante dépendante de ε . Quand N tend vers l'infini, on retrouve la limite (5.20).

Pour estimer la borne inférieure (5.21) sur le temps de mélange, il suffit de remarquer que pour un temps $\ell_N \leq (1 - \varepsilon)N \log N$ moins de $N^{1-\frac{\varepsilon}{2}}$ cartes auront été déplacées. Ceci veut dire que le paquet contiendra un nombre important de cartes ayant conservé leur ordre initial et la convergence complète n'aura pas eu lieu.

En utilisant la notation (5.23), on peut estimer le temps nécessaire pour que $N^{1-\frac{\varepsilon}{2}}$ cartes aient été choisies. On notera ce temps

$$\hat{\tau} = \tau_1 + \dots + \tau_{N^{1-\frac{\varepsilon}{2}}}.$$

L'espérance et la variance de ce temps se comportent asymptotiquement en N comme

$$\mathbb{E}(\hat{\tau}) \simeq (1 - \frac{\varepsilon}{2})N \log N \quad \text{et} \quad \mathbb{E}(\hat{\tau}^2) - \mathbb{E}(\hat{\tau})^2 \leq cN^2 \quad (5.24)$$

Nous allons montrer que l'espérance est une bonne estimation du temps nécessaire pour découvrir $N^{1-\frac{\varepsilon}{2}}$ et en particulier qu'avec grande probabilité $\hat{\tau}$ sera plus grand que $(1 - \varepsilon)N \log N$. Pour cela définissons

$$\tilde{\tau} = \hat{\tau} - (1 - \varepsilon)\mathbb{E}(\hat{\tau})$$

et utilisons la borne du second moment (prouvée dans le lemme 5.13 ci-dessous) qui implique pour $\lambda = 0$ que

$$\mathbb{P}(\tilde{\tau} > 0) \geq \frac{\mathbb{E}(\tilde{\tau})^2}{\mathbb{E}(\tilde{\tau}^2)} = \frac{\varepsilon^2 \mathbb{E}(\hat{\tau})^2}{\varepsilon^2 \mathbb{E}(\hat{\tau})^2 + \mathbb{E}(\hat{\tau}^2) - \mathbb{E}(\hat{\tau})^2}.$$

Par les estimations (5.24), le terme de droite tend vers 1 quand N tend vers l'infini et on en déduit que

$$\lim_{N \rightarrow \infty} \mathbb{P}\left(\hat{\tau} > (1 - \varepsilon)(1 - \frac{\varepsilon}{2})N \log N\right) = 1.$$

Supposons qu'exactement K cartes n'aient pas été choisies, alors les $N - K$ autres ont été déplacées en haut du paquet. Ces K cartes ont conservé leur ordre initial $i_1 < i_2 < \dots < i_K$ et se trouvent au bas du paquet. Le nombre d'arrangements possibles pour un tel évènement est

$$\binom{N}{K} (N - K)! = \frac{N!}{K!}.$$

Par conséquent la probabilité d'un tel évènement est $\frac{1}{K!}$. Soit A_N l'ensemble des permutations telles qu'au moins $N/2$ cartes soient ordonnées au bas du paquet. La probabilité $\pi(A_N)$ d'un tel évènement tend donc vers 0 quand N tend vers l'infini.

Il suffit maintenant d'associer les deux résultats précédents pour en déduire la borne inférieure. Pour toute suite de temps $\{\ell_N\}$ telle que $\ell_N \leq (1 - \varepsilon)N \log N$, il y aura eu au plus $N^{1-\varepsilon/2}$ cartes déplacées. Par conséquent avec grande probabilité, la chaîne de Markov Σ_{ℓ_N} sera dans l'ensemble A_N et

$$\lim_{N \rightarrow \infty} \mathbb{P}_{\text{Id}}(\Sigma_{\ell_N} \in A_N) - \pi(A_N) = 1.$$

Par l'identité (5.12)

$$\|P^{\ell_N}(\text{Id}, \cdot) - \pi\| \geq \mathbb{P}_{\text{Id}}(\Sigma_{\ell_N} \in A_N) - \pi(A_N).$$

On en déduit la borne inférieure (5.21). □

Lemme 5.13 (Inégalité du second moment). *Soit X une variable aléatoire telle que $\mathbb{E}(X) \geq 0$, alors*

$$0 \leq \lambda \leq 1, \quad \mathbb{P}(X > \lambda \mathbb{E}(X)) \geq (1 - \lambda)^2 \frac{\mathbb{E}(X)^2}{\mathbb{E}(X^2)}.$$

Démonstration. Pour le voir, on remarque que

$$\mathbb{E}(X) = \mathbb{E}(X 1_{\{X > \lambda \mathbb{E}(X)\}}) + \mathbb{E}(X 1_{\{X \leq \lambda \mathbb{E}(X)\}}) \leq \mathbb{E}(X 1_{\{X > \lambda \mathbb{E}(X)\}}) + \lambda \mathbb{E}(X).$$

Par conséquent

$$(1 - \lambda)\mathbb{E}(X) \leq \mathbb{E}(X 1_{\{X > \lambda \mathbb{E}(X)\}}).$$

On conclut par l'inégalité de Cauchy-Schwarz

$$(1 - \lambda)^2 \mathbb{E}(X)^2 \leq \mathbb{E}(X^2) \mathbb{E}(1_{\{X > \lambda \mathbb{E}(X)\}}).$$

□

Chapitre 6

Application aux algorithmes stochastiques

6.1 Optimisation

Dans de nombreuses applications, on souhaite minimiser une fonction $V : \mathbb{R}^K \rightarrow \mathbb{R}$ dont la structure est souvent complexe et dépend d'un grand nombre de paramètres $K \gg 1$ selon le problème à modéliser. Cette fonction sert par exemple à quantifier un coût en économie ou un rendement dans une réaction chimique, à optimiser des échanges dans un réseau informatique ou à déterminer des estimateurs en statistique (maximum de vraisemblance). On cherche aussi à identifier les valeurs où cette fonction prend son minimum

$$\text{Argmin } V = \{x \in \mathbb{R}^K; \quad V(x) = \inf_y V(y)\}.$$

Ce problème d'optimisation est purement déterministe et il peut être résolu par des méthodes analytiques. En particulier, les méthodes de *programmation linéaire* sont optimales pour résoudre des problèmes avec des contraintes linéaires. Dans le cas d'une fonction V strictement convexe, une méthode de descente de gradient [4] permet de déterminer le point x^* où la fonction atteint son minimum en suivant le flot de l'équation

$$\forall t \geq 0, \quad \dot{x}_t = -V'(x_t) \quad \text{alors} \quad \lim_{t \rightarrow \infty} x_t = x^*.$$

Par contre si la fonction V possède de nombreux minima locaux une telle méthode ne permettra pas de déterminer le minimum global facilement car la limite de x_t dépendra de l'état initial x_0 (cf. figure 6.1).

De nombreux problèmes d'optimisation nécessitent d'étudier des fonctions V particulièrement complexes, dépendant de multiples paramètres. Pour fixer les idées, considérons le cas d'école du *problème du voyageur de commerce*. Un voyageur de commerce doit visiter K clients dans K villes différentes et revenir à son point de départ en ne visitant chaque ville qu'une seule fois. Étant données les distances entre toutes les villes $\{d(i, j)\}_{1 \leq i \leq K, 1 \leq j \leq K}$, l'objectif est de minimiser le trajet à parcourir, c'est-à-dire

$$\min_{\sigma \in S_K} \{V(\sigma)\} \quad \text{avec} \quad V(\sigma) = \sum_{i=1}^K d(\sigma(i), \sigma(i+1)) \quad (6.1)$$

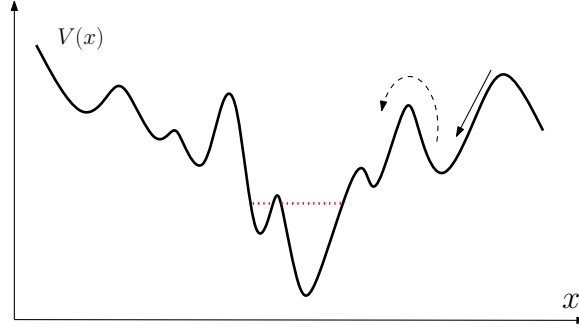


FIGURE 6.1 – Le schéma représente un potentiel $V(x)$ avec plusieurs minima locaux qui rendent la méthode de descente de gradient inefficace. Les algorithmes stochastiques permettent de franchir les barrières de potentiel (cf. la flèche en pointillés) pour atteindre le minimum global. Quand T est petit la mesure μ_T va se concentrer principalement autour des valeurs les plus basses de V par exemple sur le schéma sur les points situés sous la droite en pointillés.

où σ appartient à l'ensemble $\mathcal{S}_K$ des permutations de $\{1, \dots, K\}$. Chaque permutation σ correspond à un trajet entre les villes selon un certain ordre et comme le voyageur revient à son point de départ, on pose $\sigma(K+1) = \sigma(1)$.

Une méthode de recherche systématique du minimum consisterait à explorer tous les chemins possibles et à comparer leurs longueurs. Ceci conduirait à une complexité numérique gigantesque car l'ensemble des trajets entre les K villes est équivalent aux permutations d'un ensemble à K éléments. Il faudrait donc évaluer les longueurs des $K!$ trajets, ce qui est impossible numériquement dès que K devient grand. Ce problème appartient à la classe des problèmes *NP-complets* et il n'existe pas d'algorithme qui permette de le résoudre en un temps polynomial en K (sous l'hypothèse $P \neq NP$). Le problème du voyageur de commerce est un problème théorique qui sert souvent de référence pour tester des stratégies d'optimisation. Dans la pratique, il existe de nombreux problèmes d'optimisation similaires pour lesquels il est impossible d'obtenir la solution exacte en un temps raisonnable mais qui peuvent être résolus de manière approchée par des méthodes stochastiques. Nous reviendrons sur le problème du voyageur de commerce section 6.4.1.

Ce chapitre décrit des méthodes probabilistes pour déterminer le minimum d'une fonction V sur un espace discret E fini mais de cardinal très grand. Pour construire une solution approchée à ce problème déterministe, nous définissons μ_T la *mesure de Gibbs* associée au potentiel V et au paramètre $T > 0$

$$\forall x \in E, \quad \mu_T(x) = \frac{1}{Z_T} \exp\left(-\frac{1}{T}V(x)\right) \quad \text{avec} \quad Z_T = \sum_{y \in E} \exp\left(-\frac{1}{T}V(y)\right). \quad (6.2)$$

La mesure μ_T attribue une probabilité à chaque site de E et se concentre sur les minima de V quand T tend vers 0. Le résultat suivant est attribué à Pierre-Simon Laplace.

Lemme 6.1. Si $\mathcal{M}$ désigne l'ensemble des points de E où V atteint son minimum, on a

$$\forall x \in E, \quad \lim_{T \rightarrow 0} \mu_T(x) = \begin{cases} \frac{1}{\text{Card}(\mathcal{M})} & \text{si } x \in \mathcal{M} \\ 0 & \text{si } x \notin \mathcal{M} \end{cases}$$

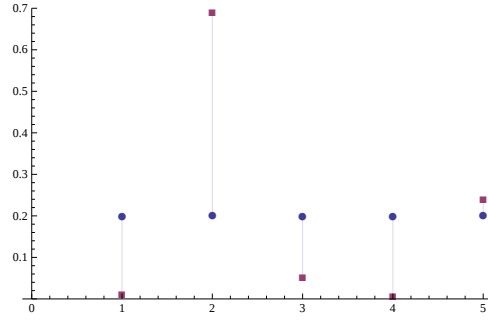


FIGURE 6.2 – On considère l’espace $E = \{1, \dots, 5\}$ et la fonction $V(1) = 87, V(2) = 4, V(3) = 55, V(4) = 99, V(5) = 25$. La distribution de la probabilité μ_T est représentée par des cercles pour $T = 10^4$ et des carrés pour $T = 2$. Quand T est très grand la mesure est presque distribuée uniformément, par contre pour T plus petit les valeurs les plus basses de V deviennent prépondérantes (cf. lemme 6.1).

Démonstration. Soit V^* le minimum de V , on peut réécrire la mesure de Gibbs (6.2)

$$\forall x \in E, \quad \mu_T(x) = \frac{1}{\sum_{y \in E} \exp\left(-\frac{1}{T}[V(y) - V^*]\right)} \exp\left(-\frac{1}{T}[V(x) - V^*]\right).$$

Dès que $x \notin \mathcal{M}$, on a $V(x) - V^* > 0$ et comme E est fini, seuls les termes dans $\mathcal{M}$ contribuent quand T tend vers 0. $\square$

La figure 6.2 illustre ce lemme et montre que pour T proche de 0, la mesure μ_T se concentre sur les points où V est minimum. Par conséquent en simulant des réalisations de la mesure μ_T pour T proche de 0, on obtiendra avec une grande probabilité une approximation de l’ensemble $\mathcal{M}$ où V atteint son minimum. La simulation de la mesure de probabilité μ_T sera l’objet de la section suivante.

6.2 Algorithmes stochastiques

6.2.1 Algorithme de Metropolis-Hastings

À première vue, la simulation de la mesure de Gibbs (6.2) suppose de calculer la distribution μ_T et donc d’évaluer $Z_T = \sum_{y \in E} \exp\left(-\frac{1}{T}V(y)\right)$. Dans la pratique, ceci est impossible à implémenter car il faudrait calculer toutes les valeurs de V pour un ensemble E de cardinal trop important. La méthode proposée en 1953 dans l’article [20] et améliorée par W. Hastings [15] en 1970 permet d’éviter cet écueil en simulant la mesure de Gibbs à l’aide d’une chaîne de Markov.

L’algorithme de Metropolis-Hastings permet de simuler une variable aléatoire sous une mesure de probabilité quelconque sur E . On note π cette mesure et on suppose que $\pi(x) > 0$ pour tout x de E . Pour réaliser la simulation, il faut se donner une matrice de transition Q irréductible sur E satisfaisant pour tous x, y de E

$$Q(x, y) > 0 \quad \Rightarrow \quad Q(y, x) > 0 \quad (6.3)$$

et une fonction croissante $h :]0, \infty[\rightarrow]0, 1]$ vérifiant $h(u) = uh(1/u)$. Par exemple on peut choisir

$$h(u) = \inf\{1, u\} \quad \text{ou} \quad h(u) = \frac{u}{1+u}.$$

Pour $x \neq y$, on pose

$$R(x, y) = \begin{cases} h\left(\frac{\pi(y)Q(y, x)}{\pi(x)Q(x, y)}\right) & \text{si } Q(x, y) \neq 0 \\ 0 & \text{sinon} \end{cases} \quad (6.4)$$

Ceci permet de construire la matrice de transition P définie par

$$\begin{cases} P(x, y) = Q(x, y)R(x, y) & \text{si } x \neq y \\ P(x, x) = 1 - \sum_{y \neq x} P(x, y) \end{cases} \quad (6.5)$$

L'algorithme de Metropolis-Hastings, décrit ci-dessous, permet de simuler une chaîne de Markov $\{X_n\}_{n \geq 0}$ de matrice de transition P :

Étape 0. Initialiser X_0

Étape $n + 1$.

Choisir y selon la loi $Q(X_n, y)$

Choisir U_{n+1} uniformément dans $[0, 1]$ (et indépendamment du passé)

Si $U_{n+1} < R(X_n, y)$ poser $X_{n+1} = y$, sinon poser $X_{n+1} = X_n$

Supposons que $\pi(x) > 0$ pour tous les états x de E , on montre alors

Théorème 6.2. *La matrice de transition P définie en (6.5) est irréductible et réversible pour la mesure π qui est donc son unique mesure invariante. Si de plus $h < 1$ alors P est apériodique.*

Démonstration. L'irréductibilité de Q implique immédiatement celle de P . Pour montrer que P est réversible, il suffit d'utiliser l'identité $h(u) = uh(1/u)$

$$\begin{aligned} x \neq y, \quad \pi(x)P(x, y) &= \pi(x)Q(x, y)h\left(\frac{\pi(y)Q(y, x)}{\pi(x)Q(x, y)}\right) \\ &= \pi(y)Q(y, x) \frac{\pi(x)Q(x, y)}{\pi(y)Q(y, x)} h\left(\frac{\pi(y)Q(y, x)}{\pi(x)Q(x, y)}\right) \\ &= \pi(y)Q(y, x)h\left(\frac{\pi(x)Q(x, y)}{\pi(y)Q(y, x)}\right) = \pi(y)P(y, x). \end{aligned}$$

Le théorème 3.11 permet d'en déduire que π est bien la mesure invariante.

Si $h < 1$, alors $P(x, x) > 0$ pour tout x de E et la matrice P est bien apériodique. On peut aussi vérifier facilement que si Q est apériodique alors P le sera même si $h \leq 1$. $\square$

L'intérêt de l'algorithme de Métropolis est évident pour simuler la mesure de Gibbs μ_T (6.2), en effet la matrice de transition P s'écrit pour $x \neq y$

$$P(x, y) = Q(x, y) h\left(\exp\left(\frac{1}{T}[V(x) - V(y)]\right) \frac{Q(y, x)}{Q(x, y)}\right)$$

et la normalisation Z_T n'a plus besoin d'être calculée. Comme h est une fonction croissante, la matrice de transition P pondère les probabilités de transition et favorise les sauts de x vers y si $V(x) > V(y)$ c'est-à-dire si le potentiel V décroît après le saut. Considérons le potentiel représenté figure 6.1 indexé par $E = \{1, \dots, L\}$ et supposons que la matrice Q corresponde à la marche aléatoire symétrique sur E . Si T est très faible, la chaîne de Markov aura tendance à évoluer vers les minima de V . Cependant, l'évolution étant aléatoire certaines transitions (assez rares) peuvent aller à l'encontre de cette tendance et éviter à la chaîne de Markov de rester bloquée dans un minimum local. Contrairement à l'approche déterministe de la descente de gradient, les fluctuations aléatoires permettent d'explorer le paysage de potentiel. Nous reviendrons sur le choix optimal du paramètre T section 6.4.

6.2.2 Modèle d'Ising

Les mesures de Gibbs (6.2) s'utilisent aussi dans des contextes très différents des méthodes d'optimisation. Elles ont été introduites initialement en physique statistique pour rendre compte de la statistique de systèmes microscopiques. La théorie de Gibbs est présentée en détail dans le cours de physique statistique [12] et nous nous contenterons ici de l'illustrer dans le cas particulier du modèle d'Ising.

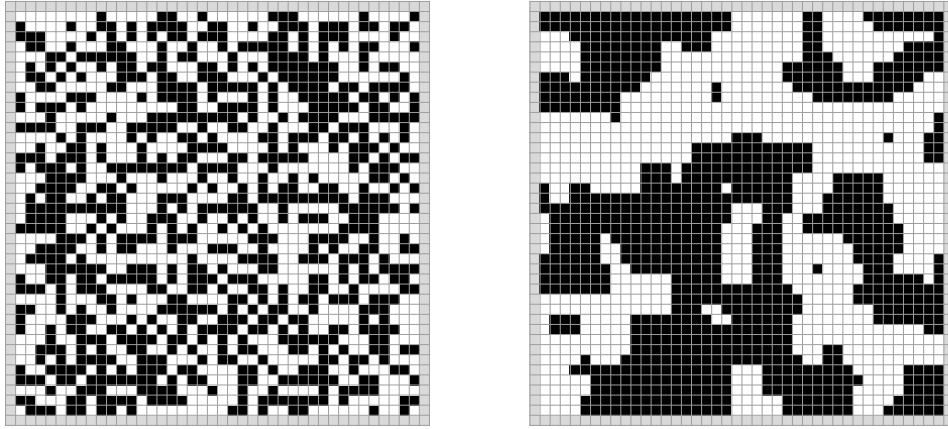


FIGURE 6.3 – Deux réalisations du modèle d'Ising (obtenues par l'algorithme de Metropolis-Hastings) pour différentes températures sur le domaine $\Lambda = \{1, \dots, 40\}^2$. La simulation de gauche montre un état très désordonné associé à une température très haute. Sur la simulation de droite, les spins de même signe se regroupent dans des régions assez larges car la température est plus basse.

Le modèle d'Ising offre un cadre théorique très simple pour décrire les transitions de phase de l'aimantation d'un métal ferromagnétique. À chaque site i du réseau $\Lambda = \{1, \dots, L\}^d$, on associe un spin s_i prenant les valeurs ± 1 et on note $S_\Lambda = \{s_i\}_{i \in \Lambda}$ une configuration de spins. Les spins interagissent avec leurs plus proches voisins et une énergie est attribuée à chaque configuration S_Λ

$$V(S_\Lambda) = - \sum_{\substack{i,j \in \Lambda \\ i \sim j}} s_i s_j$$

où $i \sim j$ signifie que les sites i et j sont à distance 1 sur le réseau Λ . Un système physique a tendance à minimiser son énergie ce qui permet de distinguer deux configurations privilégiées (les *états fondamentaux*) : les spins sont tous égaux à 1 ou tous égaux à -1 . Pour tenir compte des fluctuations thermiques, on définit la mesure de Gibbs qui attribue à la configuration S_Λ la probabilité

$$\mu_T(S_\Lambda) = \frac{1}{Z_T} \exp\left(-\frac{1}{T}V(S_\Lambda)\right)$$

où la *fonction de partition* Z_T sert à normaliser la mesure de Gibbs. Le paramètre T s'interprète comme une température : quand T est grand les fluctuations thermiques dominent et le système est désordonné, par contre pour T proche de 0 les configurations de basse énergie sont privilégiées et les spins ont tendance à s'aligner (cf. figure 6.3).

Ce modèle très simple de spins en interaction permet de mettre en évidence l'existence d'une transition de phase quand la taille du domaine L tend vers l'infini. Les transitions de phase constituent une source de questions fascinantes dont certaines seront évoquées au chapitre 7. Pour le moment, contentons-nous d'implémenter l'algorithme de Metropolis-Hasting afin de simuler le modèle d'Ising.

Retraduit dans le formalisme des chaînes de Markov, une configuration S_Λ correspond à un état et l'espace d'états est $E = \{-1, 1\}^\Lambda$. Pour un domaine bi-dimensionnel de taille $L = 40$ comme dans la figure 6.3, le cardinal de E est $2^{40 \times 40} \simeq 10^{481}$. Il est donc impossible d'énumérer toutes les configurations pour calculer la distribution μ_T . Pour simplifier les notations, nous allons omettre la dépendance en Λ et poser $S = S_\Lambda$. Pour tout i dans Λ , on note $S^{(i)}$ la configuration déduite de S en changeant simplement le signe du spin en i

$$\forall j \in \Lambda, \quad S_j^{(i)} = \begin{cases} -s_i, & \text{si } j = i \\ s_j, & \text{si } j \neq i \end{cases}$$

La matrice de référence Q décrit une évolution sur l'espace des configurations

$$\forall i \in \Lambda, \quad Q(S, S^{(i)}) = \frac{1}{\text{Card}(\Lambda)}.$$

Elle correspond au mécanisme suivant : un site i est choisi au hasard dans Λ et son spin est retourné. Ce sont les seules transitions autorisées. Ces transitions modifient les configurations seulement localement, par conséquent la variation de l'énergie correspondant au changement du spin en i ne dépend que de la moyenne des spins autour de i

$$\delta V(i, S) = V(S^{(i)}) - V(S) = 2s_i \sum_{j \sim i} s_j.$$

Étant donnée une fonction h satisfaisant $h(u) = uh(1/u)$, l'algorithme de Metropolis-Hastings s'écrit

Étape 0. Initialiser X_0 avec une configuration S quelconque

Étape $n + 1$.

Choisir i uniformément dans Λ

Choisir U_{n+1} uniformément dans $[0, 1]$ (et indépendamment du passé)

Si $U_{n+1} < h\left(\exp\left(-\frac{1}{T}\delta V(i, X_n)\right)\right)$ poser $X_{n+1} = X_n^{(i)}$, sinon poser $X_{n+1} = X_n$

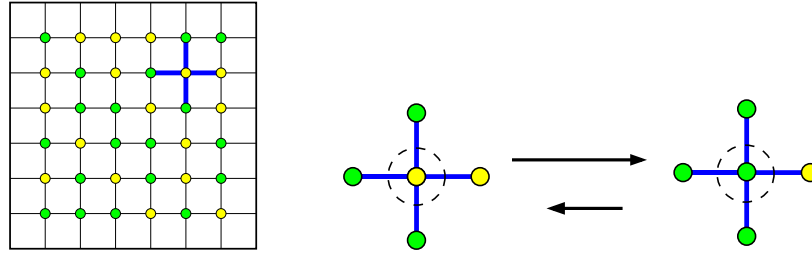


FIGURE 6.4 – L’algorithme de Metropolis-Hastings consiste à choisir un site au hasard (figure de gauche) et à mettre à jour ce spin en fonction de la moyenne de ses 4 voisins. Quand T est proche de 0, le spin aura ensuite tendance à s’aligner avec ses voisins (figure de droite).

6.3 Simulation parfaite : algorithme de Propp-Wilson ★

Nous avons vu au chapitre 5 qu’une chaîne de Markov irréductible, apériodique converge vers sa mesure invariante. En particulier, le théorème 5.5 garantit la convergence de l’algorithme de Metropolis-Hastings quand le temps tend vers l’infini. Mais en pratique, la simulation doit être arrêtée à un temps fini et il est donc important d’estimer l’erreur faite. Cette question a motivé de nombreuses études théoriques pour quantifier la vitesse de convergence et il s’agit toujours d’un sujet de recherche très actif en probabilités et en statistique bayésienne. Dans le cadre ce cours, nous avons défini au théorème 5.11 le critère de Doeblin qui permet de déterminer pour tout $\varepsilon > 0$ un temps n_ε au-delà duquel l’erreur est contrôlée

$$\forall n \geq n_\varepsilon, \quad \sup_{x \in E} \|P^n(x, \cdot) - \pi\|_{VT} \leq \varepsilon.$$

Il n’est pas toujours possible d’obtenir une estimation théorique qui fournit des bornes suffisamment précises. Ainsi dans la pratique, la durée de simulation est souvent déterminée par l’intuition ou calibrée à partir d’expérimentations.

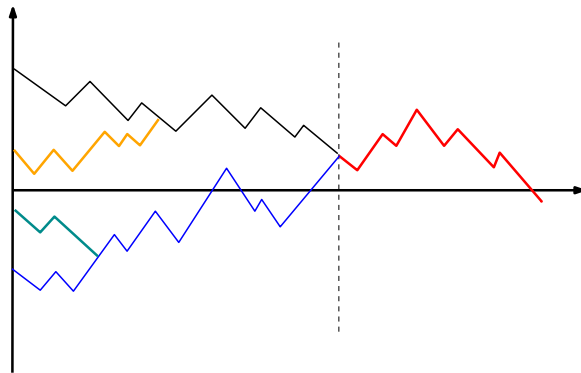


FIGURE 6.5 – Le schéma représente un couplage entre des trajectoires issues de différents états initiaux. Au-delà de la ligne en pointillés, toutes les trajectoires ont fusionné.

Nous allons décrire maintenant l’algorithme de Propp-Wilson [24] qui permet de simuler de façon exacte la mesure invariante par une méthode de *couplage par le passé*.

Avant cela, revenons sur la preuve du théorème 5.5 où la convergence était estimée en fonction du temps de couplage entre différentes trajectoires (cf. figure 6.5). Dans le cas d'une récurrence aléatoire (définie au théorème 2.2) le couplage se construit de la façon suivante. On rappelle que la chaîne de Markov $\{X_n\}_{n \geq 0}$ à valeurs dans E est obtenue par récurrence

$$n \geq 0, \quad X_{n+1} = f(X_n, \xi_{n+1})$$

en fonction d'une suite $\{\xi_n\}_{n \geq 1}$ de variables aléatoires indépendantes et identiquement distribuées sur un espace F et d'une fonction f de $E \times F$ dans E . Afin de coupler les $\text{Card}(E)$ trajectoires, on définit $\theta = \{\xi^{(x)}\}_{x \in E}$ une collection de $\text{Card}(E)$ variables aléatoires indépendantes et de même loi que ξ_0 et on note

$$\forall x \in E, \quad F_\theta(x) = f(x, \xi^{(x)}).$$

Ceci permet de construire le couplage après un pas de temps simultanément pour toutes les données initiales dans E . Pour itérer, il suffit de choisir une suite $\{\theta_n\}_{n \geq 1}$ de variables indépendantes et la chaîne de Markov partant de x s'obtient en composant les applications

$$\Phi_n(x) = F_{\theta_n} \circ F_{\theta_{n-1}} \circ \cdots \circ F_{\theta_1}(x). \quad (6.6)$$

Le couplage s'effectue au premier temps (aléatoire) T où Φ_T est constante, i.e. que Φ_T ne dépend plus de l'état initial dans E . Comme toutes les trajectoires ont fusionné au temps T , la chaîne de Markov a perdu toute la mémoire du passé et il est tentant de croire que la position de chaîne à l'instant T (i.e. Φ_T) est distribuée selon la mesure invariante. Ce n'est pas le cas comme le montre l'exemple de la figure 6.6, néanmoins une modification simple mais astucieuse permet de rendre cette idée rigoureuse.

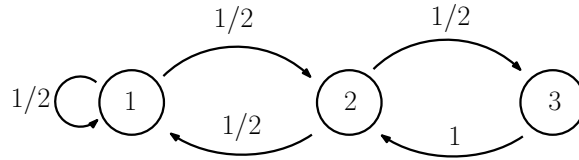


FIGURE 6.6 – Pour ce graphe de transition, l'équilibre ne peut pas correspondre au moment où les différentes trajectoires fusionnent. En effet, l'état 3 ne peut être atteint qu'en venant de l'état 2 et il n'est donc pas possible que les trajectoires se touchent pour la première fois à l'état 3.

Le point de vue Propp et Wilson consiste à coupler non pas vers le futur mais vers le passé en remontant le temps et en inversant l'ordre dans (6.6)

$$\mathcal{G}_n(x) = F_{\theta_1} \circ \cdots \circ F_{\theta_{n-1}} \circ F_{\theta_n}(x).$$

L'indexation F_{θ_k} correspond ici à la transition entre les instants $-k$ et $-k + 1$. $\mathcal{G}_n(x)$ s'interprète comme la valeur à l'instant 0 de la chaîne de Markov partie de x au temps $-n$ (cf. figure 6.7). On définit le temps d'arrêt T correspondant au premier instant où $\mathcal{G}_n$ est constante

$$T = \inf \{n \geq 0; \quad \forall x, y \in E, \quad \mathcal{G}_n(x) = \mathcal{G}_n(y)\}. \quad (6.7)$$

Dans ce cas le résultat de l'algorithme est l'état $\mathcal{G}_T = \mathcal{G}_T(x)$ obtenu au temps 0 (qui ne dépend pas de x).

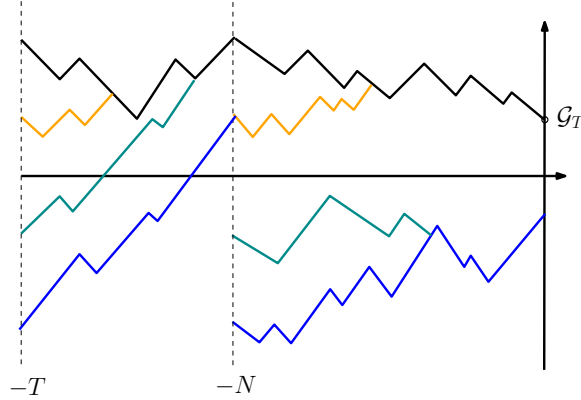


FIGURE 6.7 – Sur le schéma, les différentes trajectoires partant au temps $-N$ n'ont pas coalescé au temps 0. En remontant le temps jusqu'à $-T$ les trajectoires issues des différents états de E ont coalescé et l'état $\mathcal{G}_T$ obtenu au temps 0 est distribué exactement selon la mesure invariante. À mesure que l'algorithme remonte dans le temps, la simulation doit conserver la mémoire des trajectoires déjà utilisées.

Théorème 6.3. *On considère une chaîne de Markov irréductible, apériodique de mesure invariante π . Si le temps de coalescence T défini en (6.7) est fini presque sûrement, alors l'état $\mathcal{G}_T$ est distribué selon la mesure invariante π .*

Démonstration. Une fois que les trajectoires ont coalescé, l'état de la chaîne en 0 ne va plus varier

$$\forall n \geq T, \quad \mathcal{G}_n = \mathcal{G}_T.$$

Comme on a supposé que le temps T est fini presque sûrement, on a par le théorème de convergence dominée

$$\forall x, y \in E, \quad \lim_{n \rightarrow \infty} \mathbb{P}(\mathcal{G}_n(x) = y) = \mathbb{P}(\mathcal{G}_T(x) = y).$$

Les variables $\mathcal{G}_n$ et Φ_n (6.6) ont la même loi et par conséquent

$$\lim_{n \rightarrow \infty} \mathbb{P}(\mathcal{G}_n(x) = y) = \lim_{n \rightarrow \infty} \mathbb{P}(\Phi_n(x) = y) = \pi(y)$$

où la dernière égalité est obtenue par le théorème 5.5. On conclut ainsi la preuve de ce théorème

$$\forall x, y \in E, \quad \mathbb{P}(\mathcal{G}_T(x) = y) = \pi(y).$$

□

L'algorithme de Propp-Wilson suppose de suivre toutes les trajectoires issues de E ce qui est impossible si E a un cardinal trop grand. Cependant pour certaines chaînes de Markov (par exemple pour la dynamique de Metropolis associée au modèle d'Ising), on peut implémenter une variante de cet algorithme pour laquelle il suffit de suivre un petit nombre de trajectoires. Ceci dépasse le cadre de ce cours mais on pourra trouver une étude détaillée dans [24].

6.4 Algorithme de recuit simulé ★

Le problème initial évoqué section 6.1 est de caractériser les configurations minimisant une fonction V . Le lemme 6.1 permet d'estimer de telles configurations comme des états privilégiés d'une mesure de Gibbs μ_T en fonction d'un paramètre $T > 0$ pouvant être interprété comme une température. Cette mesure peut ensuite être simulée par un algorithme de Metropolis-Hastings. La difficulté consiste à calibrer le paramètre T de façon optimale car deux effets contraires se conjuguent dans cette procédure. T doit être proche de 0 pour que les minima soient correctement estimés, mais si T est trop faible, l'aléa disparaît et la chaîne de Markov se comporte comme les modèles déterministes de descente de gradient. Dans ce cas la convergence n'aura jamais lieu sur des échelles de temps raisonnables. Inversement si T est trop grand, les fluctuations stochastiques seront suffisantes pour empêcher la chaîne de Markov d'être piégée dans des minima locaux (cf. figure 6.1) et la convergence vers la mesure d'équilibre μ_T sera plus rapide, cependant μ_T ne décrira pas correctement les solutions cherchées.

Le *recuit simulé* consiste à simuler une chaîne de Markov par l'algorithme de Metropolis-Hastings en baissant la température à chaque pas de temps. L'idée est de permettre à la chaîne de Markov de sortir rapidement des minima locaux qu'elle pourrait rencontrer initialement, mais de baisser la température progressivement pour qu'elle se stabilise sur le minimum global recherché. On choisit une suite de températures $\{T_n\}_{n \geq 0}$ décroissant vers 0 et au temps n la chaîne de Markov (non homogène) évolue en suivant la dynamique associée à la mesure μ_{T_n} .

Étape 0. Initialiser X_0

Étape $n + 1$.

Choisir y selon la loi $Q(X_n, y)$

Choisir U_{n+1} uniformément dans $[0, 1]$ (et indépendamment du passé)

Si $U_{n+1} < h \left(\exp \left(\frac{1}{T_n} [V(X_n) - V(y)] \right) \frac{Q(y, X_n)}{Q(X_n, y)} \right)$ poser $X_{n+1} = y$

Sinon poser $X_{n+1} = X_n$

Le choix de la suite de températures est crucial. On peut démontrer que

Théorème 6.4. *Étant donnée une fonction V , il existe une constante $C(V, Q)$ qui dépend de V et Q , telle que l'algorithme de recuit simulé appliqué avec la suite de températures $T_n = \frac{C(V)}{\log n}$ sélectionne l'ensemble $\mathcal{M}$ des minima de V , c'est-à-dire*

$$\lim_{n \rightarrow \infty} \mathbb{P}(X_n \in \mathcal{M}) = 1.$$

La preuve de ce théorème pourra être trouvée dans le livre [9] et nous nous contenterons de justifier le choix de la décroissance en $\frac{1}{\log n}$ par un exemple. Soit $E = \{1, 2, 3\}$ et $V(1) = 0, V(2) = 1, V(3) = -1$. Le minimum est atteint à l'état 3 et l'état 1 constitue un minimum local (cf. figure 6.8). On pose $h(u) = \min\{u, 1\}$ et Q suit le graphe des transitions de la figure 6.8.

On suppose que l'état initial $X_0 = 1$ et on veut calculer la probabilité que la chaîne de Markov soit dans l'état 3 à un temps n donné. Pour cela il faut que la chaîne soit passée

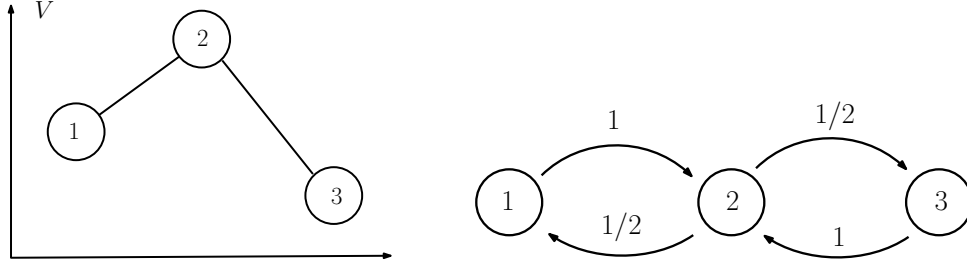


FIGURE 6.8 – Le potentiel V de l'exemple avec 3 sites est tracé à gauche. Le graphe de transition de la matrice Q est représenté à droite.

par l'état 2 avant n

$$\mathbb{P}(X_n = 3) \leq \sum_{k=0}^{n-2} \mathbb{P}(X_k = 1, X_{k+1} = 2).$$

Avec les paramètres choisis, la probabilité de transition entre 1 et 2 au temps $k + 1$ est donnée par

$$\mathbb{P}\left(U_{k+1} < \frac{1}{2} \exp\left(-\frac{1}{T_k}\right)\right) = \frac{1}{2} \exp\left(-\frac{1}{T_k}\right).$$

On en déduit donc une borne supérieure uniforme en n

$$\mathbb{P}(X_n = 3) \leq \frac{1}{2} \sum_{k=0}^{\infty} \exp\left(-\frac{1}{T_k}\right).$$

Il est facile de voir que si la constante c est suffisamment petite et $T_k \leq \frac{c}{\log k}$ alors pour tout temps n

$$\mathbb{P}(X_n = 3) < 1.$$

Par conséquent si la température tend trop vite vers 0, la chaîne de Markov restera indéfiniment piégée au point 1 avec une probabilité positive. Cet exemple très simple, montre que la décroissance de la température dans le théorème 6.4 est optimale.

Dans la pratique, la décroissance de la température en $\frac{1}{\log n}$ s'avère trop lente pour implémenter des algorithmes performants. On préfère donc souvent utiliser des décroissances polynomiales de la forme $\frac{1}{n^\gamma}$ qui ne sont pas justifiées d'un point de vue théorique mais qui donnent quand même de très bons résultats...

6.4.1 Problème du voyageur de commerce

Le recuit simulé permet d'obtenir une solution approchée du problème du voyageur de commerce défini en (6.1). On choisit comme état initial X_0 une permutation au hasard dans l'ensemble $\mathcal{S}_K$ des permutations de $\{1, \dots, K\}$. Étant donné un parcours $\sigma = (\sigma(1), \dots, \sigma(K))$, on note pour $i \neq j$ le parcours $\sigma^{(i,j)}$ obtenu en permutant l'ordre de visite des villes $\sigma(i)$ et $\sigma(j)$. Ceci permet de définir une matrice de transition Q sur les différents parcours en autorisant uniquement ce type de transitions

$$1 \leq i < j \leq K, \quad \forall \sigma \in \mathcal{S}_K, \quad Q(\sigma, \sigma^{(i,j)}) = \frac{2}{K(K-1)}.$$

On vérifie facilement que la matrice de transition Q définit une chaîne de Markov irréductible sur $\mathcal{S}_K$ et qu'elle satisfait (6.3). On peut donc implémenter le recuit simulé pour obtenir une estimation de la distance minimale à parcourir. Les résultats de l'algorithme de recuit simulé sont représentés figure 6.9.

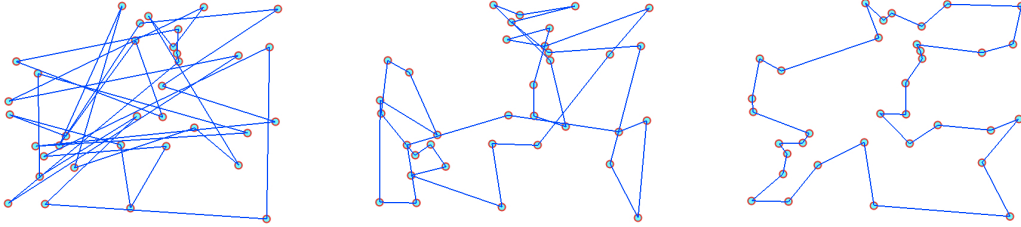


FIGURE 6.9 – Les simulations ci-dessus représentent des étapes de l'algorithme de recuit simulé pour résoudre le problème du voyageur de commerce. Les positions de 35 villes sont choisies au hasard ainsi que le circuit initial tracé à gauche. Après 2000 itérations, la longueur du parcours s'est réduite (figure du centre) et le chemin tend à converger après 10000 itérations vers une solution (presque) optimale (figure de droite).

6.4.2 Traitement d'images

Dans cette section, nous allons montrer comment le recuit simulé permet de traiter des images perturbées par un bruit aléatoire et d'identifier des formes. Pour simplifier, on suppose que l'image dégradée par le bruit est en noir et blanc. On peut indexer chaque pixel de cette image de taille $L \times L$ comme un site i de $\Lambda = \{1, \dots, L\}^2$ et lui attribuer la valeur $\sigma_i = \pm 1$ selon sa couleur. L'image est donc une collection de pixels $\Sigma = \{\sigma_i\}_{i \in \Lambda}$. On veut éliminer le bruit pour reconnaître une forme sur cette image. Pour cela on fait l'hypothèse que cette forme a des contours réguliers et est constituée de blocs de 1. On va donc modifier l'image en enlevant des pixels -1 dans les régions où les pixels égaux à 1 sont denses et inversement. Ce mécanisme rappelle celui de la dynamique de Metropolis-Hasting pour le modèle d'Ising (cf. section 6.2.2) et nous allons nous en inspirer.

L'image initiale $\Sigma = \{\sigma_i\}_\Lambda$ va être modifiée en une nouvelle image $S = \{s_i\}_\Lambda$ obtenue comme le minimum de l'énergie

$$V(S) = -\alpha \sum_{\substack{i,j \in \Lambda \\ i \sim j}} s_i s_j + \beta \sum_{i \in \Lambda} (s_i - \sigma_i)^2$$

où $\alpha, \beta \geq 0$ sont deux constantes. Si $\alpha = 0$ et $\beta > 0$, le minimum est donné par l'image initiale. Inversement si $\alpha > 0$ et $\beta = 0$, on retrouve l'énergie du modèle d'Ising et le minimum est atteint pour les 2 configurations extrêmes où tous les pixels valent 1 ou -1 . Il faut donc ajuster les paramètres α et β pour réaliser un compromis entre deux effets : l'image restaurée doit rester fidèle à l'image initiale mais les contours doivent être le plus net possible et les fluctuations dues au bruit doivent être éliminées.

Le recuit simulé est une méthode adaptée pour minimiser la fonction V qui est composée de nombreux minima locaux et indexée par L^2 variables. L'état initial est donné



FIGURE 6.10 – Pour illustrer l’algorithme de recuit simulé, l’image de gauche représentant un carré blanc sur fond noir a été perturbée par un bruit aléatoire pour donner l’image au centre. L’image de gauche est obtenue après traitement de l’image bruitée par l’algorithme de Metropolis-Hasting.

par l’image Σ et on utilise ensuite la dynamique de Metropolis-Hastings en abaissant progressivement la température (cf. figure 6.10).

Cette application du recuit simulé au traitement d’images avait simplement pour but d’illustrer les possibilités offertes par cette méthode. Pour traiter des problèmes concrets, une théorie plus sophistiquée est nécessaire et ses fondements sont décrits dans [29]. La segmentation d’images, i.e. l’identification de composantes dans des images, est particulièrement utilisée en imagerie médicale ou en cartographie.

Chapitre 7

Un exemple de modélisation en physique : la percolation ★

Les chaînes de Markov fournissent un cadre théorique très développé pour étudier les comportements asymptotiques de variables aléatoires corrélées. La dépendance des chaînes de Markov est indexée par la variable de temps. Ce chapitre constitue une introduction aux systèmes où l'indexation des variables aléatoires n'est plus linéaire et où la géométrie joue un rôle. Nous décrirons le modèle de percolation et montrerons que la structure spatiale induit des propriétés très intéressantes qui font de la percolation un modèle clef en physique statistique. Le cours de W. Werner [27] est une excellente référence sur la théorie de la percolation (en particulier on pourra y retrouver les résultats présentés dans ce chapitre).

Ce chapitre peut être omis dans le cadre du cours de MAP432, il sert simplement à présenter des développements actuels en théorie des probabilités.

7.1 Description du modèle

Imaginons une pierre poreuse immergée dans de l'eau. Peut-on déterminer en fonction de la porosité si le centre de la pierre est mouillé ? Cette question a été posée par Broadbent et Hammersley en 1957 et formulée dans le cadre mathématique suivant. On considère $\mathbb{Z}^d$ et on note $\mathcal{E}$ l'ensemble des arêtes

$$\mathcal{E} = \{(i, j) \mid i, j \in \mathbb{Z}^d, \quad \|i - j\|_2 = 1\}. \quad (7.1)$$

Pour simplifier les notations, une arête typique sera souvent notée $b = (i, j) \in \mathcal{E}$. À chaque arête b , on associe une variable aléatoire de Bernoulli ω_b de paramètre $p \in [0, 1]$ indépendamment des autres arêtes

$$\forall b \in \mathcal{E}, \quad \mathbb{P}(\omega_b = 1) = 1 - \mathbb{P}(\omega_b = 0) = p. \quad (7.2)$$

Par analogie avec la pierre poreuse, on dira qu'une arête b est ouverte si $\omega_b = 1$ (l'eau peut passer à travers l'arête) et fermée sinon (cf. figure 7.1). Un chemin de k à ℓ dans $\mathbb{Z}^d$ est une suite $\{i_0 = k, i_1, \dots, i_n = \ell\}$ de sites distincts tels que (i_{j-1}, i_j) soit dans $\mathcal{E}$. On dit

que deux sites k et ℓ sont reliés dans le modèle de percolation par un chemin ouvert s'il existe un chemin $\{i_0 = k, i_1, \dots, i_n = \ell\}$ tel que $\omega_{(i_{j-1}, i_j)} = 1$ pour tout $j \leq n$. On notera $\{k \leftrightarrow \ell\}$ l'évènement que k soit relié à ℓ par un chemin ouvert et $\{O \leftrightarrow \infty\}$ l'évènement qu'il existe un chemin infini d'arêtes ouvertes partant de l'origine.

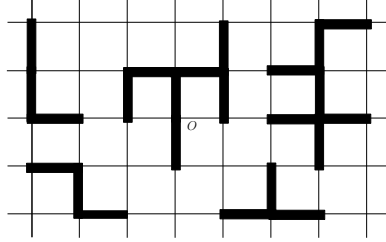


FIGURE 7.1 – Exemple d'une configuration de percolation dans un sous-ensemble de $\mathbb{Z}^2$. Les arêtes ouvertes sont représentées en gras.

On peut interpréter l'existence d'un chemin infini en disant que le centre de la pierre sera mouillé. Si $p = 1$, l'origine est toujours connectée à l'infini par un chemin d'arêtes ouvertes et inversement si $p = 0$, l'origine est toujours déconnectée de l'infini. Le problème est donc de déterminer pour quelles valeurs de $p \in (0, 1)$ un tel chemin existe avec probabilité positive.

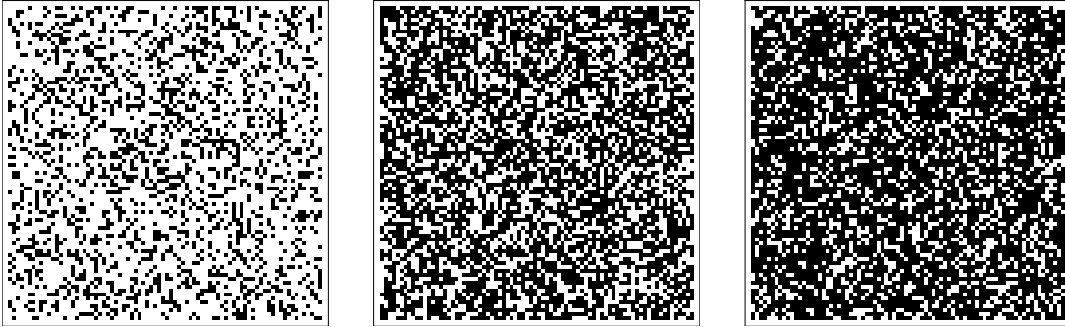


FIGURE 7.2 – Dans ces simulations, les sites de $\mathbb{Z}^2$ sont coloriés en noir avec probabilité p et en blanc avec probabilité $1 - p$. La question de la percolation peut donc se reformuler en terme de chemins de sites noirs adjacents. Chaque grille a 80×80 sites et l'intensité de p est successivement $p = 0,3$, $p = 0,59$ et $p = 0,7$. Existe-t-il, dans l'image du milieu, un chemin noir reliant l'origine (située au centre) au bord du carré ?

7.2 Transition de phase

Pour $p \in [0, 1]$, on définit

$$\theta(p) = \mathbb{P}(\{O \leftrightarrow \infty\}). \quad (7.3)$$

Le théorème suivant montre l'existence d'une transition de phase pour la percolation (cf. figure 7.3).

Théorème 7.1. Pour $d \geq 2$, il existe un point critique $p_c \in]0, 1[$ tel que

$$\forall p < p_c, \quad \theta(p) = 0, \quad \forall p > p_c, \quad \theta(p) > 0.$$

La structure spatiale est très importante et dans le cas unidimensionnel, il n'existe pas de transition de phase à une valeur non triviale. En effet si l'origine est connectée à l'infini alors pour tout n , l'origine est connectée à n ou $-n$. On a donc

$$\mathbb{P}(\{O \leftrightarrow \infty\}) \leq \mathbb{P}(\{O \leftrightarrow n\}) + \mathbb{P}(\{O \leftrightarrow -n\}) \leq 2p^n.$$

Pour tout $p < 1$, on voit que le membre de droite de l'équation tend vers 0 quand n tend vers l'infini. Par conséquent $p_c = 1$ si $d = 1$.

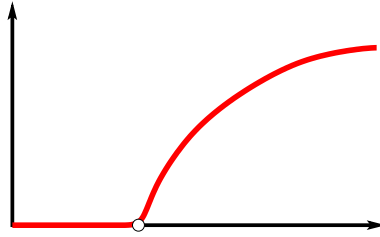


FIGURE 7.3 – Graphe de $p \rightarrow \theta(p)$ pour $d \geq 2$. La continuité de la courbe au point critique est évoquée section 7.2.3.

La preuve du théorème 7.1 se décompose en 3 étapes :

1. Il existe $p_0 > 0$ tel que pour tout $p < p_0$ alors $\theta(p) = 0$.
2. Il existe $p_1 < 1$ tel que pour tout $p > p_1$ alors $\theta(p) > 0$.
3. La fonction $p \rightarrow \theta(p)$ est croissante.

Ces trois assertions permettent donc de démontrer qu'il existe un unique point critique p_c dans $]0, 1[$. On pose

$$p_c = \inf\{p \mid \theta(p) > 0\}.$$

Chaque étape sera l'objet d'une des propositions démontrées ci-dessous.

7.2.1 Absence de percolation pour p petit

Proposition 7.2. Il existe $p_0 > 0$ tel que $\theta(p) = 0$ pour tout $p < p_0$.

Démonstration. L'idée de la preuve est similaire à l'argument utilisé en dimension 1, mais cette fois il faut tenir compte de l'entropie des chemins, i.e. des choix multiples des chemins possibles. Si $\{O \leftrightarrow \infty\}$ a lieu alors il existe au moins un chemin $\gamma = \{i_0 = O, i_1, \dots, i_n\}$ de longueur n et partant de l'origine O dont toutes les arêtes sont ouvertes (on rappelle qu'un chemin ne s'intersecte pas). On notera Γ_n l'ensemble des chemins de

longueur n partant de l'origine. On a donc

$$\begin{aligned} \mathbb{P}(\{O \leftrightarrow \infty\}) &\leq \mathbb{P}(\{\exists \text{ un chemin de longueur } n \text{ ouvert}\}) , \\ &\leq \mathbb{P}\left(\bigcup_{\gamma \in \Gamma_n} \{\gamma \text{ est un chemin ouvert}\}\right) , \\ &\leq \sum_{\gamma \in \Gamma_n} \mathbb{P}(\{\gamma \text{ est un chemin ouvert}\}) = \text{Card}(\Gamma_n) p^n . \end{aligned}$$

En dimension d , le cardinal de Γ_n est toujours trivialement inférieur à $(2d)^n$. Par conséquent si $p_0 = \frac{1}{2d}$, l'inégalité ci-dessous est valable pour tout n

$$\mathbb{P}(\{O \leftrightarrow \infty\}) \leq (2d)^n p^n = \left(\frac{p}{p_0}\right)^n .$$

Il suffit de laisser n tendre vers l'infini pour conclure. $\square$

7.2.2 Percolation pour p proche de 1

Proposition 7.3. *Il existe $p_1 < 1$ tel que $\theta(p) > 0$ pour tout $p > p_1$.*

Démonstration. Si l'origine est connectée à l'infini dans $\mathbb{Z}^2$ alors elle le sera a fortiori dans $\mathbb{Z}^d$ où les chemins ouverts sont plus nombreux. Il suffit donc de prouver la Proposition pour $d = 2$. On note $\mathcal{C}(O)$ la composante connexe contenant l'origine et formée par les liens ouverts. On va montrer que pour p proche de 1

$$\lim_{n \rightarrow \infty} \mathbb{P}(\text{Card}(\mathcal{C}(O)) \leq n) < 1 .$$

Ceci implique qu'il y a une probabilité positive pour que le cardinal de $\mathcal{C}(O)$ soit infini et donc que l'origine soit connectée à l'infini.

Une simplification importante du cas $d = 2$ est la notion de dualité. On définit le réseau dual dont les arêtes sont $\mathcal{E}^* = \{(u + \frac{1}{2}, v + \frac{1}{2}), (u, v) \in \mathbb{Z}^2\}$. Chaque arête b de $\mathcal{E}$ est associée à l'arête b^* de $\mathcal{E}^*$ qui l'intersecte. A toute réalisation aléatoire $\{\omega_b\}_{b \in \mathcal{E}}$, on peut faire correspondre une configuration de percolation dans le réseau dual $\{\omega_{b^*}^* = 1 - \omega_b\}_{b \in \mathcal{E}}$ (cf. figure 7.4). Une arête ouverte dans $\mathcal{E}$ est associée à une arête fermée dans le réseau dual et inversement. La percolation associée au réseau dual a donc pour paramètre $1 - p$.

Si la composante $\mathcal{C}(O)$ est finie, alors elle est nécessairement entourée par un chemin ouvert dans le dual (cf. figure 7.4). On notera Γ_n^* l'ensemble des chemins γ^* de longueur n dans le dual entourant l'origine. On remarque que pour entourer l'origine il faut au moins 4 arêtes duales.

$$\begin{aligned} \mathbb{P}(\text{Card}(\mathcal{C}(O)) < \infty) &= \mathbb{P}(\{\exists \text{ un chemin dual ouvert entourant l'origine}\}) , \\ &\leq \sum_{n \geq 4} \mathbb{P}\left(\bigcup_{\gamma^* \in \Gamma_n^*} \{\gamma^* \text{ est un chemin dual ouvert}\}\right) , \\ &\leq \sum_{n \geq 4} \sum_{\gamma^* \in \Gamma_n^*} \mathbb{P}(\{\gamma^* \text{ est un chemin dual ouvert}\}) . \end{aligned}$$

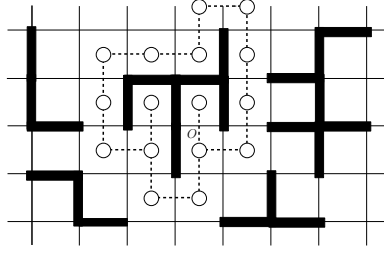


FIGURE 7.4 – La composante connexe $\mathcal{C}(0)$ des liens ouverts contenant l'origine est entourée par un contour ouvert dans le réseau dual (représenté en pointillés). Pour ne pas alourdir le dessin, seuls les points du réseau dual autour de l'origine sont représentés, mais il faut imaginer une configuration duale plus étendue avec toutes les arêtes $\{\omega_{b^*}^*\}_{b^* \in \mathcal{E}^*}$.

La probabilité qu'un chemin dual de longueur n soit ouvert est $(1 - p)^n$, on a donc

$$\mathbb{P}\left(\text{Card}(\mathcal{C}(O)) < \infty\right) \leq \sum_{n \geq 4} \text{Card}(\Gamma_n^*) (1 - p)^n.$$

Il ne reste plus qu'à estimer le cardinal de Γ_n^* . Tout chemin $\gamma^* \in \Gamma_n^*$ va croiser l'axe $\{(x, 0)\}_{x \in [0, n]}$ au moins une fois. Si on fixe une arête duale intersectant $\{(x, 0)\}_{x \in [0, n]}$, le nombre de chemins de longueur n est au plus 3^n (cette borne est loin d'être optimale). Ceci peut s'obtenir en remarquant qu'un chemin dans le dual forme une boucle sans intersections et qu'à chaque pas un chemin a au plus 3 directions possibles pour évoluer. On en déduit que pour $p > 2/3$

$$\mathbb{P}\left(\text{Card}(\mathcal{C}(O)) < \infty\right) \leq \sum_{n \geq 4} n 3^n (1 - p)^n.$$

On peut donc choisir p_1 pour que la probabilité $\mathbb{P}\left(\text{Card}(\mathcal{C}(O)) < \infty\right)$ soit strictement inférieure à 1 si $p > p_1$. $\square$

7.2.3 Point critique

Proposition 7.4. *La fonction $p \rightarrow \theta(p)$ est croissante.*

Démonstration. Une façon simple de simuler une variable aléatoire de Bernoulli ω de paramètre p est de tirer au hasard une variable aléatoire uniforme U sur $[0, 1]$ et de poser $\omega = 1_{\{U \leq p\}}$. On peut donc comparer 2 configurations de percolation de paramètres $p < q$ en les couplant, i.e. en les construisant simultanément. On se donne une collection $\{U_b\}_{b \in \mathcal{E}}$ de variables aléatoires indépendantes et uniformément distribuées sur $[0, 1]$ et on pose

$$\forall b \in \mathcal{E}, \quad \omega_b = 1_{\{U_b \leq p\}}, \quad \eta_b = 1_{\{U_b \leq q\}}.$$

Les variables $\{\omega_b\}_{b \in \mathcal{E}}$ sont des variables de Bernoulli indépendantes de paramètre p et $\{\eta_b\}_{b \in \mathcal{E}}$ sont des variables de Bernoulli indépendantes de paramètre q . Par contre $\{\omega_b\}_{b \in \mathcal{E}}$ et $\{\eta_b\}_{b \in \mathcal{E}}$ sont corrélées car $\omega_b \leq \eta_b$ pour tout $b \in \mathcal{E}$. Par conséquent tout chemin ouvert dans la configuration $\{\omega_b\}_{b \in \mathcal{E}}$ sera aussi ouvert dans la configuration $\{\eta_b\}_{b \in \mathcal{E}}$. On en déduit que si l'origine est connectée à l'infini pour la réalisation $\{\omega_b\}_{b \in \mathcal{E}}$ elle le sera aussi pour la réalisation $\{\eta_b\}_{b \in \mathcal{E}}$. Ceci permet de conclure que $\theta(p) \leq \theta(q)$. $\square$

Nous avons démontré que la fonction $p \rightarrow \theta(p)$ est croissante et il est naturel de se demander si elle est continue. Par définition, $p \rightarrow \theta(p)$ est nulle pour $p < p_c$ et donc continue. On peut montrer (avec un peu d'efforts) qu'elle est aussi continue pour $p > p_c$. On conjecture la continuité au point p_c en toute dimension, mais elle n'a été démontrée que pour $d = 2$ et $d \geq 19$. Du point de vue de la physique, une discontinuité en p_c s'interpréterait comme une transition de phase du premier ordre et impliquerait l'existence d'une composante ouverte infinie de densité macroscopique à p_c . Personne ne s'attend à un tel scénario et on conjecture que la transition de phase est du second ordre ($\theta(p_c) = 0$), cependant le cas physiquement intéressant de la dimension $d = 3$ reste un problème mathématique ouvert !

7.2.4 Dimension 2

On remarquera que la preuve du théorème 7.1 établit l'existence de p_c sans en déterminer la valeur. En général, cette valeur n'est pas connue mais dans certains cas particuliers des symétries permettent de la deviner.

En dimension 2, on peut démontrer que $p_c = 1/2$. La preuve est délicate (cf. [27]) et on se contentera de justifier l'intuition du résultat. La dualité implique l'alternative suivante pour tout domaine de la forme $\Lambda_n = \{-n, \dots, n\}^2$:

- ou Λ_n est traversé par un chemin ouvert reliant le bord droit au bord gauche
- ou il existe un chemin dual ouvert reliant le haut et le bas de Λ_n (de Λ_n^* si on veut être précis).

Ces deux événements ne peuvent pas arriver simultanément (cf. figure 7.5).

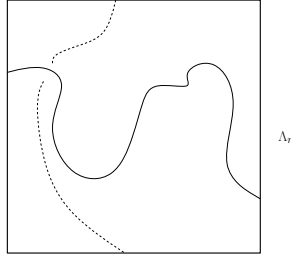


FIGURE 7.5 – Un chemin ouvert reliant le bord droit au bord gauche de la boîte Λ_n . Ce chemin coupe la boîte en 2 morceaux et empêche tout chemin dual (représenté en pointillés) de la traverser du haut vers le bas.

On peut démontrer que $p > p_c$ si la probabilité qu'un chemin ouvert relie les bords gauche et droit de Λ_n tend vers 1 quand n tend vers l'infini. Par conséquent si $p > p_c$ la percolation dans $\mathcal{E}$ va empêcher la percolation dans le réseau dual $\mathcal{E}^*$. Mais les deux types de percolation ont par symétrie un comportement identique. L'absence de percolation dans le réseau dual implique donc $1 - p < p_c$. Ceci conduit à la relation $p_c = 1 - p_c$ et justifie heuristiquement $p_c = 1/2$. On peut aussi démontrer la continuité au point critique $\theta(1/2) = 0$.

La percolation au point critique est très étudiée en physique. Par exemple, on conjecture un comportement universel de $\theta(p)$ proche de p_c pour une grande classe de modèles

bidimensionnels

$$p > 1/2, \quad \theta(p) \simeq (p - 1/2)^{5/36}.$$

L'exposant $5/36$ ne devrait pas dépendre de la structure microscopique du réseau. Pour le moment cette relation a été établie "uniquement" pour le réseau triangulaire par S. Smirnov et W. Werner. Le cas de $\mathbb{Z}^2$ considéré dans ces notes reste un problème ouvert.

L'objet mathématique caché derrière ce comportement universel est l'évolution de *Schramm-Loewner* (Schramm-Loewner evolution). Ce processus stochastique est la limite du processus d'exploration discret représenté figure 7.6. Il encode la structure limite de l'interface et de façon implicite les exposants critiques du modèle. Ce processus dépasse largement le cadre de la percolation car il apparaît comme la limite universelle des modèles critiques bidimensionnels (invariants par transformations conformes) : le modèle d'Ising, le modèle de Potts, la marche auto-évitante, le champ libre gaussien...

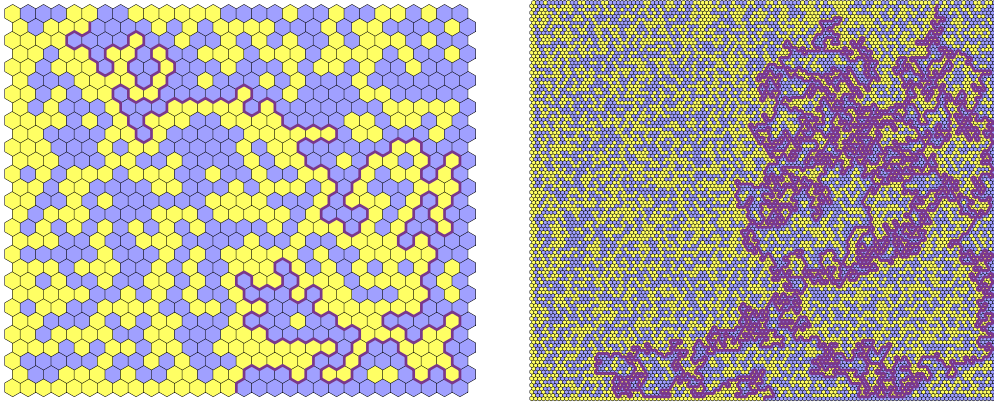


FIGURE 7.6 – On considère la percolation sur le réseau hexagonal. Les sites du bord droit sont coloriés en bleu et ceux du bord gauche en jaune, les autres couleurs sont choisies aléatoirement avec la probabilité p_c . On construit un chemin en partant du bas de l'image et en explorant l'interface entre le bleu et le jaune. Les conditions aux bords forcent le chemin à traverser le domaine du bas vers le haut. Quand la taille du domaine augmente la trajectoire revient sur elle même et forme des boucles. Les simulations ci-dessus ont été réalisées par V. Beffara.

Deuxième partie

Martingales

Chapitre 8

De la marche aléatoire aux stratégies optimales

Les martingales font souvent référence à des méthodes secrètes et mystérieuses destinées à gagner aux jeux de hasard. La version mathématique est moins romanesque et en particulier, nous allons démontrer que de telles méthodes n'existent pas. Avant de formaliser le concept de martingales, commençons par décrire les gains d'un joueur, qui à chaque étape n d'un jeu de hasard, gagne une somme ξ_n si $\xi_n \geq 0$ ou perd la somme $-\xi_n$ si $\xi_n < 0$. La fortune du joueur au temps n sera alors notée

$$\forall n \geq 1, \quad S_n = S_0 + \sum_{i=1}^n \xi_i, \quad (8.1)$$

où S_0 représente la fortune initiale. Si chaque pas de temps correspond à un tirage au sort indépendant, comme à la roulette, les $\{\xi_n\}_{n \geq 1}$ forment une suite de variables aléatoires indépendantes à valeurs dans $\mathbb{R}$. Comme l'issue de chaque partie est totalement incertaine, on pourra supposer que le gain moyen à chaque étape est nul, i.e. $\mathbb{E}(\xi_n) = 0$.

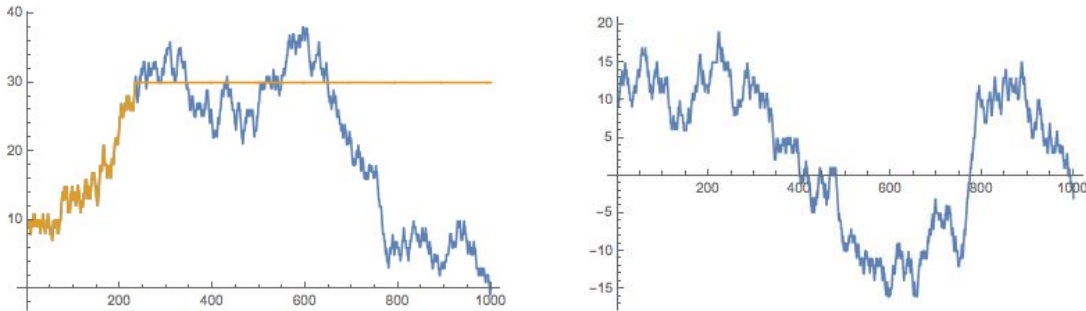


FIGURE 8.1 – À gauche, une réalisation de la marche aléatoire S_n qui représente la fortune du joueur après 1000 pas en partant initialement de $S_0 = 10$. Le graphe orange correspond à la fortune d'un joueur qui décide de s'arrêter de jouer dès qu'il a atteint la valeur 30 (dans cette réalisation, ceci se fait après 233 pas). Le reste de la courbe, en bleu, montre les gains que ce joueur aurait pu obtenir dans le futur. À droite, une autre réalisation de la marche aléatoire S_n . Cette fois, le joueur est ruiné avant d'avoir pu atteindre la valeur 30 !

Le processus $\{S_n\}_{n \geq 0}$ est une marche aléatoire dans $\mathbb{R}$ et il constitue notre premier

exemple de martingale. Comme les gains à chaque étape sont de moyenne nulle, l'espérance du gain reste constante au cours du temps $\mathbb{E}(S_n) = S_0$. Le joueur aimerait trouver une stratégie qui lui permette de gagner à coup sûr sans avoir à tricher, i.e. sans connaître ou influencer l'issue des parties futures. Par exemple au lieu de s'arrêter à un temps n fixé, une stratégie possible serait de jouer jusqu'à accumuler une fortune supérieure à la mise initiale, puis de quitter le jeu pour ne pas perdre ce gain (cf. figure 8.1). La notion de temps d'arrêt, déjà utilisée pour les chaînes de Markov (cf. section 2.4), jouera un rôle clef dans la construction des stratégies. La figure 8.1 montre cependant que la stratégie qui consiste à s'arrêter à partir d'un certain seuil présente des risques et que le joueur n'est pas à l'abri de finir ruiné. Dans la suite de ce cours, nous prouverons qu'il n'existe pas de stratégie pour gagner à coup sûr. Nous verrons aussi comment les martingales peuvent servir en finance pour calculer le prix d'une option et pour déterminer des stratégies de couverture.

La marche aléatoire définie en (8.1) est à la fois une chaîne de Markov et une martingale. Ces 2 points de vue sont complémentaires et permettent d'obtenir des informations différentes sur les processus aléatoires. Plus généralement les martingales servent à modéliser des fluctuations dans des séries temporelles, comme les cours de la bourse ou les signaux bruités, avec une dépendance par rapport au passé d'une autre nature que celle imposée par la propriété de Markov. La théorie des martingales permettra aussi de caractériser des comportements asymptotiques différents de ceux rencontrés dans les chaînes de Markov. Pour illustrer ces nouveaux résultats, nous allons définir la notion de série aléatoire.

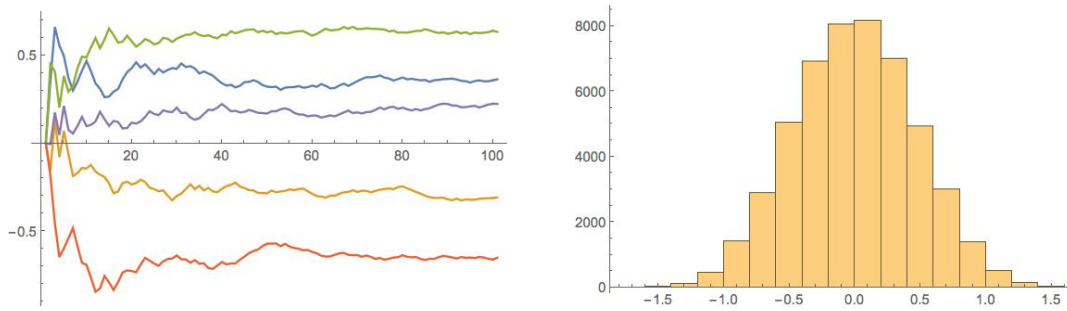


FIGURE 8.2 – Le graphe de gauche représente 5 trajectoires de longueur 100 de la martingale (8.2) avec $\alpha = 1$. On remarque des fluctuations importantes initialement, puis chaque trajectoire semble converger vers une limite aléatoire. Le graphe de droite est la distribution de M_{100} obtenue à partir de 50000 tirages aléatoires.

Considérons une suite $\{\zeta_i\}_{i \geq 1}$ de variables aléatoires indépendantes et distribuées uniformément dans $[-1, 1]$. Pour tout $\alpha > 0$, nous montrerons que le processus

$$\forall n \geq 1, \quad M_n = \sum_{i=1}^n \frac{\zeta_i}{i^\alpha}. \quad (8.2)$$

est une martingale. Comme $\mathbb{E}(\zeta_1) = 0$, l'espérance de M_n est nulle pour tout n . Les variables ζ_i étant bornées, la martingale M_n va converger si $\alpha > 1$ pour toute réalisation de la suite $\{\zeta_i\}_{i \geq 1}$. Par contre la limite M_∞ sera une variable aléatoire car elle dépend

de chaque réalisation des $\{\zeta_i\}_{i \geq 1}$ (cf. figure 8.2). On peut se demander si la convergence reste valable quand α est inférieur à 1. En effet, les variables ζ_i sont de moyenne nulle et des compensations devraient s'opérer entre les différents termes de la série. Au chapitre 11, nous prouverons par la théorie des martingales que la série converge presque sûrement dès que $\alpha > 1/2$.

Le formalisme des martingales va nous permettre d'étudier de nouveaux types de processus aléatoires et d'établir des résultats complémentaires à ceux obtenus par les chaînes de Markov. À l'aide des martingales, nous modéliserons l'effet du renforcement qui a des applications en apprentissage statistique et dans les réseaux sociaux. Nous analyserons aussi les fluctuations des processus aléatoires car elles peuvent jouer un rôle déterminant, par exemple dans la dérive génétique ou dans la modélisation de la turbulence. Au chapitre 13, nous construirons des stratégies pour optimiser les prises de décision.

La théorie des martingales nécessite de considérer des processus aléatoires à valeurs dans $\mathbb{R}$. Ceci suppose de généraliser la notion de conditionnement aux variables continues et un chapitre préliminaire sera donc nécessaire pour définir l'espérance conditionnelle avant d'introduire les martingales.

Chapitre 9

Espérance conditionnelle

9.1 Espérance conditionnelle sur un espace d'états discret

À l'occasion d'un sondage, des gens sont interrogés et ils doivent attribuer une note de 1 à 100 pour le nouveau produit qu'ils viennent de tester. On peut modéliser l'ensemble des réponses par une variable aléatoire X à valeurs dans $E = \{1, \dots, 100\}$. Le résultat du sondage sera la moyenne des réponses et il donnera une bonne estimation de $\mathbb{E}(X)$. Pour affiner le sondage, on voudrait classer les réponses en fonction de la personne sondée selon son genre, son âge ou la couleur de ses cheveux. Par exemple si Y correspond à l'âge de la personne sondée, la *probabilité conditionnelle* permet de déterminer la probabilité qu'une personne d'âge y attribue la note x

$$\mathbb{P}(X = x|Y = y) = \frac{\mathbb{P}(X = x, Y = y)}{\mathbb{P}(Y = y)}. \quad (9.1)$$

Cette expression n'a un sens que si $\mathbb{P}(Y = y) > 0$. On a ainsi défini une nouvelle mesure de probabilité sur E appelée *probabilité conditionnelle*

$$\sum_{x \in E} \mathbb{P}(X = x|Y = y) = 1$$

et le résultat du sondage pour la classe d'âge y correspondra à la note moyenne attribuée sous cette probabilité, i.e. l'espérance sous la probabilité conditionnelle

$$\mathbb{E}(X|Y = y) = \sum_{x \in E} x \mathbb{P}(X = x|Y = y).$$

La fonction $\mathbb{E}(X|Y = y)$ ne dépend que de y et elle fournit (en général) une information plus précise que $\mathbb{E}(X)$. On définit l'*espérance conditionnelle* $\mathbb{E}(X|Y)$ comme la fonction de Y qui prend la valeur $\mathbb{E}(X|Y = y)$ quand $Y = y$. Il existe donc une fonction $h(Y)$ telle que

$$h(Y) = \mathbb{E}(X|Y) \quad \text{et} \quad \forall y, \quad h(y) = \mathbb{E}(X|Y = y).$$

Si Y prend ses valeurs dans l'espace discret E' , on peut retrouver $\mathbb{E}(X)$ en intégrant sur Y

$$\begin{aligned}\mathbb{E}(\mathbb{E}(X|Y)) &= \sum_{y \in E'} \mathbb{P}(Y = y) \mathbb{E}(X|Y = y) = \sum_{x \in E} x \sum_{y \in E'} \mathbb{P}(Y = y) \mathbb{P}(X = x|Y = y) \\ &= \sum_{x \in E} x \sum_{y \in E'} \mathbb{P}(X = x, Y = y) = \sum_{x \in E} x \mathbb{P}(X = x) = \mathbb{E}(X).\end{aligned}\quad (9.2)$$

Si $\mathbb{E}(X^2)$ est finie, l'espérance $\mathbb{E}(X)$ peut être vue comme la meilleure approximation de X par une constante car elle minimise la distance quadratique

$$\mathbb{E}\left((X - \mathbb{E}(X))^2\right) = \inf_{c \in \mathbb{R}} \mathbb{E}\left((X - c)^2\right) = \inf_{c \in \mathbb{R}} \left\{ \mathbb{E}(X^2) - 2c\mathbb{E}(X) + c^2 \right\}.$$

Le minimum du polynôme est atteint en $c = \mathbb{E}(X)$. L'espérance conditionnelle $\mathbb{E}(X|Y)$ s'interprète aussi comme la meilleure approximation (pour la distance quadratique) de la variable X par la variable Y . Si Y prend ses valeurs dans l'espace discret E' , on cherche à déterminer la fonction $h : E' \rightarrow \mathbb{R}$ qui réalise le minimum

$$\inf_h \mathbb{E}\left((X - h(Y))^2\right) = \inf_h \left\{ \mathbb{E}(X^2) - 2\mathbb{E}(h(Y)X) + \mathbb{E}(h(Y)^2) \right\}.$$

La distance quadratique est minimisée par le projeté orthogonal de X sur l'espace des variables $\mathcal{H} = \{h(Y); h : E' \rightarrow \mathbb{R}\}$ (cf. figure 9.1). Pour déterminer cette projection, reprenons le calcul (9.2)

$$\begin{aligned}\mathbb{E}(h(Y)\mathbb{E}(X|Y)) &= \sum_{y \in E'} \mathbb{P}(Y = y) h(y) \mathbb{E}(X|Y = y) \\ &= \sum_{x \in E} \sum_{y \in E'} x h(y) \mathbb{P}(X = x, Y = y) = \mathbb{E}(h(Y)X).\end{aligned}$$

On se restreint aux fonctions $\mathbb{E}(h(Y)^2) < \infty$ pour que les espérances ci-dessus soient bien définies. Ceci prouve la relation d'orthogonalité pour toute fonction h de E' dans $\mathbb{R}$

$$\mathbb{E}\left(h(Y)(X - \mathbb{E}(X|Y))\right) = 0$$

et permet de résoudre le problème variationnel

$$\begin{aligned}\inf_h \mathbb{E}\left((X - h(Y))^2\right) &= \inf_h \left\{ \mathbb{E}(X^2) - 2\mathbb{E}(h(Y)X) + \mathbb{E}(h(Y)^2) \right\} \\ &= \inf_h \left\{ \mathbb{E}(X^2) - 2\mathbb{E}(h(Y)\mathbb{E}(X|Y)) + \mathbb{E}(h(Y)^2) \right\} \\ &= \inf_h \left\{ \mathbb{E}(X^2) - \mathbb{E}(\mathbb{E}(X|Y)^2) + \mathbb{E}\left((h(Y) - \mathbb{E}(X|Y))^2\right) \right\} \\ &= \mathbb{E}(X^2) - \mathbb{E}(\mathbb{E}(X|Y)^2).\end{aligned}\quad (9.3)$$

L'unique minimum est donc atteint pour $h(Y) = \mathbb{E}(X|Y)$. L'interprétation de l'espérance conditionnelle en terme de projection orthogonale sera particulièrement utile par la suite pour les variables à valeurs dans $\mathbb{R}$.

Supposons maintenant que le couple (X, Y) soit à valeurs dans $\mathbb{R}^2$, on aimerait définir l'espérance conditionnelle en s'inspirant du cas discret

$$\mathbb{P}(X = x | Y = y) = \lim_{\varepsilon \rightarrow 0} \mathbb{P}(X \in [x - \varepsilon, x + \varepsilon] \mid Y \in [y - \varepsilon, y + \varepsilon])$$

en général il n'est pas facile de donner un sens à ces expressions et nous allons renoncer à cette approche intuitive pour utiliser le point de vue de l'approximation quadratique illustré dans le cas discret.

9.2 Définition de l'espérance conditionnelle

Cette section résume quelques résultats nécessaires pour généraliser l'espérance conditionnelle aux variables à valeurs dans $\mathbb{R}^n$. Les preuves complètes figurent dans les annexes A et B.

Soit $X : \Omega \rightarrow \mathbb{R}$ une variable aléatoire à valeurs dans $\mathbb{R}$. On souhaite attribuer une probabilité aux événements de la forme $\{X \leq a\}$ pour tout a de $\mathbb{R}$, mais aussi évaluer la probabilité des complémentaires et des intersections entre tous ces ensembles. Ceci nous amène à définir la notion de σ -algèbre.

Définition 9.1 (σ -algèbre). Une σ -algèbre $\mathcal{A}$ sur un espace Ω est une famille d'événements satisfaisant les trois propriétés suivantes :

- (i) Ω appartient à $\mathcal{A}$.
- (ii) Si A est dans $\mathcal{A}$ alors A^c est dans $\mathcal{A}$.
- (iii) Toute réunion dénombrable d'événements de $\mathcal{A}$ appartient à $\mathcal{A}$.

Si $\mathcal{C}$ est une collection d'événements on notera $\sigma(\mathcal{C})$ la plus petite σ -algèbre contenant $\mathcal{C}$. On dira que $\sigma(\mathcal{C})$ est la σ -algèbre engendrée par $\mathcal{C}$.

Dans $\mathbb{R}$, la σ -algèbre engendrée par les intervalles de la forme $] - \infty, a]$ est la σ -algèbre borélienne et sera notée $\mathcal{B}_{\mathbb{R}}$ (voir aussi (A.1) en annexe). Pour $\mathbb{R}^n$, la σ -algèbre borélienne $\mathcal{B}_{\mathbb{R}^n}$ est engendrée par les ensembles de la forme $\otimes_{i=1}^n] - \infty, a_i]$.

Une σ -algèbre constitue le bon cadre théorique pour définir une mesure de probabilité (cf. théorème A.7). On appelle alors *espace de probabilité* le triplet $(\Omega, \mathcal{A}, \mathbb{P})$ où $\mathcal{A}$ est une σ -algèbre sur Ω et $\mathbb{P} : \mathcal{A} \rightarrow [0, 1]$ est une mesure de probabilité. Si $\mathcal{A}$ est une σ -algèbre sur Ω , on dira que la variable $X : \Omega \rightarrow \mathbb{R}^n$ est mesurable par rapport à $\mathcal{A}$ (on abrège souvent par $\mathcal{A}$ -mesurable), si tous les événements $\{X \in B\}$ pour $B \in \mathcal{B}_{\mathbb{R}^n}$ appartiennent à $\mathcal{A}$. On peut ainsi mesurer $\mathbb{P}(\{X \in B\})$.

Les σ -algèbres considérées dans ce cours seront souvent construites à partir de variables aléatoires et serviront à décrire l'information codée par ces variables. Si X est une variable aléatoire à valeurs dans $\mathbb{R}$, il est naturel de considérer la σ -algèbre engendrée par les ensembles de la forme $\{X \leq a\}$. On la notera $\sigma(X)$. On remarquera que $\sigma(X)$ contient aussi tous les événements de la forme $\{X \in B\}$ pour B appartenant à $\mathcal{B}_{\mathbb{R}}$. Plus généralement, si $X = (X_1, \dots, X_n)$ est une variable aléatoire à valeurs dans $\mathbb{R}^n$, alors la σ -algèbre $\sigma(X)$ associée est engendrée par les événements de la forme $\{X \in B\}$ pour B appartenant à $\mathcal{B}_{\mathbb{R}^n}$. Cette σ -algèbre est aussi notée $\sigma(X_1, \dots, X_n)$ car elle décrit la collection des n variables aléatoires $X_1, \dots, X_n$.

Le résultat suivant permet de représenter facilement toutes les variables aléatoires mesurables par rapport à $\sigma(X)$.

Lemme 9.2. *On considère deux variables aléatoires sur $(\Omega, \mathcal{A}, \mathbb{P})$: Y à valeurs dans $(E, \mathcal{E})$ et W à valeurs dans $\mathbb{R}$. Alors W est mesurable par rapport à $\sigma(Y)$ si et seulement s'il existe une fonction mesurable $f : E \mapsto \mathbb{R}$ telle que $W = f(Y)$.*

Cette caractérisation des variables $\sigma(X)$ -mesurables sera fondamentale pour la suite et sa démonstration se trouve au lemme B.4.

Le formalisme précédent va nous permettre de définir la notion d'espérance conditionnelle pour des variables aléatoires à valeurs dans $\mathbb{R}$ comme une projection orthogonale en s'inspirant des variables aléatoires à valeurs dans un espace discret. Considérons un espace de probabilité $(\Omega, \mathcal{A}, \mathbb{P})$ et deux variables X et Y à valeurs dans $\mathbb{R}$, mesurables par rapport à $\mathcal{A}$. On suppose que X est dans $\mathbb{L}^2$, i.e. que $\mathbb{E}(X^2) < \infty$. Le lemme 9.2 permet d'affirmer que les variables mesurables par rapport à $\sigma(Y)$ sont de la forme $h(Y)$ où $h : \mathbb{R} \mapsto \mathbb{R}$ est une fonction mesurable. On cherche à approcher X en fonction des variables du sous-espace

$$\mathcal{H} = \left\{ h(Y); \quad h \text{ fonction mesurable telle que } \mathbb{E}(h(Y)^2) < \infty \right\}. \quad (9.4)$$

Comme $\mathcal{H}$ est un sous-espace vectoriel fermé, on peut définir la projection orthogonale de X sur $\mathcal{H}$ pour le produit scalaire $\langle Z, W \rangle = \mathbb{E}(ZW)$. On note cette projection $\mathbb{E}(X|Y)$ (cf. figure 9.1) et elle satisfait la relation d'orthogonalité pour tout $h(Y)$ dans $\mathcal{H}$

$$\mathbb{E}\left((X - \mathbb{E}(X|Y))h(Y)\right) = 0 \quad \Leftrightarrow \quad \mathbb{E}(Xh(Y)) = \mathbb{E}(\mathbb{E}(X|Y)h(Y)). \quad (9.5)$$

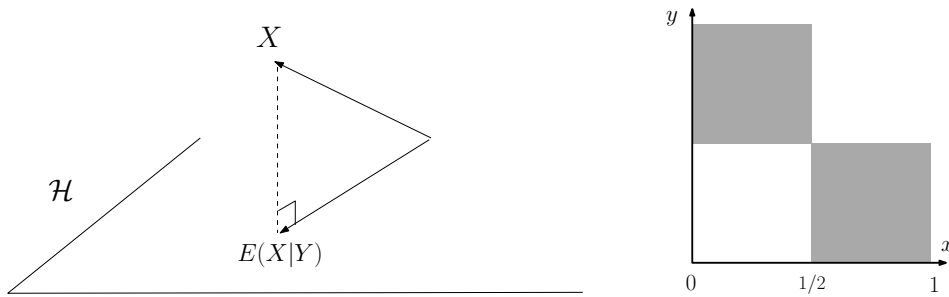


FIGURE 9.1 – L'espérance conditionnelle $\mathbb{E}(X|Y)$ s'interprète comme la projection orthogonale au sens $\mathbb{L}^2$ de X sur $\mathcal{H}$ (schéma de gauche). Le schéma de droite représente la densité $f_{X,Y}$ de l'exemple (9.7).

Un calcul identique à celui fait en (9.3) dans le cas discret montre que $\mathbb{E}(X|Y)$ est la meilleure prédiction possible de X (au sens $\mathbb{L}^2$) par la variable Y

$$\inf_{h \in \mathcal{H}} \mathbb{E}\left((X - h(Y))^2\right) = \mathbb{E}(X^2) - \mathbb{E}(\mathbb{E}(X|Y)^2).$$

Par ailleurs toutes les variables mesurables par rapport à $\sigma(Y)$ s'écrivent sous la forme $h(Y)$; la variable $\mathbb{E}(X|Y)$ est donc $\sigma(Y)$ -mesurable. On utilisera aussi la notation équivalente $\mathbb{E}(X|\sigma(Y)) = \mathbb{E}(X|Y)$ pour insister sur le fait que le conditionnement par rapport à Y revient à conditionner par rapport à toute l'information dans la σ -algèbre $\sigma(Y)$.

Cette stratégie s'étend aux variables aléatoires intégrables.

Théorème 9.3. *Soit X une variable aléatoire $\mathcal{A}$ -mesurable appartenant à $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$, i.e. telle que $\mathbb{E}(|X|) < \infty$. On considère $\mathcal{F} \subset \mathcal{A}$ une autre σ -algèbre. Il existe une unique variable aléatoire Z (définie presque sûrement) telle que*

- (a) Z est $\mathcal{F}$ -mesurable,
- (b) $\mathbb{E}(|Z|) < \infty$,
- (c) Pour toute variable aléatoire W bornée et $\mathcal{F}$ -mesurable alors $\mathbb{E}(XW) = \mathbb{E}(ZW)$.

On définit l'espérance conditionnelle de X sachant $\mathcal{F}$ par $\mathbb{E}(X|\mathcal{F}) := Z$.

La propriété (c) est l'analogue de la condition d'orthogonalité (9.5), mais l'espace $\mathcal{H}$ est réduit aux variables aléatoires bornées et $\mathcal{F}$ -mesurables. Le théorème précédent se déduit facilement du cas $\mathbb{L}^2$ en approchant la variable X par des variables dans $\mathbb{L}^2$, la preuve est faite en annexe (cf. théorème B.25).

Dans la pratique, on possède une information sur des variables aléatoires $Y_1, \dots, Y_k$ et on cherche à déterminer l'espérance conditionnelle sachant $\{Y_1, \dots, Y_k\}$. On note

$$\mathbb{E}(X|Y_1, \dots, Y_k) := \mathbb{E}(X|\mathcal{F}),$$

où $\mathcal{F} = \sigma(Y_1, \dots, Y_k)$ est la σ -algèbre engendrée par $Y_1, \dots, Y_k$. Dans ce cas, on déduit du lemme 9.2 que la propriété (c) du théorème 9.3 est équivalente à

$$\forall h \in \mathbb{L}^\infty(\mathcal{B}_{\mathbb{R}^k}), \quad \mathbb{E}(Xh(Y_1, \dots, Y_k)) = \mathbb{E}(\mathbb{E}(X|Y_1, \dots, Y_k)h(Y_1, \dots, Y_k)) \quad (9.6)$$

où $\mathbb{L}^\infty(\mathcal{B}_{\mathbb{R}^k})$ est l'ensemble des fonctions boréliennes bornées.

L'espérance conditionnelle $\mathbb{E}(X|Y)$ fournit une prédiction sur X sachant Y . On distingue deux cas extrêmes :

- Si $Y = c$ est constante. La σ -algèbre associée se réduit à $\sigma(Y) = \{\emptyset, \Omega\}$, c'est la plus petite σ -algèbre et elle correspond à l'absence totale d'information. La condition (a) du théorème 9.3 dit que Z est déterministe, i.e. $Z = \mathbb{E}(Z)$, et la condition d'orthogonalité (c) permet d'identifier cette constante $\mathbb{E}(X) = \mathbb{E}(Z) = Z$. Dans ce cas, l'espérance conditionnelle se confond avec l'espérance $\mathbb{E}(X|Y) = \mathbb{E}(X)$.
- Si $Y = X$, toute l'information sur X est connue et $\mathbb{E}(X|X) = X$. En effet, la variable $Z = X$ vérifie les propriétés (a), (b) et (c) du théorème 9.3 avec $\mathcal{F} = \sigma(X)$. Il s'agit de la meilleure prédiction possible de X (sachant X !).

Espérance conditionnelle pour des variables à densité.

Revenons maintenant sur le cas particulier des variables à densité. Supposons que le couple (X, Y) soit à valeurs dans $\mathbb{R}^2$ et admette une distribution absolument continue par rapport à la mesure de Lebesgue dans $\mathbb{R}^2$, de densité $f_{(X,Y)}(x,y)dxdy$. Définissons la loi

marginale de Y obtenue par intégration par rapport à la variable x

$$f_Y(y) = \int f_{(X,Y)}(x', y) dx'$$

et pour $f_Y(y) > 0$, la probabilité conditionnelle par

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

que l'on interprète comme l'analogie de (9.1) dans ce contexte. Pour tout y fixé tel que $f_Y(y) > 0$, la fonction $f_{X|Y=y}(x)$ définit une densité sur $\mathbb{R}$ et on peut lui associer

$$\mathbb{E}(X|Y=y) = \int x f_{X|Y=y}(x) dx.$$

Comme dans le cas discret, $\mathbb{E}(X|Y) = \varphi(Y)$ est une fonction de Y et elle définit une variable aléatoire $\mathbb{E}(X|Y)$ appelée espérance conditionnelle de X sachant Y . La condition d'orthogonalité (c) du théorème 9.3 se vérifie aussi par un calcul direct : pour toute fonction $h : \mathbb{R} \rightarrow \mathbb{R}$ bornée

$$\begin{aligned} \mathbb{E}(\mathbb{E}(X|Y) h(Y)) &= \int dy f_Y(y) h(y) \int x f_{X|Y=y}(x) dx \\ &= \iint x h(y) f_{(X,Y)}(x, y) dx dy = \mathbb{E}(X h(Y)) \end{aligned}$$

où on a utilisé le théorème de Fubini. On retrouve ainsi la propriété (9.6).

Pour conclure prenons l'exemple de la densité

$$\forall x, y \in [0, 1]^2, \quad f_{X,Y}(x, y) = 2(\mathbf{1}_{\{x \leq 1/2, y \geq 1/2\}} + \mathbf{1}_{\{x > 1/2, y \leq 1/2\}}) \quad (9.7)$$

représentée figure 9.1. Le calcul des marginales permet de voir que les variables considérées séparément sont uniformément distribuées sur $[0, 1]$

$$\forall x, y \in [0, 1], \quad f_X(x) = 1, \quad f_Y(y) = 1.$$

Par contre, la probabilité conditionnelle de X sachant Y est ou bien uniforme sur $[0, 1/2]$ ou bien sur $[1/2, 1]$

$$f_{X|Y=y}(x) = 2(\mathbf{1}_{\{x \leq 1/2, y \geq 1/2\}} + \mathbf{1}_{\{x > 1/2, y \leq 1/2\}}).$$

L'espérance conditionnelle est donnée par

$$\mathbb{E}(X|Y) = \frac{1}{4} \mathbf{1}_{\{Y \geq 1/2\}} + \frac{3}{4} \mathbf{1}_{\{Y < 1/2\}}.$$

9.3 Propriétés de l'espérance conditionnelle

Commençons par les propriétés déjà esquissées dans le paragraphe précédent. Soit $\mathcal{F}$ une σ -algèbre, par exemple $\mathcal{F} = \sigma(Y_1, \dots, Y_k)$.

Proposition 9.4. *L'espérance conditionnelle $\mathbb{E}(\cdot|\mathcal{F})$ est linéaire et pour tout $X \in \mathbb{L}^1(\mathcal{A}, \mathbb{P})$*

- (i) $\mathbb{E}(\mathbb{E}(X|\mathcal{F})) = \mathbb{E}(X)$,
- (ii) Si $X \geq 0$, alors $\mathbb{E}(X|\mathcal{F}) \geq 0$ presque sûrement,
- (iii) Si X est $\mathcal{F}$ -mesurable, alors $\mathbb{E}(X|\mathcal{F}) = X$ presque sûrement.

L'espérance conditionnelle jouit des mêmes propriétés de passage à la limite que l'espérance. Le résultat suivant est démontré proposition B.26.

Proposition 9.5. *Soient $\{X_n\}_{n \geq 1}$ et X des variables aléatoires dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$. L'espérance conditionnelle satisfait les résultats de convergence suivants :*

- (i) (Convergence monotone)
Si $X_n \geq 0$ converge presque sûrement vers X en croissant, alors

$$\lim_{n \rightarrow \infty} \uparrow \mathbb{E}(X_n|\mathcal{F}) = \mathbb{E}(X|\mathcal{F}).$$

- (ii) (Lemme de Fatou)
Si $X_n \geq 0$, alors

$$\mathbb{E}\left(\liminf_{n \rightarrow \infty} X_n \mid \mathcal{F}\right) \leq \liminf_{n \rightarrow \infty} \mathbb{E}(X_n \mid \mathcal{F}).$$

- (iii) (Convergence dominée)
S'il existe Y dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ tel que $\sup_n |X_n| \leq Y$ et $\{X_n\}_{n \geq 1}$ converge vers X presque sûrement alors

$$\lim_{n \rightarrow \infty} \mathbb{E}(X_n|\mathcal{F}) = \mathbb{E}(X|\mathcal{F}).$$

Les inégalités classiques pour l'espérance (cf. la section B.2.2 de l'appendice) restent valables pour les espérances conditionnelles. En particulier, les généralisations de l'inégalité de Jensen (théorème B.8) et de l'inégalité de Hölder (proposition B.5) sont écrites ci-dessous.

Proposition 9.6 (Inégalité de Jensen). *Soit X dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ et $g : \mathbb{R}^n \rightarrow \mathbb{R} \cup \{\infty\}$ une fonction convexe telle que $\mathbb{E}(|g(X)|) < \infty$. Alors $\mathbb{E}(g(X)|\mathcal{F}) \geq g(\mathbb{E}(X|\mathcal{F}))$.*

Proposition 9.7 (Inégalité de Hölder). *Soient X, Y deux variables aléatoires à valeurs dans $\mathbb{R}$ telles que*

$$\mathbb{E}(|X|^p) < \infty \quad \text{et} \quad \mathbb{E}(|Y|^q) < \infty \quad \text{avec} \quad p > 1 \quad \text{et} \quad \frac{1}{p} + \frac{1}{q} = 1.$$

L'inégalité de Hölder pour l'espérance conditionnelle par rapport à la σ -algèbre $\mathcal{F}$ s'écrit

$$\mathbb{E}(|XY| \mid \mathcal{F}) \leq \mathbb{E}(|X|^p \mid \mathcal{F})^{1/p} \mathbb{E}(|Y|^q \mid \mathcal{F})^{1/q}. \quad (9.8)$$

On considère maintenant des conditionnements successifs appliqués à deux σ -algèbres $\mathcal{F}, \mathcal{G}$. Commençons par décrire les liens possibles entre ces deux σ -algèbres.

- Si $\mathcal{F} \subset \mathcal{G}$, on dira alors que $\mathcal{G}$ contient plus d'information que $\mathcal{F}$. C'est le cas par exemple si $\mathcal{G} = \sigma(Y_1, \dots, Y_k)$ et $\mathcal{F} = \sigma(Y_1, \dots, Y_{k'})$ avec $k > k'$. En effet, la σ -algèbre $\mathcal{G}$ possède plus d'ensembles mesurables que $\mathcal{F}$.
- Si X et Y sont deux variables aléatoires indépendantes, alors les σ -algèbres $\mathcal{F} = \sigma(X)$ et $\mathcal{G} = \sigma(Y)$ sont dites indépendantes car les ensembles de $\mathcal{F}$ sont indépendants de ceux de $\mathcal{G}$. Plus généralement, la notion de σ -algèbres indépendantes est définie dans la section B.4.1.

La proposition suivante résume les propositions B.27, B.28 et B.29 prouvées en annexe.

Proposition 9.8. Soit $X \in \mathbb{L}^1(\mathcal{A}, \mathbb{P})$ et $\mathcal{F}, \mathcal{G}$ des sous- σ -algèbres de $\mathcal{A}$.

- (i) Si $\mathcal{F} \subset \mathcal{G}$, alors $\mathbb{E}(\mathbb{E}(X|\mathcal{G})|\mathcal{F}) = \mathbb{E}(X|\mathcal{F})$.
- (ii) Si $\mathcal{G}$ est indépendante de $\sigma(\sigma(X), \mathcal{F})$, alors $\mathbb{E}(X|\sigma(\mathcal{F}, \mathcal{G})) = \mathbb{E}(X|\mathcal{F})$.
- (iii) Si Y est une variable aléatoire mesurable par rapport à $\mathcal{F}$ et $\mathbb{E}(|XY|) < \infty$, alors

$$\mathbb{E}(XY|\mathcal{F}) = Y\mathbb{E}(X|\mathcal{F}).$$

Pour illustrer cette proposition, supposons que X mesure le rendement d'une réaction chimique qui dépend de nombreux paramètres codés par la σ -algèbre $\mathcal{A}$: la température T , la pression P , l'habileté des expérimentateurs, etc. Dans (i), si les paramètres T et P sont connus le résultat moyen sachant $\mathcal{G} = \sigma(T, P)$ est $\mathbb{E}(X|\mathcal{G})$. Si seul T est déterminé, il faut intégrer sur toutes les valeurs possibles de la pression P pour obtenir $\mathbb{E}(X|\mathcal{F})$ l'espérance conditionnelle sachant $\mathcal{F} = \sigma(T)$. La relation (ii) dit que rajouter une information $\mathcal{G}$ qui n'a rien à voir avec cette expérience ne permet pas d'améliorer la prédiction $\mathbb{E}(X|\mathcal{F})$. Le dernier point (iii) est l'analogue du (iii) dans la proposition 9.4. Si $Y = h(T)$ ne dépend que de T , la meilleure prédiction de Y sachant $\mathcal{F} = \sigma(T)$ est Y .

9.4 Processus aléatoire

Nous avons déjà rencontré la notion de processus aléatoire sous la forme d'une chaîne de Markov $\{X_n\}_{n \geq 0}$ à valeurs discrètes. Plus généralement, si $(\Omega, \mathcal{A}, \mathbb{P})$ désigne un espace de probabilité, un *processus aléatoire* $\{X_n\}_{n \geq 0}$ est une suite de variables aléatoires sur $(\Omega, \mathcal{A})$ à valeurs dans un espace mesurable $(E, \mathcal{E})$. Dans la suite, nous nous intéresserons principalement à des processus à valeurs dans $(\mathbb{R}, \mathcal{B}_{\mathbb{R}})$. L'indice n indique la date à laquelle la variable aléatoire X_n est observée. Afin de quantifier le déroulement du temps et la structure de l'information qui en découle, on introduit la notion de filtration.

Définition 9.9. (Filtration) Une filtration de $\mathcal{A}$ est une suite croissante $\mathbb{F} = \{\mathcal{F}_n\}_{n \geq 0}$ de sous- σ -algèbres de $\mathcal{A}$

$$\mathcal{F}_0 \subset \mathcal{F}_1 \subset \mathcal{F}_2 \subset \dots \subset \mathcal{A}.$$

On dit que $(\Omega, \mathcal{A}, \mathbb{F}, \mathbb{P})$ un espace de probabilité filtré.

En particulier, si $\{X_n\}_{n \geq 0}$ est un processus aléatoire. Alors la suite

$$\mathcal{F}_n = \sigma(X_i, \quad i \leq n), \quad n \geq 0,$$

est appelée la filtration naturelle du processus.

Pour chaque entier n , la sous- σ -algèbre $\mathcal{F}_n$ représente l'information disponible à la date n . La croissance de la suite $\{\mathcal{F}_n\}_{n \geq 0}$ traduit l'idée que l'information s'accumule au fil du temps et qu'il n'y a pas de possibilité d'oublier des informations passées. De la même façon qu'une variable aléatoire est associée à une σ -algèbre, un processus aléatoire est codé par une filtration.

Dès le chapitre 2, nous avons utilisé la notion de filtration, sans en évoquer le formalisme, pour décrire les chaînes de Markov. Si $\{X_n\}_{n \geq 0}$ est une chaîne de Markov sur un espace dénombrable E , alors $\mathcal{F}_n = \sigma(X_i, i \leq n)$ sera la filtration associée. Comme E est un espace discret, la σ -algèbre $\mathcal{F}_n$ est engendrée par les ensembles de la forme $\{X_0 = x_0, \dots, X_n = x_n\}$ pour toute famille $x_0, \dots, x_n$ d'éléments de E . Étant donné un événement $B \subset E$, la probabilité conditionnelle $\mathbb{P}(X_{n+1} \in B \mid \mathcal{F}_n)$ est une variable aléatoire qui dépend des valeurs prises par $X_0, \dots, X_n$. Par la propriété de Markov, on sait que cette dépendance se réduit à la variable X_n . On peut ainsi écrire

$$\mathbb{P}(X_{n+1} \in B \mid \mathcal{F}_n) = \mathbb{P}(X_{n+1} \in B \mid X_n).$$

Dans le cas discret, le formalisme de la section 9.1 s'applique : la probabilité conditionnelle $\mathbb{P}(X_{n+1} \in B \mid X_n)$ est une variable aléatoire qui s'exprime comme une fonction de X_n dont la valeur est $\mathbb{P}(X_{n+1} \in B \mid X_n = x_n)$ si $X_n = x_n$. L'espérance conditionnelle sachant X_n vérifie pour toute fonction $h : E \mapsto \mathbb{R}$

$$\mathbb{E}(h(X_{n+1}) \mid X_n) = \sum_{y \in E} h(y) \mathbb{P}(X_{n+1} = y \mid X_n).$$

Les notions suivantes seront utilisées au chapitre 13 pour définir des stratégies.

Définition 9.10. Soit $\mathbb{F} = \{\mathcal{F}_n\}_{n \geq 0}$ une filtration de $\mathcal{A}$.

- (i) Le processus $\{X_n\}_{n \geq 0}$ est dit adapté à la filtration $\mathbb{F}$ si X_n est $\mathcal{F}_n$ -mesurable pour tout $n \geq 0$.
- (ii) Le processus $\{X_n\}_{n \geq 0}$ est dit prévisible pour la filtration $\mathbb{F}$ si X_n est $\mathcal{F}_{n-1}$ -mesurable pour tout $n \geq 0$. Par convention, on note $\mathcal{F}_{-1} = \{\emptyset, \Omega\}$.

Les temps d'arrêt ont déjà été introduits section 2.4 car ils jouent un rôle central dans l'analyse des processus aléatoires. Ils peuvent être redéfinis en utilisant le formalisme des filtrations.

Définition 9.11 (Temps d'arrêt). Un temps d'arrêt T (pour la filtration $\mathbb{F} = \{\mathcal{F}_n\}_{n \geq 0}$) est une variable aléatoire à valeurs dans $\mathbb{N} \cup \{\infty\}$ telle que

$$\{T = n\} \in \mathcal{F}_n \quad \text{pour tout } n \geq 0.$$

Si $\mathcal{F}_n = \sigma(X_i, i \leq n)$ est la filtration naturelle, cette notion est équivalente à celle décrite section 2.4. L'événement $\{T = n\}$ étant mesurable par rapport à $\mathcal{F}_n$, il peut donc être exprimé en fonction des $(n+1)$ premières observations

$$1_{\{T=n\}} = \varphi_n(X_0, \dots, X_n)$$

où φ_n est une fonction mesurable.

On rappelle qu'une classe importante de temps d'arrêt correspond au premier temps d'atteinte d'un ensemble A

$$T_A = \inf \{n \geq 0; \quad X_n \in A\}$$

avec la convention $\inf \emptyset = \infty$.

Chapitre 10

Martingales et stratégies

Dans tout ce chapitre, on considérera des processus aléatoires à valeurs dans $\mathbb{R}$.

10.1 Martingales

Une martingale est un processus aléatoire dont l'espérance conditionnelle par rapport au passé reste constante. La notion de filtration, introduite dans la section 9.4, permet de donner une définition mathématique précise d'une martingale et de ses variantes.

Définition 10.1. (*Martingale*)

Soit $X = \{X_n\}_{n \geq 0}$ un processus aléatoire adapté sur l'espace de probabilité filtré $(\Omega, \mathcal{A}, \mathbb{F}, \mathbb{P})$. Si X_n est intégrable pour tout n (i.e. $\mathbb{E}(|X_n|) < \infty$), on dit que X est

— une martingale si

$$\mathbb{E}(X_n | \mathcal{F}_{n-1}) = X_{n-1} \quad \text{pour tout } n \geq 1,$$

— une surmartingale si

$$\mathbb{E}(X_n | \mathcal{F}_{n-1}) \leq X_{n-1} \quad \text{pour tout } n \geq 1,$$

— une sous-martingale si

$$\mathbb{E}(X_n | \mathcal{F}_{n-1}) \geq X_{n-1} \quad \text{pour tout } n \geq 1.$$

Il existe de nombreux exemples de tels processus.

Marche aléatoire symétrique.

Soit $\{\xi_n\}_{n \geq 1}$ une suite de variables aléatoires dans $\mathbb{R}$ indépendantes et de moyenne nulle $\mathbb{E}(\xi_n) = 0$. Alors la marche aléatoire $S = \{S_n\}_{n \geq 0}$

$$\forall n \geq 1, \quad S_n = \sum_{i=1}^n \xi_i \quad \text{et} \quad S_0 = 0 \quad (10.1)$$

est une martingale pour la filtration $\mathcal{F}_n = \sigma(\xi_i, i \leq n)$. Les variables aléatoires $\{\xi_n\}_{n \geq 1}$ peuvent représenter le gain d'un joueur (ou la perte) à la $n^{\text{ième}}$ partie d'un jeu de hasard.

La propriété de martingale se démontre facilement en utilisant la linéarité de l'espérance conditionnelle

$$\mathbb{E}(S_{n+1}|\mathcal{F}_n) = \mathbb{E}(S_n|\mathcal{F}_n) + \mathbb{E}(\xi_{n+1}|\mathcal{F}_n) = S_n + \mathbb{E}(\xi_{n+1}) = S_n$$

où on a utilisé que S_n est mesurable par rapport à $\mathcal{F}_n$ (Proposition 9.4 (iii)) et que ξ_{n+1} est indépendant de $\mathcal{F}_n$ (Proposition 9.8 (ii)).

Contrairement aux hypothèses faites sur les marches aléatoires dans le cadre des chaînes de Markov, nous n'avons pas supposé que les variables $\{\xi_n\}_{n \geq 1}$ prennent des valeurs entières, ni qu'elles soient identiquement distribuées.

Produit de variables aléatoires.

La structure multiplicative permet aussi d'obtenir des martingales en posant

$$\forall n \geq 1, \quad M_n = \prod_{i=1}^n \xi_i \quad \text{et} \quad M_0 = 1, \quad (10.2)$$

où les $\{\xi_n\}_{n \geq 0}$ sont des variables aléatoires indépendantes de moyenne $\mathbb{E}(\xi_n) = 1$. Le processus $\{M_n\}_{n \geq 0}$ est une martingale pour la filtration $\mathcal{F}_n = \sigma(\xi_i, i \leq n)$. Pour le voir il suffit d'utiliser que M_n est mesurable par rapport à $\mathcal{F}_n$ (Proposition 9.8 (iii)) et que ξ_{n+1} est indépendant de $\mathcal{F}_n$ (Proposition 9.8 (ii))

$$\mathbb{E}(M_{n+1}|\mathcal{F}_n) = \mathbb{E}(\xi_{n+1} M_n|\mathcal{F}_n) = M_n \mathbb{E}(\xi_{n+1}|\mathcal{F}_n) = M_n.$$

Chaînes de Markov.

Nous allons maintenant construire une martingale à partir d'une chaîne de Markov $\{X_n\}_{n \geq 0}$ de matrice de transition P sur un espace d'états E dénombrable. Rappelons la définition (3.7) d'une fonction harmonique h pour P

$$\forall x \in E, \quad h(x) = \sum_{y \in E} P(x, y) h(y).$$

Si $\mathbb{E}(|h(X_n)|)$ est fini pour tout n , alors $\{h(X_n)\}_{n \geq 0}$ est une martingale pour la filtration $\mathcal{F}_n = \sigma(X_1, \dots, X_n)$. En effet

$$\begin{aligned} \mathbb{E}(h(X_{n+1})|\mathcal{F}_n) &= \mathbb{E}(h(X_{n+1})|X_0, \dots, X_n) = \mathbb{E}(h(X_{n+1})|X_n) \\ &= \sum_{y \in E} P(X_n, y) h(y) = h(X_n), \end{aligned}$$

où la propriété de Markov a été utilisée dans la première ligne.

De la même façon

- si h est *surharmonique*, i.e. $h \geq Ph$, alors $\{h(X_n)\}_{n \geq 0}$ est une surmartingale,
- si h est *sous-harmonique*, i.e. $h \leq Ph$, alors $\{h(X_n)\}_{n \geq 0}$ est une sous-martingale.

Inégalité de Jensen.

Proposition 10.2. Soient $\{M_n\}_{n \geq 0}$ une martingale et $g : \mathbb{R} \rightarrow \mathbb{R}$ une application convexe telle que $\mathbb{E}(|g(M_n)|) < \infty$, alors le processus aléatoire $\{g(M_n)\}_{n \geq 0}$ est une sous-martingale.

En particulier, $\{|M_n|\}_{n \geq 0}$ est une sous-martingale.

Démonstration. Cette proposition est une conséquence de l'inégalité de Jensen pour l'espérance conditionnelle (proposition 9.6)

$$\mathbb{E}(g(M_{n+1}) | \mathcal{F}_n) \geq g(\mathbb{E}(M_{n+1} | \mathcal{F}_n)) = g(M_n). \quad (10.3)$$

□

La proposition 10.2 appliquée à la marche aléatoire $\{S_n\}_{n \geq 0}$ définie en (10.1) et à la fonction $g(x) = x^2$ montre que $\{S_n^2\}_{n \geq 0}$ est une sous-martingale. Par conséquent, son espérance croît au cours du temps (en supposant que $\mathbb{E}(\xi_1^2) < \infty$). Cette croissance peut être compensée et le processus

$$n \geq 0, \quad M_n = S_n^2 - n\mathbb{E}(\xi_1^2) \quad (10.4)$$

est alors une martingale. Pour le voir, décomposons l'espérance conditionnelle par la propriété de linéarité (proposition 9.4)

$$\begin{aligned} \mathbb{E}(M_{n+1} | \mathcal{F}_n) &= \mathbb{E}(S_n^2 + 2\xi_{n+1}S_n + \xi_{n+1}^2 | \mathcal{F}_n) - (n+1)\mathbb{E}(\xi_1^2) \\ &= \mathbb{E}(S_n^2 | \mathcal{F}_n) - n\mathbb{E}(\xi_1^2) + 2\mathbb{E}(\xi_{n+1}S_n | \mathcal{F}_n) + \mathbb{E}(\xi_{n+1}^2 | \mathcal{F}_n) - \mathbb{E}(\xi_1^2). \end{aligned}$$

En utilisant que S_n est mesurable par rapport à $\mathcal{F}_n$ et que ξ_{n+1} est indépendant de $\mathcal{F}_n$ on obtient

$$\begin{aligned} \mathbb{E}(S_n^2 | \mathcal{F}_n) - n\mathbb{E}(\xi_1^2) &= S_n^2 - n\mathbb{E}(\xi_1^2) = M_n, && \text{par la proposition 9.4 (iii),} \\ \mathbb{E}(\xi_{n+1}^2 | \mathcal{F}_n) &= \mathbb{E}(\xi_1^2), && \text{par la proposition 9.8 (ii),} \\ \mathbb{E}(\xi_{n+1}S_n | \mathcal{F}_n) &= S_n\mathbb{E}(\xi_{n+1} | \mathcal{F}_n) = S_n\mathbb{E}(\xi_{n+1}) = 0, && \text{par la proposition 9.8 (iii).} \end{aligned}$$

Par conséquent $\mathbb{E}(M_{n+1} | \mathcal{F}_n) = M_n$ et la sous-martingale $S_n^2 = M_n + n\mathbb{E}(\xi_1^2)$ se décompose comme la somme d'une martingale et d'un processus croissant.

Concluons cette section avec une propriété qui s'avérera très utile dans la suite.

Proposition 10.3. Soit $\{M_n\}_{n \geq 0}$ une martingale, alors pour tout $n \geq 0$

$$\forall k \geq 1, \quad \mathbb{E}(M_{n+k} | \mathcal{F}_n) = M_n. \quad (10.5)$$

Des inégalités similaires sont vérifiées pour des surmartingales et des sous-martingales.

Démonstration. On va montrer le résultat par récurrence dans le cas des martingales. Pour $k = 1$, il s'agit de la définition. Supposons que la propriété soit vraie au rang k . Comme la filtration $\mathcal{F}_{n+k}$ contient $\mathcal{F}_n$, on peut utiliser les espérances conditionnelles emboîtées (cf. proposition 9.8 (i)) pour écrire que

$$\begin{aligned} \mathbb{E}(M_{n+k+1} | \mathcal{F}_n) &= \mathbb{E}(\mathbb{E}(M_{n+k+1} | \mathcal{F}_{n+k}) | \mathcal{F}_n) \\ &= \mathbb{E}(M_{n+k} | \mathcal{F}_n) = M_n. \end{aligned}$$

La seconde égalité s'obtient par la propriété de martingale $M_{n+k} = \mathbb{E}(M_{n+k+1} | \mathcal{F}_{n+k})$, et la troisième par l'hypothèse de récurrence. □

On en déduit que l'espérance d'une martingale $\{M_n\}_{n \geq 0}$ est constante

$$\forall n \geq 0, \quad \mathbb{E}(M_n) = \mathbb{E}(M_0). \quad (10.6)$$

La réciproque est fautive, un processus d'espérance constante n'est pas toujours une martingale. Il suffit par exemple de considérer $S_n^3 = (\sum_{i=1}^n \xi_i)^3$ où les incréments ξ_i sont des variables aléatoires symétriques, indépendantes et vérifiant $\mathbb{E}(|\xi_1|^3) < \infty$. On verra en Proposition 10.8 une condition plus élaborée pour obtenir une forme de réciproque.

10.2 Stratégies et théorème d'arrêt

Considérons un joueur au casino qui à l'instant n va miser à la roulette la somme Φ_n sur le numéro 13. Si le 13 sort, le joueur empoche $\xi_n = 35$ fois sa mise sinon il perd et on pose $\xi_n = -1$. Son gain $\Phi_n \xi_n$ est alors proportionnel à sa mise Φ_n et ξ_n représente le résultat aléatoire du jeu. La filtration naturelle associée à ce jeu est $\mathcal{F}_n = \sigma(\xi_i, i \leq n)$ et le processus $\{\xi_n\}_{n \geq 1}$ des résultats est adapté à $\mathcal{F}_n$. Au temps n , le joueur mise avant de connaître le résultat ξ_n , son choix ne dépend que des résultats précédents. Le processus $\{\Phi_n\}_{n \geq 1}$ décrit la *stratégie* du joueur et il est prévisible, i.e. Φ_n est $\mathcal{F}_{n-1}$ -mesurable pour tout $n \geq 1$ (cf. définition 9.10). La fortune du joueur au temps n est alors

$$X_n = \sum_{k=1}^n \Phi_k \xi_k.$$

On peut généraliser cette structure

Proposition 10.4. Soit $M = \{M_n\}_{n \geq 0}$ une martingale et $\{\Phi_n\}_{n \geq 1}$ un processus prévisible borné, alors le processus défini par

$$X_0 = 0 \quad \text{et} \quad X_n = \sum_{k=1}^n \Phi_k (M_k - M_{k-1}), \quad n \geq 1 \quad (10.7)$$

est une martingale.

Si $\Phi_k \geq 0$ pour tout k et $\{M_n\}_{n \geq 0}$ est une surmartingale (resp. sous-martingale), alors $\{X_n\}_{n \geq 0}$ est aussi une surmartingale (resp. sous-martingale).

Démonstration. Par construction $\{X_n\}_{n \geq 0}$ est mesurable par rapport à $\mathcal{F}_n$. En remarquant que Φ_{n+1} est mesurable par rapport à $\mathcal{F}_n$, on obtient

$$\begin{aligned} \mathbb{E}(X_{n+1} | \mathcal{F}_n) &= \mathbb{E}(X_n | \mathcal{F}_n) + \mathbb{E}(\Phi_{n+1}(M_{n+1} - M_n) | \mathcal{F}_n) \\ &= X_n + \Phi_{n+1} \mathbb{E}(M_{n+1} - M_n | \mathcal{F}_n) = X_n \end{aligned}$$

où on a utilisé dans la dernière égalité que $\mathbb{E}(M_{n+1} - M_n | \mathcal{F}_n) = 0$ car $\{M_n\}_{n \geq 0}$ est une martingale. $\square$

Cette proposition montre que si $\{M_n\}_{n \geq 0}$ est une martingale, il n'existe aucune stratégie (dont les mises restent majorées) qui puisse transformer un jeu équitable en un jeu profitable. Quelle que soit la stratégie $\{\Phi_n\}_{n \geq 1}$ adoptée, la moyenne du gain est constante $\mathbb{E}(X_n) = \mathbb{E}(X_0)$. La proposition 10.4 jouera un rôle majeur dans la suite de ce

cours. On pourra notamment trouver des exemples de stratégies appliquées à la gestion d'actifs financiers dans la section 10.3.

Revenons à l'exemple du joueur décrit précédemment et supposons que le gain de la $k^{\text{ième}}$ partie est donné par $M_k - M_{k-1}$ où $\{M_n\}_{n \geq 0}$ est une martingale. Le joueur décide de miser à chaque fois 1 euro jusqu'à un temps d'arrêt T après lequel il s'arrête définitivement de jouer (cf. figure 10.1), sa fortune s'écrit alors

$$X_0 = 0 \quad \text{et} \quad X_n = \sum_{k=1}^n 1_{\{T \geq k\}} (M_k - M_{k-1}), \quad n \geq 1.$$

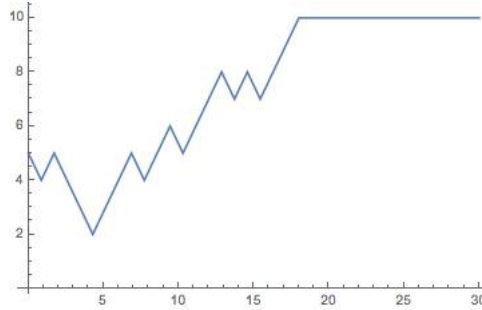


FIGURE 10.1 – Cet exemple représente un jeu de pile ou face dont les gains aléatoires suivent la loi $\mathbb{P}(\zeta_k = \pm 1) = 1/2$. Le joueur commence avec un capital égal à 5 et il arrête de jouer dès que sa fortune atteint la valeur 10. Ceci revient à utiliser le temps d'arrêt $T = \inf\{n \geq 0, M_n = 10\}$ dans la formule (10.8).

Comme le temps d'arrêt est mesurable par rapport à la σ -algèbre $\mathcal{F}_n$ (cf. définition 9.11) le processus $\Phi_n = 1_{\{T \geq n\}}$ est prévisible. En effet, l'évènement $\{T \geq n\}$ est mesurable par rapport à $\mathcal{F}_{n-1}$ car il se décompose uniquement en fonction d'évènements $\mathcal{F}_{n-1}$ -mesurables

$$\{T \geq n\} = \bigcap_{k=1}^{n-1} \{T \neq k\}.$$

La proposition 10.4 implique que $\{X_n\}_{n \geq 0}$ est donc une martingale. Il s'agit de la martingale M arrêtée au temps T que l'on notera $M^T = \{M_n^T\}_{n \geq 0}$

$$X_n = M_n^T = \begin{cases} M_n, & \text{si } n \leq T, \\ M_T, & \text{si } n \geq T. \end{cases} \quad (10.8)$$

Plus généralement, pour un processus aléatoire $Y = (Y_n)_{n \geq 0}$, on définit Y^T le processus arrêté au temps d'arrêt T par

$$Y_n^T = Y_{n \wedge T} \quad \text{pour tout } n \geq 0 \quad \text{où } n \wedge T = \inf\{n, T\}.$$

On déduit du calcul précédent et de la proposition 10.4

Proposition 10.5. Soient X une surmartingale (resp. sous-martingale, martingale) et T un temps d'arrêt sur $(\Omega, \mathcal{A}, \mathbb{F}, \mathbb{P})$. Alors le processus arrêté X^T est une surmartingale (resp. sous-martingale, martingale).

Le théorème suivant est une conséquence importante de la Proposition 10.5.

Théorème 10.6. (Théorème d'arrêt de Doob)

Soit X une martingale (resp. sous-martingale, surmartingale) et T un temps d'arrêt. Si une des trois propriétés est satisfaite

- (i) T est borné
- (ii) X est borné et T est fini presque sûrement
- (iii) $\mathbb{E}(T) < \infty$ et il existe une constante c telle que $\sup_n |X_n(\omega) - X_{n-1}(\omega)| \leq c$ presque sûrement

alors

$$\mathbb{E}(X_T) = \mathbb{E}(X_0). \quad (10.9)$$

Dans le cas d'une sous-martingale, on a alors $\mathbb{E}(X_T) \geq \mathbb{E}(X_0)$ et $\mathbb{E}(X_T) \leq \mathbb{E}(X_0)$ pour une surmartingale.

Démonstration. Nous montrons le résultat pour les martingales car le cas des sous-martingales et des surmartingales se traite de manière identique.

Par la proposition 10.5, le processus arrêté X^T est une martingale et son espérance reste constante pour tout n

$$\mathbb{E}(X_{T \wedge n}) = \mathbb{E}(X_0).$$

Pour prendre la limite $n \rightarrow \infty$, on examine successivement chaque condition :

- (i) T étant uniformément borné par c , il suffit de choisir $n \geq c$ pour conclure.
- (ii) T étant fini presque sûrement, on en déduit la convergence presque sûre

$$X_{n \wedge T} \xrightarrow{n \rightarrow \infty} X_T.$$

Comme X est borné, la suite de variables $\{X_{n \wedge T}\}_{n \geq 0}$ l'est aussi et le théorème de convergence dominée permet de conclure.

- (iii) Sous l'hypothèse $|X_n(\omega) - X_{n-1}(\omega)| \leq c$, on peut majorer $X_{n \wedge T}$ par

$$|X_{n \wedge T}| \leq \sum_{k=1}^{n \wedge T} |X_k - X_{k-1}| \leq cT.$$

Par conséquent la suite de variables $\{X_{n \wedge T}\}_n$ est dominée par une variable intégrable et elle converge presque sûrement vers X_T . Le théorème de convergence dominée permet une nouvelle fois de conclure.

□

Le théorème précédent est valable sous des hypothèses moins restrictives que les trois conditions mentionnées, cependant le contre-exemple ci-dessous montre que le résultat

ne peut pas être généralisé systématiquement. Considérons S la marche aléatoire symétrique (10.1)

$$S_n = \sum_{k=1}^n \xi_k \quad \text{avec} \quad \mathbb{P}(\xi_i = \pm 1) = \frac{1}{2} \quad (10.10)$$

et T_1 le premier temps d'atteinte de 1 pour cette marche. Comme S est une martingale, le processus arrêté S^{T_1} est aussi une martingale par la proposition 10.5. On a donc

$$\forall n \geq 0, \quad \mathbb{E}(S_n^{T_1}) = \mathbb{E}(S_0) = 0. \quad (10.11)$$

D'après le théorème de Polya 4.4, la marche aléatoire en dimension 1 est récurrente et T_1 est fini presque sûrement. On en déduit que $n \wedge T_1$ converge vers T_1 quand n tend vers l'infini. Cependant on ne peut pas passer à la limite dans l'espérance car

$$0 = \lim_{n \rightarrow \infty} \mathbb{E}(S_n^{T_1}) \neq \mathbb{E}(S_{T_1}) = 1. \quad (10.12)$$

En effet, le processus $\{S_n^{T_1}\}_{n \geq 0}$ a de rares fluctuations très négatives qui suffisent pour préserver la moyenne $\mathbb{E}(S_n^{T_1}) = 0$ (cf. figure 10.2).

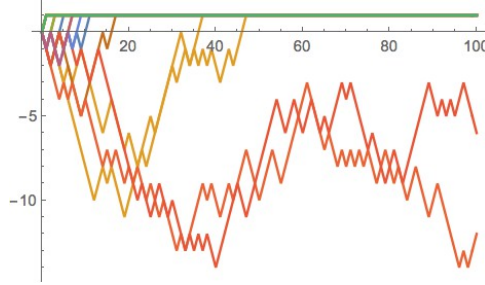


FIGURE 10.2 – Cette simulation représente 30 réalisations de marches aléatoires arrêtées au niveau 1. Après 100 pas, seules 2 marches n'ont pas atteint le niveau 1 mais leurs valeurs sont très négatives. Ces rares fluctuations permettent d'expliquer le comportement asymptotique (10.12).

Remarque 10.7. Dans le jeu de loterie un phénomène similaire s'opère. Si N personnes misent 1 euro chacune et qu'un unique gagnant, tiré au sort parmi les N personnes, reçoit la somme totale de N euros, alors le gain moyen vaut 1. Pourtant quand N est grand, un joueur perd avec une probabilité proche de 1, mais en moyenne ceci est compensé par un gain très important qui a lieu rarement.

Par définition une martingale a une espérance constante puisque $\mathbb{E}(X_n) = \mathbb{E}(X_0)$ pour tout entier n . Le résultat suivant donne une sorte de réciproque grâce à la notion de temps d'arrêt.

Proposition 10.8. Soit $X = \{X_n\}_{n \geq 0}$ un processus aléatoire adapté à la filtration $\mathbb{F}$ tel que $\mathbb{E}(|X_n|) < \infty$ pour tout entier n . Alors, X est une martingale si et seulement si

$$\mathbb{E}(X_T) = \mathbb{E}(X_0) \quad \text{pour tout temps d'arrêt } T \text{ borné.} \quad (10.13)$$

Démonstration. Par le théorème d'arrêt de Doob, la relation (10.13) est toujours satisfaite si X est une martingale, il s'agit donc d'une condition nécessaire.

Pour démontrer que (10.13) est aussi une condition suffisante, commençons par fixer un entier n et un événement $A \in \mathcal{F}_n$ arbitraire. La variable aléatoire

$$T = n\mathbf{1}_A + (n+1)\mathbf{1}_{A^c}$$

est un temps d'arrêt borné car

$$\{T = n\} = A \in \mathcal{F}_n, \quad \{T = n+1\} = A^c \in \mathcal{F}_n \subset \mathcal{F}_{n+1}$$

et pour tout $k \neq n, n+1$, on voit que $\{T = k\} = \emptyset \in \mathcal{F}_k$. Par conséquent si la condition (10.13) est satisfaite, on a

$$\mathbb{E}(X_0) = \mathbb{E}(X_T) = \mathbb{E}(X_n\mathbf{1}_A + X_{n+1}\mathbf{1}_{A^c}).$$

En prenant $T = n+1$, la condition (10.13) impose aussi que $\mathbb{E}(X_0) = \mathbb{E}(X_{n+1})$. On déduit des 2 identités précédentes que

$$0 = \mathbb{E}(X_T - X_{n+1}) = \mathbb{E}(X_n\mathbf{1}_A - X_{n+1}\mathbf{1}_A) = \mathbb{E}((X_{n+1} - X_n)\mathbf{1}_A).$$

Comme cette relation est valable pour tout événement A dans $\mathcal{F}_n$, elle caractérise l'espérance conditionnelle

$$\mathbb{E}(X_{n+1} - X_n \mid \mathcal{F}_n) = 0.$$

Ceci montre que le processus X est bien une martingale. □

Application du théorème d'arrêt.

Nous allons maintenant revisiter le problème de la ruine du joueur, déjà étudié par les chaînes de Markov dans la section 2.5.2, et montrer comment le théorème d'arrêt permet de le résoudre simplement.

Soient $a, b \geq 1$ des entiers. On considère un jeu équilibré

$$X_n = a + \sum_{i=1}^n \xi_i$$

où $\{\xi_i\}_{i \geq 1}$ est une suite de variables aléatoires indépendantes de loi $\mathbb{P}(\xi_i = \pm 1) = 1/2$. La fortune initiale du joueur est donc $X_0 = a$ et on définit les temps d'arrêt

$$T_0 = \inf\{n; \quad X_n = 0\}, \quad T_{a+b} = \inf\{n; \quad X_n = a+b\}, \quad \tau = \inf\{T_0, T_{a+b}\}.$$

Comme la marche aléatoire symétrique $\{X_n\}_{n \geq 0}$ est récurrente par le théorème 4.4, elle atteindra toujours 0 ou $a+b$. Le temps d'arrêt τ est donc fini presque sûrement.

La marche aléatoire $\{X_n\}_{n \geq 0}$ étant une martingale, le processus arrêté $\{X_n^\tau\}_{n \geq 0}$ est aussi une martingale et il vérifie

$$\forall n \geq 0, \quad \mathbb{E}(X_{n \wedge \tau}) = \mathbb{E}(X_0) = a.$$

Comme τ est fini presque sûrement et $X_{n \wedge \tau}$ est dans l'intervalle $[0, a + b]$, le théorème de convergence dominée permet de passer à la limite et d'obtenir

$$a = \mathbb{E}(X_\tau) = \mathbb{E}((a + b) 1_{\{T_0 > T_{a+b}\}}) = (a + b) \mathbb{P}(T_0 > T_{a+b}).$$

car le processus arrêté ne peut prendre que les valeurs 0 et $a + b$. On retrouve donc le résultat de la section 2.5.2.

$$u(a) = \mathbb{P}_a(T_0 < T_{a+b}) = \frac{b}{a + b}.$$

Pour calculer $\mathbb{E}(\tau)$, nous allons utiliser la martingale $M = \{(X_n - a)^2 - n\}_{n \geq 0}$ définie en (10.4). Comme $\mathbb{E}(M_{n \wedge \tau}) = 0$, on a donc

$$\mathbb{E}((X_{\tau \wedge n} - a)^2) = \mathbb{E}(\tau \wedge n).$$

Quand n tend vers l'infini, le théorème de convergence dominée permet de justifier la convergence du terme de gauche et le théorème de convergence monotone, celle du terme de droite

$$\mathbb{E}((X_\tau - a)^2) = \mathbb{E}(\tau).$$

En utilisant que $\mathbb{P}(T_0 < T_{a+b}) = \frac{b}{a+b}$, on obtient donc

$$\mathbb{E}(\tau) = \mathbb{E}(a^2 1_{\{T_0 < T_{a+b}\}}) + \mathbb{E}(b^2 1_{\{T_0 > T_{a+b}\}}) = a^2 \frac{b}{a+b} + b^2 \frac{a}{a+b} = ab.$$

Ceci permet de retrouver le résultat (2.31).

10.3 Stratégies de gestion et évaluation des options ★

La finance offre un cadre d'application à la théorie des martingales que nous allons illustrer dans cette section par quelques exemples sur la gestion d'actifs financiers (cette section peut être omise en première lecture). On considère d actifs financiers dont les prix, le jour n (à l'heure de la clôture du marché), sont représentés par le vecteur $S_n = \{S_n^1, \dots, S_n^d\}$. Ces prix évoluent quotidiennement et on notera $\mathcal{F}_n = \sigma(S_i, i \leq n)$ la filtration engendrée par les prix au cours du temps. L'échelle de temps de cet exemple est arbitraire et on pourrait suivre les cotations à l'année ou à la minute (voire à la seconde). Parallèlement aux actifs boursiers dont l'évolution est incertaine, il existe aussi des placements non risqués à un taux fixe $r > 0$. Ainsi, après avoir investi initialement la somme Φ_0^0 , un placement à taux fixe permet d'obtenir au temps n la somme $(1 + r)^n \Phi_0^0$.

Une *stratégie de gestion* est un processus $\Phi_n = \{\Phi_n^0, \Phi_n^1, \dots, \Phi_n^d\} \in \mathbb{R}^{d+1}$ décrivant au cours du temps la quantité des différents actifs détenus, y compris du placement à taux fixe par la coordonnée Φ_n^0 . Nous supposons que le processus $\{\Phi_n\}_{n \geq 1}$ est prévisible, i.e. que Φ_n est $\mathcal{F}_{n-1}$ -mesurable (avec $\mathcal{F}_0 = \{\emptyset, \Omega\}$ la σ -algèbre triviale). Au temps 0, le prix de tous les actifs est connu et une somme initiale $V_0(\Phi)$ est investie dans un portefeuille dont la valeur totale s'écrit

$$V_0(\Phi) = \Phi_1^0 + \sum_{i=1}^d \Phi_1^i S_0^i.$$

Après une journée, les valeurs des actifs ont changé et la valeur du portefeuille est devenue

$$V_1(\Phi) = (1+r)\Phi_1^0 + \sum_{i=1}^d \Phi_1^i S_1^i.$$

L'actif 0 correspond au placement à taux fixe $1+r$. En fonction de l'évolution des cours boursiers, c'est-à-dire de la connaissance de $\mathcal{F}_1$, le portefeuille peut être réorganisé selon la nouvelle répartition Φ_2

$$(1+r)\Phi_1^0 + \sum_{i=1}^d \Phi_1^i S_1^i = (1+r)\Phi_2^0 + \sum_{i=1}^d \Phi_2^i S_1^i.$$

Comme la valeur totale reste inchangée, la stratégie est dite *autofinancée* (il n'y a ni apport, ni retrait de fonds). On peut cependant recourir à l'emprunt et dans ce cas $\Phi_n^0 < 0$. De même considérer des valeurs $\Phi_n^i < 0$ revient à vendre à découvert. Une fois l'argent placé, la valeur du portefeuille après une seconde journée de cotation devient

$$V_2(\Phi) = (1+r)^2\Phi_2^0 + \sum_{i=1}^d \Phi_2^i S_2^i = \Phi_2 \cdot S_2,$$

où le produit scalaire est pris entre Φ_2 et le vecteur des actifs $S_2 = \{(1+r)^2, S_2^1, \dots, S_2^d\} \in \mathbb{R}^{d+1}$ auquel le placement à taux fixe a été intégré. Ainsi l'investisseur réajuste son portefeuille à chaque pas de temps et la valeur sera

$$\forall n \geq 1, \quad V_n(\Phi) = \Phi_n \cdot S_n \quad \text{avec} \quad \Phi_n \cdot S_n = \Phi_{n+1} \cdot S_n. \quad (10.14)$$

La dernière égalité impose que la stratégie soit autofinancée. Cette condition implique la relation suivante

$$V_{n+1}(\Phi) - V_n(\Phi) = \Phi_{n+1} \cdot S_{n+1} - \Phi_n \cdot S_n = \Phi_{n+1} \cdot (S_{n+1} - S_n).$$

La valeur du portefeuille au cours du temps peut donc se réécrire en fonction du vecteur des incréments $\Delta S_k = S_k - S_{k-1}$

$$V_n(\Phi) = V_0(\Phi) + \sum_{k=1}^n \Phi_k \cdot \Delta S_k. \quad (10.15)$$

Le processus $V_n(\Phi)$ a une structure très proche de celui défini en (10.7), cependant les incréments ΔS_k ne proviennent pas d'une martingale. En particulier, la première coordonnée de S_n croît exponentiellement comme $(1+r)^n$ (cf. Figure 10.3). Pour compenser cet effet, nous allons normaliser les processus et définir la *valeur actualisée* du portefeuille

$$\tilde{S}_n = \frac{1}{(1+r)^n} S_n, \quad \tilde{V}_n(\Phi) = \frac{1}{(1+r)^n} V_n(\Phi). \quad (10.16)$$

En utilisant l'identité $\Phi_n \cdot \tilde{S}_n = \Phi_{n+1} \cdot \tilde{S}_n$, on retrouve une relation analogue à (10.15)

$$\tilde{V}_n(\Phi) = V_0(\Phi) + \sum_{k=1}^n \Phi_k \cdot \Delta \tilde{S}_k, \quad (10.17)$$

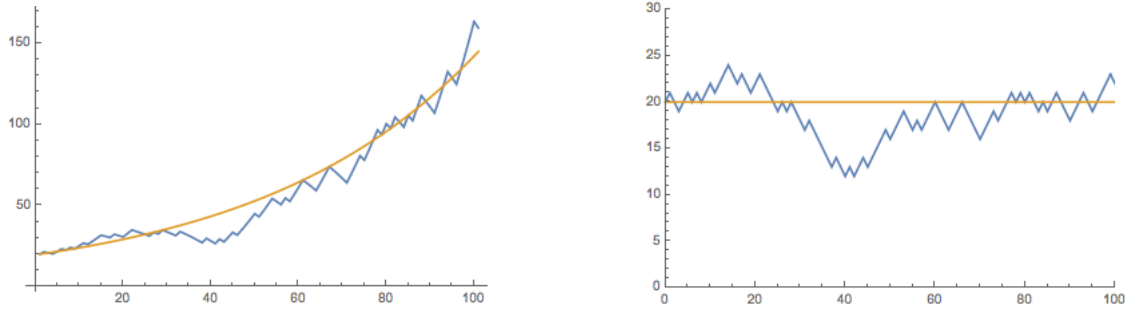


FIGURE 10.3 – Sur la figure de gauche, les fluctuations d'un actif risqué $\{S_n^1\}_{n \leq 100}$ sont représentées en bleu et la courbe orange correspond au placement non risqué qui croît exponentiellement. Les mêmes courbes sont tracées sur la figure de droite après normalisation par le facteur $(1+r)^{-n}$ pour tenir compte de l'actualisation (10.16).

en notant les incréments $\Delta \tilde{S}_k = \tilde{S}_k - \tilde{S}_{k-1}$.

Dans le processus actualisé $\tilde{S}_n$ défini en (10.16), la première coordonnée est constante et représente un placement sans risque, par contre les autres coordonnées évoluent selon une mesure de probabilité $\mathbb{P}$ et correspondent aux actifs risqués. Dans la suite, la théorie des martingales va s'avérer très utile pour étudier les processus actualisés.

Nous allons maintenant nous intéresser à la notion d'arbitrage qui joue un rôle clef en finance. Étant donné un temps N , une *stratégie d'arbitrage* est une stratégie autofinancée telle que $V_0(\Phi) = 0$ et

$$\forall n \leq N, \quad V_n(\Phi) \geq 0 \quad \text{et} \quad \mathbb{P}(V_N(\Phi) > 0) > 0. \quad (10.18)$$

Si elle existe, une telle stratégie permet de gagner de l'argent sans prendre de risques. On dira que le marché est *viable*, s'il n'existe pas de stratégie d'arbitrage. Ceci impose des contraintes sur les actifs que nous allons détailler dans le cadre simplifié du *modèle binomial*. Considérons à partir de maintenant le cas d'un seul actif risqué

$$S_n = \{(1+r)^n, S_n^1\}$$

avec la valeur initiale $S_0^1 = 1$. On suppose qu'il existe deux valeurs $0 < b < h$ et qu'à chaque pas de temps, le prix de l'actif change de la façon suivante

$$S_{n+1}^1 = \zeta_{n+1} S_n^1, \quad (10.19)$$

où les variables aléatoires $\{\zeta_n\}_{n \geq 1}$ sont indépendantes et identiquement distribuées selon la mesure de probabilité $\mathbb{P}^*$ sur $\{b, h\}^N$ telle que

$$\forall n \leq N, \quad \mathbb{P}^*(\zeta_n = h) = 1 - \mathbb{P}^*(\zeta_n = b) = p^* > 0. \quad (10.20)$$

Pour que le marché soit viable, il faut que

$$b < 1 + r < h. \quad (10.21)$$

Il s'agit d'une condition nécessaire. En effet si $1 + r \leq b$, alors on peut emprunter 1 euro au temps 0 et placer cet argent sur l'actif risqué qui rapportera au moins autant car

$$V_n(\Phi) = -(1+r)^n + S_n^1 \geq -(1+r)^n + b^n \geq 0$$

en utilisant la stratégie $\Phi_n = (-1, 1)$ pour tout n . Comme $\mathbb{P}^*(\zeta_n = h) = p^*$, la probabilité de gagner est strictement positive. Inversement si $1 + r \geq h$, alors il suffit de vendre l'actif risqué à découvert et de placer l'argent à taux fixe r pour gagner avec une probabilité positive en suivant la stratégie $\Phi_n = (1, -1)$.

Montrons maintenant que la condition (10.21) est suffisante pour que le marché soit viable. Pour cela, commençons par supposer que le paramètre p^* introduit en (10.20) vaut

$$p^* = \frac{(1+r) - b}{h - b} \in]0, 1[. \quad (10.22)$$

Avec ce choix, l'espérance sous $\mathbb{P}^*$ s'écrit

$$\forall n \geq 1, \quad \mathbb{E}^* \left(\frac{\zeta_n}{1+r} \right) = 1.$$

Par conséquent, si on considère l'actif actualisé défini en (10.16)

$$\tilde{S}_n^1 = \frac{1}{(1+r)^n} S_n^1 \quad \text{avec} \quad S_n^1 = \zeta_n S_{n-1}^1 = \prod_{i=1}^n \zeta_i, \quad (10.23)$$

alors le processus $\{\tilde{S}_n^1\}_{n \geq 0}$ est une martingale sous la loi $\mathbb{P}^*$ comme dans l'exemple (10.2). En particulier, la proposition 10.4 permet d'affirmer que la valeur du portefeuille

$$\tilde{V}_n(\Phi) = V_0(\Phi) + \sum_{k=1}^n \Phi_k \cdot \Delta \tilde{S}_k, \quad (10.24)$$

est aussi une martingale pour toute stratégie de gestion Φ et que son espérance reste constante au cours du temps

$$\forall n \leq N, \quad \mathbb{E}^* \left(\tilde{V}_n(\Phi) \right) = \mathbb{E}^* \left(\tilde{V}_0(\Phi) \right) = \mathbb{E}^* (V_0(\Phi)). \quad (10.25)$$

Il ne peut donc pas exister de stratégie d'arbitrage Φ satisfaisant (10.18) car sinon on aurait

$$\mathbb{E}^* \left(\tilde{V}_N(\Phi) \right) = \frac{\mathbb{E}^* (V_N(\Phi))}{(1+r)^N} > 0 \quad \text{et} \quad \mathbb{E}^* \left(\tilde{V}_N(\Phi) \right) = \mathbb{E}^* (V_0(\Phi)) = 0.$$

Ainsi nous avons prouvé que pour toute stratégie de gestion *admissible* Φ c'est-à-dire telle que

$$\forall n \leq N, \quad V_n(\Phi) \geq 0 \quad \mathbb{P}^* - \text{presque sûrement}, \quad (10.26)$$

alors

$$\mathbb{P}^* (V_N(\Phi) > 0) = 0 \quad \text{si} \quad V_0(\Phi) = 0. \quad (10.27)$$

Cette preuve repose sur le choix très spécifique de la loi $\mathbb{P}^*$ et du paramètre p^* défini en (10.22). Cette mesure de probabilité n'a probablement rien à voir avec la véritable

loi $\mathbb{P}$ des actifs boursiers car la distribution des variables aléatoires $\{\zeta_n\}_{n \leq N}$ n'a aucune raison d'être indépendante à chaque pas de temps. Dans un marché financier, les fluctuations sont corrélées au cours du temps. Néanmoins, nous allons montrer que l'absence de stratégie d'arbitrage établie en (10.27) reste valable pour des distributions beaucoup plus générales. Il suffit de supposer que la loi $\mathbb{P}$, qui régit le marché, est équivalente à $\mathbb{P}^*$, c'est-à-dire que pour toute réalisation $\omega \in \{h, b\}^N$ le rapport des probabilités satisfait

$$0 < \frac{\mathbb{P}(\omega)}{\mathbb{P}^*(\omega)}. \quad (10.28)$$

Cette hypothèse implique que toute stratégie de gestion Φ admissible sous la loi $\mathbb{P}$

$$\forall n \leq N, \quad V_n(\Phi) \geq 0 \quad \mathbb{P} - \text{presque sûrement},$$

sera aussi admissible sous la loi $\mathbb{P}^*$ (10.26). Par l'identité (10.27), pour toute stratégie admissible, l'évènement $\{\tilde{V}_N(\Phi) > 0\}$ est négligeable sous la loi $\mathbb{P}^*$ si $V_0(\Phi) = 0$. Un changement de mesure suffit alors à transférer cette propriété à la loi $\mathbb{P}$

$$\begin{aligned} \mathbb{P}(V_N(\Phi) > 0) &= \sum_{\omega \in \{h, b\}^N} \mathbb{P}(\omega) \mathbf{1}_{\{V_N(\Phi) > 0\}} = \sum_{\omega \in \{h, b\}^N} \mathbb{P}^*(\omega) \frac{\mathbb{P}(\omega)}{\mathbb{P}^*(\omega)} \mathbf{1}_{\{V_N(\Phi) > 0\}} \\ &= \mathbb{E}^* \left(\frac{\mathbb{P}(\omega)}{\mathbb{P}^*(\omega)} \mathbf{1}_{\{V_N(\Phi) > 0\}} \right) = 0. \end{aligned}$$

Nous avons donc démontré que la condition $b < 1 + r < h$ est nécessaire¹ et suffisante pour que le marché soit viable pour toute loi $\mathbb{P}$ satisfaisant (10.28). La mesure $\mathbb{P}^*$ sert simplement d'artifice de calcul pour appliquer la théorie des martingales. L'intérêt de cette approche est de fournir un résultat général sans avoir à estimer les véritables paramètres du marché pour déterminer la loi $\mathbb{P}$. Ce résultat s'étend à des distributions plus générales à valeurs dans $\mathbb{R}$ [17].

Nous allons montrer maintenant que la théorie des martingales permet aussi de calculer le prix d'options et de déterminer des stratégies de couverture. Une *option européenne* est un titre qui donne le droit d'acheter un actif à une échéance fixée et à un prix déterminé à l'avance. Par exemple, une compagnie aérienne peut décider d'acquérir une option d'achat (dite *call*) afin de pouvoir acheter, l'année suivante, une certaine quantité de carburant à un prix fixé. Cette option garantit à la compagnie aérienne un prix maximum pour le carburant et la préserve ainsi des aléas du marché. Concrètement si le prix du carburant est indexé au temps n par l'actif S_n^1 et que l'échéance de temps est notée N (par exemple $N = 365$ jours), alors moyennant un coût initial C_N perçu par le vendeur de l'option, la compagnie pourra acheter à l'échéance N le carburant au prix K quel que soit le cours du carburant au temps N . Ainsi si $S_N^1 > K$, la compagnie aérienne n'aura à déboursier que la valeur K et le vendeur de l'option devra s'acquitter de la somme $S_N^1 - K$ pour compenser la différence. Par contre si $K \geq S_N^1$, la compagnie aérienne n'a pas intérêt à exercer son option d'achat et elle achètera le carburant au prix du marché S_N^1 . Dans ce cas, la compagnie aérienne aura perdu la somme initiale C_N qu'elle avait versée au

1. On remarquera que la condition d'équivalence des mesures (10.28) suffit pour appliquer l'argument (10.3) à la loi $\mathbb{P}$.

vendeur de l'option et cette fois c'est le vendeur qui sera gagnant. Une *option européenne* est donc une forme d'assurance contre les aléas du marché qui sont alors assumés par le vendeur de l'option moyennant une commission C_N .

La difficulté consiste à établir le juste prix C_N de l'option. Définissons le prix d'exercice de l'option par $(S_N^1 - K)_+$ ², i.e. la somme que le vendeur devra fournir si le cours de l'action dépasse le seuil K au temps N . Pour cela le vendeur va devoir adopter une stratégie autofinancée Φ qui lui permette de s'acquitter de cette somme au temps N en utilisant la prime initiale $V_0(\Phi) = C_N$

$$\forall n \leq N, \quad V_n(\Phi) \geq 0 \quad \text{et} \quad V_N(\Phi) = (S_N^1 - K)_+. \quad (10.29)$$

Cette stratégie est dite admissible car V_n doit rester positif à tout temps. La théorie des martingales permet de construire de telles stratégies de couverture sous des hypothèses générales (cf. [17, 26]). Dans la suite, nous traiterons uniquement le cas simplifié du modèle de Cox, Ross, Rubinstein (parfois appelé modèle CRR) où tous les calculs sont explicites.

Considérons le modèle binomial (10.20) avec un seul actif dont l'évolution est donnée par $S_{n+1}^1 = \zeta_{n+1} S_n^1$. Supposons, comme dans l'exemple (10.22), que les variables $\{\zeta_n\}_{n \geq 1}$ sont indépendantes et identiquement distribuées selon la loi

$$\mathbb{P}^*(\zeta_n = h) = 1 - \mathbb{P}^*(\zeta_n = b) = p^* = \frac{(1+r) - b}{h - b} \in]0, 1[. \quad (10.30)$$

Rappelons que sous l'hypothèse $b < 1 + r < h$, le choix de p^* implique que les valeurs actualisées des actifs $\{\tilde{S}_n^1\}_{n \geq 0}$ forment une martingale (cf. (10.23)). Par conséquent, l'espérance de la valeur actualisée du portefeuille est constante d'après (10.25)

$$\forall n \leq N, \quad \mathbb{E}^*(\tilde{V}_n(\Phi)) = \mathbb{E}^*(\tilde{V}_0(\Phi)).$$

S'il existe une stratégie de couverture satisfaisant (10.29), alors

$$C_N = \mathbb{E}^*(\tilde{V}_0(\Phi)) = \mathbb{E}^*(\tilde{V}_N(\Phi)) = \mathbb{E}^*\left(\frac{(S_N^1 - K)_+}{(1+r)^N}\right), \quad (10.31)$$

en utilisant que $V_N(\Phi) = (S_N^1 - K)_+$. Par la propriété de martingale de $\{\tilde{V}_n(\Phi)\}_{n \geq 0}$ (10.24), on peut non seulement calculer le prix initial C_N de l'option, mais aussi la valeur du portefeuille à chaque instant $n \leq N$

$$\frac{V_n(\Phi)}{(1+r)^n} = \tilde{V}_n(\Phi) = \mathbb{E}^*(\tilde{V}_N(\Phi) \mid \mathcal{F}_n) = \mathbb{E}^*\left(\frac{(S_N^1 - K)_+}{(1+r)^N} \mid \mathcal{F}_n\right). \quad (10.32)$$

Il suffit alors de remarquer que $S_N^1 = S_n^1 \prod_{i=n+1}^N \zeta_i$ pour en déduire une expression explicite de la valeur du portefeuille à tout temps $n \leq N$

$$V_n(\Phi) = c_n(S_n^1) \quad \text{avec} \quad c_n(s) = \frac{(1+r)^n}{(1+r)^N} \mathbb{E}^*\left(\left(s \prod_{i=n+1}^N \zeta_i - K\right)_+\right) \quad (10.33)$$

2. La partie positive est notée $(x)_+ = \max\{x, 0\}$

car le conditionnement au temps n se réduit à calculer une espérance par rapport à $N - n$ variables aléatoires indépendantes. L'aléa ayant disparu au temps N , la formule (10.33) se réduit à $c_N(s) = (s - K)_+$. Ainsi non seulement $V_n(\Phi)$ a une valeur explicite, mais nous allons pouvoir calculer une stratégie de couverture parfaite. En effet à chaque pas de temps, $\Phi_n = (\Phi_n^0, \Phi_n^1)$ devra satisfaire l'identité (10.33)

$$\Phi_n^0(1+r)^n + \Phi_n^1 S_n^1 = c_n(S_n^1).$$

La stratégie Φ_n étant mesurable par rapport à $\mathcal{F}_{n-1}$, cette identité doit donc être vraie quelle que soit la valeur de $\zeta_n \in \{h, b\}$. Comme $S_n^1 = S_{n-1}^1 \zeta_n$, la stratégie doit satisfaire les contraintes suivantes

$$\begin{aligned} \Phi_n^0(1+r)^n + \Phi_n^1 S_{n-1}^1 h &= c_n(S_{n-1}^1 h), \\ \Phi_n^0(1+r)^n + \Phi_n^1 S_{n-1}^1 b &= c_n(S_{n-1}^1 b). \end{aligned} \quad (10.34)$$

En résolvant ces équations, on construit la stratégie de couverture

$$\forall n \leq N, \quad \Phi_n^1 = \frac{c_n(S_{n-1}^1 h) - c_n(S_{n-1}^1 b)}{S_{n-1}^1(h-b)} \quad \text{et} \quad \Phi_n^0 = \frac{hc_n(S_{n-1}^1 b) - bc_n(S_{n-1}^1 h)}{h-b}. \quad (10.35)$$

La stratégie de gestion trouvée est la seule possible, car la solution du système (10.34) est unique. L'hypothèse (10.31) sur l'existence d'une stratégie de couverture est donc justifiée a posteriori. Étant donnée la valeur de l'actif S_{n-1}^1 , les relations (10.35) déterminent une stratégie prévisible pour réinvestir au temps n . Insistons sur le fait que cette stratégie est déterministe et qu'elle permet de produire exactement la valeur $c_N(s) = (s - K)_+$ au temps N . À l'aide de la loi $\mathbb{P}^*$, nous avons pu utiliser la théorie des martingales et construire une stratégie explicite qui peut s'appliquer même si la loi du marché n'est pas $\mathbb{P}^*$. Comme dans l'exemple (10.28), la mesure $\mathbb{P}^*$ sert d'outil de calcul pour déterminer des stratégies sans avoir à connaître explicitement la loi $\mathbb{P}$ du marché.

Pour conclure cette section, nous allons retrouver la célèbre formule de Black-Scholes à partir de la valeur de l'option C_N calculée en (10.31)

$$C_N = \mathbb{E}^* \left(\left(\frac{\prod_{i=1}^N \zeta_i}{(1+r)^N} - \frac{K}{(1+r)^N} \right)_+ \right). \quad (10.36)$$

On peut décomposer l'échéance de l'option (disons une année) en fonction d'un pas de temps très court (de l'ordre de la minute). Dans cette nouvelle échelle de temps, le paramètre N est très grand et le taux d'intérêt $r = \frac{R}{N}$ très faible. On interprétera le nouveau paramètre $R > 0$ comme le rendement instantané. Pour N très grand, le rendement du placement sans risque à l'échéance sera approximativement

$$(1+r)^N = \left(1 + \frac{R}{N}\right)^N \simeq \exp(R). \quad (10.37)$$

À chaque pas de temps, les variations de l'actif risqué sont aussi très faibles et les valeurs des cours $\{b, h\}$ seront maintenant indexées en fonction d'un seul paramètre $\sigma > 0$ tel que

$$\log\left(\frac{b}{1+r}\right) = -\frac{\sigma}{\sqrt{N}} \quad \text{et} \quad \log\left(\frac{h}{1+r}\right) = \frac{\sigma}{\sqrt{N}}.$$

La probabilité p^* , définie en (10.30), peut aussi se réécrire en fonction de σ

$$p^* = \frac{(1+r) - b}{h - b} = \frac{1}{1 + \exp\left(\frac{\sigma}{\sqrt{N}}\right)} \simeq \frac{1}{2} - \frac{\sigma}{4\sqrt{N}},$$

où le dernier terme décrit le comportement asymptotique de p^* quand N tend vers l'infini. Les variables aléatoires $X_i = \log \frac{\zeta_i}{1+r}$ sont indépendantes et identiquement distribuées sous $\mathbb{P}^*$ avec une moyenne et une variance données par

$$\mathbb{E}^*(X_1) = \frac{\sigma}{\sqrt{N}}(2p^* - 1) \quad \text{et} \quad \mathbb{E}^*(X_1^2) - \mathbb{E}^*(X_1)^2 = \frac{\sigma^2}{N}(4p^* - 4p^{*2}).$$

On peut ainsi montrer qu'après renormalisation, la moyenne et la variance convergent quand N tend vers l'infini

$$\lim_{N \rightarrow \infty} N \mathbb{E}^*(X_1) = -\frac{\sigma^2}{2} \quad \text{et} \quad \lim_{N \rightarrow \infty} N \left(\mathbb{E}^*(X_1^2) - \mathbb{E}^*(X_1)^2 \right) = \sigma^2.$$

En utilisant une preuve identique à celle du théorème central limite, on déduit la convergence en loi

$$\prod_{i=1}^N \frac{\zeta_i}{1+r} = \exp\left(\sum_{i=1}^N X_i\right) \xrightarrow[N \rightarrow \infty]{(loi)} \exp(\gamma),$$

où γ est une variable aléatoire gaussienne de moyenne $-\frac{\sigma^2}{2}$ et de variance σ^2 . Ceci démontre donc que le prix de l'option (10.36) converge, quand N tend vers l'infini, vers

$$\begin{aligned} \lim_{N \rightarrow \infty} C_N &= \mathbb{E} \left(\left(\exp(\gamma) - K \exp(-R) \right)_+ \right) \\ &= \frac{1}{\sqrt{2\pi\sigma^2}} \int_{\mathbb{R}} dx \exp\left(-\frac{x^2}{2\sigma^2}\right) \left(\exp\left(-\frac{\sigma^2}{2} + x\right) - K \exp(-R) \right)_+, \end{aligned}$$

où nous avons aussi utilisé la formule asymptotique du rendement (10.37). Cette expression asymptotique est la formule de Black-Scholes qui permet d'estimer le prix d'une option. Le paramètre σ correspond à la *volatilité stochastique*, c'est-à-dire à l'amplitude des fluctuations du prix de l'actif. Ce paramètre n'est pas directement observable sur le marché et il doit être estimé par des méthodes statistiques.

Il existe plusieurs variantes des options européennes, citons notamment les options américaines qui peuvent être exercées à n'importe quel temps $n \leq N$ avant l'échéance N . Par analogie avec la formule (10.31), le prix d'une option américaine se calcule en optimisant sur toutes les stratégies mesurables possibles, i.e. sur tous les temps d'arrêt τ bornés par N

$$c_N = \sup_{\tau} \mathbb{E}^* \left(\frac{(S_{\tau}^1 - K)_+}{(1+r)^{\tau}} \right). \quad (10.38)$$

Dans le chapitre 13 consacré à l'arrêt optimal, nous verrons comment calculer (10.38) en déterminant la stratégie optimale par la théorie des martingales. Pour en savoir plus sur les mathématiques financières, on pourra consulter le polycopié du cours de N. Touzi [26] et le livre [17].

10.4 Inégalités de martingales

Le résultat suivant permet de contrôler le maximum de la trajectoire d'un processus aléatoire en fonction de la valeur finale de ce processus, ce qui n'est pas toujours intuitif sur une simulation (cf. figure 10.4). Ainsi l'estimation des fluctuations de toute la trajectoire d'un processus peut se réduire à l'étude d'une seule variable aléatoire (la valeur finale) qui est une quantité beaucoup plus simple à évaluer.

Théorème 10.9 (Inégalité maximale de Doob). *Soit $\{M_n\}_{n \geq 0}$ une sous-martingale et*

$$M_n^* = \max\{M_k, k \leq n\}$$

le maximum de la trajectoire jusqu'au temps n .

(i) *Pour tout $c > 0$, les fluctuations du maximum sont bornées par*

$$c \mathbb{P}(M_n^* \geq c) \leq \mathbb{E}(M_n \mathbf{1}_{\{M_n^* \geq c\}}) \quad \text{pour tout entier } n.$$

(ii) *Soit $p > 1$. Supposons que la sous-martingale $\{M_n\}_{n \geq 0}$ est positive et appartient à $\mathbb{L}^p$, i.e. que*

$$\forall n \geq 0, \quad M_n \geq 0 \quad \text{et} \quad \mathbb{E}(|M_n|^p) < \infty.$$

Alors pour tout entier n , le maximum M_n^ appartient aussi à $\mathbb{L}^p$ et sa norme vérifie*

$$\|M_n^*\|_p \leq \frac{p}{p-1} \|M_n\|_p$$

avec la notation $\|X\|_p = (\mathbb{E}(|X|^p))^{1/p}$.

Par la proposition 10.2, si $\{M_n\}_{n \geq 0}$ est une martingale de signe quelconque, les inégalités de Doob s'appliquent à la sous-martingale $\{|M_n|\}_{n \geq 0}$.

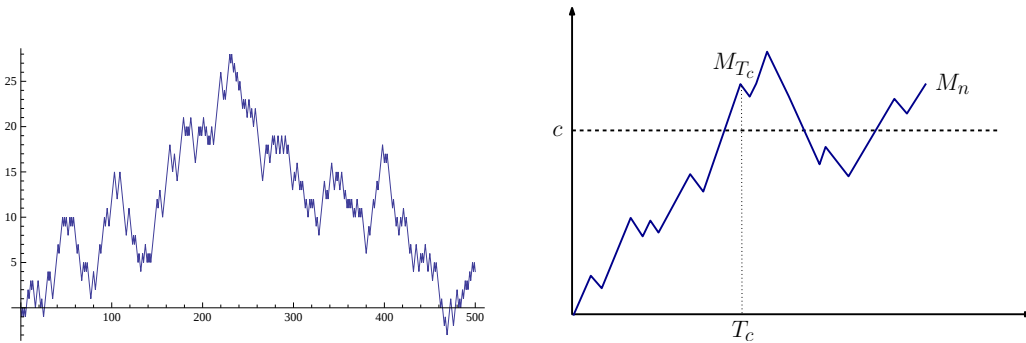


FIGURE 10.4 – Une simulation d'une marche aléatoire de longueur 500 est représentée sur la figure de gauche. Dans cette réalisation, la valeur finale de la marche aléatoire est beaucoup plus petite que le maximum de la trajectoire. Le schéma de droite décrit la décomposition d'une trajectoire en fonction du premier temps de passage T_c au-dessus du niveau c .

Démonstration. (i) On définit le temps d'arrêt $T_c = \inf\{k > 0; M_k \geq c\}$ comme le premier temps où la sous-martingale passe au-dessus du niveau c (cf. figure 10.4). La

trajectoire du processus peut se décomposer en fonction du temps d'arrêt T_c

$$\begin{aligned} \mathbb{E}\left(M_n \mathbf{1}_{\{M_n^* \geq c\}}\right) &= \mathbb{E}\left(M_n \mathbf{1}_{\{T_c \leq n\}}\right) = \sum_{k=1}^n \mathbb{E}\left(M_n \mathbf{1}_{\{T_c=k\}}\right) \\ &= \sum_{k=1}^n \mathbb{E}\left((M_n - M_k) \mathbf{1}_{\{T_c=k\}}\right) + \mathbb{E}\left(M_k \mathbf{1}_{\{T_c=k\}}\right). \end{aligned} \quad (10.39)$$

Comme $\{T_c = k\}$ est mesurable par rapport à $\mathcal{F}_k$, on peut conditionner par la trajectoire jusqu'au temps k pour obtenir

$$\mathbb{E}\left((M_n - M_k) \mathbf{1}_{\{T_c=k\}}\right) = \mathbb{E}\left(\mathbf{1}_{\{T_c=k\}} \mathbb{E}(M_n - M_k \mid \mathcal{F}_k)\right) \geq 0,$$

où la dernière inégalité vient du fait que $\{M_n\}_{n \geq 0}$ est une sous-martingale. Par conséquent la contribution de la trajectoire au-delà du temps T_c peut être minorée. Comme la sous-martingale passe au-dessus du niveau c en T_c , elle vérifie

$$M_k \mathbf{1}_{\{T_c=k\}} \geq c \mathbf{1}_{\{T_c=k\}}.$$

L'inégalité de l'assertion (i) se déduit donc de (10.39)

$$\mathbb{E}\left(M_n \mathbf{1}_{\{M_n^* \geq c\}}\right) \geq c \mathbb{E}\left(\mathbf{1}_{\{T_c \leq n\}}\right) = c \mathbb{P}\left(M_n^* \geq c\right).$$

(ii) Prouvons maintenant la seconde partie du théorème. Soit $q = \frac{p}{p-1}$. L'inégalité établie en (i) implique

$$\int_0^\infty p c^{p-1} \mathbb{P}(M_n^* \geq c) dc \leq \int_0^\infty p c^{p-2} \mathbb{E}(M_n \mathbf{1}_{\{M_n^* \geq c\}}) dc.$$

Comme la martingale est positive, le théorème de Fubini permet de réécrire le membre de gauche de l'inégalité précédente sous la forme

$$\int_0^\infty p c^{p-1} \mathbb{P}(M_n^* \geq c) dc = \mathbb{E}\left(\int_0^{M_n^*} p c^{p-1} dc\right) = \mathbb{E}\left((M_n^*)^p\right)$$

et le membre de droite comme

$$\begin{aligned} \int_0^\infty p c^{p-2} \mathbb{E}(M_n \mathbf{1}_{\{M_n^* \geq c\}}) dc &= \mathbb{E}\left(M_n \int_0^{M_n^*} p c^{p-2} dc\right) = q \mathbb{E}\left(M_n (M_n^*)^{p-1}\right) \\ &\leq q \|M_n\|_p \|(M_n^*)^{p-1}\|_q = q \|M_n\|_p \mathbb{E}((M_n^*)^p)^{1/q}, \end{aligned}$$

en utilisant successivement l'inégalité de Hölder et l'identité $\frac{1}{p} + \frac{1}{q} = 1$. On obtient ainsi

$$\mathbb{E}((M_n^*)^p) \leq q \|M_n\|_p \mathbb{E}((M_n^*)^p)^{1/q}.$$

Ceci conclut la preuve de l'assertion (ii). □

Chapitre 11

Convergence des martingales

Au chapitre 8, l'exemple des séries aléatoires (8.2) a permis de mettre en évidence la convergence d'une martingale vers un objet limite aléatoire. Plus généralement, les martingales offrent un cadre théorique pour démontrer des théorèmes de convergence dans $\mathbb{L}^p$, presque sûre ou en loi. Ce chapitre présente ces différents modes de convergence et les hypothèses correspondantes.

11.1 Convergence des martingales dans $\mathbb{L}^2$

Une martingale $\{M_n\}_{n \geq 0}$ est bornée dans $\mathbb{L}^2$ si elle vérifie

$$\sup_{n \geq 0} \mathbb{E}(M_n^2) < \infty.$$

Le cadre $\mathbb{L}^2$ joue un rôle privilégié car les accroissements

$$n \geq 1, \quad \Delta M_n = M_n - M_{n-1}$$

sont orthogonaux pour le produit scalaire de $\mathbb{L}^2$

$$\mathbb{E}((\Delta M_i)(\Delta M_j)) = \mathbb{E}(\Delta M_i \mathbb{E}(\Delta M_j | \mathcal{F}_{j-1})) = 0 \quad \text{pour } 1 \leq i < j, \quad (11.1)$$

où la dernière égalité est une conséquence de la propriété de martingale

$$\mathbb{E}(\Delta M_j | \mathcal{F}_{j-1}) = \mathbb{E}(M_j | \mathcal{F}_{j-1}) - M_{j-1} = 0.$$

En particulier, l'identité précédente implique que pour tout $n \geq 0$

$$\begin{aligned} \mathbb{E}(M_n^2) &= \mathbb{E} \left(\left(M_0 + \sum_{i=1}^n \Delta M_i \right)^2 \right) \\ &= \mathbb{E}(M_0^2) + \sum_{i=1}^n \mathbb{E}((\Delta M_i)^2) + 2 \sum_{i < j} \mathbb{E}((\Delta M_i)(\Delta M_j)) \\ &= \mathbb{E}(M_0^2) + \sum_{i=1}^n \mathbb{E}((\Delta M_i)^2). \end{aligned} \quad (11.2)$$

Ainsi la suite de réels $\{\mathbb{E}(M_n^2)\}_{n \geq 0}$ est croissante. Comme elle est bornée, cette suite converge vers une valeur positive finie. Cette remarque est la clef du théorème suivant.

Théorème 11.1. Soit $\{M_n\}_{n \geq 0}$ une martingale bornée dans $\mathbb{L}^2$, i.e. telle que

$$\sup_{n \geq 0} \mathbb{E}(M_n^2) < \infty. \quad (11.3)$$

Alors il existe une variable aléatoire limite M_∞ appartenant à $\mathbb{L}^2$ et la convergence de $\{M_n\}_{n \geq 0}$ vers M_∞ peut être quantifiée de différentes façons :

(i) dans $\mathbb{L}^2$

$$\lim_{n \rightarrow \infty} \mathbb{E}((M_n - M_\infty)^2) = 0,$$

(ii) presque sûrement, $M_n \xrightarrow[n \rightarrow \infty]{p.s.} M_\infty$.

Pour tout $p \in [1, 2]$, l'inégalité (B.11)

$$\|M_n - M_\infty\|_p \leq \|M_n - M_\infty\|_2,$$

permet de déduire la convergence dans $\mathbb{L}^p$ de la convergence dans $\mathbb{L}^2$. Une conséquence immédiate du théorème 11.1 est la convergence de la série (8.2) dès que $\alpha > 1/2$.

Démonstration. (i) Commençons par démontrer la convergence dans $\mathbb{L}^2$. Comme la martingale est bornée dans $\mathbb{L}^2$, la suite $\{\mathbb{E}(M_n^2)\}_{n \geq 0}$ est convergente dans $\mathbb{R}$ d'après la relation (11.2). Par la propriété de martingale, on voit que pour $n, p \geq 0$

$$\begin{aligned} \mathbb{E}((M_{n+p} - M_n)^2) &= \mathbb{E}(M_{n+p}^2 + M_n^2 - 2M_n \mathbb{E}(M_{n+p} | \mathcal{F}_n)) \\ &= \mathbb{E}(M_{n+p}^2) - \mathbb{E}(M_n^2) \xrightarrow{n \rightarrow \infty} 0, \end{aligned}$$

où la convergence vers 0 est une conséquence de la convergence de la suite $\{\mathbb{E}(M_n^2)\}_{n \geq 0}$. Ainsi, $\{M_n\}_{n \geq 0}$ forme une suite de Cauchy dans l'espace de Hilbert $\mathbb{L}^2$ et on en déduit par le théorème B.9 l'existence d'une variable aléatoire M_∞ dans $\mathbb{L}^2$ telle que

$$\lim_{n \rightarrow \infty} \mathbb{E}((M_n - M_\infty)^2) = 0.$$

(ii) Nous allons montrer maintenant la convergence presque sûre. Soit $\varepsilon > 0$, en appliquant l'inégalité de Chebychev, on obtient

$$\begin{aligned} \mathbb{P}\left(\sup_{k \geq n} |M_k - M_\infty| \geq \varepsilon\right) &\leq \frac{1}{\varepsilon^2} \mathbb{E}\left(\sup_{k \geq n} |M_k - M_\infty|^2\right) \\ &\leq \frac{2}{\varepsilon^2} \left(\mathbb{E}(|M_n - M_\infty|^2) + \mathbb{E}\left(\sup_{k \geq n} |M_k - M_n|^2\right)\right). \end{aligned}$$

Par ailleurs, le théorème de convergence monotone implique

$$\mathbb{E}\left(\sup_{k \geq n} |M_k - M_n|^2\right) = \lim_{N \rightarrow \infty} \mathbb{E}\left(\sup_{n \leq k \leq N} |M_k - M_n|^2\right).$$

L'application $x \rightarrow |x|$ est convexe et l'inégalité de Jensen de la proposition 10.2 montre que le processus $\{|M_k - M_n|\}_{k \geq n}$ est une sous-martingale positive. Il ne reste plus qu'à appliquer l'inégalité maximale de Doob établie dans le théorème 10.9 pour obtenir

$$\mathbb{E} \left(\sup_{k \geq n} |M_k - M_n|^2 \right) \leq \lim_{N \rightarrow \infty} 4 \mathbb{E} (|M_N - M_n|^2) = 4 \mathbb{E} (|M_\infty - M_n|^2).$$

En utilisant le résultat de convergence dans $\mathbb{L}^2$ du cas (i), on montre que

$$\mathbb{P} \left(\sup_{k \geq n} |M_k - M_\infty| \geq \varepsilon \right) \leq \frac{10}{\varepsilon^2} \mathbb{E} (|M_n - M_\infty|^2) \xrightarrow{n \rightarrow \infty} 0. \quad (11.4)$$

Pour retrouver la convergence presque sûre de M_n vers M_∞ , il suffit alors d'utiliser (11.4) pour choisir une suite d'entiers $(n_\ell)_{\ell \geq 1}$ tendant vers l'infini suffisamment vite pour que la série suivante converge

$$\sum_{\ell \geq 1} \mathbb{P} \left(\sup_{k \geq n_\ell} |M_k - M_\infty| \geq \frac{1}{\ell} \right) < \infty.$$

La convergence presque sûre se déduit alors du théorème de Borel-Cantelli. $\square$

Comme le montre le théorème précédent, les accroissements jouent un rôle clef pour les martingales bornées dans $\mathbb{L}^2$. La notion de *crochet* d'une martingale définie ci-dessous repose sur la structure des accroissements et elle permettra de généraliser le théorème central limite (cf. théorème 11.14).

Définition 11.2 (Crochet). Soit $\{M_n\}_{n \geq 0}$ une martingale bornée dans $\mathbb{L}^2$. On appelle *crochet* de $\{M_n\}_{n \geq 0}$, le processus croissant défini pour tout $n \geq 1$ par

$$\langle M \rangle_n = \sum_{k=1}^n \mathbb{E} ((M_k - M_{k-1})^2 | \mathcal{F}_{k-1}), \quad (11.5)$$

avec $\langle M \rangle_0 = 0$. Par construction $\{\langle M \rangle_n\}_{n \geq 0}$ est un processus prévisible.

En utilisant la relation

$$\mathbb{E} ((M_k - M_{k-1})^2 | \mathcal{F}_{k-1}) = \mathbb{E} (M_k^2 - M_{k-1}^2 | \mathcal{F}_{k-1}),$$

on vérifie facilement que pour toute martingale $\{M_n\}_{n \geq 0}$ bornée dans $\mathbb{L}^2$, le processus $N_n = M_n^2 - \langle M \rangle_n$ est aussi une martingale. Ce résultat étend la propriété vue en (10.4) pour une somme de variables indépendantes.

11.2 Application : loi des grands nombres

Comme application du théorème 11.1, nous allons maintenant montrer la loi des grands nombres pour les suites de variables aléatoires *intégrables*, indépendantes et identiquement distribuées. Ceci renforce le résultat vu dans le cours de première année où la loi des grands nombres a été établie pour les suites de variables aléatoires de carré intégrable. Nous commençons par prouver la loi des grands nombres dans le cadre des martingales.

Théorème 11.3. Soit $\{M_n\}_{n \geq 0}$ une martingale vérifiant

$$\sum_{n \geq 1} \frac{1}{n^2} \mathbb{E}(|\Delta M_n|^2) < \infty. \quad (11.6)$$

Alors

$$\lim_{n \rightarrow \infty} \frac{1}{n} M_n = 0 \quad \text{presque sûrement.}$$

Démonstration. Le processus $X_n = \sum_{k=1}^n \frac{1}{k} \Delta M_k$ est une martingale bornée dans $\mathbb{L}^2$

$$\mathbb{E}(X_n^2) \leq \sum_{k, \ell=1}^n \frac{1}{k\ell} \mathbb{E}(\Delta M_k \Delta M_\ell) = \sum_{k=1}^n \frac{1}{k^2} \mathbb{E}((\Delta M_k)^2) + \sum_{k < \ell} \frac{2}{k\ell} \mathbb{E}(\Delta M_k \mathbb{E}(\Delta M_\ell | \mathcal{F}_k)).$$

Par l'hypothèse (11.6) le premier terme est borné uniformément en n tandis que le second est nul par la propriété de martingale. D'après le théorème 11.1, il existe donc une variable aléatoire X_∞ appartenant à $\mathbb{L}^2$ telle que X_n converge vers X_∞ presque sûrement.

Pour conclure, il suffit de reproduire l'argument classique du lemme de Kronecker pour les suites déterministes

$$\begin{aligned} \frac{1}{n}(M_n - M_0) &= \frac{1}{n} \sum_{i=1}^n i(X_i - X_{i-1}) = \frac{1}{n} \left(\sum_{i=1}^n iX_i - \sum_{i=1}^n (i-1)X_{i-1} - \sum_{i=1}^n X_{i-1} \right) \\ &= X_n - \frac{1}{n} \sum_{i=1}^n X_{i-1}. \end{aligned}$$

Comme X_n converge vers X_∞ presque sûrement, on en déduit que $\frac{1}{n} M_n$ tend vers 0. $\square$

Le résultat suivant utilise le théorème précédent pour montrer la version la plus forte de la loi des grands nombres.

Théorème 11.4 (Loi forte des grands nombres). Soit $\{X_n\}_{n \geq 1}$ une suite de variables aléatoires dans $\mathbb{R}$ indépendantes, identiquement distribuées et intégrables $\mathbb{E}(|X_1|) < \infty$. Alors

$$\frac{1}{n} \sum_{i=1}^n X_i \xrightarrow[n \rightarrow \infty]{} \mathbb{E}(X_1) \quad \text{presque sûrement.}$$

Démonstration. Sans perte de généralité, on suppose $\mathbb{E}(X_1) = 0$. La marche aléatoire $\sum_{i=1}^n X_i$ est une martingale à laquelle on voudrait appliquer le théorème 11.3, cependant l'hypothèse (11.6) de majoration dans $\mathbb{L}^2$ n'est pas vérifiée et il faut considérer la martingale modifiée

$$n \geq 1, \quad M_n = \sum_{i=1}^n X_i \mathbf{1}_{\{|X_i| \leq i\}} - \mathbb{E}(X_i \mathbf{1}_{\{|X_i| \leq i\}}).$$

Les accroissements étant tronqués à des échelles de plus en plus grandes, on s'attend à ce que le comportement asymptotique de M_n reste proche de celui de $\sum_{i=1}^n X_i$ quand n tend vers l'infini.

Vérifions que $\{M_n\}_{n \geq 1}$ satisfait l'hypothèse (11.6)

$$\begin{aligned} \sum_{k=1}^n \frac{1}{k^2} \mathbb{E}(|\Delta M_k|^2) &\leq \sum_{k=1}^n \frac{1}{k^2} \mathbb{E}(|X_1|^2 \mathbf{1}_{\{|X_1| \vee 1 \leq k\}}) = \mathbb{E} \left(|X_1|^2 \sum_{k=1}^n \frac{1}{k^2} \mathbf{1}_{\{k \geq |X_1| \vee 1\}} \right) \\ &\leq 2 \mathbb{E} \left(|X_1|^2 \int_{|X_1| \vee 1}^{\infty} \frac{dt}{t^2} \right) = 2 \mathbb{E} \left(\frac{|X_1|^2}{|X_1| \vee 1} \right) \leq 2 \mathbb{E}(|X_1|) < \infty. \end{aligned}$$

Le théorème 11.3 permet de conclure que $\frac{1}{n} M_n$ tend vers 0 presque sûrement. De plus le théorème de convergence dominée implique

$$\lim_{n \rightarrow \infty} \mathbb{E}(X_1 \mathbf{1}_{\{|X_1| \leq n\}}) = \mathbb{E}(X_1) = 0 \quad \text{et donc} \quad \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n \mathbb{E}(X_i \mathbf{1}_{\{|X_i| \leq i\}}) = 0.$$

On en déduit la convergence presque sûre

$$\frac{1}{n} \sum_{i=1}^n X_i \mathbf{1}_{\{|X_i| \leq i\}} \xrightarrow[n \rightarrow \infty]{} 0.$$

Il reste à prouver que les trop grandes valeurs de X_i ne contribuent pas à la moyenne quand n tend vers l'infini. Les variables étant intégrables, on remarque que

$$\sum_{i \geq 1} \mathbb{P}(|X_i| \geq i) = \sum_{i \geq 1} i \mathbb{P}(i \leq |X_1| < i+1) \leq \mathbb{E}(|X_1|) < \infty.$$

Ceci se réécrit

$$\mathbb{E} \left(\sum_{i \geq 1} \mathbf{1}_{\{|X_i| \geq i\}} \right) < \infty \quad \text{et donc} \quad \sum_{i \geq 1} \mathbf{1}_{\{|X_i| \geq i\}} < \infty \quad \text{presque sûrement.}$$

On en déduit que presque sûrement, il existe un entier $N(\omega)$ tel que pour tout $i \geq N(\omega)$, alors $|X_i(\omega)| < i$. Ceci implique la convergence

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i(\omega) \mathbf{1}_{\{|X_i(\omega)| \geq i\}} = 0$$

car pour presque tout ω , la somme n'a qu'un nombre fini de termes.

En additionnant les deux limites, on conclut le théorème 11.4. □

11.3 Convergence des sous-martingales

Le résultat suivant généralise le théorème 11.1 pour obtenir la convergence de martingales bornées uniformément dans $\mathbb{L}^1$ et non plus dans $\mathbb{L}^2$.

Théorème 11.5 (Théorème de convergence de Doob). *Soit $\{X_n\}_{n \geq 0}$ une sous-martingale satisfaisant*

$$\sup_{n \geq 0} \mathbb{E}(|X_n|) < \infty. \tag{11.7}$$

Alors il existe une variable aléatoire X_∞ appartenant à $\mathbb{L}^1$ telle que

$$X_n \xrightarrow[n \rightarrow \infty]{} X_\infty \quad \text{presque sûrement.}$$

L'hypothèse (11.7) est moins restrictive que l'hypothèse (11.3) utilisée au théorème 11.1. Pour s'en convaincre, il suffit d'utiliser l'inégalité de Jensen qui implique que

$$\|X_n\|_1 = \mathbb{E}(|X_n|) \leq \mathbb{E}(|X_n|^2)^{1/2} = \|X_n\|_2.$$

Par contre la convergence établie au théorème 11.5 est moins forte que celle du théorème 11.1 car la limite n'est valable que presque sûrement. La remarque 11.8 permettra de préciser ce point.

Un corollaire très utile pour les applications est le suivant.

Corollaire 11.6. *Si une des trois propriétés suivantes est satisfaite :*

- $\{X_n\}_{n \geq 0}$ est une martingale positive
- $\{X_n\}_{n \geq 0}$ est une sous-martingale majorée uniformément par une constante
- $\{X_n\}_{n \geq 0}$ est une sur-martingale minorée uniformément par une constante

alors il existe une variable aléatoire X_∞ appartenant à $\mathbb{L}^1$ et $\{X_n\}_{n \geq 0}$ converge presque sûrement vers X_∞ .

La démonstration de ce corollaire sera faite page 160.

La preuve du théorème 11.5 repose sur un résultat préliminaire énoncé ci-dessous. Pour $a < b$, on définit la suite de temps d'arrêt

$$\tau_0 = 0, \quad \theta_{n+1} = \inf \{i \geq \tau_n; \quad X_i \leq a\} \quad \text{et} \quad \tau_{n+1} = \inf \{i \geq \theta_{n+1}; \quad X_i \geq b\}.$$

Alors pour tout $n \geq 0$, la variable aléatoire

$$U_n^{a,b} = \max \{j; \quad \tau_j \leq n\} \tag{11.8}$$

représente le nombre de traversées du niveau b en partant en dessous du niveau a avant la date n (cf. figure 11.1). Nous dirons plus simplement que $U_n^{a,b}$ est le nombre de traversées montantes de l'intervalle $[a, b]$ avant la date n .

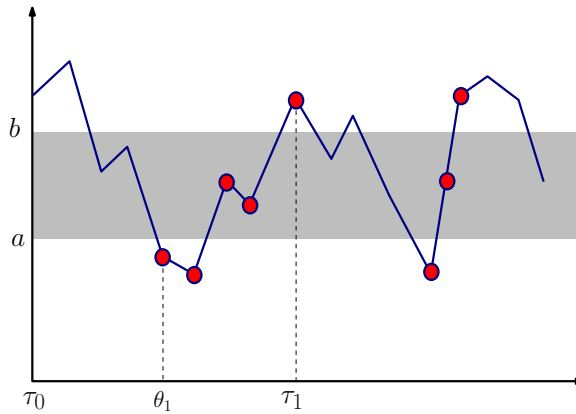


FIGURE 11.1 – Les traversées montantes de l'intervalle $[a, b]$ correspondent aux états marqués par des cercles pendant les intervalles de temps $[\theta_i, \tau_i]$.

Lemme 11.7. Soient $\{X_n\}_{n \geq 0}$ une sous-martingale et $a < b$. Alors la moyenne du nombre de traversées montantes de l'intervalle $[a, b]$ défini en (11.8) vérifie

$$\mathbb{E} \left(U_n^{a,b} \right) \leq \frac{1}{b-a} \mathbb{E} \left((X_n - a)^+ \right)$$

avec la notation $Z^+ = \sup\{0, Z\}$.

Démonstration. Posons $Y_n = (X_n - a)^+$ pour $n \geq 0$ et décomposons la trajectoire $\{Y_n\}_{n \geq 0}$ en fonction des traversées montantes de l'intervalle $[a, b]$ (cf. figure 11.1). Comme $\theta_{n+1} \geq n$, on a

$$\begin{aligned} Y_n &= Y_{n \wedge \theta_1} + \sum_{i=1}^n (Y_{n \wedge \tau_i} - Y_{n \wedge \theta_i}) + \sum_{i=1}^n (Y_{n \wedge \theta_{i+1}} - Y_{n \wedge \tau_i}) \\ &\geq \sum_{i=1}^n (Y_{n \wedge \tau_i} - Y_{n \wedge \theta_i}) + \sum_{i=1}^n (Y_{n \wedge \theta_{i+1}} - Y_{n \wedge \tau_i}). \end{aligned}$$

Par définition des traversées montantes, on a pour $i \geq 1$

$$Y_{\theta_i} = 0, \quad Y_{\tau_i} \geq b - a \quad \text{sur} \quad \{\tau_i \leq n\} \quad \text{et} \quad Y_{n \wedge \tau_i} - Y_{n \wedge \theta_i} \geq 0 \quad \text{presque sûrement.}$$

Par la convexité de $x \mapsto (x - a)^+$ et l'inégalité de Jensen (cf. proposition 10.2), on remarque aussi que le processus $\{Y_n\}_{n \geq 0}$ est une sous-martingale. Le théorème d'arrêt implique donc $\mathbb{E} (Y_{n \wedge \theta_{i+1}} - Y_{n \wedge \tau_i}) \geq 0$. On en déduit l'inégalité cherchée

$$(b-a) \mathbb{E} \left(U_n^{a,b} \right) \leq \sum_{i=1}^n \mathbb{E} (Y_{n \wedge \tau_i} - Y_{n \wedge \theta_i}) \leq \mathbb{E} (Y_n) - \sum_{i=1}^n \mathbb{E} (Y_{n \wedge \theta_{i+1}} - Y_{n \wedge \tau_i}) \leq \mathbb{E} (Y_n).$$

□

Démonstration du théorème 11.5. Le processus comptant le nombre de traversées montantes $\{U_n^{a,b}\}_{n \geq 0}$ est croissant, on note alors $U^{a,b} = \lim_{n \rightarrow \infty} U_n^{a,b}$. D'après le théorème de convergence monotone et le lemme 11.7

$$\begin{aligned} \mathbb{E} \left(U^{a,b} \right) &= \lim_{n \rightarrow \infty} \mathbb{E} \left(U_n^{a,b} \right) \leq \frac{1}{b-a} \sup_{n \geq 0} \mathbb{E} \left((X_n - a)^+ \right) \\ &\leq \frac{1}{b-a} \left(\sup_{n \geq 0} \mathbb{E} (X_n^+) + |a| \right) < \infty. \end{aligned}$$

La dernière inégalité est une conséquence de l'hypothèse du théorème.

En particulier, ceci prouve que $U^{a,b}$ est fini presque sûrement. L'événement

$$\mathcal{N}^{a,b} = \left\{ \liminf_{n \rightarrow \infty} X_n \leq a < b \leq \limsup_{n \rightarrow \infty} X_n \right\}$$

est de probabilité nulle car il correspond aux trajectoires qui oscillent un nombre infini de fois de part et d'autre de l'intervalle $[a, b]$. En prenant l'union sur les rationnels, on voit que

$$\mathcal{N} = \left\{ \liminf_{n \rightarrow \infty} X_n < \limsup_{n \rightarrow \infty} X_n \right\} = \bigcup \left\{ \mathcal{N}^{a,b}; \quad a, b \in \mathbb{Q}, \quad a < b \right\}$$

est négligeable, comme union dénombrable d'ensembles négligeables. Ceci montre bien que $X_\infty = \lim_n X_n$ existe presque sûrement.

Pour montrer que X_∞ appartient à $\mathbb{L}^1$, il suffit d'utiliser le lemme de Fatou et la borne uniforme dans $\mathbb{L}^1$

$$\mathbb{E}(|X_\infty|) = \mathbb{E}\left(\liminf_{n \rightarrow \infty} |X_n|\right) \leq \liminf_{n \rightarrow \infty} \mathbb{E}(|X_n|) \leq \sup_{n \geq 0} \mathbb{E}(|X_n|) < \infty.$$

□

Démonstration du corollaire 11.6. Si $\{X_n\}_{n \geq 0}$ est une martingale positive alors pour tout n

$$\mathbb{E}(|X_n|) = \mathbb{E}(X_n) = \mathbb{E}(X_0) < \infty$$

et l'hypothèse du théorème 11.5 est bien satisfaite.

Plus généralement le théorème 11.5 s'applique à des sous-martingales satisfaisant l'hypothèse $\sup_{n \geq 0} \mathbb{E}(X_n^+) < \infty$. En effet, une sous-martingale $\{X_n\}_{n \geq 0}$ est bornée dans $\mathbb{L}^1$ si et seulement si $\sup_{n \geq 0} \mathbb{E}(X_n^+) < \infty$

$$\mathbb{E}(|X_n|) = \mathbb{E}(X_n^+) + \mathbb{E}(X_n^-) = 2\mathbb{E}(X_n^+) - \mathbb{E}(X_n) \leq 2\mathbb{E}(X_n^+) - \mathbb{E}(X_0).$$

En particulier, si la sous-martingale est bornée supérieurement on a bien $\sup_{n \geq 0} \mathbb{E}(X_n^+) < \infty$.

Si $\{X_n\}_{n \geq 0}$ est une surmartingale bornée inférieurement, il suffit d'appliquer le théorème à la sous-martingale $\{-X_n\}_{n \geq 0}$. □

Remarque 11.8. Bien que la limite X_∞ dans le théorème 11.5 soit dans $\mathbb{L}^1$, la convergence n'a pas toujours lieu dans $\mathbb{L}^1$. Pour le voir, considérons $S = \{S_n\}_{n \geq 0}$ la marche aléatoire symétrique définie en (10.1)

$$S_0 = 1, \quad n \geq 1, \quad S_n = 1 + \sum_{k=1}^n \xi_k \quad \text{avec} \quad \mathbb{P}(\xi_i = \pm 1) = \frac{1}{2}$$

et T_0 est le premier temps d'atteinte de 0 pour cette marche. Comme S est une martingale, le processus arrêté S^{T_0} est aussi une martingale par la proposition 10.5. Il s'agit d'une martingale positive et donc elle converge presque sûrement par le corollaire 11.6. De plus

$$\forall n \geq 0, \quad \mathbb{E}(S_n^{T_0}) = \mathbb{E}(S_0) = 1.$$

D'après le théorème 4.4 de Polya, la marche aléatoire en dimension 1 est récurrente et T_0 est fini presque sûrement. On en déduit que $n \wedge T_0$ converge vers T_0 quand n tend vers l'infini et

$$S_n^{T_0} \xrightarrow[n \rightarrow \infty]{p.s.} S_{T_0} = 0.$$

Cependant, on ne peut pas passer à la limite dans l'espérance

$$1 = \lim_{n \rightarrow \infty} \mathbb{E}(S_n^{T_0}) \neq \mathbb{E}(S_{T_0}) = 0. \quad (11.9)$$

Par conséquent la convergence n'a pas lieu dans $\mathbb{L}^1$

$$\lim_{n \rightarrow \infty} \mathbb{E}(|S_n^{T_0} - S_{T_0}|) \neq 0.$$

Quand n est grand, une trajectoire atteint 0 avant le temps n avec grande probabilité, cependant avec une faible probabilité certaines trajectoires ne vont pas toucher 0 et vont prendre de très grandes valeurs au temps n . Ces rares fluctuations de la marche aléatoire suffisent à expliquer la différence entre les 2 expressions dans (11.9) (cf. figure 10.2).

11.4 Application : le modèle de Wright-Fisher

La dérive génétique correspond au changement de la proportion d'un allèle au sein d'une population par l'effet du hasard, indépendamment de la sélection naturelle ou des migrations. Le modèle de Wright-Fisher propose une interprétation simplifiée du rôle de l'aléatoire dans l'évolution pour illustrer le mécanisme de dérive génétique.

L'enjeu est de décrire l'évolution de la répartition de deux allèles A et a au sein d'une population asexuée de taille constante N . Le mécanisme de reproduction est supposé aléatoire et chaque individu de la génération $n + 1$ choisit son parent uniformément dans la génération n de manière indépendante. Ainsi, si X_n est le nombre d'allèles A à la génération n , alors X_{n+1} sachant X_n a une distribution binomiale

$$\mathbb{P}(X_{n+1} = k | X_n) = \binom{N}{k} \left(\frac{X_n}{N}\right)^k \left(1 - \frac{X_n}{N}\right)^{N-k}.$$

On suppose qu'initialement X_0 est fixé dans $\{0, \dots, N\}$. Une façon simple de représenter X_{n+1} en fonction de X_n consiste à définir les variables aléatoires

$$\xi_{i,n} = \mathbf{1}_{\{U_{i,n} \leq \frac{X_n}{N}\}}$$

où les $U_{i,n}$ sont des variables aléatoires indépendantes et distribuées uniformément sur $[0, 1]$. On peut ainsi réécrire

$$X_{n+1} = \sum_{i=1}^N \xi_{i,n} \quad \text{avec} \quad \mathbb{P}(\xi_{i,n} = 1 | X_n) = \frac{X_n}{N} = 1 - \mathbb{P}(\xi_{i,n} = 0 | X_n) \quad (11.10)$$

Les variables $\xi_{i,n}$ prennent la valeur 1 si l'allèle du parent est A et 0 sinon. On pose $\mathcal{F}_n = \sigma(X_i, i \leq n)$. Le processus $\{X_n\}_{n \geq 0}$ est une martingale car

$$\mathbb{E}(X_{n+1} | \mathcal{F}_n) = \sum_{k=1}^N \mathbb{E}(\xi_{k,n} | X_n) = N \frac{X_n}{N} = X_n.$$

Les variables X_n sont bornées uniformément. Elles convergent donc, par le théorème 11.1, dans $\mathbb{L}^2$ et presque sûrement vers une limite X_∞ .

Pour déterminer cette limite, nous définissons une nouvelle martingale

$$n \geq 0, \quad M_n = \left(\frac{N}{N-1} \right)^n X_n (N - X_n).$$

La propriété de martingale se vérifie facilement en utilisant la décomposition (11.10)

$$\begin{aligned} \left(\frac{N-1}{N} \right)^{n+1} \mathbb{E}(M_{n+1} | \mathcal{F}_n) &= \mathbb{E}(X_{n+1}(N - X_{n+1}) | \mathcal{F}_n) = N\mathbb{E}(X_{n+1} | \mathcal{F}_n) - \mathbb{E}(X_{n+1}^2 | \mathcal{F}_n) \\ &= NX_n - \sum_{i \neq j} \mathbb{E}(\xi_{i,n} \xi_{j,n} | X_n) - \sum_{i=1}^N \mathbb{E}(\xi_{i,n}^2 | X_n) \\ &= NX_n - N(N-1) \left(\frac{X_n}{N} \right)^2 - X_n \\ &= \frac{N-1}{N} X_n (N - X_n) = \left(\frac{N-1}{N} \right)^{n+1} M_n. \end{aligned}$$

On en déduit que

$$\mathbb{E}(X_n(N - X_n)) = \left(\frac{N-1}{N} \right)^n \mathbb{E}(M_n) = \left(\frac{N-1}{N} \right)^n \mathbb{E}(M_0) = \left(\frac{N-1}{N} \right)^n X_0(N - X_0).$$

Comme X_n converge vers X_∞ et est uniformément bornée par N , le théorème de convergence dominée implique

$$\begin{aligned} 0 &\leq \mathbb{E}(X_\infty(N - X_\infty)) = \lim_n \mathbb{E}(X_n(N - X_n)) \\ &= \lim_n \left(\frac{N-1}{N} \right)^n X_0(N - X_0) = 0. \end{aligned}$$

On conclut que $\mathbb{E}(X_\infty(N - X_\infty)) = 0$ et donc X_∞ vaut 0 ou N presque sûrement. La suite X_n converge vers $X_\infty \in \{0, N\}$ et ne prend que des valeurs discrètes, par conséquent il existe un temps aléatoire presque sûrement fini

$$\tau = \inf\{n \geq 0; \quad X_n = 0 \text{ ou } X_n = N\}$$

au-delà duquel $X_n = 0$ ou $X_n = N$ pour tout $n \geq \tau$. Ceci montre que presque sûrement un des allèles disparaît.

Nous allons montrer que la probabilité de disparition de l'allèle A vaut

$$\mathbb{P}(X_\tau = 0) = 1 - \frac{X_0}{N}.$$

La martingale $\{X_n\}_{n \geq 0}$ est bornée et τ est fini presque sûrement, par conséquent ce résultat s'obtient en appliquant le théorème 10.6 d'arrêt de Doob (cas (ii))

$$X_0 = \mathbb{E}(X_\tau) = \mathbb{E}(\mathbf{1}_{X_\tau=0} X_\tau) + \mathbb{E}(\mathbf{1}_{X_\tau=N} X_\tau) = N(1 - \mathbb{P}(X_\tau = 0)).$$

Revenons maintenant sur la nature de la convergence de la martingale $\{M_n\}_{n \geq 0}$. Les résultats précédents ont montré que

$$\mathbb{P}(\exists n \text{ tel que } \forall k \geq n, \quad X_k(N - X_k) = 0) = 1$$

et donc

$$\mathbb{P}(\exists n \text{ tel que } \forall k \geq n, \quad M_k = 0) = 1.$$

Cela signifie que M_n converge vers 0 presque sûrement. Cependant si $X_0 \notin \{0, N\}$, alors M_n ne peut pas converger vers 0 dans $\mathbb{L}^1$ car

$$\mathbb{E}(|M_n - M_\infty|) = \mathbb{E}(M_n) = M_0 \neq 0$$

où $M_0 = X_0(N - X_0) \neq 0$.

11.5 Martingales fermées

Pour que la convergence dans le théorème 11.5 ait lieu dans $\mathbb{L}^1$, il est nécessaire de considérer une classe particulière de martingales.

Définition 11.9 (Martingales fermées). *Un processus aléatoire $\{M_n\}_{n \geq 0}$ est une martingale fermée s'il existe une variable aléatoire intégrable M_∞ telle que $M_n = \mathbb{E}(M_\infty | \mathcal{F}_n)$ pour tout $n \geq 0$.*

Toute l'information d'une martingale fermée est codée dans une variable aléatoire limite. Les séries aléatoires illustrent cette notion car pour $\alpha > 1/2$

$$\forall n \geq 1, \quad M_n = \sum_{i=1}^n \frac{\zeta_i}{i^\alpha} \xrightarrow[n \rightarrow \infty]{\mathbb{L}^2} M_\infty := \sum_{i=1}^{\infty} \frac{\zeta_i}{i^\alpha}.$$

On vérifie alors que la martingale $\{M_n\}_{n \geq 1}$ est fermée pour la filtration $\mathcal{F}_n = \sigma(\zeta_i, i \leq n)$, car $M_n = \mathbb{E}(M_\infty | \mathcal{F}_n)$.

Dans le cas du théorème 11.5, si la convergence n'a pas lieu dans $\mathbb{L}^1$, une partie de l'information sur le processus a disparu dans le passage à la limite. Ainsi, la martingale arrêtée définie dans la remarque 11.8 ne converge pas dans $\mathbb{L}^1$ et elle n'est pas fermée. Un autre exemple est donné par la martingale

$$M_n = \exp \left(\sum_{i=1}^n \zeta_i - \frac{1}{2}n \right),$$

où $\{\zeta_i\}_{i \geq 1}$ est une suite de variables aléatoires gaussiennes $\mathcal{N}(0, 1)$ indépendantes et identiquement distribuées. Par la loi des grands nombres $\frac{1}{n} \sum_{i=1}^n \zeta_i$ converge vers 0 et donc M_n converge vers $M_\infty = 0$ presque sûrement. La limite ne contient plus d'information sur le processus et la martingale $\{M_n\}_{n \geq 1}$ n'est pas fermée.

Il faut donc renforcer l'hypothèse

$$\sup_{n \geq 0} \mathbb{E}(|M_n|) < \infty$$

du théorème 11.5 pour obtenir une convergence dans $\mathbb{L}^1$ et la propriété de martingale fermée. Nous allons voir maintenant qu'il suffit de modifier très légèrement cette hypothèse et de supposer qu'il existe $p > 1$ tel que $\sup_{n \geq 0} \mathbb{E}(|M_n|^p) < \infty$ pour que $\{M_n\}_{n \geq 0}$ soit une martingale fermée (cf. remarque 11.11). La convergence a alors lieu dans $\mathbb{L}^p$ et donc aussi dans $\mathbb{L}^1$.

Théorème 11.10 (Convergence dans $\mathbb{L}^p$). *Soit $p > 1$ et $\{M_n\}_{n \geq 0}$ une martingale telle que*

$$\sup_n \mathbb{E}(|M_n|^p) < \infty. \quad (11.11)$$

Alors $\{M_n\}_{n \geq 0}$ est fermée et il existe M_∞ dans $\mathbb{L}^p$, mesurable pour $\mathcal{F}_\infty = \sigma(\bigcup_{n \geq 0} \mathcal{F}_n)$ telle que

$$\forall n \geq 0, \quad M_n = \mathbb{E}(M_\infty | \mathcal{F}_n).$$

De plus, $\{M_n\}_{n \geq 0}$ converge vers M_∞ presque sûrement et dans $\mathbb{L}^p$

$$\lim_{n \rightarrow \infty} \mathbb{E}((M_n - M_\infty)^p) = 0.$$

Ce résultat généralise la convergence $\mathbb{L}^2$ du théorème 11.1.

Démonstration. Par l'inégalité de Hölder, la borne uniforme (11.11) dans $\mathbb{L}^p$ implique que $\sup_n \mathbb{E}(|M_n|) < \infty$. Par conséquent le théorème 11.5 s'applique et on en déduit que M_n converge presque sûrement vers une variable M_∞ mesurable pour $\mathcal{F}_\infty$ (cf. Proposition A.13).

La borne uniforme (11.11) permet de retrouver la convergence dans $\mathbb{L}^p$ car l'inégalité maximale de Doob (théorème 10.9) implique que $M_n^* = \sup_{k \leq n} |M_k|$ satisfait

$$\|M_n^*\|_p \leq \frac{p}{p-1} \|M_n\|_p \quad \text{pour tout } n \geq 0,$$

avec la notation $\|X\|_p = \mathbb{E}(|X|^p)^{1/p}$. Comme la suite $n \rightarrow M_n^*$ est croissante, la limite $M_\infty^* = \sup_k M_k^*$ existe. Par le théorème de convergence monotone et l'hypothèse (11.11), on vérifie que M_∞^* est dans $\mathbb{L}^p$

$$\|M_\infty^*\|_p = \lim_{n \rightarrow \infty} \|M_n^*\|_p \leq \frac{p}{p-1} \sup_{n \geq 0} \|M_n\|_p < \infty.$$

On en déduit une borne uniforme sur la martingale et sa limite

$$n \geq 0, \quad |M_n| \leq M_\infty^* \quad \text{et} \quad |M_\infty| \leq M_\infty^*.$$

Par conséquent, le théorème de convergence dominée s'applique et la convergence a aussi lieu dans $\mathbb{L}^p$

$$M_n \xrightarrow[n \rightarrow \infty]{\mathbb{L}^p} M_\infty.$$

Il ne reste plus à vérifier que $M_n = \mathbb{E}(M_\infty | \mathcal{F}_n)$. Soit A un événement arbitraire dans $\mathcal{F}_n$. Comme $\{M_n\}_{n \geq 0}$ est une martingale, on observe que pour tout $N \geq n$

$$|\mathbb{E}((M_n - M_\infty)\mathbf{1}_A)| = |\mathbb{E}((M_N - M_\infty)\mathbf{1}_A)| \leq \mathbb{E}(|M_N - M_\infty|) \xrightarrow[N \rightarrow \infty]{} 0.$$

Par conséquent, pour tout $A \in \mathcal{F}_n$

$$\mathbb{E}((M_n - M_\infty)\mathbf{1}_A) = 0$$

ce qui montre bien que $M_n = \mathbb{E}(M_\infty | \mathcal{F}_n)$.

□

Remarque 11.11. Pour mieux comprendre le rôle de l'hypothèse (11.11) sur la norme $\mathbb{L}^p$, reprenons l'exemple de la loterie vu dans la remarque 10.7. Si le gain est modélisé par une variable aléatoire prenant 2 valeurs

$$\mathbb{P}(X_N = N) = \frac{1}{N}, \quad \mathbb{P}(X_N = 0) = 1 - \frac{1}{N},$$

alors le gain moyen est constant $\mathbb{E}(X_N) = 1$, même si X_N tend vers 0 presque sûrement quand N tend vers l'infini. Nous avons déjà interprété ce phénomène par l'existence de fluctuations très rares. La norme $\mathbb{L}^p$ permet de mettre en évidence ces fluctuations

$$\mathbb{E}(X_N^p) = \frac{1}{N} N^p \xrightarrow[n \rightarrow \infty]{} \infty \quad \text{dès que } p > 1.$$

L'hypothèse (11.11) est donc naturelle car si elle est satisfaite, le processus aléatoire ne pourra pas trop fluctuer et on s'attend à une convergence dans un sens très fort.

L'hypothèse (11.11) sur la borne $\mathbb{L}^p$ peut être améliorée en considérant la classe des processus uniformément intégrables définie ci-dessous (voir aussi la définition B.12). Cette notion peut être omise en première lecture.

Définition 11.12 (Uniforme intégrabilité). Un processus $\{X_n\}_{n \geq 0}$ est dit uniformément intégrable si

$$\lim_{c \rightarrow \infty} \sup_n \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| \geq c\}}) = 0.$$

L'uniforme intégrabilité permet de renforcer la convergence en probabilité pour obtenir l'équivalence

$$X_n \xrightarrow[n \rightarrow \infty]{\mathbb{L}^1} X \quad \Leftrightarrow \quad \begin{cases} X_n \xrightarrow[n \rightarrow \infty]{\text{probabilité}} X \\ \{X_n\}_{n \geq 0} \text{ est uniformément intégrable} \end{cases}$$

démontrée dans le théorème B.13.

Le résultat suivant caractérise la convergence des martingales dans $\mathbb{L}^1$. On notera $\mathcal{F}_\infty = \sigma(\cup_{n \geq 0} \mathcal{F}_n)$ la σ -algèbre limite.

Théorème 11.13 (*). Soit $M = \{M_n\}_{n \geq 0}$ une martingale. Les deux assertions suivantes sont équivalentes :

- (i) M est fermée, i.e. il existe M_∞ dans $\mathbb{L}^1$, mesurable pour $\mathcal{F}_\infty$, telle que pour tout $n \geq 0$

$$M_n = \mathbb{E}(M_\infty | \mathcal{F}_n)$$

et la convergence $M_n \rightarrow M_\infty$ a lieu dans $\mathbb{L}^1$ et presque sûrement.

- (ii) M est uniformément intégrable.

Démonstration. La preuve se décompose en deux étapes.

(i) $\implies$ (ii) :

Comme M_∞ appartient à $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$, on déduit du lemme A.20 que pour $\varepsilon > 0$, il existe $\delta > 0$ tel que

$$\mathbb{E}(|M_\infty| \mathbf{1}_A) \leq \varepsilon \quad \text{pour tout } A \in \mathcal{A} \text{ de mesure } \mathbb{P}(A) \leq \delta. \quad (11.12)$$

D'après l'inégalité de Jensen, on a

$$|M_n| \leq \mathbb{E}(|M_\infty| \mid \mathcal{F}_n) \quad \text{pour tout } n \geq 0. \quad (11.13)$$

L'inégalité de Chebychev et la propriété des projections itérées de l'espérance conditionnelle impliquent donc

$$\mathbb{P}(|M_n| \geq c) \leq \frac{\mathbb{E}(|M_n|)}{c} \leq \frac{\mathbb{E}(\mathbb{E}(|M_\infty| \mid \mathcal{F}_n))}{c} = \frac{\mathbb{E}(|M_\infty|)}{c} \leq \delta \quad \text{pour } c \text{ assez grand.}$$

On peut donc appliquer (11.12) à l'événement $A = \{|M_n| \geq c\}$ pour obtenir

$$\varepsilon \geq \mathbb{E}(|M_\infty| \mathbf{1}_{\{|M_n| \geq c\}}) = \mathbb{E}(\mathbb{E}(|M_\infty| \mid \mathcal{F}_n) \mathbf{1}_{\{|M_n| \geq c\}}) \geq \mathbb{E}(|M_n| \mathbf{1}_{\{|M_n| \geq c\}})$$

où nous avons utilisé la propriété des projections itérées de l'espérance conditionnelle et (11.13).

(ii) $\implies$ (i) :

Pour la réciproque supposons maintenant que M soit uniformément intégrable. Alors $\sup_n \mathbb{E}(|M_n|) < \infty$ et on retrouve le cadre du théorème 11.5 qui assure l'existence d'une variable aléatoire M_∞ dans $\mathbb{L}^1$ telle que $M_n \rightarrow M_\infty$ presque sûrement. Dans le reste de cette preuve, nous allons montrer que cette convergence a lieu dans $\mathbb{L}^1$ et que $M_n = \mathbb{E}(M_\infty \mid \mathcal{F}_n)$ pour tout $n \geq 0$.

Pour $c > 0$, on définit la fonction

$$x \in \mathbb{R}, \quad f_c(x) = x \mathbf{1}_{|x| \leq c} + c \mathbf{1}_{x > c} - c \mathbf{1}_{x < -c}$$

et on décompose

$$\begin{aligned} \mathbb{E}(|M_n - M_\infty|) &\leq \mathbb{E}(|f_c(M_n) - M_n|) + \mathbb{E}(|f_c(M_\infty) - M_\infty|) \\ &\quad + \mathbb{E}(|f_c(M_n) - f_c(M_\infty)|). \end{aligned}$$

On fixe $\varepsilon > 0$. Comme

$$|f_c(x) - x| \leq |x| \mathbf{1}_{|x| \geq c}$$

l'uniforme intégrabilité assure l'existence d'une constante $c_0 > 0$ telle que pour $c \geq c_0$ et pour tout $n \geq 0$

$$\mathbb{E}(|f_c(M_n) - M_n|) \leq \varepsilon/2.$$

De plus (11.12) implique que pour c assez grand

$$\mathbb{E}(|f_c(M_\infty) - M_\infty|) \leq \varepsilon/2.$$

Comme la fonction f_c est continue bornée, on déduit du théorème de convergence dominée que

$$\mathbb{E} (|f_c(M_n) - f_c(M_\infty)|) \leq \varepsilon \quad \text{pour } n \text{ assez grand.}$$

On a finalement $\mathbb{E} (|M_n - M_\infty|) \leq 2\varepsilon$. Ceci prouve la convergence dans $\mathbb{L}^1$ de M_n vers M_∞ .

Finalement, on peut démontrer que $M_n = \mathbb{E}(M_\infty | \mathcal{F}_n)$ en suivant la même stratégie qu'à la fin de la preuve du théorème 11.10. $\square$

11.6 Théorème central limite

Dans cette section, nous montrons d'abord un théorème central limite pour des martingales que nous étendrons ensuite aux chaînes de Markov.

11.6.1 Théorème central limite pour les martingales

Soient $\{X_n\}_{n \geq 0}$ des variables aléatoires indépendantes, identiquement distribuées telles que

$$\mathbb{E}(X_1) = 0, \quad \mathbb{E}(X_1^2) = \sigma^2.$$

Le théorème central limite (rappelé théorème B.23) implique, après normalisation, la convergence en loi de $M_n = \sum_{i=1}^n X_i$ vers une variable gaussienne centrée de variance σ^2 notée $\mathcal{N}(0, \sigma^2)$

$$\frac{1}{\sqrt{n}} M_n \xrightarrow[n \rightarrow \infty]{\text{loi}} \mathcal{N}(0, \sigma^2).$$

Dans cet énoncé, $\{M_n\}_{n \geq 0}$ est une martingale dont les accroissements sont indépendants. Le résultat suivant généralise le théorème central limite aux martingales (dont les accroissements ne sont plus indépendants mais satisfont certaines propriétés).

Théorème 11.14. *Soit $\{M_n\}_{n \geq 0}$ une martingale dont les accroissements $\Delta M_n = M_n - M_{n-1}$ vérifient les propriétés suivantes :*

(i) *il existe une constante σ^2 pour laquelle la convergence ci-dessous a lieu en probabilité*

$$\frac{1}{n} \langle M \rangle_n := \frac{1}{n} \sum_{k=1}^n \mathbb{E} ((\Delta M_k)^2 | \mathcal{F}_{k-1}) \xrightarrow[n \rightarrow \infty]{\mathbb{P}} \sigma^2. \quad (11.14)$$

Rappelons que le crochet $\langle M \rangle_n$ de la martingale a déjà été introduit dans la définition 11.2.

(ii) *il existe une constante $\alpha > 2$ telle que la convergence ci-dessous a lieu en probabilité*

$$\frac{1}{n^{\alpha/2}} \sum_{k=1}^n \mathbb{E} (|\Delta M_k|^\alpha | \mathcal{F}_{k-1}) \xrightarrow[n \rightarrow \infty]{\mathbb{P}} 0. \quad (11.15)$$

Sous ces deux conditions, la martingale renormalisée converge en loi

$$\frac{1}{\sqrt{n}} M_n \xrightarrow[n \rightarrow \infty]{\text{loi}} \mathcal{N}(0, \sigma^2). \quad (11.16)$$

Si la martingale $M_n = \sum_{i=1}^n X_i$ est une somme de variables indépendantes, la condition (11.14) se résume au calcul de la variance de X_1 car $\langle M \rangle_n = n\mathbb{E}(X_1^2)$. Le critère (11.15), dit de Lyapounov, est automatiquement vérifié dès que les variables X_i ont un moment d'ordre $\alpha > 2$ car dans ce cas

$$\frac{1}{n^{\alpha/2}} \sum_{k=1}^n \mathbb{E}(|\Delta M_k|^\alpha | \mathcal{F}_{k-1}) = n^{1-\alpha/2} \mathbb{E}(|X_1|^\alpha) \xrightarrow{n \rightarrow \infty} 0.$$

Pour des variables indépendantes, la condition $\alpha > 2$ ne permet pas de retrouver exactement le théorème central limite B.23 qui s'applique sous l'hypothèse plus faible de variance finie $\mathbb{E}(X_1^2) < \infty$. Cependant le critère (11.15) peut être généralisé par la condition de Lindeberg (cf. [3] théorème 3.9.1)

$$\forall \varepsilon > 0, \quad \frac{1}{n} \sum_{k=1}^n \mathbb{E} \left(|\Delta M_k|^2 \mathbf{1}_{\{|\Delta M_k|^2 > \varepsilon n\}} | \mathcal{F}_{k-1} \right) \xrightarrow[n \rightarrow \infty]{\mathbb{P}} 0$$

qui implique aussi le théorème central limite B.23 pour des variables indépendantes. La normalisation par $\sqrt{n}$ dans (11.16) provient du comportement asymptotique des incréments (11.14). D'autres types de normalisation existent sous des hypothèses plus générales (cf. [3] théorème 3.9.1).

Démonstration. Pour simplifier la preuve, nous allons renforcer l'hypothèse (ii) et remplacer la condition (11.15) par la borne suivante

$$\forall k \geq 1, \quad \mathbb{E}(|\Delta M_k|^3 | \mathcal{F}_{k-1}) \leq K, \quad (11.17)$$

avec K une constante fixée. Cette hypothèse simplificatrice permet en particulier de garantir, en utilisant l'inégalité de Hölder, la majoration suivante

$$\forall n \geq 1, \quad \mathbb{E}((\Delta M_n)^2 | \mathcal{F}_{n-1}) \leq \mathbb{E}(|\Delta M_n|^3 | \mathcal{F}_{n-1})^{2/3} \leq K^{2/3}, \quad (11.18)$$

et une borne uniforme sur $\frac{1}{n}\langle M \rangle_n$

$$\forall n \geq 1, \quad \left| \frac{\langle M \rangle_n}{n} \right| \leq \frac{1}{n} \sum_{k=1}^n \mathbb{E}((\Delta M_k)^2 | \mathcal{F}_{k-1}) \leq K^{2/3}. \quad (11.19)$$

Comme dans le cas des sommes de variables indépendantes, pour prouver le théorème 11.14, il suffit de montrer la convergence des fonctions caractéristiques

$$\forall u \in \mathbb{R}, \quad \lim_{n \rightarrow \infty} \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n \right) \right] = \exp \left(-\sigma^2 \frac{u^2}{2} \right) \quad (11.20)$$

pour en déduire ensuite la convergence vers une loi gaussienne par le théorème B.15 de Lévy. L'étape principale vers ce résultat sera d'établir la limite suivante

$$\forall u \in \mathbb{R}, \quad \lim_{n \rightarrow \infty} \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n + \frac{u^2}{2n} \langle M \rangle_n \right) \right] = 1. \quad (11.21)$$

En effet, supposons pour le moment le résultat (11.21) et montrons qu'il implique la convergence (11.20). Pour tout u dans $\mathbb{R}$ et ε dans $]0, 1[$, on peut écrire

$$\begin{aligned} & \left| \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n \right) \right] \exp \left(\sigma^2 \frac{u^2}{2} \right) - \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n + \frac{u^2}{2n} \langle M \rangle_n \right) \right] \right| \\ & \leq \mathbb{E} \left[\left| \exp \left(\sigma^2 \frac{u^2}{2} \right) - \exp \left(\frac{u^2}{2n} \langle M \rangle_n \right) \right| \right] \\ & \leq C_1 \varepsilon \mathbb{P} \left[\left| \frac{\langle M \rangle_n}{n} - \sigma^2 \right| \leq \varepsilon \right] + C_2 \mathbb{P} \left[\left| \frac{\langle M \rangle_n}{n} - \sigma^2 \right| > \varepsilon \right], \end{aligned}$$

où la dernière inégalité est obtenue en utilisant la régularité de la fonction $x \mapsto \exp(\frac{u^2}{2}x)$ et la borne uniforme (11.19). Les constantes C_1, C_2 dépendent de K, u, σ^2 mais pas de ε . En utilisant la convergence en probabilité de l'hypothèse (i) pour contrôler la limite $n \rightarrow \infty$, puis en laissant tendre ε vers 0, on conclut que

$$\lim_{n \rightarrow \infty} \left| \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n \right) \right] \exp \left(\sigma^2 \frac{u^2}{2} \right) - \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n + \frac{u^2}{2n} \langle M \rangle_n \right) \right] \right| = 0.$$

Il ne reste plus qu'à appliquer le résultat (11.21) pour en déduire la convergence (11.20) de la fonction caractéristique.

Nous allons maintenant prouver la convergence (11.21) en décomposant la contribution de chaque accroissement. Pour cela, fixons n et posons pour tout k dans $\{1, \dots, n\}$

$$Z_k = \exp \left(i \frac{u}{\sqrt{n}} \Delta M_k + \frac{u^2}{2n} \mathbb{E}((\Delta M_k)^2 | \mathcal{F}_{k-1}) \right). \quad (11.22)$$

Un développement de Taylor permet d'obtenir l'inégalité suivante pour tout z dans $\mathbb{C}$

$$\left| \exp(z) - 1 - z - \frac{1}{2}z^2 \right| \leq |z|^3 \exp(|\operatorname{Re}(z)|),$$

où $\operatorname{Re}(z)$ représente la partie réelle de z . Cette inégalité implique

$$\begin{aligned} \mathbb{E}(Z_k | \mathcal{F}_{k-1}) &= 1 + \mathbb{E} \left(i \frac{u}{\sqrt{n}} \Delta M_k + \frac{u^2}{2n} \mathbb{E}((\Delta M_k)^2 | \mathcal{F}_{k-1}) \mid \mathcal{F}_{k-1} \right) \\ &+ \frac{1}{2} \mathbb{E} \left(-\frac{u^2}{n} (\Delta M_k)^2 + i \frac{u^3}{n^{3/2}} \Delta M_k \mathbb{E}((\Delta M_k)^2 | \mathcal{F}_{k-1}) \mid \mathcal{F}_{k-1} \right) + R_k, \end{aligned} \quad (11.23)$$

avec un reste R_k majoré par

$$\begin{aligned} |R_k| &\leq \mathbb{E} \left(\frac{u^4}{8n^2} \mathbb{E}((\Delta M_k)^2 | \mathcal{F}_{k-1})^2 \mid \mathcal{F}_{k-1} \right) \\ &+ \mathbb{E} \left(\left| i \frac{u}{\sqrt{n}} \Delta M_k + \frac{u^2}{2n} \mathbb{E}((\Delta M_k)^2 | \mathcal{F}_{k-1}) \right|^3 \exp \left(\frac{u^2}{2n} \mathbb{E}((\Delta M_k)^2 | \mathcal{F}_{k-1}) \right) \mid \mathcal{F}_{k-1} \right) \\ &\leq \frac{C}{n^{3/2}}. \end{aligned} \quad (11.24)$$

La dernière majoration est une conséquence de l'estimation (11.18) et la constante C ne dépend que de K et u .

La structure de martingale permet d'obtenir les identités suivantes

$$\mathbb{E} \left(\Delta M_k \mid \mathcal{F}_{k-1} \right) = 0 \quad \text{et} \quad \mathbb{E} \left(\Delta M_k \mathbb{E} \left((\Delta M_k)^2 \mid \mathcal{F}_{k-1} \right) \mid \mathcal{F}_{k-1} \right) = 0$$

et ainsi de simplifier (11.23)

$$1 \leq k \leq n, \quad \mathbb{E} \left(Z_k \mid \mathcal{F}_{k-1} \right) = 1 + R_k. \quad (11.25)$$

Nous sommes maintenant en mesure de prouver (11.21). Comme $\langle M \rangle_n - \langle M \rangle_{n-1} = \mathbb{E} \left((\Delta M_n)^2 \mid \mathcal{F}_{n-1} \right)$, on peut décomposer en fonction du dernier incrément

$$\begin{aligned} \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n + \frac{u^2}{2n} \langle M \rangle_n \right) \right] &= \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_{n-1} + \frac{u^2}{2n} \langle M \rangle_{n-1} \right) Z_n \right] \\ &= \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_{n-1} + \frac{u^2}{2n} \langle M \rangle_{n-1} \right) \mathbb{E} \left(Z_n \mid \mathcal{F}_{n-1} \right) \right] \\ &= \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_{n-1} + \frac{u^2}{2n} \langle M \rangle_{n-1} \right) (1 + R_n) \right], \end{aligned}$$

où la dernière égalité repose sur l'identité (11.25) avec $k = n$. Il suffit alors d'utiliser la majoration (11.24) de R_n et le fait que $\frac{1}{n-1} \langle M \rangle_{n-1}$ soit borné (11.19) pour conclure que

$$\left| \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n + \frac{u^2}{2n} \langle M \rangle_n \right) \right] - \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_{n-1} + \frac{u^2}{2n} \langle M \rangle_{n-1} \right) \right] \right| \leq \frac{C}{n^{3/2}}.$$

En itérant n fois le calcul précédent, on en déduit

$$\begin{aligned} &\left| \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_n + \frac{u^2}{2n} \langle M \rangle_n \right) \right] - \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_1 + \frac{u^2}{2n} \langle M \rangle_1 \right) \right] \right| \\ &\leq \sum_{k=2}^n \left| \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_k + \frac{u^2}{2n} \langle M \rangle_k \right) \right] - \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_{k-1} + \frac{u^2}{2n} \langle M \rangle_{k-1} \right) \right] \right| \leq \frac{C}{\sqrt{n}}. \end{aligned} \quad (11.26)$$

Dans (11.26), le terme exponentiel impliquant M_1 et $\langle M \rangle_1$ converge vers 1 presque sûrement quand n tend vers l'infini. Comme $\langle M \rangle_1$ est borné d'après (11.19), il suffit d'appliquer le théorème de convergence dominée pour en déduire

$$\lim_{n \rightarrow \infty} \mathbb{E} \left[\exp \left(i \frac{u}{\sqrt{n}} M_1 + \frac{u^2}{2n} \langle M \rangle_1 \right) \right] = 1.$$

Par l'inégalité (11.26), ceci démontre la limite (11.21) et conclut la preuve du théorème central limite. $\square$

11.6.2 Inégalité de Hoeffding

Le théorème central limite 11.14 fournit un résultat asymptotique de l'écart entre M_n et sa moyenne. Le théorème suivant permet d'estimer les fluctuations pour des valeurs de n fixées.

Théorème 11.15 (Inégalité de Hoeffding). *Soit $\{M_n\}_{n \geq 0}$ une martingale telle que $M_0 = 0$ et dont les accroissements $\Delta M_n = M_n - M_{n-1}$ sont majorés par une suite $\{K_n\}_{n \geq 1}$*

$$\forall n \geq 1, \quad |\Delta M_n| \leq K_n.$$

Alors, pour tout $x \geq 0$ et $n \geq 1$

$$\mathbb{P}(|M_n| \geq x) \leq 2 \exp \left(-\frac{x^2}{2 \sum_{i=1}^n K_i^2} \right). \quad (11.27)$$

L'hypothèse du théorème sur les accroissements rappelle la condition (11.14).

Appliquons cette inégalité au cas d'une marche aléatoire symétrique $S_n = \sum_{i=1}^n \xi_i$ où les $\{\xi_i\}_{i \geq 1}$ sont des variables aléatoires indépendantes centrées et identiquement distribuées à valeurs dans $[-1, 1]$. L'hypothèse sur les accroissements est donc satisfaite avec $K_n = 1$ et par conséquent

$$\forall u > 0, \quad \mathbb{P}(|S_n| \geq u) \leq 2 \exp \left(-\frac{u^2}{2n} \right).$$

En particulier, on peut remplacer u par $n\varepsilon$ avec $\varepsilon > 0$

$$\mathbb{P} \left(\left| \frac{S_n}{n} \right| \geq \varepsilon \right) \leq 2 \exp \left(-\frac{n\varepsilon^2}{2} \right).$$

La probabilité de s'écarter de l'espérance est donc exponentiellement petite en n .

On peut aussi remplacer u par $x\sqrt{n}$ pour tout $x \geq 0$

$$\mathbb{P} \left(\left| \frac{S_n}{\sqrt{n}} \right| \geq x \right) \leq 2 \exp \left(-\frac{x^2}{2} \right).$$

Cette borne supérieure rappelle l'asymptotique du théorème central limite, mais elle est cette fois valable pour tout n . Ceci est intéressant en pratique car le nombre de données disponibles est parfois faible et cette borne théorique permet de quantifier les écarts dus aux fluctuations.

Démonstration. La preuve se décompose en deux étapes.

Étape 1. Soit Z une variable aléatoire à valeurs dans $[a, b] \subset \mathbb{R}$ telle que $\mathbb{E}(Z) = 0$. Nous allons montrer que

$$\forall \lambda \in \mathbb{R}, \quad \mathbb{E} \left(e^{\lambda Z} \right) \leq e^{\frac{\lambda^2}{8} (b-a)^2}.$$

Posons $\psi(\lambda) = \log \mathbb{E} (e^{\lambda Z})$. On vérifie aisément les faits suivants

$$\psi(0) = 0$$

$$\psi'(\lambda) = \frac{\mathbb{E} (Z e^{\lambda Z})}{\mathbb{E} (e^{\lambda Z})}, \quad \psi'(0) = 0$$

$$\psi''(\lambda) = \frac{\mathbb{E} (Z^2 e^{\lambda Z}) \mathbb{E} (e^{\lambda Z}) - (\mathbb{E} (Z e^{\lambda Z}))^2}{(\mathbb{E} (e^{\lambda Z}))^2} = \text{Var}_{\mathbb{P}_\lambda}(Z)$$

où la mesure $\mathbb{P}_\lambda$ est définie par

$$\mathbb{P}_\lambda(\cdot) = \frac{\mathbb{P}(\cdot e^{\lambda Z})}{\mathbb{E}(e^{\lambda Z})}.$$

Un développement de Taylor en 0 à l'ordre deux implique

$$\begin{aligned} \psi(\lambda) &= \psi(0) + \lambda \psi'(0) + \int_0^\lambda (\lambda - s) \psi''(s) ds = \int_0^\lambda (\lambda - s) \psi''(s) ds \\ &= \int_0^\lambda (\lambda - s) \text{Var}_{\mathbb{P}_s}(Z) ds. \end{aligned}$$

Pour estimer $\text{Var}_{\mathbb{P}_\lambda}(Z)$, on utilise l'identité

$$\text{Var}_{\mathbb{P}_\lambda}(Z) = \inf_{m \in \mathbb{R}} \{ \mathbb{E}_{\mathbb{P}_\lambda}((Z - m)^2) \}.$$

Comme Z est à valeurs dans $[a, b]$, on a

$$\left| Z - \frac{a+b}{2} \right| \leq \frac{b-a}{2}.$$

Par conséquent, il suffit de choisir $m = \frac{a+b}{2}$ pour obtenir une borne supérieure sur la variance

$$\text{Var}_{\mathbb{P}_\lambda}(Z) \leq \mathbb{E}_{\mathbb{P}_\lambda} \left(\left(Z - \frac{a+b}{2} \right)^2 \right) \leq \frac{(b-a)^2}{4}.$$

Cette borne uniforme sur la variance permet d'obtenir l'inégalité cherchée

$$\psi(\lambda) \leq \int_0^\lambda (\lambda - s) \frac{(b-a)^2}{4} ds = \frac{(b-a)^2}{8} \lambda^2.$$

Étape 2. Supposons que

$$\mathbb{E}(\exp(\lambda M_n)) \leq \exp \left(\frac{\lambda^2}{2} \sum_{i=1}^n K_i^2 \right). \quad (11.28)$$

On peut en déduire (11.27) en appliquant l'inégalité de Markov pour tout $\lambda \geq 0$

$$\mathbb{P}(M_n \geq x) \leq \mathbb{E}(\exp(\lambda M_n)) \exp(-\lambda x) \leq \exp \left(\frac{\lambda^2}{2} \sum_{i=1}^n K_i^2 - \lambda x \right).$$

Comme $x \geq 0$, l'optimisation de l'inégalité en λ conduit à choisir la valeur $\lambda = x / (\sum_{i=1}^n K_i^2) \geq 0$ pour obtenir

$$\mathbb{P}(M_n \geq x) \leq \exp \left(-\frac{x^2}{2 \sum_{i=1}^n K_i^2} \right).$$

Par symétrie, cette inégalité est aussi vraie pour $-M_n$ et on retrouve ainsi l'inégalité (11.27).

Il reste donc à démontrer (11.28). Pour cela, on écrit M_n comme la somme télescopique des accroissements

$$M_n = \sum_{i=1}^n \Delta M_i \quad \text{avec} \quad \Delta M_i = M_i - M_{i-1}$$

qui vérifient $\mathbb{E}(\Delta M_i \mid \mathcal{F}_{i-1}) = 0$ et qui sont à valeurs dans l'intervalle $[-K_i, K_i]$. En appliquant à

$$\mathbb{E} \left(e^{\lambda M_n} \right) = \mathbb{E} \left(e^{\lambda \sum_{i=1}^n \Delta M_i} \right) = \mathbb{E} \left(e^{\lambda \sum_{i=1}^{n-1} \Delta M_i} \mathbb{E} \left(e^{\lambda \Delta M_n} \mid \mathcal{F}_{n-1} \right) \right)$$

l'inégalité de la première étape (avec $b - a = 2K_n$) pour les espérances conditionnelles, on en déduit

$$\mathbb{E} \left(e^{\lambda M_n} \right) \leq \mathbb{E} \left(e^{\lambda \sum_{i=1}^{n-1} \Delta M_i} \right) \exp \left(\frac{\lambda^2}{2} K_n^2 \right) \leq \exp \left(\frac{\lambda^2}{2} \sum_{i=1}^n K_i^2 \right)$$

où la dernière inégalité s'obtient par itération. $\square$

11.6.3 Théorème central limite pour les chaînes de Markov

On considère $\{X_n\}_{n \geq 0}$ une chaîne de Markov irréductible sur un espace d'états E fini. On notera P sa matrice de transition, π sa mesure invariante et $\mathbb{E}_\pi$ l'espérance sous π . Le théorème central limite 11.14 s'étend aux chaînes de Markov

Théorème 11.16. *Soit f une fonction de E dans $\mathbb{R}$ alors la convergence en loi ci-dessous a lieu quelle que soit la distribution de X_0*

$$\frac{1}{\sqrt{n}} \sum_{j=0}^n [f(X_j) - \mathbb{E}_\pi(f)] \xrightarrow[n \rightarrow \infty]{\text{loi}} \gamma \quad \text{et} \quad \gamma \stackrel{\text{loi}}{=} \mathcal{N}(0, \sigma^2)$$

où la variance σ^2 est définie en (11.29).

Démonstration. Quitte à changer f en $f - \mathbb{E}_\pi(f)$, on se ramène au cas $\mathbb{E}_\pi(f) = 0$. L'idée de la preuve consiste à comparer la somme $\sum_{j=0}^n f(X_j)$ à une martingale afin d'utiliser le théorème central limite 11.14 pour les martingales.

Étape 1 : Décomposition en martingale.

L'espace d'états E étant fini le critère de Doeblin est satisfait (5.18) et le théorème 5.11 s'applique

$$\forall n \geq 0, \quad \frac{1}{2} \sup_{x \in E} \left| \sum_{y \in E} P^n(x, y) f(y) - \pi(y) f(y) \right| \leq C \exp(-cn) \|f\|_\infty$$

où $c, C > 0$ sont des constantes indépendantes de f . Comme $\mathbb{E}_\pi(f) = 0$, on en déduit

$$\forall n \geq 0, \quad \sup_{x \in E} |P^n f(x)| \leq 2C \exp(-cn) \|f\|_\infty$$

et la convergence de la série

$$\forall x \in E, \quad u(x) = \sum_{k=0}^{\infty} P^k f(x)$$

avec la notation $P^0 f(x) = f(x)$. La fonction u est bornée car l'espace d'états E est fini. On remarque que

$$Pu(x) = \sum_{k=0}^{\infty} P^{k+1} f(x) = u(x) - f(x).$$

Par conséquent f satisfait

$$\forall x \in E, \quad f(x) = u(x) - Pu(x).$$

La somme se décompose donc sous la forme

$$\sum_{j=0}^n f(X_j) = u(X_0) + \left[\sum_{j=1}^n u(X_j) - Pu(X_{j-1}) \right] - Pu(X_n).$$

Soit $Z_j = u(X_j) - Pu(X_{j-1})$. En appliquant la propriété de Markov, on obtient

$$\mathbb{E}(Z_j | \mathcal{F}_{j-1}) = \mathbb{E}(u(X_j) - Pu(X_{j-1}) | \mathcal{F}_{j-1}) = \mathbb{E}(u(X_j) | X_{j-1}) - Pu(X_{j-1}) = 0.$$

Par conséquent $M_n = \sum_{j=1}^n Z_j$ est une martingale et nous avons prouvé la décomposition

$$\sum_{j=0}^n f(X_j) = u(X_0) - Pu(X_n) + M_n.$$

Les termes u et Pu sont bornés et ils ne contribuent pas à la limite $\sum_{j=0}^n f(X_j) / \sqrt{n}$. Il suffit donc de vérifier la convergence en loi de $M_n / \sqrt{n}$.

Étape 2 : Théorème central limite.

Pour montrer que l'hypothèse (11.14) du théorème 11.14 est satisfaite, on remarque que les accroissements de la martingale sont donnés par $\Delta M_j = u(X_j) - Pu(X_{j-1})$. L'espace E étant fini, la fonction u est bornée et par conséquent les accroissements ΔM_j aussi.

On peut réécrire $\mathbb{E}((\Delta M_j)^2 | \mathcal{F}_{j-1}) = \psi(X_{j-1})$ pour une certaine fonction ψ . Il suffit donc d'appliquer le théorème ergodique 5.1 pour les chaînes de Markov afin d'en déduire la convergence presque sûre

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{j=1}^n \mathbb{E}((\Delta M_j)^2 | \mathcal{F}_{j-1}) = \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{j=1}^n \psi(X_{j-1}) = \mathbb{E}_\pi(\psi) = \mathbb{E}_\pi((u(X_1) - Pu(X_0))^2).$$

Le théorème 11.14 permet de conclure que $\frac{1}{\sqrt{n}} \sum_{j=0}^n f(X_j)$ converge en loi vers une variable gaussienne centrée de variance

$$\sigma^2 = \mathbb{E}_\pi((u(X_1) - Pu(X_0))^2). \quad (11.29)$$

□

Le théorème central limite peut être généralisé au cas des espaces d'états dénombrables (cf. le livre [6]). En général, il n'est pas facile d'estimer la variance σ^2 et il faut utiliser les données de plusieurs trajectoires pour obtenir une estimation pertinente.

Chapitre 12

Applications des martingales

12.1 Cascades de Mandelbrot

Le phénomène de turbulence, observé fréquemment dans les écoulements de fluides (liquides ou gaz), se caractérise par l'apparition de tourbillons de différentes tailles dont la localisation et l'orientation varient très rapidement. Dans un écoulement turbulent la dissipation de l'énergie met en jeu des échelles multiples : les grands tourbillons se divisent en plus petits tourbillons et transmettent ainsi l'énergie aux petites échelles. Les mécanismes de la turbulence sont extrêmement compliqués et ils restent encore très mal compris, du point de vue mathématique, si on cherche à décrire analytiquement l'écoulement d'un fluide à partir des équations de Navier-Stokes.

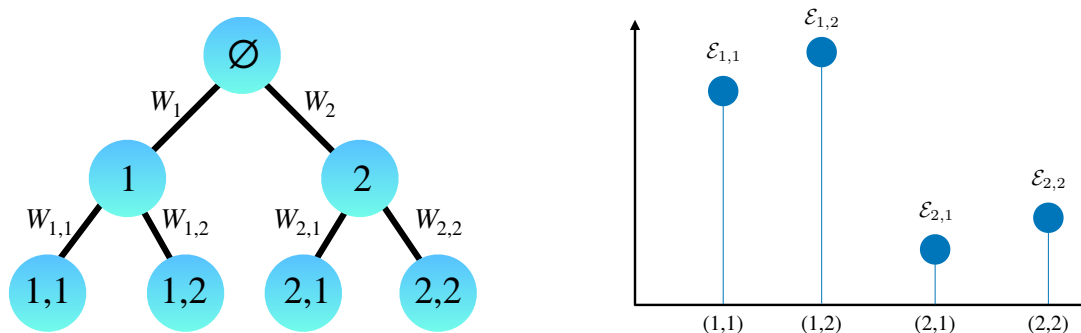


FIGURE 12.1 – Le schéma de gauche représente un arbre binaire de hauteur $k = 2$. Les énergies se transmettent en fonction des variables aléatoires W . Sur le graphique de droite, les énergies à l'échelle $k = 2$ sont tracées en énumérant les nœuds de l'arbre de gauche à droite. Chaque point correspond à l'énergie d'une cellule. La même convention sera adoptée figure 12.2.

Pour modéliser simplement les transferts d'énergie dans un écoulement turbulent, B. Mandelbrot a proposé une description où les différentes échelles du fluide sont codées par un arbre. La racine de l'arbre, notée $\emptyset$, représente l'échelle macroscopique et les petites échelles sont indexées en fonction de la hauteur $k \geq 1$. Étant donnée une hauteur k , un nœud de l'arbre binaire est déterminé par les indices $\{a_1, a_2, \dots, a_k\} \in \{1, 2\}^k$ qui décrivent le chemin entre la racine et ce nœud : on passe du niveau $k - 1$ au niveau k en suivant la branche de gauche si $a_k = 1$ et la branche de droite si $a_k = 2$ (cf. figure 12.1).

On considère aussi une collection de variables aléatoires indépendantes et identiquement distribuées $\{W_{a_1, a_2, \dots, a_k}\}$ indexées par les nœuds de l'arbre telles que

$$W_{a_1, a_2, \dots, a_k} \geq 0 \quad \text{et} \quad \mathbb{E}(W_{a_1, a_2, \dots, a_k}) = 1.$$

L'énergie $\mathcal{E}_\emptyset = 1$ est injectée à la racine, puis elle se répartit aléatoirement à l'échelle $k = 1$ sous la forme $\mathcal{E}_1 = \frac{1}{2}W_1\mathcal{E}_\emptyset$ et $\mathcal{E}_2 = \frac{1}{2}W_2\mathcal{E}_\emptyset$ en fonction des variables aléatoires W_1, W_2 . Ces énergies se divisent encore au niveau $k = 2$ de la façon suivante

$$\mathcal{E}_{1,1} = \frac{1}{2}W_{1,1}\mathcal{E}_1, \quad \mathcal{E}_{1,2} = \frac{1}{2}W_{1,2}\mathcal{E}_1, \quad \mathcal{E}_{2,1} = \frac{1}{2}W_{2,1}\mathcal{E}_2, \quad \mathcal{E}_{2,2} = \frac{1}{2}W_{2,2}\mathcal{E}_2.$$

En itérant cette procédure, on construit une cascade d'énergie, dite *cascade de Mandelbrot*, qui décrit la transmission de l'énergie de l'échelle macroscopique vers les plus petites échelles (cf. figure 12.1). À la hauteur k , le nœud $(a_1, a_2, \dots, a_k)$ recevra l'énergie

$$\mathcal{E}_{a_1, a_2, \dots, a_k} = \frac{1}{2^k} W_{a_1} W_{a_1, a_2} W_{a_1, a_2, a_3} \dots W_{a_1, a_2, \dots, a_k}, \quad (12.1)$$

dont la moyenne vaut $\mathbb{E}(\mathcal{E}_{a_1, a_2, \dots, a_k}) = \frac{1}{2^k}$. Il s'agit d'une représentation très idéalisée du transfert d'énergie dans un fluide. On peut imaginer que $\mathcal{E}_\emptyset$ représente l'énergie moyenne dans une boîte de taille 1. À l'échelle k , cette boîte est subdivisée en 2^k petites cellules de taille $\frac{1}{2^k}$ indexées par les nœuds à la hauteur k et l'énergie de la cellule $(a_1, \dots, a_k)$ vaut $\mathcal{E}_{a_1, a_2, \dots, a_k}$. L'énergie totale $\hat{\mathcal{E}}_k$ de l'échelle k a une espérance constante

$$\hat{\mathcal{E}}_k = \sum_{a_1, a_2, \dots, a_k} \mathcal{E}_{a_1, a_2, \dots, a_k} \quad \text{et} \quad \mathbb{E}(\hat{\mathcal{E}}_k) = 1, \quad (12.2)$$

où la somme porte sur les 2^k cellules. En moyenne, l'énergie totale est intégralement conservée d'une échelle à l'autre, cependant le mécanisme de transmission conduit à une répartition très irrégulière de l'énergie entre les différents nœuds de l'arbre comme on peut le voir sur la figure 12.2. Les cascades de Mandelbrot permettent ainsi de reproduire simplement le phénomène d'intermittence qui caractérise la turbulence.

Pour comprendre les transferts d'énergie dans l'arbre, nous allons d'abord suivre la variation de l'énergie le long d'une branche. Considérons la branche dont tous les indices a_i valent 1 et notons l'énergie du nœud situé à l'échelle k par

$$\mathbf{e}_k = \frac{1}{2^k} W_1 W_{1,1} W_{1,1,1} \dots W_{1, \dots, 1} \quad \text{où on a choisi} \quad a_i = 1, \quad i \leq k. \quad (12.3)$$

La suite des énergies $\{2^k \mathbf{e}_k\}_{k \geq 1}$ est une martingale car il s'agit d'un processus multiplicatif comme dans l'exemple (10.2). Cette martingale étant positive, elle converge presque sûrement d'après le corollaire 11.6. Pour déterminer sa limite, commençons par réécrire l'énergie en utilisant de nouvelles variables aléatoires

$$2^k \mathbf{e}_k = \exp \left(\sum_{i=1}^k X_i \right) \quad \text{avec} \quad X_i = \log \left(\underbrace{W_{1, \dots, 1}}_{i \text{ fois}} \right),$$

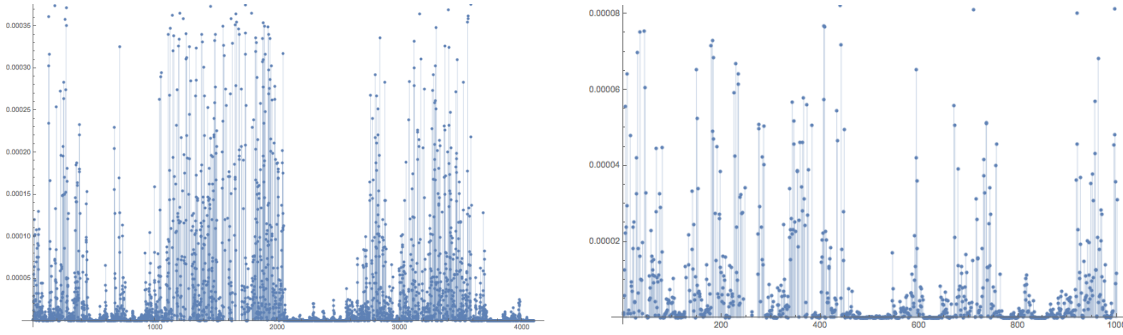


FIGURE 12.2 – Sur le graphique de gauche, les énergies d’une cascade de Mandelbrot à l’échelle $k = 12$ sont représentées en énumérant les nœuds de l’arbre de gauche à droite comme dans la figure 12.1. On remarque une variabilité extrême dans la répartition de l’énergie qui se concentre principalement dans certaines zones. Ceci rappelle un phénomène caractéristique de la turbulence, nommé *intermittence*, où l’énergie est localisée dans des régions qui varient au cours du temps et dont la répartition a des propriétés fractales. Les cascades de Mandelbrot possèdent aussi une structure fractale et le schéma de droite correspond simplement à un agrandissement des 1000 premières énergies du graphique de gauche.

où les variables aléatoires X_i sont indépendantes et identiquement distribuées. La fonction log étant strictement concave, l’inégalité de Jensen¹ implique que

$$\mathbb{E}(X_i) < \log \mathbb{E}(W_{1,\dots,1}) = 0.$$

Il suffit d’appliquer la loi des grands nombres à la suite $\{X_k\}_{k \geq 1}$ pour en déduire la convergence de la martingale $\{2^k \mathbf{e}_k\}_{k \geq 1}$ presque sûrement vers 0 car

$$\frac{1}{k} \sum_{i=1}^k X_i \xrightarrow[k \rightarrow \infty]{p.s.} \mathbb{E}(X_i) < 0 \quad \text{et donc} \quad 2^k \mathbf{e}_k \xrightarrow[k \rightarrow \infty]{p.s.} 0.$$

Ainsi le long de chaque branche de l’arbre, l’énergie renormalisée $\{2^k \mathbf{e}_k\}_{k \geq 1}$ tend vers 0, même si sa moyenne reste constante

$$\forall k \geq 0, \quad \mathbb{E}(2^k \mathbf{e}_k) = 1.$$

Nous avons déjà rencontré des phénomènes semblables au chapitre 10 (cf. par exemple la figure 10.2) et les mécanismes en jeu dans les cascades de Mandelbrot sont identiques. De grandes fluctuations des variables $2^k \mathbf{e}_k$ compensent le comportement typique et suffisent à maintenir une espérance constante à chaque échelle k . Ces fluctuations sont très difficiles à observer en suivant simplement l’énergie le long d’une branche donnée car la suite $\{2^k \mathbf{e}_k\}_{k \geq 1}$ ne prendra que rarement de grandes valeurs. Les cascades de Mandelbrot possèdent une structure beaucoup plus riche qui code en même temps 2^k processus aléatoires (fortement corrélés par l’organisation arborescente) à chaque hauteur k . En simulant une cascade de Mandelbrot, on construit simultanément un nombre exponentiellement grand de processus ce qui permet de voir les fluctuations de l’énergie car

1. Nous avons utilisé une amélioration de l’inégalité de Jensen présentée Proposition B.8 qui permet d’obtenir une inégalité stricte car les variables $W_{1,\dots,1}$ prennent au moins 2 valeurs distinctes.

parmi tous ces processus, certains ne vont pas suivre le comportement typique (cf. figure 12.2). Pour illustrer ce phénomène, nous allons montrer maintenant que l'énergie totale $\hat{\mathcal{E}}_k$ définie en (12.2) se comporte asymptotiquement très différemment de l'énergie $2^k \mathbf{e}_k$ associée à une cellule.

Théorème 12.1. *Le processus aléatoire $\{\hat{\mathcal{E}}_k\}_{k \geq 1}$ représentant l'énergie totale à l'échelle k est une martingale pour la filtration engendrée par les σ -algèbres*

$$\mathcal{F}_k = \sigma\left(W_{a_1, a_2, \dots, a_k}; \quad a_i \in \{1, 2\}, i \leq k\right).$$

Si $\mathbb{E}(W_1^2) < 2$, alors le processus $\{\hat{\mathcal{E}}_k\}_{k \geq 1}$ converge dans $\mathbb{L}^2$ et presque sûrement vers une variable aléatoire $\hat{\mathcal{E}}$ telle que $\mathbb{E}(\hat{\mathcal{E}}) = 1$. Sous l'hypothèse supplémentaire $\mathbb{P}(W_1 > 0) = 1$, la limite $\hat{\mathcal{E}}$ ne s'annule jamais, i.e. $\mathbb{P}(\hat{\mathcal{E}} > 0) = 1$.

Sous les conditions du théorème 12.1, l'énergie totale $\hat{\mathcal{E}}_k$ reste strictement positive quand k tend vers l'infini. Par conséquent, l'énergie totale n'est pas gouvernée par le comportement moyen des cellules mais par de grandes fluctuations de l'énergie qui sont localisées dans des zones restreintes de l'arbre (cf. figure 12.2).

Démonstration du théorème 12.1. Par la définition (12.2), l'énergie totale $\hat{\mathcal{E}}_k$ est mesurable par rapport à $\mathcal{F}_k$ et elle peut se réécrire sous la forme

$$\hat{\mathcal{E}}_k = \sum_{a_1, \dots, a_k} \frac{W_{a_1, a_2, \dots, a_k}}{2} \mathcal{E}_{a_1, a_2, \dots, a_{k-1}}.$$

Pour vérifier que le processus $\{\hat{\mathcal{E}}_k\}_{k \geq 1}$ est une martingale, il suffit d'utiliser l'indépendance des variables aléatoires W et leurs moyennes $\mathbb{E}(W_{a_1, a_2, \dots, a_k}) = 1$

$$\begin{aligned} \mathbb{E}(\hat{\mathcal{E}}_k | \mathcal{F}_{k-1}) &= \sum_{a_1, \dots, a_k} \mathbb{E}\left(\frac{W_{a_1, a_2, \dots, a_k}}{2} \mathcal{E}_{a_1, a_2, \dots, a_{k-1}} \middle| \mathcal{F}_{k-1}\right) = \sum_{a_1, \dots, a_k} \mathbb{E}\left(\frac{W_{a_1, a_2, \dots, a_k}}{2}\right) \mathcal{E}_{a_1, a_2, \dots, a_{k-1}} \\ &= \sum_{a_1, \dots, a_k} \frac{1}{2} \mathcal{E}_{a_1, a_2, \dots, a_{k-1}} = \sum_{a_1, \dots, a_{k-1}} \mathcal{E}_{a_1, a_2, \dots, a_{k-1}} = \hat{\mathcal{E}}_{k-1}, \end{aligned}$$

où le facteur $\frac{1}{2}$ a été compensé par les 2 choix possibles pour $a_k \in \{1, 2\}$.

La martingale $\{\hat{\mathcal{E}}_k\}_{k \geq 1}$ étant positive, elle converge presque sûrement d'après le corollaire 11.6. Nous allons maintenant prouver la convergence dans $\mathbb{L}^2$ de l'énergie totale en évaluant $\mathbb{E}(\hat{\mathcal{E}}_k^2)$ par récurrence. La stratégie consiste à décomposer un arbre de hauteur k en 2 arbres indépendants reliés par la racine comme dans la figure 12.3. L'énergie totale du sous-arbre de gauche $\hat{\mathcal{E}}_{k-1}^{(1)}$ ne dépend que des variables W avec l'indice $a_1 = 1$

$$\hat{\mathcal{E}}_{k-1}^{(1)} = \sum_{a_2, \dots, a_k} \frac{1}{2^{k-1}} W_{1, a_2} W_{1, a_2, a_3} \dots W_{1, a_2, \dots, a_k}.$$

Comme la pondération par $\frac{1}{2} W_1$ ne participe pas à l'énergie du sous-arbre, les variables $\hat{\mathcal{E}}_{k-1}^{(1)}$ et W_1 sont indépendantes. L'énergie $\hat{\mathcal{E}}_{k-1}^{(2)}$ se définit de façon analogue en se restreignant aux indices $a_1 = 2$. L'énergie totale $\hat{\mathcal{E}}_k$ peut donc se réécrire en fonction de l'énergie des sous-arbres

$$\hat{\mathcal{E}}_k = \frac{W_1}{2} \hat{\mathcal{E}}_{k-1}^{(1)} + \frac{W_2}{2} \hat{\mathcal{E}}_{k-1}^{(2)}. \quad (12.4)$$

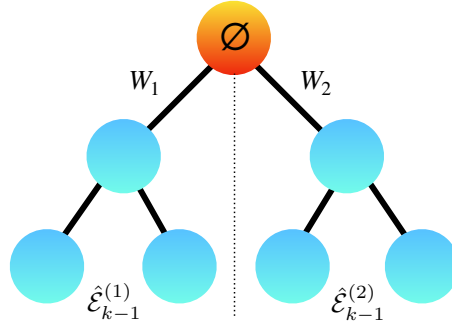


FIGURE 12.3 – Un arbre de hauteur k contient 2 arbres disjoints de hauteur $k - 1$ représentés en bleu. Ce schéma illustre l'identité (12.4) où $\hat{\mathcal{E}}_{k-1}^{(1)}$ correspond à l'énergie totale du sous-arbre de gauche et $\hat{\mathcal{E}}_{k-1}^{(2)}$ celle de celui de droite. Les énergies $\hat{\mathcal{E}}_{k-1}^{(1)}$ et $\hat{\mathcal{E}}_{k-1}^{(2)}$ sont indépendantes.

Les énergies $\hat{\mathcal{E}}_{k-1}^{(1)}$ et $\hat{\mathcal{E}}_{k-1}^{(2)}$ sont associées à des arbres disjoints et elles sont donc indépendantes. L'identité (12.4) et l'indépendance des variables $W_1, W_2, \hat{\mathcal{E}}_{k-1}^{(1)}, \hat{\mathcal{E}}_{k-1}^{(2)}$ permettent d'écrire

$$\begin{aligned} \mathbb{E}(\hat{\mathcal{E}}_k^2) &= \frac{1}{2} \mathbb{E} \left(W_1^2 (\hat{\mathcal{E}}_{k-1}^{(1)})^2 \right) + \frac{1}{2} \mathbb{E} \left(W_1 W_2 \hat{\mathcal{E}}_{k-1}^{(1)} \hat{\mathcal{E}}_{k-1}^{(2)} \right) \\ &= \frac{1}{2} \mathbb{E} (W_1^2) \mathbb{E} \left((\hat{\mathcal{E}}_{k-1}^{(1)})^2 \right) + \frac{1}{2} \mathbb{E} (W_1) \mathbb{E} (W_2) \mathbb{E} \left(\hat{\mathcal{E}}_{k-1}^{(1)} \right) \mathbb{E} \left(\hat{\mathcal{E}}_{k-1}^{(2)} \right). \end{aligned}$$

Il suffit d'utiliser que $\mathbb{E}(W_1) = 1$ et le résultat $\mathbb{E}(\hat{\mathcal{E}}_{k-1}^{(1)}) = \mathbb{E}(\hat{\mathcal{E}}_{k-1}^{(2)}) = 1$ prouvé en (12.2) pour obtenir une récurrence

$$k \geq 1, \quad \mathbb{E}(\hat{\mathcal{E}}_k^2) = \frac{\mathbb{E}(W_1^2)}{2} \mathbb{E} \left((\hat{\mathcal{E}}_{k-1}^{(1)})^2 \right) + \frac{1}{2} = \frac{\mathbb{E}(W_1^2)}{2} \mathbb{E}(\hat{\mathcal{E}}_{k-1}^2) + \frac{1}{2}.$$

Pour la dernière égalité, il suffit de remarquer que $\hat{\mathcal{E}}_{k-1}^{(1)}$ a la même loi que $\hat{\mathcal{E}}_{k-1}$. La suite $\{\mathbb{E}(\hat{\mathcal{E}}_k^2)\}_{k \geq 0}$ converge dès que $\mathbb{E}(W_1^2) < 2$. Elle est donc bornée dans $\mathbb{L}^2$

$$\sup_k \mathbb{E}(\hat{\mathcal{E}}_k^2) < \infty$$

et il ne reste plus qu'à appliquer le théorème 11.1 pour en déduire la convergence de la martingale $\{\hat{\mathcal{E}}_k\}_{k \geq 1}$ dans $\mathbb{L}^2$ (et presque sûrement) vers une variable aléatoire $\hat{\mathcal{E}}$ de moyenne égale à 1.

Supposons maintenant que $\mathbb{E}(W_1^2) < 2$ et $\mathbb{P}(W_1 > 0) = 1$, alors $\{\hat{\mathcal{E}}_k\}_{k \geq 1}$ converge presque sûrement et on peut passer à la limite dans l'identité (12.4) pour obtenir

$$\hat{\mathcal{E}} = \frac{W_1}{2} \hat{\mathcal{E}}^{(1)} + \frac{W_2}{2} \hat{\mathcal{E}}^{(2)},$$

où $\hat{\mathcal{E}}^{(1)}$ et $\hat{\mathcal{E}}^{(2)}$ sont deux variables indépendantes, de même loi que $\hat{\mathcal{E}}$, obtenues comme

la limite des variables $\hat{\mathcal{E}}_{k-1}^{(1)}$ et $\hat{\mathcal{E}}_{k-1}^{(2)}$. On en déduit que

$$\begin{aligned}\mathbb{P}(\hat{\mathcal{E}} = 0) &= \mathbb{P}\left(\{W_1 \hat{\mathcal{E}}^{(1)} = 0\} \cap \{W_2 \hat{\mathcal{E}}^{(2)} = 0\}\right) = \mathbb{P}\left(W_1 \hat{\mathcal{E}}^{(1)} = 0\right)^2 \\ &= \left(1 - \mathbb{P}\left(W_1 \hat{\mathcal{E}}^{(1)} \neq 0\right)\right)^2 = \left(1 - \mathbb{P}(W_1 \neq 0) \mathbb{P}\left(\hat{\mathcal{E}}^{(1)} \neq 0\right)\right)^2,\end{aligned}$$

en utilisant l'indépendance entre W_1 et $\hat{\mathcal{E}}^{(1)}$ (cf. figure 12.3) dans la dernière égalité. Comme $\hat{\mathcal{E}}$ et $\hat{\mathcal{E}}^{(1)}$ ont la même loi, l'égalité précédente se réécrit

$$\mathbb{P}(\hat{\mathcal{E}} = 0) = \left(1 - \mathbb{P}(W_1 > 0) \mathbb{P}(\hat{\mathcal{E}} \neq 0)\right)^2.$$

Sous l'hypothèse $\mathbb{P}(W_1 > 0) = 1$, la probabilité que l'énergie totale s'annule doit satisfaire l'équation suivante

$$\mathbb{P}(\hat{\mathcal{E}} = 0) = \mathbb{P}(\hat{\mathcal{E}} \neq 0)^2 \quad \text{et donc} \quad \mathbb{P}(\hat{\mathcal{E}} = 0) \in \{0, 1\}.$$

La variable aléatoire $\hat{\mathcal{E}}$ est positive, de moyenne $\mathbb{E}(\hat{\mathcal{E}}) = 1$. Par conséquent, la seule solution possible de l'équation ci-dessus est $\mathbb{P}(\hat{\mathcal{E}} = 0) = 0$. Ceci conclut le théorème 12.1. $\square$

Le théorème 12.1 démontre la convergence de l'énergie totale dans $\mathbb{L}^2$ sous la condition $\mathbb{E}(W_1^2) < 2$. Une étude plus précise des corrélations de l'énergie et de sa répartition dans l'arbre permettrait de découvrir une structure fractale extrêmement riche qui dépasse le cadre de ce cours. Les cascades de Mandelbrot ont joué un rôle important pour modéliser des phénomènes variés et elles ont permis de poser les bases de l'analyse multifractale et de la théorie du chaos multiplicatif. On pourra consulter l'article de Y. Heurteaux [16] pour une introduction à ces deux domaines des mathématiques.

12.2 Mécanismes de renforcement

Les mécanismes de renforcement permettent de comprendre comment des comportements individuels aléatoires peuvent générer des effets collectifs dans des contextes variés. Le modèle de Wright-Fisher défini section 11.4 pour décrire la dérive génétique peut être aussi interprété comme une forme de renforcement car l'évolution aléatoire finit par sélectionner un des allèles qui devient dominant. Nous allons illustrer d'autres phénomènes de renforcement par quelques modèles probabilistes simples.

12.2.1 Urne de Pólya

À la suite d'une avancée technologique, un nouveau marché se développe avec différents produits concurrents. Assez rapidement, on observe souvent l'émergence de standards : certains produits sont systématiquement délaissés par les consommateurs et d'autres deviennent la norme. Ces choix ne sont pas toujours dictés par la supériorité technologique d'un produit sur un autre. Le consommateur suit l'avis de ses proches et les tendances du moment. Par exemple lors de l'achat d'un nouveau téléphone, un consommateur va privilégier la marque qui propose le plus grand nombre d'applications, parallèlement, l'offre des applications va s'amplifier si le marché se développe. Les économistes

parlent d'*externalité de réseau* pour décrire le fait que plusieurs individus en achetant le même produit augmentent la valeur de ce produit en l'érigeant en standard.

L'urne de Pólya permet d'appréhender le phénomène précédent en le modélisant par un choix aléatoire. Initialement, on considère une urne avec r boules rouges et v boules vertes. On choisit une boule au hasard et on la replace dans l'urne en ajoutant en plus une autre boule de la même couleur. Ceci revient à dire que le consommateur acquiert le produit rouge ou le produit vert en fonction de la proportion des produits déjà vendus.

La proportion X_n des boules vertes après le $n^{\text{ième}}$ -tirage est une martingale. En effet, si il existe i boules rouges et j boules vertes au temps n alors

$$\mathbb{P}\left(X_{n+1} = \frac{j+1}{i+j+1}\right) = \frac{j}{i+j} \quad \text{et} \quad \mathbb{P}\left(X_{n+1} = \frac{j}{i+j+1}\right) = \frac{i}{i+j}.$$

On a donc $\mathbb{E}(X_{n+1}|\mathcal{F}_n) = X_n$ où $\mathcal{F}_n$ représente la σ -algèbre des choix jusqu'à l'instant n . La martingale $\{X_n\}_{n \geq 1}$ étant positive elle converge presque sûrement, par le corollaire 11.6, vers une variable aléatoire limite X_∞ que nous allons déterminer maintenant. On remarquera aussi que $\{X_n\}_{n \geq 1}$ est bornée dans tout $\mathbb{L}^p$ et la convergence a donc aussi lieu dans $\mathbb{L}^p$ avec $p \geq 1$ (cf. le théorème 11.10).

La probabilité que les m premiers tirages soient des boules vertes et les $\ell = n - m$ suivants soient des boules rouges vaut

$$\frac{v}{v+r} \cdot \frac{v+1}{v+r+1} \cdots \frac{v+(m-1)}{v+r+(m-1)} \cdot \frac{r}{v+r+m} \cdots \frac{r+(\ell-1)}{v+r+(n-1)}.$$

Tout autre tirage au sort de m boules vertes et ℓ boules rouges aura la même probabilité car le dénominateur sera inchangé et les coefficients du numérateur seront simplement permutés. Considérons le cas particulier $r = 1$ et $v = 1$, alors la probabilité qu'il y ait $m+1$ boules vertes au temps n s'écrit

$$\mathbb{P}\left(X_n = \frac{m+1}{n+2}\right) = \binom{n}{m} \frac{m!(n-m)!}{(n+1)!} = \frac{1}{n+1}.$$

Le nombre de boules vertes est donc uniformément distribué à tout temps et X_∞ aura la loi uniforme sur $[0, 1]$. La proportion finale X_∞ des boules vertes est aléatoire et est déterminée principalement par l'aléa des choix initiaux. On observe un comportement asymptotique stable et le nombre de boules vertes croît linéairement comme nX_∞ quand n devient grand : la proportion X_n reste figée quand n est grand (cf. figure 12.4).

Pour des données initiales générales, X_∞ suit une loi beta sur $[0, 1]$ de paramètres (v, r) dont la densité s'écrit

$$f_{X_\infty}(x) = \frac{(v+r-1)!}{(v-1)!(r-1)!} (1-x)^{r-1} x^{v-1}. \quad (12.5)$$

12.2.2 Graphes aléatoires de Barabási-Albert

Les processus de renforcement sont multiples et ils façonnent de nombreux aspects de la vie courante. Par exemple, les réseaux sociaux ou le réseau internet peuvent être

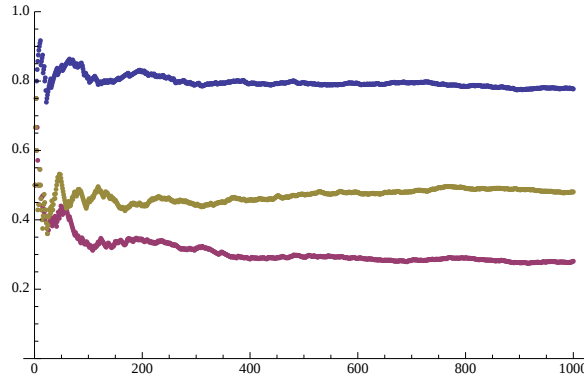


FIGURE 12.4 – La proportion X_n de boules vertes est représentée dans trois réalisations de l'urne de Pólya avec 1000 tirages au sort. Les fluctuations initiales sont importantes, mais la proportion se stabilise très vite et reste ensuite asymptotiquement constante.

interprétés comme des graphes aléatoires dont les structures ont de nombreuses similarités. Les interconnexions dans ces graphes sont très différentes du modèle d'Erdős-Rényi évoqué section 4.4.2. En particulier, il existe des sites avec un très grand nombre de connections et la statistique des degrés de ces graphes est souvent régie par des lois de puissance. Ces graphes se sont constitués au fil du temps sans suivre un dessein préétabli. De nombreux modèles ont été proposés pour essayer de décrire cette "auto-organisation" et les lois correspondantes. Barabási et Albert ont proposé dans leur article [1] un mécanisme de renforcement, que nous allons présenter ci-dessous, pour construire dynamiquement des graphes dont les degrés ont des propriétés statistiques similaires à celles observées en pratique.

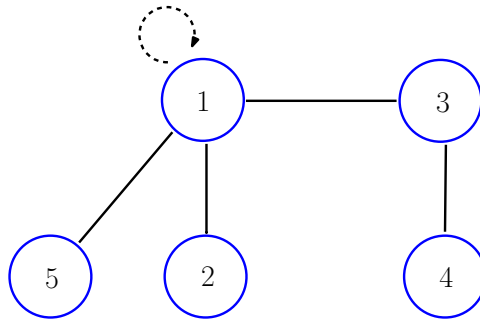


FIGURE 12.5 – Un exemple de graphe de Barabási-Albert avec 5 sites. La convention consistant à relier le site 1 à lui même lui confère le degré 1 initialement. Les autres sites ont été ajoutés aléatoirement selon l'ordre indiqué sur le dessin.

Au temps initial $n = 1$, le graphe $\mathcal{G}_1$ est constitué par un unique site relié à lui même. À chaque pas de temps, un site est ajouté et on notera $\mathcal{G}_n$ le graphe correspondant. La règle est la suivante : au temps n , le nouveau site n est connecté à un site dans le graphe $\mathcal{G}_{n-1}$ choisi proportionnellement à son degré (cf. figure 12.5). Ce graphe est construit dynamiquement par analogie au *world wide web* où les sites les plus importants ont tendance à attirer le plus grand nombre de liens, les nouveaux sites se connectant de façon privi-

légée aux serveurs principaux. Deux exemples de graphes aléatoires de Barabási-Albert sont représentés figure 12.6.

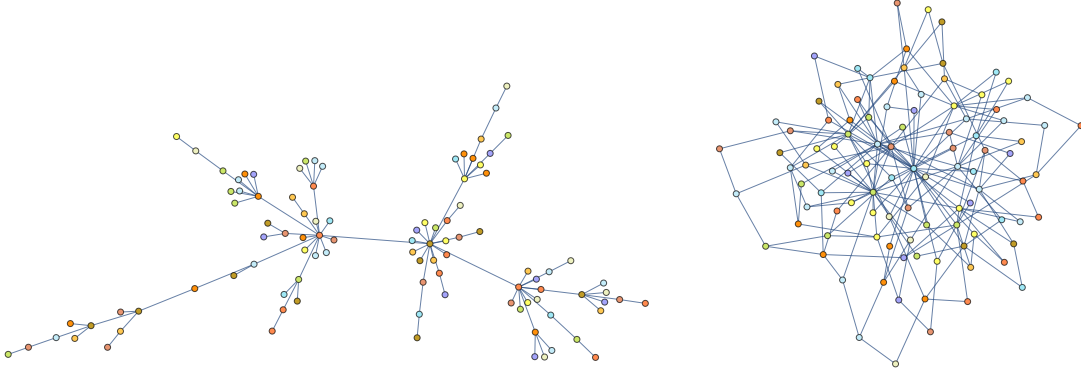


FIGURE 12.6 – Deux graphes de Barabási-Albert avec 100 sites sont représentés ci-dessus. Celui de droite est construit en ajoutant un site à chaque étape et celui de gauche deux sites. On remarquera que certains sites privilégiés sont fortement connectés.

La construction du graphe peut s'interpréter à l'aide d'une urne de Pólya. Initialement en $n = 1$, on considère une urne contenant 1 boule avec le label 1. Supposons qu'au temps n , il existe $2n - 1$ boules dans l'urne chacune étant associée à un label entre 1 et n . Au temps $n + 1$, on choisit une boule au hasard et on note $k \in \{1, \dots, n\}$ son label. On replace alors dans l'urne deux boules de label k et une nouvelle de label $n + 1$. Retraduit en terme de graphe, ceci correspond à ajouter le site $n + 1$ et à le relier au site k par une arête. Le nombre de boules avec le label k croît exactement comme le degré du site k , i.e. le nombre d'arêtes du site k .

Pour comprendre l'évolution du nombre d'arêtes du site $k > 1$, on considère l'urne au temps k et on colorie les $2k - 2$ boules de label strictement inférieur à k en rouge et la boule k en vert. On continue les tirages au sort. Au temps $n > k$, si on choisit une boule de label strictement supérieur à k on ignore ce tirage car il correspond à la création d'une arête sur un site qui n'est pas dans $\{1, \dots, k\}$. Si une boule de label $j < k$ est tirée, cela revient à ajouter une boule rouge et si une boule de label k est choisie cela correspond à l'ajout d'une boule verte. Par conséquent, pour $k > 1$, la distribution relative des boules vertes et rouges suit celle d'une urne de Pólya de donnée initiale $r = 2k - 2$ et $v = 1$ et le comportement asymptotique est donné par la loi (12.5) qui vaut $f_{X_\infty}(x) = r(1 - x)^{r-1}$.

Au lieu de suivre le degré d'un site fixé, on peut aussi chercher une information plus globale et déterminer l'espérance $\mathcal{N}(d, n)$ du nombre de sites de degré d au temps n . Comme pour les chaînes de Markov, on obtient en conditionnant par rapport au passé

$$\mathcal{N}(d, n+1) = \begin{cases} \left(1 - \frac{d}{2n-1}\right) \mathcal{N}(d, n) + \frac{d-1}{2n-1} \mathcal{N}(d-1, n), & \text{si } d > 1 \\ \left(1 - \frac{d}{2n-1}\right) \mathcal{N}(d, n) + 1, & \text{si } d = 1 \end{cases}$$

En effet, un nouveau site de degré $d > 1$ ne peut être créé que si une arête a été ajoutée à un site de degré $d - 1$ et inversement un site de degré d disparaît si on lui ajoute une

arête. Un calcul (simple mais douloureux) permet de montrer que pour $d \geq 1$

$$\lim_{n \rightarrow \infty} \frac{1}{n} \mathcal{N}(d, n) = \frac{4}{d(d+1)(d+2)}.$$

Pour de très grands graphes, la probabilité qu'un site choisi au hasard ait le degré d décroît comme $1/d^3$ quand d diverge. La répartition des degrés selon une loi de puissance se retrouve dans de nombreux réseaux dont la structure n'est pas dictée par le degré moyen mais est caractérisée par quelques nœuds fortement connectés. La construction de Barabási-Albert conduit à une structure de graphes très stable, insensible aux variations aléatoires de la construction. Par contre ces graphes sont très vulnérables aux attaques des sites fortement connectés. Une étude approfondie des graphes aléatoires pourra être trouvée dans le livre de R. Durrett [11].

12.3 Descente de gradient stochastique

12.3.1 Algorithme du gradient stochastique

Dans cette section, nous allons décrire une méthode, différente de celle présentée au chapitre 6, pour déterminer les valeurs où une fonction $V : \mathbb{R}^d \rightarrow \mathbb{R}$ atteint son minimum

$$\text{Argmin } V = \{\theta^* \in \mathbb{R}^d; \quad V(\theta^*) = \inf_{\theta \in \mathbb{R}^d} V(\theta)\}.$$

Si la fonction V est dérivable, ce problème revient à identifier les valeurs θ^* telles que $\nabla V(\theta^*) = 0$. On rappelle que le gradient de V correspond au vecteur de $\mathbb{R}^d$

$$\nabla V(\theta) = (\partial_1 V(\theta), \dots, \partial_d V(\theta))$$

obtenu en dérivant selon chacune des coordonnées de $\theta \in \mathbb{R}^d$.

Dans de nombreuses applications, l'expression de la fonction V n'est pas explicite, mais on suppose qu'elle s'écrit comme une espérance $V(\theta) = \mathbb{E}(W(\theta, X))$ par rapport à une variable aléatoire X dont la loi est fixée et ne dépend pas de θ . On suppose de plus que le gradient de V est donné par

$$\nabla V(\theta) = \mathbb{E}(\nabla W(\theta, X)) = \left(\mathbb{E}(\partial_1 W(\theta, X)), \dots, \mathbb{E}(\partial_d W(\theta, X)) \right) \in \mathbb{R}^d, \quad (12.6)$$

où les dérivées portent uniquement sur la variable θ . Cette représentation de la fonction V se justifie pour des raisons pratiques, car souvent, seules des observations aléatoires du type $\{W(\theta, X_i)\}_i$ sont disponibles et la fonction $V(\theta) = \mathbb{E}(W(\theta, X))$ décrit leur comportement moyen. Dans la section 12.3.2, nous verrons que cette représentation sert pour l'apprentissage des réseaux de neurones artificiels car l'ampleur des données à traiter est trop importante pour calculer numériquement ∇V . Inversement cette représentation de V s'avère aussi très utile quand le nombre des données est faible. Par exemple, si on cherche à calibrer le dosage d'un médicament en fonction d'un paramètre θ pour produire un effet optimal quantifié par la fonction V , on ne pourra mesurer que les effets $\{W(\theta, X_1), \dots, W(\theta, X_n)\}$ du médicament sur les n patients impliqués dans une phase de test. L'effet moyen $V(\theta)$ ne sera pas directement accessible et il faudrait de nombreux patients pour l'évaluer avec précision. On cherche donc à déterminer θ^* en testant

seulement un petit nombre de patients, sans avoir à connaître V . L'algorithme du gradient stochastique, que nous allons décrire, permet de réaliser cette tâche et d'estimer θ^* récursivement en ajustant progressivement le paramètre θ au fur et à mesure des observations.

Dans la suite, nous supposons que la fonction V admet un unique minimum θ^* ($\nabla V(\theta^*) = 0$) et qu'elle satisfait l'inégalité

$$\forall \theta \in \mathbb{R}^d, \quad \langle \theta - \theta^*, \nabla V(\theta) \rangle > 0, \quad (12.7)$$

avec la notation $\langle \cdot, \cdot \rangle$ pour représenter le produit scalaire et nous utiliserons aussi $\| \cdot \|$ pour la distance euclidienne dans $\mathbb{R}^d$

$$\forall x = (x_1, \dots, x_d), y = (y_1, \dots, y_d), \quad \langle x, y \rangle = \sum_{i=1}^d x_i y_i, \quad \|x\| = \sqrt{\sum_{i=1}^d x_i^2}.$$

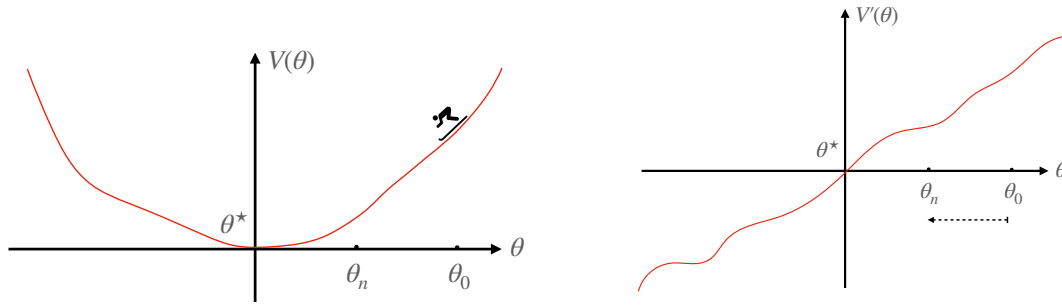


FIGURE 12.7 – Ci-dessus un exemple de potentiel V satisfaisant la condition (12.7) avec un unique minimum en θ^* . La méthode de descente de gradient consiste à suivre la pente pour trouver le minimum.

Dans certaines applications, le paramètre θ peut prendre ses valeurs dans un espace de dimension d très grande. En particulier, dans le cas des réseaux de neurones artificiels, il est très courant de considérer des dimensions $d \gg 10^6$. Commençons cependant par décrire le cas de la dimension $d = 1$ pour fixer les idées. Étant donnée une fonction $V : \mathbb{R} \rightarrow \mathbb{R}$ (cf. figure 12.7), la condition (12.7) signifie simplement

$$\begin{cases} V'(\theta) > 0 & \text{si } \theta > \theta^*, \\ V'(\theta) < 0 & \text{si } \theta < \theta^*. \end{cases}$$

Par ailleurs, la condition (12.6) s'écrit

$$V'(\theta) = \mathbb{E}(W'(\theta, X)) \quad (12.8)$$

où $W'(\theta, X)$ est la dérivée par rapport à la variable θ .

Pour déterminer la valeur θ^* , la procédure la plus simple consiste à choisir une donnée initiale θ_0 et à suivre le flot de l'équation

$$t \geq 0, \quad \partial_t \theta_t = -V'(\theta_t).$$

Sous les hypothèses considérées, la fonction $t \mapsto \theta_t$ converge vers θ^* . Pour l'implémentation numérique, il est préférable d'utiliser des pas de temps discrets

$$\theta_{n+1} = \theta_n - \gamma_n V'(\theta_n) \quad (12.9)$$

où les incréments $\gamma_n > 0$ sont choisis tels que

$$\lim_{n \rightarrow \infty} \gamma_n = 0 \quad \text{et} \quad \sum_n \gamma_n = \infty. \quad (12.10)$$

La dérivée $\partial_t \theta_t$ s'interprète alors comme $(\theta_{n+1} - \theta_n) / \gamma_n$. Ces conditions suffisent pour montrer que θ_n converge vers θ^* quand n tend vers l'infini. Pour le comprendre, considérons le cas particulier $V(\theta) = \alpha \theta^2$. La récurrence (12.9) s'écrit

$$\theta_{n+1} = (1 - 2\alpha\gamma_n)\theta_n = \beta_n \theta_0 \quad \text{avec} \quad \beta_n = \prod_{k=0}^n (1 - 2\alpha\gamma_k). \quad (12.11)$$

La condition (12.7) impose $\alpha > 0$ ce qui est une condition nécessaire pour que le produit β_n converge vers $\theta^* = 0$ quand n tend vers l'infini, en effet

$$\log \beta_n = \log \beta_K + \sum_{k=K}^n \log(1 - 2\alpha\gamma_k) \leq \log \beta_K - 2\alpha \sum_{k=K}^n \gamma_k \xrightarrow{n \rightarrow \infty} -\infty, \quad (12.12)$$

où K est choisi, grâce à (12.10), afin que $1 > 2\alpha\gamma_k$ pour tout $k \geq K$. La divergence de la série $\sum_n \gamma_n$ permet donc de déduire que θ_n converge vers $\theta^* = 0$ quelle que soit la donnée initiale θ_0 . De la même façon, en dimension d , la descente de gradient s'écrit

$$\theta_{n+1} = \theta_n - \gamma_n \nabla V(\theta_n) \in \mathbb{R}^d \quad (12.13)$$

et l'hypothèse (12.7) garantit la convergence vers θ^* .

Si le gradient de V est de la forme $\nabla V(\theta) = \mathbb{E}(\nabla W(\theta, X))$, le théorème suivant généralise la procédure déterministe de descente de gradient (12.13) et permet de trouver le point critique θ^* en utilisant uniquement des observations du type $\{\nabla W(\theta, X_i)\}_i$ sans avoir à calculer V explicitement.

Théorème 12.2 (Algorithme du gradient stochastique).

On considère une fonction $V : \mathbb{R}^d \mapsto \mathbb{R}$ satisfaisant l'hypothèse (12.7) dont le gradient est continu et de la forme $\nabla V(\theta) = \mathbb{E}(\nabla W(\theta, X))$ comme dans (12.6). On suppose, de plus, qu'il existe une constante K telle que

$$\forall \theta \in \mathbb{R}^d, \quad \mathbb{E}(\|\nabla W(\theta, X)\|^2) \leq K(1 + \|\theta - \theta^*\|^2). \quad (12.14)$$

Soit $\{X_n\}_{n \geq 1}$ une suite de variables indépendantes, distribuées selon la même loi que X . L'algorithme du gradient stochastique consiste à construire une récurrence aléatoire

$$\forall n \geq 0, \quad \theta_{n+1} = \theta_n - \gamma_n \nabla W(\theta_n, X_{n+1}) \quad (12.15)$$

partant initialement de la valeur θ_0 et dépendant de la suite $\{\gamma_n\}_{n \geq 0}$ d'incrément positifs vérifiant

$$\sum_n \gamma_n = \infty \quad \text{et} \quad \sum_n \gamma_n^2 < \infty. \quad (12.16)$$

Alors le processus $\{\theta_n\}_{n \geq 0}$ converge presque sûrement, quelle que soit la donnée initiale θ_0 , vers θ^* la solution de $\nabla V(\theta^*) = \mathbb{E}(\nabla W(\theta^*, X)) = 0$.

Pour comprendre intuitivement cet algorithme, on peut reformuler la récurrence (12.15) et l'interpréter comme une perturbation de la descente de gradient déterministe (12.13)

$$\theta_{n+1} = \theta_n - \gamma_n \nabla W(\theta_n, X_{n+1}) = \theta_n - \gamma_n \nabla V(\theta_n) - \gamma_n \varepsilon_n$$

par une variable aléatoire $\varepsilon_n = \nabla W(\theta_n, X_{n+1}) - \nabla V(\theta_n)$ dont la moyenne est nulle grâce à (12.6). À chaque étape, une petite erreur $\gamma_n \varepsilon_n$ est commise mais sous les hypothèses du théorème, nous allons montrer que ces erreurs successives se moyennent à 0 et que le processus $\{\theta_n\}_{n \geq 0}$ converge.

Démonstration. Pour quantifier la convergence de la suite $\{\theta_n\}_{n \geq 0}$, définissons le processus

$$Z_n = \|\theta_n - \theta^*\|^2 \in \mathbb{R}^+ \quad (12.17)$$

qui mesure la distance euclidienne entre θ_n et θ^* . En appliquant la récurrence (12.15) et en développant le produit scalaire, on en déduit

$$\begin{aligned} Z_{n+1} &= \|\theta_n - \gamma_n \nabla W(\theta_n, X_{n+1}) - \theta^*\|^2 \\ &= Z_n - 2\gamma_n \langle \theta_n - \theta^*, \nabla W(\theta_n, X_{n+1}) \rangle + \gamma_n^2 \|\nabla W(\theta_n, X_{n+1})\|^2. \end{aligned}$$

Le processus $\{\theta_n\}_{n \geq 0}$, défini par une récurrence aléatoire, est mesurable par rapport à la filtration $\mathcal{F}_n = \sigma(X_1, \dots, X_n)$. Comme θ_n, Z_n sont mesurables par rapport à $\mathcal{F}_n$, l'espérance conditionnelle s'écrit

$$\mathbb{E}(Z_{n+1} | \mathcal{F}_n) = Z_n - 2\gamma_n \langle \theta_n - \theta^*, \mathbb{E}(\nabla W(\theta_n, X_{n+1}) | \mathcal{F}_n) \rangle + \gamma_n^2 \mathbb{E}(\|\nabla W(\theta_n, X_{n+1})\|^2 | \mathcal{F}_n).$$

Comme X_{n+1} est indépendant de $\mathcal{F}_n$, l'espérance conditionnelle se réduit à un simple calcul d'espérance

$$\begin{aligned} \mathbb{E}(\nabla W(\theta_n, X_{n+1}) | \mathcal{F}_n) &= \mathbb{E}(\nabla W(\theta_n, X)) = \nabla V(\theta_n), \\ \mathbb{E}(\|\nabla W(\theta_n, X_{n+1})\|^2 | \mathcal{F}_n) &= \mathbb{E}(\|\nabla W(\theta_n, X)\|^2) \leq K(1 + Z_n), \end{aligned}$$

où la dernière inégalité vient de l'hypothèse (12.14). En appliquant les résultats ci-dessus, on en déduit

$$\begin{aligned} \mathbb{E}(Z_{n+1} | \mathcal{F}_n) &\leq Z_n - 2\gamma_n \langle \theta_n - \theta^*, \nabla V(\theta_n) \rangle + \gamma_n^2 K(1 + Z_n) \\ &\leq (1 + \gamma_n^2 K)Z_n + \gamma_n^2 K - 2\gamma_n \langle \theta_n - \theta^*, \nabla V(\theta_n) \rangle. \end{aligned} \quad (12.18)$$

En utilisant la notation $\beta'_n = \prod_{\ell=0}^{n-1} (1 + \gamma_\ell^2 K)$ ainsi que la mesurabilité de θ_n (et donc de Z_n) par rapport à $\mathcal{F}_n$, cette inégalité se réécrit

$$\mathbb{E}\left(\frac{Z_{n+1}}{\beta'_{n+1}} + \left(2\frac{\gamma_n}{\beta'_{n+1}} \langle \theta_n - \theta^*, \nabla V(\theta_n) \rangle - \frac{\gamma_n^2 K}{\beta'_{n+1}}\right) \middle| \mathcal{F}_n\right) \leq \frac{Z_n}{\beta'_n}.$$

On en déduit que le processus

$$\forall n \geq 1, \quad U_n = \frac{1}{\beta'_n} Z_n + \sum_{\ell=1}^{n-1} \left(2\frac{\gamma_\ell}{\beta'_{\ell+1}} \langle \theta_\ell - \theta^*, \nabla V(\theta_\ell) \rangle - \frac{\gamma_\ell^2 K}{\beta'_{\ell+1}}\right) \quad (12.19)$$

est une surmartingale car il vérifie

$$\mathbb{E}(U_{n+1} | \mathcal{F}_n) \leq U_n.$$

L'hypothèse $\sum_n \gamma_n^2 < \infty$ (12.16) et le même raisonnement que dans (12.12) permettent d'affirmer que

$$\lim_{n \rightarrow \infty} \beta'_n = \beta > 0 \quad \text{et donc} \quad \sum_{\ell=1}^{\infty} \frac{\gamma_{\ell}^2 K}{\beta'_{\ell+1}} < \infty. \quad (12.20)$$

Comme $\inf_{\theta} \langle \theta - \theta^*, \nabla V(\theta) \rangle \geq 0$, la surmartingale $\{U_n\}_{n \geq 1}$ est bornée inférieurement

$$U_n \geq - \sum_{\ell=1}^{\infty} \frac{\gamma_{\ell}^2 K}{\beta'_{\ell+1}}$$

et elle converge, par le corollaire 11.6, presque sûrement vers une limite U_{∞} qui est finie presque sûrement.

En utilisant (12.19), on définit la suite

$$S_n := \sum_{\ell=1}^{n-1} \frac{2\gamma_{\ell}}{\beta'_{\ell+1}} \langle \theta_{\ell} - \theta^*, \nabla V(\theta_{\ell}) \rangle = U_n - \frac{1}{\beta'_n} Z_n + \sum_{\ell=1}^{n-1} \frac{\gamma_{\ell}^2 K}{\beta'_{\ell+1}} \quad (12.21)$$

qui est bornée supérieurement par

$$S_n \leq U_n + \sum_{\ell=1}^{\infty} \frac{\gamma_{\ell}^2 K}{\beta'_{\ell+1}}$$

car $Z_n \geq 0$. Le résultat précédent sur la convergence $U_n(\omega) \xrightarrow[n \rightarrow \infty]{p.s.} U_{\infty}(\omega) < \infty$ implique que la suite $\{S_n\}_{n \geq 1}$ est majorée presque sûrement. Comme cette suite est croissante, elle converge presque sûrement. Finalement, la suite $\{Z_n\}_{n \geq 1}$ se réécrit par (12.19)

$$\|\theta_n - \theta^*\|^2 = Z_n = \beta'_n \left(U_n - S_n + \sum_{\ell=1}^{n-1} \frac{\gamma_{\ell}^2 K}{\beta'_{\ell+1}} \right), \quad (12.22)$$

en fonction de suites convergentes et elle converge donc presque sûrement vers une limite $Z_{\infty} < \infty$.

Pour montrer que $Z_{\infty} = 0$ et ainsi en déduire la convergence de θ_n vers θ^* , nous allons raisonner par l'absurde. Si la limite Z_{∞} n'est pas nulle, alors il existe $c > 0$ et un entier L tels que pour tout $\ell \geq L$

$$\frac{Z_{\infty}}{2} \leq \|\theta_{\ell} - \theta^*\|^2 \leq 2Z_{\infty} \quad \text{et donc} \quad \langle \theta_{\ell} - \theta^*, \nabla V(\theta_{\ell}) \rangle \geq c, \quad (12.23)$$

car ∇V est continu et ne s'annule qu'en θ^* par l'hypothèse (12.7). On remarquera que la constante c dépend de Z_{∞} . Comme $\sum_n \gamma_n = \infty$ par l'hypothèse (12.16), cette dernière inégalité contredirait la convergence de la suite $\{S_n\}_{n \geq 1}$. Par conséquent, $Z_{\infty} = 0$ ce qui conclut la preuve du théorème 12.2. $\square$

Remarque 12.3. La preuve du théorème 12.2 peut être simplifiée en renforçant les hypothèses sur V . Supposons qu'il existe une constante $\alpha > 0$ telle que

$$\forall \theta \in \mathbb{R}^d, \quad \langle \theta - \theta^*, \nabla V(\theta) \rangle \geq \alpha \|\theta - \theta^*\|^2. \quad (12.24)$$

Alors l'inégalité (12.18) implique

$$\mathbb{E}(Z_{n+1} | \mathcal{F}_n) \leq (1 - 2\gamma_n \alpha + \gamma_n^2 K) Z_n + \gamma_n^2 K.$$

Ceci impose une relation de récurrence sur l'espérance de Z_n

$$\mathbb{E}(Z_{n+1}) \leq (1 - 2\gamma_n \alpha + \gamma_n^2 K) \mathbb{E}(Z_n) + \gamma_n^2 K. \quad (12.25)$$

On s'est donc ramené à l'étude d'une suite à valeurs réelles. Grâce aux hypothèses (12.16) sur $\{\gamma_n\}_{n \geq 1}$, on peut en déduire (après quelques calculs) la convergence

$$\lim_{n \rightarrow \infty} \mathbb{E}(\|\theta_n - \theta^*\|_2) = 0.$$

Notons que cette méthode plus simple ne suffit pas à retrouver la convergence presque sûre.

12.3.2 Réseaux de neurones artificiels

L'algorithme du gradient stochastique sert dans de nombreux contextes. Il permet d'ajuster des paramètres de façon incrémentale à partir de données reçues au fur et à mesure du temps. Ainsi même si les données sont rares ou partielles, la récurrence aléatoire (12.15) du théorème 12.2 donne une procédure pour estimer le paramètre θ^* . Quand le nombre de données est très important, cet algorithme s'avère aussi très utile pour estimer des paramètres dans un espace de grande dimension. Nous allons illustrer ceci en montrant comment l'algorithme du gradient stochastique permet d'entraîner des réseaux de neurones artificiels afin de reconnaître des images [19].

Étant donnée une image I représentant un chiffre manuscrit, on veut construire une fonction $I \mapsto \Phi(\theta, I) \in \{0, 1, \dots, 9\}$ donnant la valeur de ce chiffre. Cette fonction dépend d'un paramètre $\theta \in \mathbb{R}^d$ (dans un espace de très grande dimension) qui doit être déterminé pour réaliser au mieux cette tâche. La construction d'un réseau de neurones artificiels dépasse le cadre de ce cours, mais nous allons en résumer les grands principes pour expliquer l'origine de la fonction Φ .

Inspiré du fonctionnement d'un neurone biologique, le *perceptron* est la brique élémentaire constituant le réseau. Il reçoit une information $I = \{a_1, \dots, a_k\}$ et restitue une valeur $\varphi(\theta, I)$ proche de 1 s'il est activé et proche de 0 sinon (cf. figure 12.8). La fonction de transfert $\varphi(\theta, I) = S(\sum_{i=1}^k \sigma_i a_i - \sigma_{k+1}) \in [0, 1]$ est simplement une combinaison linéaire des coordonnées de I pondérées par des poids synaptiques $\theta = \{\sigma_1, \dots, \sigma_k\}$ et d'une fonction d'activation S . En calibrant les poids $\theta = \{\sigma_1, \dots, \sigma_k\}$, on peut contrôler la réponse du perceptron. La fonction $I \mapsto \varphi(\theta, I)$ ainsi construite est trop simple pour classifier efficacement des entrées I très générales et il faut associer de nombreuses fonctions de ce type, en les combinant entre elles à travers des couches successives pour aboutir à une fonction $I \mapsto \Phi(\theta, I)$ susceptible d'effectuer une tâche de classification automatique performante. L'architecture des connexions est variée et fait actuellement l'objet de nombreuses recherches pour des applications multiples. Dans l'exemple de la reconnaissance

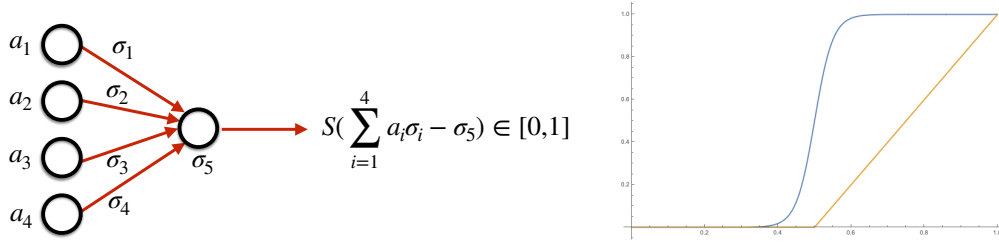


FIGURE 12.8 – Un perceptron est représenté sur la figure de gauche. Il reçoit les données $I = (a_1, \dots, a_4)$ et renvoie un signal $\varphi(\theta, I) = S(\sum_{i=1}^4 \sigma_i a_i - \sigma_5)$. La fonction $I \mapsto \varphi(\theta, I)$ dépend donc des paramètres $\theta = (\sigma_1, \dots, \sigma_5)$ et d'une fonction S qui est activée au-delà d'un certain seuil. Deux types de fonctions d'activation S sont tracés sur la figure de droite.

de chiffres (cf. figure 12.9), l'entrée $I \in [0, 1]^{28 \times 28}$ est une image de 28×28 pixels chacun avec un niveau de gris dans $[0, 1]$. Le signal I est ensuite traité par un réseau de neurones formé de plusieurs couches (de tailles variables) dont chacune prend en entrée le résultat de la précédente. La dernière couche renvoie la fonction $\Phi(\theta, I) \in \{0, 1, \dots, 9\}$ qui donne la valeur du chiffre reconnu.

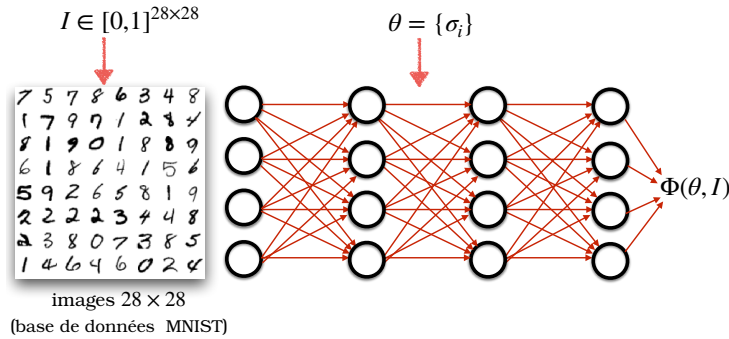


FIGURE 12.9 – L'image I d'un chiffre, sous la forme d'un vecteur représentant les niveaux de gris de 28×28 pixels, est fournie en entrée d'un réseau de neurones multicouches et la réponse est donnée par $\Phi(\theta, I) \in \{0, 1, \dots, 9\}$. Le dessin ci-dessus décrit de façon schématique les neurones répartis dans des couches successives et connectés entre eux par les poids synaptiques $\theta = \{\sigma_j\}_{j \leq d}$. Dans la pratique, la taille et le nombre des couches intermédiaires sont variables.

Pour que le réseau de neurones accomplisse sa tâche de classification, il faut ajuster le paramètre $\theta = \{\sigma_j\}_{j \leq d} \in \mathbb{R}^d$ ce qui est difficile car le nombre d des poids à calibrer est très grand. Cette étape, dite d'apprentissage supervisé, nécessite de disposer d'un immense jeu de données aux caractéristiques connues, par exemple une collection $\{X_j = (I_j, y_j)\}_{j \leq K}$ où chaque image I_j est associée au chiffre correspondant y_j . On cherche alors à ajuster $\theta = \{\sigma_j\} \in \mathbb{R}^d$ pour minimiser la fonction

$$V(\theta) = \frac{1}{K} \sum_{j=1}^K W(\theta, X_j) \quad \text{avec} \quad W(\theta, X_j) = \mathbf{1}_{\{\Phi(\theta, I_j) \neq y_j\}}. \quad (12.26)$$

Dans le cas de la simple reconnaissance de chiffres, le nombre de variables en jeu est déjà

immense. La base de données MNIST² propose $K = 60000$ exemples de chiffres manuscrits, chaque image est représentée par un vecteur de dimension $784 = 28 \times 28$ et un réseau de neurones avec 2 ou 3 couches intermédiaires nécessite d'ajuster plusieurs milliers de poids $\{\sigma_j\}_{j \leq d}$. Dans la pratique, on ne peut donc pas optimiser directement la fonction V , mais l'algorithme de descente de gradient stochastique permet de contourner cette difficulté. En effet, la fonction $V(\theta)$ définie en (12.26) peut s'interpréter comme une espérance et chaque élément X_j de la base d'exemples comme une observation aléatoire. En s'inspirant du théorème 12.2, on construit alors à l'aide de ∇W une récurrence aléatoire $(\theta_n)_{n \geq 1}$. La fonction W étant très compliquée, la fonction V définie en (12.26) ne satisfait pas nécessairement les hypothèses du théorème 12.2, en particulier il peut exister de nombreux minima locaux θ^* . Cependant la structure du réseau de neurones permet facilement de calculer le gradient ∇W et d'implémenter numériquement l'algorithme pour obtenir un résultat très convaincant même avec un réseau de neurones de "petite" taille (cf. figure 12.10).

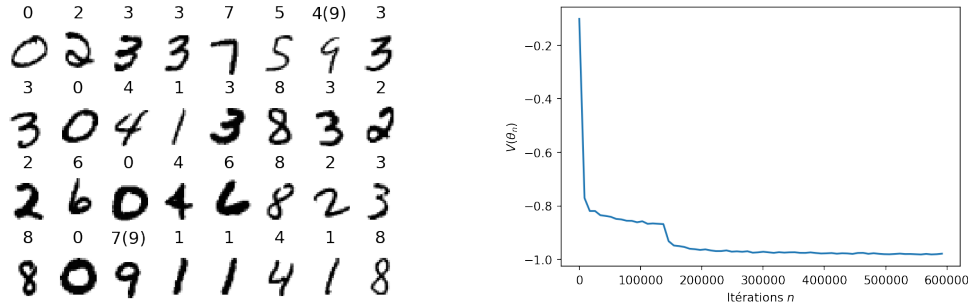


FIGURE 12.10 – Cet exemple, réalisé par C. Mantoux, montre les prédictions (indiquées au-dessus de chaque chiffre manuscrit) d'un réseau de neurones artificiels formé de deux couches cachées avec environ 30000 poids. La courbe de gauche représente $n \mapsto V(\theta_n)$.

12.4 L'algorithme de Robbins-Monro

Nous allons maintenant déterminer les solutions θ d'une équation du type $f(\theta) = a$ où f est une fonction à valeurs dans $\mathbb{R}$ qui s'écrit sous la forme $f(\theta) = \mathbb{E}(F(X, \theta))$. On s'intéresse à un cas où la fonction f ne peut pas être calculée explicitement, mais simplement estimée par des observations de la forme $\{F(X_i, \theta)\}_i$.

Considérons d'abord le cadre déterministe d'une fonction $f : \mathbb{R} \rightarrow \mathbb{R}$ continue qui admet une unique solution θ^* à l'équation $f(\theta^*) = a$ pour un niveau a donné. On suppose que f vérifie une condition analogue à (12.7)

$$\forall \theta \in \mathbb{R}, \quad (f(\theta) - a) \cdot (\theta - \theta^*) < 0. \quad (12.27)$$

Pour déterminer θ^* , il suffit d'utiliser la récurrence

$$\theta_{n+1} = \theta_n + \gamma_n (f(\theta_n) - a) \quad (12.28)$$

2. La base de données MNIST fournit 60000 images de chiffres dans un format identique destinés à entraîner des réseaux de neurones artificiels (<http://yann.lecun.com/exdb/mnist/>).

où les incréments $\gamma_n > 0$ sont choisis tels que

$$\lim_{n \rightarrow \infty} \gamma_n = 0 \quad \text{et} \quad \sum_n \gamma_n = \infty$$

afin d'assurer la convergence de θ_n converge vers θ^* .

Plus généralement, on peut montrer que des petites perturbations n'altèrent pas la convergence de la suite $\{\theta_n\}_{n \geq 0}$.

Lemme 12.4. Soit $f : \mathbb{R} \rightarrow \mathbb{R}$ une fonction continue bornée vérifiant l'hypothèse (12.27). On considère une suite $\{\gamma_n\}_{n \geq 0}$ d'incrémentes positifs et une suite $\{\varepsilon_n\}_{n \geq 0}$ à valeurs dans $\mathbb{R}$ qui modélise une perturbation. On suppose que

$$\lim_{n \rightarrow \infty} \gamma_n = 0, \quad \sum_n \gamma_n = \infty \quad \text{et que la série} \quad \sum_n \gamma_n \varepsilon_n \quad \text{converge.}$$

Alors la suite

$$\theta_{n+1} = \theta_n + \gamma_n (f(\theta_n) - a + \varepsilon_n) \quad (12.29)$$

converge vers θ^* quelle que soit la donnée initiale θ_0 .

La preuve de ce lemme est reportée à la fin de cette section et nous montrons maintenant comment une version stochastique de cet algorithme permet de traiter les fonctions de la forme $f(\theta) = \mathbb{E}(F(X, \theta))$. Dans la pratique, on ne connaît pas la fonction f , mais on peut observer des réalisations $F(X_n, \theta_n)$ et adapter le paramètre θ_n au cours du temps.

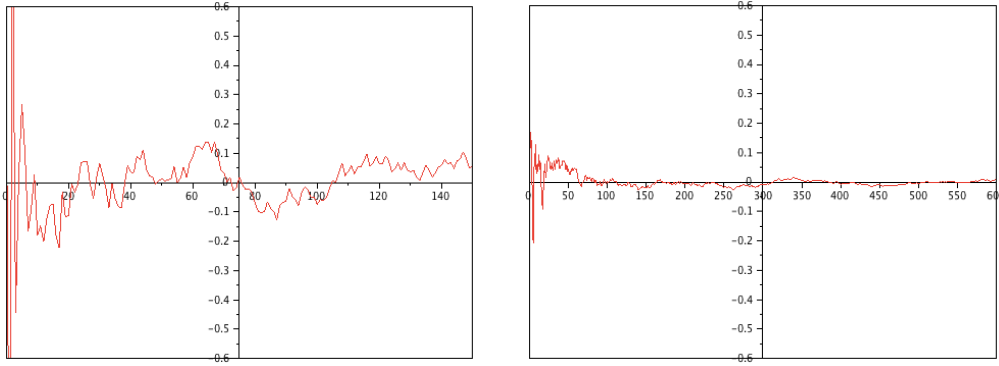


FIGURE 12.11 – L'algorithme de Robbins-Monro est utilisé dans ces 2 simulations avec la récurrence aléatoire $\theta_{n+1} = \theta_n + \gamma_n (-\arctan(\theta_n) + X_n)$ où X_n sont des variables aléatoires uniformément distribuées sur $[-1, 1]$. La simulation représentée à gauche est réalisée en choisissant $\gamma_n = \frac{1}{n^{0.7}}$ pour 150 pas de temps. La série θ_n s'approche de la solution $\theta^* = 0$ en oscillant. La convergence est améliorée sur la simulation de droite réalisée avec $\gamma_n = \frac{1}{n}$ pour 600 pas de temps.

Théorème 12.5 (Algorithme de Robbins-Monro).

On suppose que la fonction $f(\theta) = \mathbb{E}(F(X, \theta))$ est continue et vérifie l'hypothèse (12.27). De plus, on suppose que F est bornée par une constante K .

Soit $\{X_n\}_{n \geq 0}$ une suite de variables indépendantes et distribuées selon la même loi que X . L'algorithme de Robbins-Monro consiste à construire un processus aléatoire

$$\theta_{n+1} = \theta_n + \gamma_n (F(X_n, \theta_n) - a) \quad (12.30)$$

où la suite $\{\gamma_n\}_{n \geq 0}$ d'incrément positifs vérifie

$$\sum_n \gamma_n = \infty \quad \text{et} \quad \sum_n \gamma_n^2 < \infty.$$

Alors le processus $\{\theta_n\}_{n \geq 0}$ converge presque sûrement, quelle que soit la donnée initiale θ_0 , vers θ^* la solution de $f(\theta^*) = \mathbb{E}(F(X, \theta^*)) = a$.

Démonstration. La preuve consiste à réécrire la récurrence (12.30) sous la forme (12.29)

$$\theta_{n+1} = \theta_n + \gamma_n (f(\theta_n) - a + F(X_n, \theta_n) - f(\theta_n))$$

et à poser $\varepsilon_n = F(X_n, \theta_n) - f(\theta_n)$. L'idée est de considérer l'erreur faite en remplaçant $f(\theta_n)$ par $F(X_n, \theta_n)$ comme une perturbation afin d'appliquer le lemme 12.4.

Par l'indépendance des X_k , on vérifie que

$$M_n = \sum_{k=1}^n \gamma_k (F(X_k, \theta_k) - f(\theta_k)) = \sum_{k=1}^n \gamma_k \varepsilon_k$$

est une martingale pour la filtration $\mathcal{F}_n = \sigma(X_1, \dots, X_n)$. De plus elle est bornée dans $\mathbb{L}^2$

$$\mathbb{E}(M_n^2) = \sum_{k=1}^n \gamma_k^2 \mathbb{E} \left((F(X_k, \theta_k) - f(\theta_k))^2 \right) \leq K^2 \sum_{k=1}^n \gamma_k^2 \leq K^2 \sum_{k=1}^{\infty} \gamma_k^2 < \infty$$

où on a utilisé le fait que F est bornée par K et que $\sum_n \gamma_n^2 < \infty$. Le théorème 11.1 implique la convergence presque sûre de M_n et donc de la série $\sum_n \gamma_n \varepsilon_n$. Les hypothèses du lemme 12.4 étant satisfaites, il suffit de l'appliquer pour conclure la démonstration du théorème. $\square$

Démonstration du lemme 12.4. Quitte à changer f en $f(\cdot + \theta^*) - a$, on peut supposer que $a = \theta^* = 0$ et considérer la récurrence $\theta_{n+1} = \theta_n + \gamma_n (f(\theta_n) + \varepsilon_n)$.

On distingue trois comportements possibles pour la suite :

Cas 1. Supposons qu'il existe $\alpha > \beta$ tels que la suite $\{\theta_n\}_{n \geq 0}$ passe infiniment souvent au-dessus de α et au-dessous de β (cf. figure 12.12). Si $\alpha > 0$, alors il suffit de considérer le cas $\beta > 0$. En effet si $\beta \leq 0$ il est clair que la suite va osciller de part et d'autre de l'intervalle $[\alpha/2, \alpha]$ et on peut remplacer β par $\alpha/2$.

On va identifier les portions de la trajectoire où la suite passe au-dessus de β pour la dernière fois avant de remonter au-dessus du niveau α . On notera $[t_k, \tau_k]$ le $k^{\text{ième}}$ intervalle de temps correspondant (cf. figure 12.12). La fonction f est bornée par K et on a

$$|\theta_{n+1} - \theta_n| \leq \gamma_n (K + \varepsilon_n).$$

Comme γ_n et $\gamma_n \varepsilon_n$ tendent vers 0, on peut choisir n assez grand tel que juste avant de franchir le niveau $\beta > 0$ on ait $\theta_{t_k} > 0$.

Si $\theta_n > 0$ alors la condition (12.27) implique

$$\theta_{n+1} = \theta_n + \gamma_n (f(\theta_n) + \varepsilon_n) \leq \theta_n + \gamma_n \varepsilon_n.$$

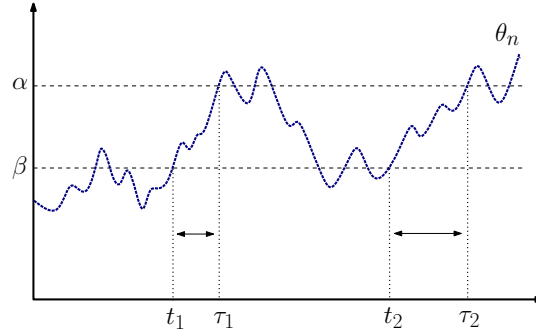


FIGURE 12.12 – La suite $\{\theta_n\}_{n \geq 0}$ est représentée et les intervalles de temps $[t_1, \tau_1]$, $[t_2, \tau_2]$ correspondent au dernier passage au-dessus du niveau β avant d'atteindre le niveau α .

Comme la suite reste positive pendant l'intervalle $[t_k, \tau_k]$, on peut écrire

$$0 < \alpha - \beta < \theta_{\tau_k} - \theta_{t_k} \leq \sum_{n=t_k}^{\tau_k} \gamma_n \varepsilon_n.$$

La suite oscille infiniment souvent entre α et β , par conséquent on peut choisir t_k arbitrairement grand et $\sum_{n=t_k}^{\tau_k} \gamma_n \varepsilon_n$ tend vers 0 car la série converge. Ceci conduit à une contradiction. On peut traiter de façon identique le cas $\alpha < 0$ et ainsi exclure d'éventuelles oscillations entre deux valeurs $\alpha > \beta$.

On se ramène donc au cas où la suite $\{\theta_n\}_{n \geq 0}$ admet une limite (éventuellement infinie).

Cas 2. Supposons que $\lim_n \theta_n = \alpha > 0$. Alors il existe n_0 et $\delta > 0$ tel que

$$\forall n \geq n_0, \quad f(\theta_n) < -\delta$$

en utilisant la condition (12.27) et la continuité de f . La relation de récurrence permet ainsi d'écrire pour $n \geq n_0$

$$\theta_n - \theta_{n_0} = \sum_{k=n_0}^{n-1} \gamma_k (f(\theta_k) + \varepsilon_k) \leq -\delta \sum_{k=n_0}^{n-1} \gamma_k + \sum_{k=n_0}^{n-1} \gamma_k \varepsilon_k.$$

Comme la suite $\sum_n \gamma_n$ diverge on obtient une contradiction. De la même façon, on peut exclure le cas $\lim_n \theta_n = \alpha < 0$.

Cas 3. Supposons que $\lim_n \theta_n = +\infty > 0$. Alors la suite est positive au-delà d'un rang n_0 et on a

$$\theta_n - \theta_{n_0} \leq \sum_{k=n_0}^n \gamma_k (f(\theta_k) + \varepsilon_k) \leq \sum_{k=n_0}^n \gamma_k \varepsilon_k.$$

Comme la suite $\sum_k \gamma_k \varepsilon_k$ converge on en déduit une contradiction. Par symétrie, on peut aussi exclure le cas $\lim_n \theta_n = -\infty$.

L'unique possibilité est que $\lim_n \theta_n = 0$ ce qui conclut le théorème. $\square$

L'algorithme de Robbins-Monro se généralise aux fonctions $f : \mathbb{R}^d \rightarrow \mathbb{R}^d$ satisfaisant l'hypothèse (12.27) et dont la croissance à l'infini est au plus linéaire.

12.5 Processus de Galton-Watson

Nous allons maintenant poursuivre l'étude des arbres aléatoires commencée section 4.4.1 en utilisant cette fois le formalisme des martingales. On rappelle que les arbres aléatoires de Galton-Watson sont définis par récurrence à l'aide d'une loi $\nu = \{p_k\}_{k \geq 0}$ sur $\mathbb{N}$ qui caractérise le nombre de descendants d'un individu. Initialement, il existe un unique ancêtre $Z_0 = 1$. Au temps t , le nombre d'individus est noté Z_t et l'évolution suit la récurrence

$$Z_{t+1} = \begin{cases} \zeta_1^{t+1} + \dots + \zeta_{Z_t}^{t+1}, & \text{si } Z_t > 0 \\ 0, & \text{si } Z_t = 0 \end{cases}$$

où $\{\zeta_i^t\}_{i \geq 1, t \geq 1}$ est une suite de variables aléatoires indépendantes et identiquement distribuées de loi ν

$$\forall k \geq 0, \quad \mathbb{P}(\zeta_i^t = k) = p_k.$$

On note aussi $\mu = \mathbb{E}(\zeta_1^1)$.

Soit $\mathcal{F}_t = \sigma(Z_k, k \leq t)$ la σ -algèbre décrivant la population jusqu'au temps t . En conditionnant par la génération précédente (comme en (4.9)), on vérifie que le processus $M_t = Z_t / \mu^t$ est une martingale

$$\mathbb{E}(M_{t+1} | \mathcal{F}_t) = \mathbb{E}\left(\frac{Z_{t+1}}{\mu^{t+1}} | \mathcal{F}_t\right) = \frac{1}{\mu^{t+1}} \mathbb{E}\left(\sum_{k=0}^{Z_t} \zeta_k^{t+1} | Z_t\right) = \frac{Z_t}{\mu^t} = M_t.$$

Cette martingale étant positive, elle converge presque sûrement, par le corollaire 11.6, vers une variable aléatoire limite M_∞ que nous allons caractériser.

Si $\mu \leq 1$ nous avons montré au théorème 4.10 que la population s'éteint presque sûrement, i.e. que $Z_t = 0$ à partir d'un certain temps (aléatoire). Par conséquent $M_\infty = 0$ presque sûrement. La convergence de M_t vers M_∞ ne peut donc pas avoir lieu dans $\mathbb{L}^1$ car

$$\forall t \geq 0, \quad \mathbb{E}(M_t) = \mathbb{E}(M_0) = 1 \neq 0 = \mathbb{E}(M_\infty).$$

Le théorème suivant permet de décrire le comportement asymptotique dans le cas sur-critique

Théorème 12.6. *Si $\mu > 1$ et la variance $\sigma^2 = \mathbb{E}(\zeta^2) - \mathbb{E}(\zeta)^2$ est finie, alors la martingale $\{M_t\}_{t \geq 0}$ converge dans $\mathbb{L}^2$ vers une limite M_∞ qui vérifie*

$$\mathbb{E}(M_\infty) = 1, \quad \mathbb{E}(M_\infty^2) - \mathbb{E}(M_\infty)^2 = \frac{\sigma^2}{\mu^2 - \mu} \quad \text{et} \quad \mathbb{P}(M_\infty = 0) = \rho$$

où ρ est la probabilité d'extinction définie au théorème 4.10.

Démonstration. Commençons par montrer que la martingale $\{M_t\}_{t \geq 0}$ est bornée dans $\mathbb{L}^2$. On calcule

$$\begin{aligned} \mathbb{E}(M_t^2 | \mathcal{F}_{t-1}) &= \mathbb{E}((M_t - M_{t-1})^2 | \mathcal{F}_{t-1}) + M_{t-1}^2 + 2M_{t-1} \mathbb{E}((M_t - M_{t-1}) | \mathcal{F}_{t-1}) \\ &= M_{t-1}^2 + \mathbb{E}((M_t - M_{t-1})^2 | \mathcal{F}_{t-1}). \end{aligned} \tag{12.31}$$

Le second terme du membre de droite dans (12.31) s'écrit

$$\begin{aligned}\mathbb{E}((M_t - M_{t-1})^2 | \mathcal{F}_{t-1}) &= \mathbb{E}\left(\left(\frac{Z_t}{\mu^t} - \frac{Z_{t-1}}{\mu^{t-1}}\right)^2 \middle| \mathcal{F}_{t-1}\right) = \frac{1}{\mu^{2t}} \mathbb{E}((Z_t - \mu Z_{t-1})^2 | \mathcal{F}_{t-1}) \\ &= \frac{1}{\mu^{2t}} \mathbb{E}\left(\left(\sum_{i=1}^{Z_{t-1}} \zeta_i^t - \mu Z_{t-1}\right)^2 \middle| \mathcal{F}_{t-1}\right) = \frac{Z_{t-1} \sigma^2}{\mu^{2t}}.\end{aligned}$$

Les identités précédentes impliquent

$$\mathbb{E}(M_t^2) = \mathbb{E}(\mathbb{E}(M_t^2 | \mathcal{F}_{t-1})) = \mathbb{E}(M_{t-1}^2) + \frac{\sigma^2}{\mu^{2t}} \mathbb{E}(Z_{t-1}) = \mathbb{E}(M_{t-1}^2) + \frac{\sigma^2}{\mu^{t+1}}$$

en utilisant que $\mathbb{E}(Z_{t-1}) = \mu^{t-1}$. Comme $\mathbb{E}(Z_0^2) = 1$, on obtient par induction

$$\mathbb{E}(M_t^2) = 1 + \sigma^2 \sum_{k=2}^{t+1} \mu^{-k} = 1 + \sigma^2 \frac{1 - \mu^{-t}}{\mu^2 - \mu}. \quad (12.32)$$

On en déduit que la martingale $\{M_t\}_{t \geq 0}$ est bornée dans $\mathbb{L}^2$ et le théorème 11.1 permet d'affirmer qu'elle converge dans $\mathbb{L}^2$ et presque sûrement vers M_∞ . Par conséquent

$$\mathbb{E}(M_\infty) = 1 = \lim_{t \rightarrow \infty} \mathbb{E}(M_t) \quad \text{et} \quad \mathbb{E}(M_\infty^2) - \mathbb{E}(M_\infty)^2 = \frac{\sigma^2}{\mu^2 - \mu} = \lim_{t \rightarrow \infty} \mathbb{E}(M_t^2) - \mathbb{E}(M_t)^2.$$

Pour calculer la probabilité $\mathbb{P}(M_\infty = 0)$, on conditionne l'évolution après la première génération Z_1

$$\mathbb{P}(M_\infty = 0) = \sum_{k=0}^{\infty} \mathbb{P}(M_\infty = 0 | Z_1 = k) p_k.$$

Si $Z_1 = k$, l'arbre se scinde en k arbres indépendants de même loi, on obtient donc

$$\mathbb{P}(M_\infty = 0) = \sum_{k=0}^{\infty} \mathbb{P}(M_\infty = 0)^k p_k = \varphi(\mathbb{P}(M_\infty = 0)).$$

Ceci permet d'identifier la probabilité ρ qui a été définie au théorème 4.10 comme l'unique solution de $\rho = \varphi(\rho)$. $\square$

Chapitre 13

Stratégies, arrêt optimal et contrôle stochastique ★

Dans la vie courante, de nombreuses circonstances nécessitent d'effectuer des choix. Nous allons formaliser le processus de décision pour construire des stratégies qui permettent d'optimiser certains critères en tenant compte des facteurs aléatoires inhérents aux problèmes rencontrés.

13.1 Arrêt optimal

Quel est le meilleur instant pour prendre une décision ? Par exemple, on souhaite vendre un stock de produits au meilleur prix avant l'échéance N et chaque jour $n \in \{0, \dots, N\}$ on démarché un acheteur potentiel qui offre la somme X_n . On peut accepter cette offre ou attendre la suivante sachant qu'on ne pourra plus bénéficier des offres passées. On cherche donc à déterminer le meilleur moment pour vendre sur la base des offres passées sans connaître le futur. Les filtrations permettent de hiérarchiser l'information, et on supposera que le processus aléatoire $X = \{X_n, \quad n = 0, \dots, N\}$ est adapté à la filtration $\mathbb{F} = \{\mathcal{F}_n, \quad n \leq N\}$ où $\mathcal{F}_n$ contient toute l'information jusqu'au temps n . Notre objectif est de construire une stratégie optimale, c'est-à-dire de définir un temps d'arrêt τ pour que l'espérance $\mathbb{E}(X_\tau)$ soit maximale. Plus précisément, si $\mathcal{T}^N$ représente l'ensemble des temps d'arrêt à valeurs dans $\{0, \dots, N\}$, on cherche à résoudre le *problème d'arrêt optimal*

$$V^N = \sup_{\tau \in \mathcal{T}^N} \mathbb{E}(X_\tau). \quad (13.1)$$

On dira qu'un temps d'arrêt τ^* dans $\mathcal{T}^N$ est optimal si $V^N = \mathbb{E}(X_{\tau^*})$. Il s'agit d'une stratégie d'arrêt optimal en *horizon fini* car la décision doit être prise avant l'instant N .

13.1.1 Enveloppe de Snell

Pour résoudre le problème (13.1), on définit le processus Y par la récurrence rétrograde

$$Y_N = X_N \quad \text{et} \quad Y_n = \max \{X_n, \mathbb{E}(Y_{n+1} | \mathcal{F}_n)\} \quad \text{pour} \quad n = 0, \dots, N-1. \quad (13.2)$$

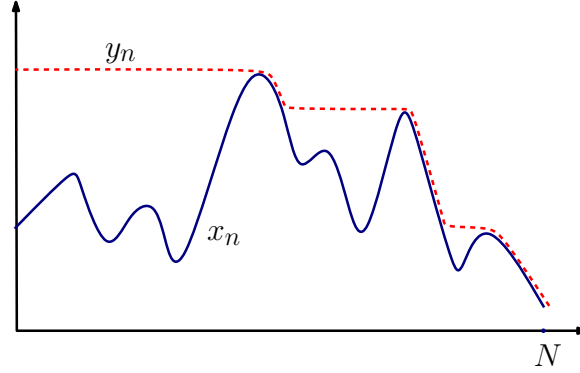


FIGURE 13.1 – Pour déterminer le maximum d’une suite de réels $\{x_n\}_{n \leq N}$, on construit récursivement une suite majorante en partant de $y_N = x_N$ et en remontant ensuite le temps $y_n = \max\{x_n, y_{n+1}\}$. La suite $\{y_n\}_{n \leq N}$, représentée en pointillés, est décroissante et $y_0 = \max\{x_n, n \leq N\}$. De plus, toute autre suite décroissante majorant $\{x_n\}_{n \leq N}$ sera aussi supérieure à $\{y_n\}_{n \leq N}$. L’enveloppe de Snell (13.2) est l’analogue de cette construction dans un cas stochastique.

La figure 13.1 illustre cette construction. Le processus Y est appelé *enveloppe de Snell* du processus X et il s’interprète de la façon suivante. Si le problème d’arrêt optimal se pose au temps N , le seul choix possible de temps d’arrêt est $\tau = N$ ce qui justifie la définition $Y_N = X_N$. À la date $N - 1$, on choisit la stratégie d’arrêt en comparant le gain X_{N-1} obtenu en s’arrêtant à $N - 1$ et le gain espéré si on continuait $\mathbb{E}(Y_N | \mathcal{F}_{N-1}) = \mathbb{E}(X_N | \mathcal{F}_{N-1})$. Ceci explique la définition de Y_{N-1} . En procédant de manière rétrograde date par date, on comprend la logique derrière l’enveloppe de Snell.

Le résultat suivant montre que l’enveloppe de Snell permet de résoudre le problème d’arrêt optimal V^N et de déterminer un temps d’arrêt optimal.

Proposition 13.1. *Supposons que X soit intégrable. Alors l’enveloppe de Snell Y est la plus petite surmartingale majorant le processus X . De plus*

— *la variable aléatoire*

$$\tau^* = \inf \left\{ n \in \{0, \dots, N\}; \quad Y_n = X_n \right\}$$

est un temps d’arrêt.

— *le processus arrêté $Y^{\tau^*} = Y_{\cdot \wedge \tau^*}$ est une martingale.*

— *le temps d’arrêt τ^* est optimal car*

$$Y_0 = \sup_{\tau \in \mathcal{T}^N} \mathbb{E}(X_\tau) = \mathbb{E}(X_{\tau^*}).$$

Cette proposition fournit une construction théorique d’un temps d’arrêt optimal. En pratique, la difficulté consiste à déduire de ce résultat une stratégie explicite. Dans les sections 13.1.2 et 13.1.3, deux exemples sont traités où τ^* peut se réécrire simplement en fonction d’une stratégie de seuil.

Démonstration.

Étape 1. *Construction de la surmartingale Y .*

Vérifions d'abord par une récurrence rétrograde que le processus Y défini en (13.2) est bien intégrable. À la date finale, on a $Y_N = X_N \in \mathbb{L}^1$. Si on suppose que Y_n appartient à $\mathbb{L}^1$, alors

$$\mathbb{E}(|Y_{n-1}|) \leq \mathbb{E}(|X_{n-1}|) + \mathbb{E}(|\mathbb{E}(Y_n|\mathcal{F}_{n-1})|) \leq \mathbb{E}(|X_{n-1}|) + \mathbb{E}(|Y_n|).$$

Ainsi Y est intégrable.

Par définition Y est une surmartingale majorant X . Soit $\tilde{Y}$ une autre surmartingale majorant X . Montrons par récurrence rétrograde que presque sûrement

$$\tilde{Y}_n \geq Y_n \quad \text{pour tout } n = 0, \dots, N.$$

Au temps N , on a $\tilde{Y}_N \geq X_N = Y_N$. Supposons maintenant que $\tilde{Y}_n \geq Y_n$. Comme $\tilde{Y}$ est une surmartingale, elle vérifie

$$\tilde{Y}_{n-1} \geq \mathbb{E}(\tilde{Y}_n|\mathcal{F}_{n-1}) \geq \mathbb{E}(Y_n|\mathcal{F}_{n-1}).$$

De plus $\tilde{Y}$ majore X et par conséquent

$$\tilde{Y}_{n-1} \geq \max\{X_{n-1}, \mathbb{E}(Y_n|\mathcal{F}_{n-1})\} = Y_{n-1}.$$

Étape 2. *Construction du temps d'arrêt optimal.*

On vérifie facilement que $\{\tau^* = n\}$ est bien $\mathcal{F}_n$ -mesurable. Comme $Y_N = X_N$, la variable τ^* définit donc un temps d'arrêt dans $\mathcal{T}^N$, comme premier temps d'atteinte du processus $Y - X$ du niveau 0. Dans le cas déterministe illustré figure 13.1, le temps d'arrêt correspond au premier temps où $y_n = x_n$. On remarque de plus que la suite déterministe $\{y_n\}$ est constante et égale à y_0 tant que $y_n \geq \max\{x_k, k \leq N\}$. Ceci suggère que dans le cas aléatoire, le processus arrêté Y^{τ^*} est une martingale. Nous allons maintenant vérifier cette propriété.

On remarque que l'évènement $\{\tau^* \geq n+1\} = \{\tau^* \leq n\}^c$ est mesurable par rapport à $\mathcal{F}_n$. Par définition, le processus $\{Y_n\}_{n \leq N}$ satisfait sur l'évènement $\{\tau^* \geq n+1\}$

$$Y_n > X_n \quad \text{et} \quad Y_n = \mathbb{E}(Y_{n+1}|\mathcal{F}_n).$$

On en déduit que

$$Y_{n+1}^{\tau^*} - Y_n^{\tau^*} = (Y_{n+1} - Y_n)\mathbf{1}_{\{\tau^* \geq n+1\}} = (Y_{n+1} - \mathbb{E}(Y_{n+1}|\mathcal{F}_n))\mathbf{1}_{\{\tau^* \geq n+1\}}.$$

En utilisant le fait que $\{\tau^* \geq n+1\} \in \mathcal{F}_n$ et en prenant l'espérance conditionnelle par rapport à $\mathcal{F}_n$, on vérifie que le processus arrêté Y^{τ^*} est une martingale

$$\mathbb{E}(Y_{n+1}^{\tau^*} | \mathcal{F}_n) - Y_n^{\tau^*} = 0.$$

Ceci implique que

$$Y_0 = \mathbb{E}(Y_N^{\tau^*}) = \mathbb{E}(Y_{N \wedge \tau^*}) = \mathbb{E}(Y_{\tau^*}) = \mathbb{E}(X_{\tau^*}).$$

Par ailleurs, par la proposition 10.5, pour tout temps d'arrêt τ dans $\mathcal{T}^N$, le processus arrêté Y^τ est une surmartingale. D'après le théorème d'arrêt de Doob, il satisfait donc

$$Y_0 \geq \mathbb{E}(Y_N^\tau) = \mathbb{E}(Y_\tau) \geq \mathbb{E}(X_\tau)$$

par définition de Y . On a ainsi montré la dernière partie de la proposition. $\square$

13.1.2 Le problème du parking

Vous conduisez le long d'une rue infinie vers votre lieu de rendez-vous qui se situe dans un quartier très fréquenté. Le stationnement dans la rue est autorisé, mais bien sûr peu de places sont libres. Alors, si vous avez la possibilité de vous garer, à quelle distance de votre lieu de rendez-vous décidez-vous de prendre la place ?

Cette question peut se modéliser de la façon suivante :

1. Vous démarrez à l'origine. Des emplacements de stationnement sont disponibles à tous les entiers. On considère une suite $\{\xi_n\}_{n \geq 0}$ de variables aléatoires indépendantes de loi de Bernoulli de paramètre p , où $\xi_n = 1$ si et seulement si l'emplacement au point n est déjà occupé. Le lieu de rendez-vous se trouve au point entier $N > 0$.
2. Si l'emplacement n est disponible, i.e. si $\xi_n = 0$, et que vous décidez de vous y garer, vous subissez le coût $|N - n|$ correspondant à l'effort de faire la distance restante en marchant.
3. Quand vous arrivez au niveau du point n , vous ne pouvez pas savoir si des places de stationnement sont disponibles au niveau $n + 1$ ou plus loin. Si vous décidez de passer au point $n + 1$, vous ne pouvez plus retourner aux points précédents.
4. Enfin, si vous arrivez en N , votre point de rendez-vous sans vous être garé, vous prenez la première place de stationnement libre qui se présente. Ainsi, si l'emplacement N est occupé, le coût moyen que vous subissez est alors

$$-X_N = C \mathbf{1}_{\{\xi_N=1\}} \quad \text{avec} \quad C = \sum_{j \geq 1} j p^{j-1} (1-p) = \frac{1}{1-p}.$$

Avant d'arriver en N le processus de coût s'écrit

$$-X_n = (N - n) \mathbf{1}_{\{\xi_n=0\}} + \infty \mathbf{1}_{\{\xi_n=1\}}$$

où le coût infini pour $\xi_n = 1$ signifie que vous ne pouvez occuper la place à aucun coût fini.

Le problème d'arrêt optimal consiste à chercher le temps d'arrêt qui minimise le coût de l'effort de l'agent, ou en inversant les signes

$$\sup_{\tau \in \mathcal{T}^N} \mathbb{E}(X_\tau).$$

L'enveloppe de Snell est donnée par $Y_N = X_N$ et

$$Y_n = \max \{X_n, \mathbb{E}(Y_{n+1} | \mathcal{F}_n)\} \quad \text{pour} \quad n < N.$$

Un simple raisonnement par récurrence rétrograde, utilisant l'indépendance des ξ_n , montre que Y_n est une fonction de ξ_n . Par conséquent

$$\mathbb{E}(Y_{n+1} | \mathcal{F}_n) = \mathbb{E}(Y_{n+1}) = f(N - (n + 1)) \quad (13.3)$$

où $f : \{0, \dots, N\} \rightarrow \mathbb{R}$ est une fonction que nous allons déterminer. Comme Y est une surmartingale, l'espérance $n \rightarrow \mathbb{E}(Y_n)$ décroît et par conséquent $n \rightarrow f(N - n)$ est décroissante (cf. figure 13.2). Le premier temps où $Y_n = X_n$ revient à déterminer le premier $n < N$ tel que

$$\xi_n = 0 \quad \text{et} \quad -(N - n) \geq f(N - n). \quad (13.4)$$

Soit $r^* \geq 0$ le premier point tel que $-(N - r^*) \geq f(N - r^*)$ (cf. figure 13.2). Si une place est disponible avant r^* alors l'inégalité de la relation (13.4) ne sera pas satisfaite, par conséquent il suffit de choisir la première place disponible après r^* .

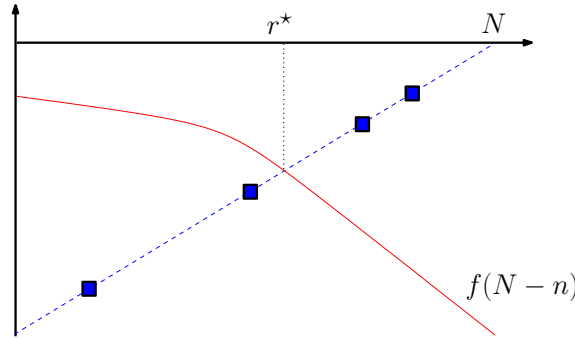


FIGURE 13.2 – Les graphes de $n \rightarrow f(N - n)$ et de $n \rightarrow -(N - n)$ sont représentés. Les carrés marquent les positions n des places disponibles ($\xi_n = 0$) et le gain correspondant est $-(N - n)$. Le seuil r^* est le point où les deux graphes s'intersectent. La place choisie est la première place libre après r^* , i.e. le troisième carré sur le schéma.

La stratégie d'arrêt optimale est nécessairement de la forme

$$\tau^* = \inf \{n \geq N - r^*; \quad \xi_n = 0\} \quad (13.5)$$

où r^* est une constante à déterminer. On note $\ell(r)$ la performance espérée en utilisant une stratégie de seuil avec paramètre r , i.e.

$$\ell(r) = \mathbb{E}(Y_{\tau(r)}) \quad \text{où} \quad \tau(r) = \inf \{n \geq N - r; \quad \xi_n = 0\}.$$

Nous allons calculer $\ell(r)$ et optimiser en fonction de r pour identifier r^* .

Pour $r = 0$, on obtient

$$\ell(0) = -(1 - p) \times 0 - pC = -\frac{p}{1 - p}.$$

Pour $r \geq 1$, on calcule par récurrence $\ell(r) = -(1 - p)r + p\ell(r - 1)$. On déduit alors que

$$-\ell(r) = r + 1 + \frac{2p^{r+1} - 1}{1 - p}, \quad r \leq N.$$

Pour maximiser $\ell(r)$, on remarque que la fonction $r \mapsto \ell(r + 1) - \ell(r) = -1 + 2p^{r+1}$ est décroissante en r . Par conséquent

$$r^* = \inf \left\{ r \geq 0; \quad -1 + 2p^{r+1} \leq 0 \right\}.$$

À titre d'exemple, on peut voir que pour $p \leq 0.5$, il faut chercher à se garer en arrivant à destination et que pour $p = 0.9$, il faut chercher la première place disponible dès qu'on arrive à 6 places de la destination.

13.1.3 Problème des secrétaires

Nous allons nous intéresser à un problème classique connu sous le nom de *problème des secrétaires*. N candidats effectuent un entretien pour un poste de secrétaire vacant et ceux-ci peuvent être tous comparés strictement (une fois rencontrés) et classés. Les candidats sont auditionnés un à un dans un ordre arbitraire, les $N!$ façons d'ordonner les candidats sont équiprobables. À l'issue de l'audition de chaque candidat et sur la base de son rang relatif par rapport aux candidats précédemment auditionnés, on doit

- soit le sélectionner pour le poste, terminant ainsi la procédure de sélection,
- soit refuser définitivement sa candidature, sans possibilité de le rappeler ultérieurement, et passer au candidat suivant.

L'objectif est de sélectionner le meilleur candidat parmi les N : le gain est défini par 1 si on sélectionne le meilleur candidat et par 0 sinon.

Modélisation.

On peut classer les N candidats selon un ordre décroissant $\{\sigma(1), \dots, \sigma(N)\}$ où σ est une permutation uniforme de $\{1, \dots, N\}$. Le candidat i a donc le rang $\sigma(i)$ et on cherche à déterminer le meilleur candidat, c'est-à-dire le candidat j tel que $\sigma(j) = 1$. L'enjeu est d'optimiser le gain $\mathbb{P}(\sigma(\tau) = 1)$ où τ est un temps d'arrêt représentant le candidat choisi.

La première difficulté est que le processus $\{\mathbf{1}_{\sigma(k)=1}\}_{k \leq N}$ que l'on cherche à optimiser n'est pas mesurable par rapport à la σ -algèbre engendrée par les observations jusqu'à l'instant k . En effet savoir que $\sigma(k) = 1$ suppose de connaître le classement des N candidats. Les variables mesurées naturellement sont les rangs relatifs $\{\xi_n\}_{1 \leq n \leq N}$. Précisément, ξ_n désigne le rang du $n^{\text{ième}}$ candidat auditionné parmi les n candidats auditionnés. Par exemple si $N = 5$ et que le classement des candidats est donné par $\sigma = \{3, 2, 5, 1, 4\}$, alors les ordres relatifs seront

$$\xi_1 = 1, \quad \xi_2 = 1, \quad \xi_3 = 3, \quad \xi_4 = 1, \quad \xi_5 = 4.$$

On aura toujours $\xi_1 = 1$ et $\xi_N = \sigma(N)$ car au temps N tous les candidats sont connus. On remarque que la connaissance de tous les ordres relatifs jusqu'à N permet de reconstituer le classement. Pour chaque $n \geq 1$, la variable aléatoire ξ_n est distribuée selon la loi uniforme sur $\{1, \dots, n\}$

$$\forall k \in \{1, \dots, n\}, \quad \mathbb{P}(\xi_n = k) = \frac{1}{n}.$$

De plus la variable ξ_n est indépendante des $\{\xi_i, i \leq n-1\}$. On note $\mathcal{F}_n = \sigma(\xi_i, i \leq n)$ la filtration canonique correspondante.

Nous allons considérer le processus

$$n \in \{1, \dots, N\}, \quad X_n = \mathbb{E}(\mathbf{1}_{\{\sigma(n)=1\}} \mid \mathcal{F}_n).$$

Pour tout temps d'arrêt τ , on vérifie que

$$\begin{aligned} \mathbb{P}(\sigma(\tau) = 1) &= \sum_{n=1}^N \mathbb{E}(\mathbf{1}_{\{\tau=n\}} \mathbf{1}_{\{\sigma(n)=1\}}) \\ &= \sum_{n=1}^N \mathbb{E}(\mathbf{1}_{\{\tau=n\}} \mathbb{E}(\mathbf{1}_{\{\sigma(n)=1\}} | \mathcal{F}_n)) \\ &= \sum_{n=1}^N \mathbb{E}(\mathbf{1}_{\{\tau=n\}} X_n) = \mathbb{E}(X_\tau). \end{aligned}$$

Pour déterminer le temps d'arrêt optimal, il est donc équivalent de travailler avec le processus $\{X_n\}_{n \leq N}$ qui est mesurable par rapport aux observations contrairement à $\{\mathbf{1}_{\{\sigma(n)=1\}}\}_{n \leq N}$. De plus le processus $\{X_n\}_{n \leq N}$ peut se réécrire sous la forme

$$X_n = \frac{n}{N} \mathbf{1}_{\{\xi_n=1\}}.$$

En effet l'évènement $\{\sigma(n) = 1\}$ correspond à $\{\xi_n = 1, \xi_{n+1} \neq 1, \dots, \xi_N \neq 1\}$ et en utilisant l'indépendance des variables ξ_i , on retrouve

$$\begin{aligned} X_n &= \mathbb{E}(\mathbf{1}_{\{\sigma(n)=1\}} | \mathcal{F}_n) = \mathbf{1}_{\{\xi_n=1\}} \mathbb{P}(\xi_{n+1} \neq 1, \dots, \xi_N \neq 1) \\ &= \mathbf{1}_{\{\xi_n=1\}} \frac{n}{n+1} \cdot \frac{n+1}{n+2} \cdots \frac{N-1}{N} = \frac{n}{N} \mathbf{1}_{\{\xi_n=1\}}. \end{aligned}$$

On peut aussi interpréter ce résultat plus intuitivement en remarquant que $\{\sigma(n) = 1\}$ équivaut à ce que $\xi_n = 1$ et que le meilleur candidat figure parmi les n premiers (ce qui a pour probabilité $\frac{n}{N}$).

La dernière étape consiste à calculer l'enveloppe de Snell du processus $\{X_n\}_{n \leq N}$

$$Y_N = \mathbf{1}_{\{\xi_N=1\}} \quad \text{et} \quad Y_n = \max \left\{ \frac{n}{N} \mathbf{1}_{\{\xi_n=1\}}, \mathbb{E}(Y_{n+1} | \mathcal{F}_n) \right\} \quad \text{pour } n = 1, \dots, N-1. \quad (13.6)$$

Comme dans le problème du parking (13.3), on utilise l'indépendance des variables $\{\xi_i\}_{i \leq N}$ pour conclure qu'il existe une fonction $f : \{0, \dots, N\} \rightarrow \mathbb{R}$ telle que

$$\mathbb{E}(Y_{n+1} | \mathcal{F}_n) = \mathbb{E}(Y_{n+1}) = f(N-n)$$

de plus $n \rightarrow f(N-n)$ est décroissante. Le temps d'arrêt optimal τ^* correspond donc au premier temps $n < N$ où

$$\frac{n}{N} \mathbf{1}_{\{\xi_n=1\}} \geq f(N-n).$$

sinon on pose $\tau^* = N$. Ceci revient à considérer une stratégie de seuil similaire à (13.5). Pour déterminer le seuil r^* optimal, on calcule le gain de chaque stratégie

$$\tau_r = N \wedge \inf \{n \geq r; \quad \xi_n = 1\} \quad \text{avec} \quad \ell(r) = \mathbb{E}(X_{\tau_r})$$

et on optimise ensuite r . En utilisant l'indépendance des variables $\{\xi_n\}_{n \leq N}$, on obtient

$$\begin{aligned} \ell(r) &= \mathbb{E}(X_{\tau_r}) = \mathbb{P}(\sigma(\tau_r) = 1) = \sum_{n=r}^N \mathbb{E}(\mathbf{1}_{\{\tau_r=n\}} \mathbf{1}_{\{\sigma(n)=1\}}) \\ &= \sum_{n=r}^N \mathbb{P}(\xi_r \neq 1, \dots, \xi_{n-1} \neq 1, \xi_n = 1, \xi_{n+1} \neq 1, \dots, \xi_N \neq 1) \\ &= \frac{r-1}{N} \sum_{n=r}^N \frac{1}{n-1} \end{aligned}$$

avec $\ell(0) = 1/N$. On remarque que

$$\ell(r+1) - \ell(r) = -\frac{1}{N} + \frac{1}{N} \sum_{n=r+1}^N \frac{1}{n-1}.$$

Le seuil optimal r^* correspond au premier point où la courbe ℓ commence à décroître

$$r^* = \inf \left\{ r; \sum_{k=r+1}^N \frac{1}{k-1} < 1 \right\}.$$

Quand N est grand, on peut utiliser l'approximation

$$\sum_{k=r+1}^N \frac{1}{k-1} \simeq \int_r^N \frac{1}{u} du = \log \frac{N}{r}$$

pour obtenir l'ordre de grandeur $r^* \sim Ne^{-1} \approx .37 N$. Ainsi, la stratégie optimale consiste à rejeter systématiquement les premiers candidats et, à partir de 37% des candidats auditionnés, à sélectionner celui qui sera classé premier parmi tous ses prédécesseurs.

13.2 Contrôle stochastique

Avant d'évoquer les aspects aléatoires, commençons par décrire les problématiques liées à la théorie du contrôle dans un cadre déterministe. Une première application consiste à guider un mobile (par exemple un satellite) dont la position X_n évolue en suivant la dynamique

$$X_0 = x \quad \text{et} \quad X_{n+1} = F(X_n, u_n) \quad \text{pour} \quad n \geq 0 \quad (13.7)$$

où $u = \{u_n\}_{n \leq N}$ est un paramètre de contrôle qui doit être ajusté pour optimiser la trajectoire du mobile. Plus généralement, on cherche à optimiser un coût de la forme

$$\mathcal{C}(x, u) = \sum_{n=1}^{N-1} c(X_n, u_n) + W(X_N) \quad (13.8)$$

en fonction de u . Les paramètres de contrôle prennent leurs valeurs dans un ensemble $\mathcal{U}$ et l'enjeu de cette section est de déterminer le coût optimal $\mathcal{C}^*(x)$ et le contrôle optimal correspondant $u^* = \{u_n^*\}_{n \leq N}$ dans $\mathcal{U}^N$ tel que

$$\mathcal{C}^*(x) = \mathcal{C}(x, u^*) = \min_{u \in \mathcal{U}^N} \mathcal{C}(x, u)$$

où x est la donnée initiale.

Les applications du contrôle sont multiples et la fonction de coût (13.8) peut s'interpréter de différentes façons. Par exemple X_n peut représenter la population d'une espèce de poissons au début de l'année n et le paramètre u_n les quotas de pêche qui permettent de contrôler l'évolution de cette population sous la forme $X_{n+1} = F(X_n, u_n)$. On cherche à ajuster les quotas de pêche durant N années afin de garantir un certain niveau d'exploitation $c(X_n, u_n)$ chaque année mais aussi la préservation de la ressource naturelle en imposant une contrainte $W(X_N)$ au temps final. D'autres questions liées aux politiques de développement durable (exploitation des forêts, détermination des quotas d'émission de CO2) sont détaillées dans le livre [7] ainsi que la forme explicite des fonctions de coût associées.

De nombreuses applications sont liées à l'économie, citons notamment la gestion d'un stock de marchandises [8]. Le paramètre X_n représente alors la quantité du stock au jour n et le paramètre u_n permet d'ajuster ce stock au cours du temps en passant commande aux fournisseurs. À chaque période de temps, $c(X_n, u_n)$ prend en compte le gain obtenu en vendant cette marchandise, les frais de stockage, etc. Les invendus au temps N induisent la pénalisation $W(X_N)$. La théorie du contrôle est aussi très utilisée en mathématiques financières.

Pour préciser la modélisation, on peut aussi tenir compte d'éventuels aléas et modifier les règles d'évolution

$$\forall n \geq 0, \quad X_{n+1} = F(X_n, u_n, \xi_n) \quad (13.9)$$

où $\{\xi_n\}_{n \leq N}$ est une suite de variables aléatoires indépendantes et identiquement distribuées. Par exemple, la reproduction d'une espèce animale peut être affectée par des facteurs climatiques qu'on modélise par les ξ_n . La théorie du *contrôle stochastique* consiste à identifier un contrôle optimal qui minimise le coût moyen défini en (13.13).

Nous allons d'abord décrire la méthode de programmation dynamique qui permet de déterminer le contrôle optimal dans le cas déterministe puis nous généraliserons cette stratégie aux évolutions aléatoires.

13.2.1 Équation de la programmation dynamique

L'algorithme proposé par R. Bellman permet d'identifier le contrôle pour une évolution déterministe (13.7). L'algorithme de la programmation dynamique procède de façon rétrograde comme pour la construction de l'enveloppe de Snell (13.2).

On note E l'espace d'états où X_n prend ses valeurs. Étant donné un contrôle $\{u_n\}_{k \leq n \leq N}$ et x dans E , on définit la trajectoire partielle $\{X_n\}_{k \leq n \leq N}$ partant de x au temps k

$$X_k = x \quad \text{et} \quad X_{n+1} = F(X_n, u_n) \quad \text{pour} \quad n \geq k$$

et la fonction de coût partielle associée

$$C_k(x, u) = \sum_{n=k}^{N-1} c(X_n, u_n) + W(X_N). \quad (13.10)$$

Supposons que le coût optimal associé est bien défini (c'est le cas, par exemple, si E et U sont finis). Alors il est donné par

$$C_k^*(x) = \inf_u \left\{ \sum_{n=k}^{N-1} c(X_n, u_n) + W(X_N) \right\}. \quad (13.11)$$

Il satisfait les équations de la programmation dynamique : pour tout x de E

$$\forall x \in E, \quad \begin{cases} \mathcal{C}_N^*(x) &= W(x), \\ \mathcal{C}_k^*(x) &= \inf_{a \in \mathcal{U}} \left\{ c(x, a) + \mathcal{C}_{k+1}^*(F(x, a)) \right\} \quad \text{pour } k \leq N-1. \end{cases} \quad (13.12)$$

Le principe qui sous-tend ces équations est que la trajectoire optimale entre 2 points sera aussi optimale entre 2 points intermédiaires. Par conséquent, si on connaît le contrôle optimal entre $k+1$ et N , il est facile d'en déduire le contrôle optimal entre k et N .

Ces équations peuvent être résolues de façon rétrograde. La valeur de a où le minimum est atteint correspond au contrôle optimal et on la notera $u_n^*(x)$ (si cette valeur n'est pas unique, on en choisit une). En remontant jusqu'au temps $k = 0$, on établit le coût optimal pour n'importe quelle donnée initiale $X_0 = x$ dans E

$$\mathcal{C}^*(x) = \min_{u \in \mathcal{U}^N} \mathcal{C}_0(x, u).$$

En utilisant les différentes valeurs $\{u_n^*(x)\}_{n \leq N, x \in E}$ déterminées au cours de cette procédure, une trajectoire optimale $\{X_n^*\}_{n \leq N}$ peut être reconstruite pour toute donnée initiale x dans E par

$$X_0^* = x \quad \text{et} \quad X_{n+1}^* = F(X_n^*, u_n^*(X_n^*)) \quad \text{pour } n \leq N-1.$$

13.2.2 Contrôle des chaînes de Markov

Nous allons maintenant généraliser les résultats précédents au contrôle stochastique. On considère l'évolution aléatoire de la forme (13.9)

$$X_0 = x \quad \text{et} \quad X_{n+1} = F(X_n, U_n, \xi_n) \quad \text{pour } n \geq 0$$

où $\{\xi_n\}_{n \leq N}$ est une suite de variables aléatoires indépendantes et identiquement distribuées. On cherche un contrôle $U_n = u_n(X_n)$ mesurable par rapport à $\sigma(X_n)$ (cette hypothèse peut être généralisée). Ce choix permet d'affirmer que $\{X_n\}_{n \leq N}$ est bien une chaîne de Markov.

L'analogie du coût (13.8) s'écrit

$$\widehat{\mathcal{C}}(x, U) = \mathbb{E} \left(\sum_{n=1}^{N-1} c(X_n, U_n) + W(X_N) \right). \quad (13.13)$$

En décomposant les coûts partiels comme en (13.11)

$$\widehat{\mathcal{C}}_k^*(x) = \inf_U \left\{ \mathbb{E}_{k,x} \left(\sum_{n=k}^{N-1} c(X_n, u_n) + W(X_N) \right) \right\} \quad (13.14)$$

où l'espérance $\mathbb{E}_{k,x}$ porte sur les trajectoires partant de x au temps k et le minimum (dont on suppose l'existence) est choisi pour les contrôles de la forme

$$U = \{u_k(x_k), \dots, u_{N-1}(x_{N-1})\}.$$

Dans le cas aléatoire, les équations de la programmation dynamique s'écrivent pour tout x de E

$$\forall x \in E, \quad \begin{cases} \hat{\mathcal{C}}_N^*(x) &= W(x), \\ \hat{\mathcal{C}}_k^*(x) &= \inf_{a \in \mathcal{U}} \left\{ c(x, a) + \mathbb{E} \left(\hat{\mathcal{C}}_{k+1}^*(F(x, a, \xi)) \right) \right\} \quad \text{pour } k \leq N-1. \end{cases} \quad (13.15)$$

On construit ainsi le contrôle optimal par étapes, en définissant $u_k^*(x)$ comme la valeur qui minimise la relation (13.15) à chaque pas de temps et pour chaque état x dans E .

Annexe A

Théorie de la mesure

Cette annexe présente des éléments de la théorie de la mesure dans un cadre fonctionnel, ceux ci seront ensuite traduits dans une formulation probabiliste dans l'annexe B.

A.1 Espaces mesurables et mesures

L'enjeu de cette première section est de définir le formalisme nécessaire pour construire des mesures sur des espaces généraux Ω . Par exemple si $\Omega = [0, 1] \subset \mathbb{R}$, on souhaite construire une mesure μ telle que l'intervalle $[a, b]$ ait pour mesure $\mu([a, b]) = b - a$. On veut aussi pouvoir mesurer la taille du complémentaire de $[a, b]$ et la réunion (éventuellement dénombrable) de plusieurs intervalles. Ainsi même en partant initialement d'une classe très simple de sous-ensembles que l'on souhaite mesurer $\{[a, b], 0 \leq a \leq b \leq 1\}$, la structure des sous-ensembles de Ω à considérer va considérablement s'enrichir quand on considère leurs réunions dénombrables et le passage au complémentaire. Il faut donc pouvoir construire la mesure μ sur des sous-ensembles de plus en plus complexes. Pour cela, nous sommes amenés à définir une classe adaptée de sous-ensembles mesurables dite σ -algèbre (section A.1.1), sur laquelle une notion de mesure peut être construite (section A.1.2).

A.1.1 σ -algèbres

Soit Ω un ensemble quelconque, commençons par définir la notion d'algèbre.

Définition A.1 (Algèbre). Une algèbre $\mathcal{A}$ sur Ω est une famille de sous-ensembles de Ω satisfaisant les trois propriétés suivantes :

- (i) Ω appartient à $\mathcal{A}$.
- (ii) Si A est dans $\mathcal{A}$ alors A^c est dans $\mathcal{A}$.
- (iii) Toute réunion finie de sous-ensembles dans $\mathcal{A}$ appartient à $\mathcal{A}$.

Ceci permet de reformuler la définition 9.1 d'une σ -algèbre qui va constituer le cadre naturel pour construire des mesures.

Définition A.2 (σ -algèbre). Une σ -algèbre $\mathcal{A}$ sur Ω est algèbre stable par union dénombrable (ceci étend la propriété (iii) à toute réunion dénombrable de sous-ensembles dans $\mathcal{A}$).

Étant donnée une σ -algèbre $\mathcal{A}$, les propriétés (i) et (ii) impliquent que $\emptyset = \Omega^c$ est dans $\mathcal{A}$. En associant (ii) et (iii), on en déduit qu'une σ -algèbre $\mathcal{A}$ est aussi stable par intersection dénombrable d'ensembles $\{A_i\}_{i \geq 1}$ de $\mathcal{A}$

$$\bigcap_{i \geq 1} A_i = \left(\bigcup_{i \geq 1} A_i^c \right)^c.$$

Une σ -algèbre est aussi stable par différence symétrique

$$A \Delta B := (A \cup B) \setminus (A \cap B) \in \mathcal{A} \quad \text{pour tous } A, B \in \mathcal{A}.$$

L'ensemble des parties de Ω est la plus grande σ -algèbre sur Ω , mais cette dernière est souvent trop complexe pour définir une mesure. En général, la notion de σ -algèbre est tellement abstraite qu'il est impossible de décrire tous les sous-ensembles contenus dans une σ -algèbre. On préfère donc se limiter à des σ -algèbres construites à partir d'une collection réduite de sous-ensembles de Ω .

Définition A.3 (σ -algèbre engendrée). *Si $\mathcal{C}$ est une collection de sous-ensembles de Ω , on notera $\sigma(\mathcal{C})$ la plus petite σ -algèbre contenant $\mathcal{C}$. On dira que $\sigma(\mathcal{C})$ est la σ -algèbre engendrée par $\mathcal{C}$.*

Nous verrons dans la section suivante que les σ -algèbres engendrées permettent de définir un embryon de mesure sur une classe $\mathcal{C}$ restreinte de sous-ensembles et ensuite de l'étendre par un théorème abstrait à toute la σ -algèbre $\sigma(\mathcal{C})$.

Si $\Omega = \mathbb{R}$, l'exemple classique est la σ -algèbre borélienne engendrée par l'ensemble des intervalles

$$\mathcal{B}_{\mathbb{R}} = \sigma(\mathcal{I}) \quad \text{avec } \mathcal{I} = \{] - \infty, x], \quad x \in \mathbb{R} \}. \quad (\text{A.1})$$

La σ -algèbre borélienne $\mathcal{B}_{\mathbb{R}}$ peut être engendrée par d'autres classes de sous-ensembles comme par exemple les ouverts de $\mathbb{R}$. Plus généralement, si Ω est un espace topologique, la σ -algèbre borélienne $\mathcal{B}_{\Omega}$ est la σ -algèbre engendrée par les ouverts de Ω .

A.1.2 Mesures

Le couple $(\Omega, \mathcal{A})$, formé d'un ensemble Ω et d'une σ -algèbre sur Ω , est appelé un *espace mesurable*. L'enjeu est maintenant d'adjoindre à $(\Omega, \mathcal{A})$ une mesure μ définie sur les sous-ensembles de $\mathcal{A}$.

Définition A.4 (Mesure). *Une mesure positive μ sur l'espace mesurable $(\Omega, \mathcal{A})$ est une application $\mu : \mathcal{A} \rightarrow [0, \infty]$ vérifiant les propriétés suivantes :*

- (i) $\mu(\emptyset) = 0$,
- (ii) la mesure μ est σ -additive, c'est-à-dire que pour toute suite $\{A_n\}_{n \geq 0} \subset \mathcal{A}$ d'ensembles disjoints alors

$$\mu\left(\bigcup_{n \geq 0} A_n\right) = \sum_{n \geq 0} \mu(A_n). \quad (\text{A.2})$$

On dit alors que $(\Omega, \mathcal{A}, \mu)$ est un *espace mesuré*.

Une mesure μ est dite σ -finie sur Ω s'il existe une suite croissante d'ensembles mesurables vérifiant $\Omega = \bigcup_n \Omega_n$ avec $\mu(\Omega_n) < \infty$ pour tout n .

Cette définition permet d'établir plusieurs propriétés importantes vérifiées par une mesure.

Proposition A.5. *Une mesure μ sur l'espace mesurable $(\Omega, \mathcal{A})$ vérifie*

1. *Pour toute famille finie $\{A_n\}_{n \leq K}$ d'ensembles disjoints dans $\mathcal{A}$*

$$\mu\left(\bigcup_{n \leq K} A_n\right) = \sum_{n \leq K} \mu(A_n).$$

2. *Soient $A \subset B$ et $\mu(A) \leq \mu(B)$ avec $\mu(A) < \infty$ alors*

$$\mu(B \setminus A) = \mu(B) - \mu(A).$$

3. *Soient $A, B \in \mathcal{A}$ alors*

$$\mu(A) + \mu(B) = \mu(A \cup B) + \mu(A \cap B).$$

4. *Soit une suite $A_n \in \mathcal{A}$ croissante $A_n \subset A_{n+1}$, alors*

$$\mu\left(\bigcup_{n \geq 0} A_n\right) = \lim_{n \rightarrow \infty} \mu(A_n). \quad (\text{A.3})$$

5. *Soit une suite $A_n \in \mathcal{A}$ décroissante $A_{n+1} \subset A_n$ avec $\mu(A_0) < \infty$, alors*

$$\mu\left(\bigcap_{n \geq 0} A_n\right) = \lim_{n \rightarrow \infty} \mu(A_n).$$

Démonstration. L'additivité de l'assertion (1) est une conséquence de la σ -additivité (A.2) en choisissant $A_i = \emptyset$ pour $i > K$.

Les autres propriétés s'obtiennent en décomposant les différents ensembles à l'aide de sous-ensembles disjoints pour lesquels la σ -additivité s'applique. Nous détaillerons simplement la preuve de la propriété (4). À partir de la famille croissante $A_n \in \mathcal{A}$, on construit la famille de sous ensembles disjoints

$$C_0 = A_0 \quad \text{et} \quad \forall n \geq 1, \quad C_n = A_n \setminus A_{n-1},$$

qui vérifie $\bigcup_{n \geq 0} A_n = \bigcup_{n \geq 0} C_n$. Comme les C_n sont disjoints, on en déduit

$$\mu\left(\bigcup_{n \geq 0} A_n\right) = \mu\left(\bigcup_{n \geq 0} C_n\right) = \sum_{n \geq 0} \mu(C_n) = \lim_{K \rightarrow \infty} \sum_{n=0}^K \mu(C_n) = \lim_{K \rightarrow \infty} \mu(A_K).$$

La propriété (5) se déduit de (4) en considérant la suite croissante $B_n = A_0 \setminus A_n$. □

Certains ensembles ne sont pas significatifs pour une mesure donnée, on dit alors qu'ils sont négligeables.

Définition A.6. *Sur un espace mesuré $(\Omega, \mathcal{A}, \mu)$, un ensemble $\mathcal{N} \in \mathcal{A}$ est dit négligeable si $\mu(\mathcal{N}) = 0$.*

D'après la propriété de σ -additivité de la mesure, toute union dénombrable d'ensembles négligeables est négligeable.

Le résultat suivant constitue la clef de voûte de la théorie de la mesure, car il permet de construire de façon unique une mesure sur toute une σ -algèbre en partant d'une classe restreinte de sous-ensembles.

Théorème A.7. (*Théorème d'extension de Carathéodory*) Soient $\mathcal{A}_0$ une algèbre sur Ω et une fonction σ -additive $\mu_0 : \mathcal{A}_0 \rightarrow \mathbb{R}_+$ au sens de (A.2). Alors il existe une mesure μ sur la σ -algèbre $\mathcal{A} := \sigma(\mathcal{A}_0)$ engendrée par $\mathcal{A}_0$ telle que $\mu = \mu_0$ sur $\mathcal{A}_0$. Si de plus μ_0 est σ -finie (en particulier si $\mu_0(\Omega) < \infty$), alors une telle extension μ est unique.

Dans la pratique, le théorème d'extension de Carathéodory sert peu car une fois qu'on connaît l'existence de la mesure μ sur $\mathcal{A}$, il n'est plus nécessaire de faire appel à ce théorème pour travailler avec μ . La démonstration de ce théorème dépasse le cadre de ce cours, mais on pourra la trouver dans l'annexe A1, page 195 du livre [28]. Un point important de la preuve mérite néanmoins d'être signalé car l'unicité de la mesure s'obtient comme conséquence de la proposition suivante qui sera utile par la suite.

Proposition A.8. Soit $\mathcal{I} \subset \mathcal{P}(\Omega)$ une famille de sous-ensembles, stable par intersection finie, i.e.

$$\forall A, B \in \mathcal{I} \quad \Rightarrow \quad A \cap B \in \mathcal{I}.$$

Soient μ, ν deux mesures définies sur $\sigma(\mathcal{I})$ telles que $\mu(\Omega), \nu(\Omega) < \infty$. Si μ, ν coïncident sur $\mathcal{I}$ alors $\mu = \nu$ sur $\sigma(\mathcal{I})$.

La preuve de cette proposition se trouve page 194 de [28]. La construction de la mesure de Lebesgue λ sur $[0, 1]$ permet d'illustrer la proposition A.8. En effet, il suffit de définir la mesure sur les segments de la forme

$$\forall x \in [0, 1], \quad \lambda([0, x]) = x \quad \text{avec} \quad \mathcal{I} = \{[0, x] : x \in [0, 1]\}, \quad (\text{A.4})$$

pour la prescrire de façon unique sur toute la σ -algèbre borélienne $\mathcal{B}_{[0,1]} = \sigma(\mathcal{I})$ introduite en (A.1). L'existence de la mesure de Lebesgue s'obtient par le théorème A.7, mais cela nécessite un travail plus conséquent (cf. [28] page 20).

A.2 Intégration

Étant donné un espace mesuré $(\Omega, \mathcal{A}, \mu)$, nous allons construire la théorie de l'intégration pour des fonctions mesurables.

A.2.1 Fonctions mesurables

La mesurabilité d'une fonction assure une compatibilité entre les structures des σ -algèbres sous-jacentes.

Définition A.9. On considère $(\Omega, \mathcal{A})$ et $(\Gamma, \mathcal{B})$ deux espaces mesurables. Une application $f : \Omega \rightarrow \Gamma$ est dite mesurable si

$$\forall B \in \mathcal{B}, \quad f^{-1}(B) \in \mathcal{A}.$$

Dans le cas de σ -algèbres $\mathcal{A}, \mathcal{B}$ boréliennes, on dira que f est borélienne.

Dans la pratique, pour vérifier la mesurabilité d'une fonction f , il suffit de considérer une classe réduite de sous-ensembles $\mathcal{C}$ engendrant $\mathcal{B} = \sigma(\mathcal{C})$ et de montrer que

$$\forall C \in \mathcal{C}, \quad f^{-1}(C) \in \mathcal{A}.$$

En particulier, si $f : \Omega \rightarrow \mathbb{R}$ est une fonction continue alors elle est mesurable pour la σ -algèbre borélienne $\mathcal{B}_\Omega$. En effet, l'ensemble $\mathcal{C}$ des ouverts de $\mathbb{R}$ engendre $\mathcal{B}_\mathbb{R}$ et l'image réciproque d'un ouvert C par l'application continue f est aussi un ouvert $f^{-1}(C) \in \mathcal{B}_\Omega$.

Les fonctions étagées forment une classe importante de fonctions mesurables car elles constituent les briques fondamentales pour construire la théorie de l'intégration.

Définition A.10 (Fonction étagée). Une fonction $f : \Omega \rightarrow \mathbb{R}$, définie sur l'espace mesurable $(\Omega, \mathcal{A})$, est dite étagée si elle ne prend qu'un nombre fini de valeurs $\{a_i\}_{i \leq n}$ qu'on peut supposer distinctes et est de la forme

$$\forall \omega \in \Omega, \quad f(\omega) = \sum_{i=1}^n a_i \mathbf{1}_{A_i}(\omega) \quad \text{avec} \quad A_i = f^{-1}(\{a_i\}) \in \mathcal{A}. \quad (\text{A.5})$$

On notera $\mathcal{E}^+$ l'ensemble des fonctions étagées positives.

En particulier, on peut approcher une fonction mesurable par des fonctions étagées.

Proposition A.11. Soit f une fonction mesurable positive sur $(\Omega, \mathcal{A})$. Il existe une suite croissante $\{f_n\}$ de fonctions étagées positives convergeant vers f .

Démonstration. Soit n un entier. Il suffit de tronquer f en fonction de "lignes" de niveau qui par construction sont mesurables dans $\mathcal{A}$

$$A_n = \{\omega \in \Omega, f(\omega) \geq n\}, \quad A_n^{(i)} = \left\{ \omega \in \Omega, \quad \frac{i-1}{n^2} \leq f(\omega) < \frac{i}{n^2} \right\} \quad \text{avec } i \leq n^2.$$

Les fonctions étagées de la forme

$$\omega \in \Omega, \quad f_n(\omega) = \sum_{i=1}^{n^2} \frac{i-1}{n^2} \mathbf{1}_{A_n^{(i)}}(\omega) + \mathbf{1}_{A_n}(\omega)$$

convergent vers f de façon croissante quand n tend vers l'infini. □

Opérations sur les fonctions mesurables

On admettra les propriétés suivantes des fonctions mesurables.

Proposition A.12.

- (i) La composition de 2 fonctions mesurables est une fonction mesurable.
- (ii) Soient f, g deux fonctions mesurables de $(\Omega, \mathcal{A})$ dans $(\mathbb{R}, \mathcal{B}_\mathbb{R})$ alors les fonctions

$$f + g, \quad fg, \quad f^+ = \sup\{f, 0\}, \quad f^- = -\inf\{f, 0\}$$

sont aussi mesurables.

On rappelle les notions de limite supérieure (notée $\limsup$) et de limite inférieure (notée $\liminf$) pour une suite réelle $(x_n)_{n \geq 1}$ de $\mathbb{R}$

$$\begin{aligned}\limsup_{n \rightarrow \infty} x_n &:= \lim_{n \rightarrow \infty} \left(\sup_{k \geq n} x_k \right) = \inf_{n \geq 1} \left(\sup_{k \geq n} x_k \right), \\ \liminf_{n \rightarrow \infty} x_n &:= \lim_{n \rightarrow \infty} \left(\inf_{k \geq n} x_k \right) = \sup_{n \geq 1} \left(\inf_{k \geq n} x_k \right).\end{aligned}\tag{A.6}$$

Proposition A.13. Soit $\{f_n\}_{n \geq 0}$ une suite de fonctions mesurables de $(\Omega, \mathcal{A})$ dans $\mathbb{R}$ alors

$$\sup_n f_n, \quad \inf_n f_n, \quad \limsup_n f_n, \quad \liminf_n f_n$$

sont aussi mesurables. Par conséquent, si f_n converge simplement, sa limite $\lim_n f_n$ est mesurable.

Démonstration. On traite le cas de $f = \sup_n f_n$. Il suffit d'analyser l'image réciproque d'intervalles ouverts du type $]x, +\infty[$

$$f^{-1}(]x, +\infty[) = \{\omega; \sup_n f_n(\omega) > x\} = \bigcup_n \{\omega; f_n(\omega) > x\}$$

qui appartiennent bien à $\mathcal{A}$ car il s'agit une réunion dénombrable d'ensembles mesurables. L'infimum se traite de manière analogue.

Par construction $\limsup_{n \rightarrow \infty} f_n = \inf_{n \geq 1} \left(\sup_{k \geq n} f_k \right)$ et les résultats sur le supremum et l'infimum de fonctions mesurables suffisent pour conclure. $\square$

A.2.2 Intégration des fonctions positives

Nous allons définir une notion d'intégration pour les fonctions sur un espace mesuré $(\Omega, \mathcal{A}, \mu)$. Étant donné μ , on ne sait mesurer que les ensembles A de $\mathcal{A}$. On définit donc

$$\int \mathbf{1}_A d\mu = \mu(A).$$

Ensuite, il est naturel de construire l'intégrale à partir des fonctions étagées de la forme (A.5)

$$\forall \omega \in \Omega, \quad f(\omega) = \sum_{i=1}^n a_i \mathbf{1}_{A_i}(\omega) \quad \text{avec} \quad A_i = f^{-1}(\{a_i\}) \in \mathcal{A}. \tag{A.7}$$

Définition A.14. (Intégrale de fonctions étagées) Soit $f \in \mathcal{E}^+$ une fonction étagée à valeurs dans $\mathbb{R}^+$ de la forme (A.7). On définit l'intégrale de f par rapport à μ comme

$$\int f d\mu = \int f(\omega) d\mu(\omega) := \sum_{i=1}^n a_i \mu(A_i),$$

avec la convention $0 \times \infty = \infty \times 0 = 0$.

Par construction, on en déduit la linéarité de l'intégrale pour les fonctions dans $\mathcal{E}^+$

$$\int (af + bg) d\mu = a \int f d\mu + b \int g d\mu, \quad (\text{A.8})$$

et une propriété de monotonie

$$f \leq g \quad \Rightarrow \quad \int f d\mu \leq \int g d\mu. \quad (\text{A.9})$$

Cette première définition permet de construire l'intégration pour toutes les fonctions positives et mesurables sur $\mathcal{A}$.

Définition A.15. (*Intégrale de fonctions positives*) Soit une fonction mesurable $f : \Omega \rightarrow [0, \infty]$, alors son intégrale par rapport à μ est définie par

$$\int f d\mu = \int f(\omega) d\mu(\omega) := \sup_{\substack{g \in \mathcal{E}^+ \\ g \leq f}} \int g d\mu.$$

Un premier théorème de convergence se déduit facilement de cette définition.

Théorème A.16. (*Convergence monotone*) Soit $\{f_n\}_{n \geq 0}$ une suite croissante de fonctions mesurables, positives, convergeant vers la limite $f = \lim_{n \rightarrow \infty} \uparrow f_n$. Alors

$$\int f d\mu = \lim_{n \rightarrow \infty} \uparrow \int f_n d\mu.$$

Le symbole $\uparrow$ sert simplement à insister sur la croissance des suites.

Démonstration. La propriété de monotonie (A.9) pour les intégrales de fonctions étagées s'étend immédiatement aux fonctions mesurables. Comme $f \geq f_n$, on en déduit ainsi une première inégalité

$$\int f d\mu \geq \lim_{n \rightarrow \infty} \uparrow \int f_n d\mu.$$

Pour établir l'inégalité inverse, nous allons montrer que pour toute fonction étagée $g = \sum_{i=1}^k a_i \mathbf{1}_{A_i}$ de $\mathcal{E}^+$ vérifiant $g \leq f$, on a

$$\lim_{n \rightarrow \infty} \uparrow \int f_n d\mu \geq \int g d\mu. \quad (\text{A.10})$$

Il suffira ensuite de prendre le supremum sur les fonctions g pour retrouver $\int f d\mu$ et conclure ainsi que $\int f d\mu \leq \lim_{n \rightarrow \infty} \uparrow \int f_n d\mu$.

Pour démontrer (A.10), on se donne $c \in [0, 1[$ et on écrit

$$\int f_n d\mu \geq \int f_n \mathbf{1}_{\{f_n \geq cg\}} d\mu \geq c \int g \mathbf{1}_{\{f_n \geq cg\}} d\mu = c \sum_{i=1}^k a_i \mu(A_i \cap \{f_n \geq ca_i\}).$$

Comme f_n converge de façon monotone vers f , les ensembles $A_i \cap \{f_n \geq ca_i\}$ sont croissants en n et la propriété (A.3) implique la convergence de leurs mesures vers $\mu(A_i)$. On obtient ainsi

$$\lim_{n \rightarrow \infty} \uparrow \int f_n d\mu \geq c \sum_{i=1}^k a_i \mu(A_i) = c \int g d\mu.$$

Il suffit ensuite de faire tendre c vers 1 pour retrouver (A.10). □

Un second résultat de convergence s'obtient comme conséquence du théorème de convergence monotone.

Lemme A.17. (Fatou) *Pour toute suite de fonctions $\{f_n\}_{n \geq 0}$ mesurables positives, on a*

$$\int (\liminf_{n \rightarrow \infty} f_n) d\mu \leq \liminf_{n \rightarrow \infty} \int f_n d\mu.$$

Démonstration. Pour tout $\ell \geq n$, on déduit de la propriété de monotonie (A.9) que

$$\inf_{k \geq n} f_k \leq f_\ell \quad \Rightarrow \quad \int \inf_{k \geq n} f_k d\mu \leq \int f_\ell d\mu.$$

En particulier, on obtient

$$\int \inf_{k \geq n} f_k d\mu \leq \inf_{k \geq n} \int f_k d\mu.$$

En utilisant la définition de la limite inférieure (A.6)

$$\liminf_{n \rightarrow \infty} f_n := \lim_{n \rightarrow \infty} \uparrow \left(\inf_{k \geq n} f_k \right)$$

et le théorème A.16 de convergence monotone, on en déduit

$$\lim_{n \rightarrow \infty} \uparrow \int \inf_{k \geq n} f_k d\mu = \int \lim_{n \rightarrow \infty} \uparrow \left(\inf_{k \geq n} f_k \right) d\mu = \int \liminf_{n \rightarrow \infty} f_n d\mu.$$

Ceci prouve le lemme A.17. □

A.2.3 Fonctions intégrables

Pour le moment, l'intégrale de Lebesgue n'a été définie que pour des fonctions mesurables positives. La définition suivante permet de généraliser l'intégration aux fonctions à valeurs dans $\mathbb{R}$.

Définition A.18. (Intégrale) *Soit une fonction mesurable $f : \Omega \rightarrow \mathbb{R}$, on dit que f est intégrable par rapport à μ si*

$$\int |f| d\mu < \infty.$$

On définit alors son intégrale par

$$\int f d\mu = \int f^+ d\mu - \int f^- d\mu,$$

où les fonctions $f^+ = \sup\{f, 0\}$, $f^- = -\inf\{f, 0\}$ sont mesurables (Proposition A.12) et positives.

On vérifiera en particulier que la linéarité (A.8) est aussi préservée pour des fonctions intégrables f, g

$$\int (af + bg) d\mu = a \int f d\mu + b \int g d\mu. \quad (\text{A.11})$$

La classe des fonctions mesurables intégrables s'avère un peu trop large car l'intégration est insensible aux valeurs prises par les fonctions sur les ensembles négligeables

introduits dans la définition A.6. On dira que deux fonctions mesurables f, g coïncident μ presque partout si

$$\mu(\{\omega \in \Omega, f(\omega) \neq g(\omega)\}) = 0.$$

On notera

$$f = g \quad \mu\text{-p.p.} \quad \text{ou} \quad f(\omega) = g(\omega) \quad \mu(d\omega)\text{-p.p.}$$

pour spécifier la variable d'intégration. La terminologie μ presque sûrement (notée $\mu\text{-p.s.}$) est équivalente et souvent employée dans le cadre des probabilités. Si f, g sont intégrables et coïncident μ presque partout, on a alors

$$\int f d\mu = \int g d\mu.$$

Ainsi, on notera $\mathbb{L}^1 = \mathbb{L}^1(\mathcal{A}, \mu)$ l'ensemble des fonctions intégrables sur l'espace mesuré $(\Omega, \mathcal{A}, \mu)$ et on identifiera à un même élément toutes les fonctions égales μ presque partout.

Pour conclure ce paragraphe, nous rappelons 2 résultats classiques.

Lemme A.19. Soit f une fonction positive, mesurable, alors

$$\int f d\mu = 0 \quad \Leftrightarrow \quad f = 0 \quad \mu\text{-p.p.}$$

Démonstration. Il suffit de démontrer l'implication $\Rightarrow$. Pour cela, on définit la suite d'ensembles croissants

$$A_n = \{\omega \in \Omega, f(\omega) \geq \frac{1}{n}\}.$$

Ces ensembles sont tous de mesure nulle car

$$\mu(A_n) = \int \mathbf{1}_{A_n} d\mu \leq n \int f \mathbf{1}_{A_n} d\mu \leq n \int f d\mu = 0.$$

Pour conclure la preuve du lemme A.19, il suffit d'appliquer la propriété (A.3) à cette suite d'évènements croissants

$$\mu\left(\{\omega \in \Omega, f(\omega) > 0\}\right) = \mu\left(\bigcup_{n \geq 0} A_n\right) = \lim_{n \rightarrow \infty} \mu(A_n) = 0.$$

On en déduit que f est nulle μ presque partout. □

Lemme A.20. Soit f une fonction intégrable dans $\mathbb{L}^1(\mathcal{A}, \mu)$ et $\varepsilon > 0$ une constante fixée. Alors, il existe $\delta > 0$ tel que pour tout ensemble $A \in \mathcal{A}$ vérifiant $\mu(A) < \delta$, on a

$$\int |f| \mathbf{1}_A d\mu < \varepsilon.$$

Démonstration. On se restreint au cas $f \geq 0$ et pour tout entier n , on considère la fonction bornée $f_n = \min(f, n)$. L'inégalité suivante est valable pour tout ensemble $A \in \mathcal{A}$

$$\mu(f \mathbf{1}_A) = \mu(f_n \mathbf{1}_A) + \mu((f - f_n) \mathbf{1}_A) \leq n\mu(A) + \mu(f - f_n).$$

Comme f_n converge vers f de façon croissante, le théorème A.16 de convergence monotone implique qu'il existe n assez grand pour que

$$\mu(f - f_n) \leq \frac{\varepsilon}{2}.$$

Il ne reste plus qu'à choisir $\delta \leq \frac{\varepsilon}{2n}$ pour conclure la démonstration du lemme A.20. □

A.2.4 Théorème de convergence dominée

Cette section est dédiée au théorème de convergence dominée et à son application aux intégrales qui dépendent d'un paramètre.

Théorème A.21 (Convergence dominée). *Soit $\{f_n\}_{n \geq 0}$ une suite de fonctions dans $\mathbb{L}^1(\mathcal{A}, \mu)$ vérifiant les deux hypothèses suivantes :*

1. *Il existe une fonction mesurable f telle que f_n converge vers f , μ presque partout, quand n tend vers l'infini.*
2. *Il existe une fonction g positive dans $\mathbb{L}^1(\mathcal{A}, \mu)$ telle que pour tout n , l'inégalité $|f_n| \leq g$ est vraie μ presque partout.*

Alors f appartient aussi à $\mathbb{L}^1(\mathcal{A}, \mu)$ et f_n converge vers f dans $\mathbb{L}^1(\mathcal{A}, \mu)$, c'est-à-dire

$$\lim_{n \rightarrow \infty} \int |f_n - f| d\mu = 0. \quad (\text{A.12})$$

En particulier, on en déduit que

$$\lim_{n \rightarrow \infty} \int f_n d\mu = \int f d\mu. \quad (\text{A.13})$$

Démonstration. Comme $|f_n| \leq g$ pour tout n , l'inégalité $|f| \leq g$ se déduit par passage à la limite et est aussi vraie μ presque partout. La limite f appartient donc à $\mathbb{L}^1(\mathcal{A}, \mu)$. Ces deux inégalités impliquent aussi que la fonction $2g - |f - f_n|$ est positive. En appliquant l'hypothèse (1), on remarque que

$$\liminf_{n \rightarrow \infty} (2g - |f - f_n|) = 2g.$$

D'après le lemme A.17 de Fatou, on en déduit que

$$\liminf_{n \rightarrow \infty} \int (2g - |f - f_n|) d\mu \geq 2 \int g d\mu.$$

L'intégrale étant linéaire, on obtient

$$2 \int g d\mu - \limsup_{n \rightarrow \infty} \int |f - f_n| d\mu \geq 2 \int g d\mu \quad \Rightarrow \quad \limsup_{n \rightarrow \infty} \int |f - f_n| d\mu \leq 0.$$

Ceci implique la convergence (A.12) dans $\mathbb{L}^1(\mathcal{A}, \mu)$. La limite (A.13) se déduit par l'inégalité

$$\left| \int f_n - f d\mu \right| \leq \int |f_n - f| d\mu.$$

□

Soit $\mathcal{U}$ un intervalle ouvert de $\mathbb{R}$. Le théorème de convergence dominée permet de contrôler des intégrales de la forme

$$F(u) = \int f(u, \omega) d\mu(\omega),$$

où $f : \mathcal{U} \times \Omega \rightarrow \mathbb{R}$ est une application telle que

$$\text{pour tout } u \in \mathcal{U} \text{ fixé, } \omega \mapsto f(u, \omega) \text{ est mesurable et appartient à } \mathbb{L}^1(\mathcal{A}, \mu). \quad (\text{A.14})$$

La régularité de f confère différentes propriétés à la fonction $u \mapsto F(u)$.

Théorème A.22. Soit $f : \mathcal{U} \times \Omega \rightarrow \mathbb{R}$ une fonction vérifiant (A.14) et $u_0 \in \mathcal{U}$.

1. Supposons que

- (a) μ presque partout en ω , l'application $u \mapsto f(u, \omega)$ est continue en u_0 ;
- (b) il existe une fonction $g \in \mathbb{L}^1(\mathcal{A}, \mu)$ telle que pour tout $u \in \mathcal{U}$

$$|f(u, \omega)| \leq g(\omega), \quad \mu \text{ presque partout en } \omega.$$

Alors la fonction $u \mapsto F(u) = \int f(u, \omega) d\mu(\omega)$ est continue en u_0 .

2. Supposons que

- (a) μ presque partout en ω , l'application $u \mapsto f(u, \omega)$ est dérivable sur $\mathcal{U}$;
- (b) il existe une fonction $g \in \mathbb{L}^1(\mathcal{A}, \mu)$ telle que pour tout $u \in \mathcal{U}$

$$\left| \frac{\partial f}{\partial u}(u, \omega) \right| \leq g(\omega), \quad \mu \text{ presque partout en } \omega.$$

Alors la fonction $u \mapsto F(u) = \int f(u, \omega) d\mu(\omega)$ est dérivable sur $\mathcal{U}$ de dérivée

$$F'(u) = \int \frac{\partial f}{\partial u}(u, \omega) d\mu(\omega).$$

Démonstration. Nous allons d'abord établir le résultat de continuité (1). Pour toute suite $\{u_n\}_{n \geq 1}$ convergeant vers u_0 , l'hypothèse (1-a) implique la convergence

$$\lim_{n \rightarrow \infty} f(u_n, \omega) = f(u_0, \omega), \quad \mu \text{ presque partout en } \omega.$$

L'hypothèse (1-b) permet d'utiliser le théorème de convergence dominée et d'en déduire la continuité en u_0 de F

$$\lim_{n \rightarrow \infty} \int f(u_n, \omega) d\mu(\omega) = \int f(u_0, \omega) d\mu(\omega).$$

Nous allons maintenant établir la dérivabilité (2) en un point u_0 de $\mathcal{U}$. On remarque que pour toute suite $\{u_n\}_{n \geq 1}$ convergeant vers u_0 , l'hypothèse (2-a) implique la convergence

$$\lim_{n \rightarrow \infty} \frac{f(u_n, \omega) - f(u_0, \omega)}{u_n - u_0} = \frac{\partial f}{\partial u}(u_0, \omega), \quad \mu \text{ presque partout en } \omega. \quad (\text{A.15})$$

Par le théorème des accroissements finis, on déduit de l'hypothèse (2-b) la borne uniforme suivante

$$|f(u, \omega) - f(u_0, \omega)| \leq g(\omega)|u - u_0|, \quad \mu \text{ presque partout en } \omega.$$

Cette inégalité et la convergence presque partout (A.15) suffisent à déduire du théorème de convergence dominée que

$$F'(u_0) = \lim_{n \rightarrow \infty} \frac{F(u_n) - F(u_0)}{u_n - u_0} = \int \lim_{n \rightarrow \infty} \frac{f(u_n, \omega) - f(u_0, \omega)}{u_n - u_0} d\mu(\omega) = \int \frac{\partial f}{\partial u}(u_0, \omega) d\mu(\omega).$$

Ceci conclut la deuxième assertion du théorème A.22. $\square$

A.3 Espaces produits

Dans cette section, nous allons aborder les aspects spécifiques de l'intégration sur des espaces produits. Ces espaces joueront un rôle important pour formaliser l'indépendance de variables aléatoires.

A.3.1 Mesurabilité sur les espaces produits

Soient 2 espaces mesurables $(\Omega_1, \mathcal{A}_1)$ et $(\Omega_2, \mathcal{A}_2)$. On associe à l'espace produit $\Omega_1 \times \Omega_2$, la σ -algèbre produit définie par

$$\mathcal{A}_1 \otimes \mathcal{A}_2 = \sigma(A_1 \times A_2; \quad A_1 \in \mathcal{A}_1, A_2 \in \mathcal{A}_2).$$

Il s'agit de la plus petite σ -algèbre engendrée par les pavés mesurables, c'est-à-dire les ensembles de la forme $A_1 \times A_2$. Si Ω_1, Ω_2 sont des espaces métriques mesurables, on peut montrer que la σ -algèbre borélienne $\mathcal{B}_{\Omega_1 \times \Omega_2}$ coïncide avec la σ -algèbre produit $\mathcal{B}_{\Omega_1} \otimes \mathcal{B}_{\Omega_2}$. En particulier, la σ -algèbre borélienne $\mathcal{B}_{\mathbb{R}^2}$ s'obtient simplement à partir d'ensembles de $\mathcal{B}_{\mathbb{R}}$. Plus généralement pour un nombre fini d'espaces mesurables $(\Omega_i, \mathcal{A}_i)$ avec $i \leq k$, la tribu produit est construite à partir des pavés

$$\mathcal{A}_1 \otimes \cdots \otimes \mathcal{A}_k = \sigma(A_1 \times \cdots \times A_k; \quad i \leq k, \quad A_i \in \mathcal{A}_i).$$

Inversement, à partir d'un ensemble C de $\mathcal{A}_1 \otimes \mathcal{A}_2$, on peut retrouver des sous ensembles de Ω_1 et Ω_2 par projection. Soient ω_1, ω_2 dans Ω_1, Ω_2 , on note

$$C_{\omega_1} = \{\omega_2 \in \Omega_2, \quad (\omega_1, \omega_2) \in C\} \quad \text{et} \quad C^{\omega_2} = \{\omega_1 \in \Omega_1, \quad (\omega_1, \omega_2) \in C\}.$$

La proposition suivante montre que ces ensembles projetés sont mesurables.

Proposition A.23.

1. Étant donné C dans $\mathcal{A}_1 \otimes \mathcal{A}_2$, alors pour tout ω_1 dans Ω_1 , la projection C_{ω_1} appartient à $\mathcal{A}_2$. De même, pour tout ω_2 dans Ω_2 , la projection C^{ω_2} est dans $\mathcal{A}_1$.
2. Plus généralement, supposons que $f : \Omega_1 \times \Omega_2 \rightarrow E$ soit mesurable pour $\mathcal{A}_1 \otimes \mathcal{A}_2$. Alors, pour tout ω_1 dans Ω_1 , l'application $\omega_2 \mapsto f(\omega_1, \omega_2)$ est mesurable pour $\mathcal{A}_2$. De plus, pour tout ω_2 dans Ω_2 , l'application $\omega_1 \mapsto f(\omega_1, \omega_2)$ est mesurable pour $\mathcal{A}_1$.

Démonstration. Étant donné ω_1 dans Ω_1 , la famille d'ensembles

$$\mathcal{C} = \{C \in \mathcal{A}_1 \otimes \mathcal{A}_2, \quad C_{\omega_1} \in \mathcal{A}_2\},$$

contient tous les pavés de la forme $A_1 \times A_2$ de $\mathcal{A}_1 \otimes \mathcal{A}_2$. De plus $\mathcal{C}$ est une σ -algèbre. Par conséquent, $\mathcal{C}$ coïncide avec $\mathcal{A}_1 \otimes \mathcal{A}_2$ et on obtient la mesurabilité des ensembles de la forme C_{ω_1} .

Pour étudier la mesurabilité de l'application $\omega_2 \mapsto f(\omega_1, \omega_2)$, on considère l'image réciproque d'un ensemble U de E

$$\{\omega_2 \in \Omega_2, \quad f(\omega_1, \omega_2) \in U\} = (f^{-1}(U))_{\omega_1}.$$

Comme cet ensemble est la projection de l'ensemble mesurable $f^{-1}(U)$, on déduit de (1) qu'il est mesurable dans $\mathcal{A}_2$. Ceci implique la mesurabilité de $\omega_2 \mapsto f(\omega_1, \omega_2)$. $\square$

On peut maintenant définir une mesure sur des espaces produits.

Théorème A.24. Soient $(\Omega_1, \mathcal{A}_1, \mu_1)$ et $(\Omega_2, \mathcal{A}_2, \mu_2)$ des espaces mesurés avec μ_1, μ_2 deux mesures σ -finies. Il existe une unique mesure, notée $\mu_1 \otimes \mu_2$, sur l'espace produit $(\Omega_1 \otimes \Omega_2, \mathcal{A}_1 \otimes \mathcal{A}_2)$ vérifiant

$$\forall A_1 \in \mathcal{A}_1, \forall A_2 \in \mathcal{A}_2, \quad \mu_1 \otimes \mu_2(A_1 \times A_2) = \mu_1(A_1) \mu_2(A_2). \quad (\text{A.16})$$

La mesure d'un ensemble C de $\mathcal{A}_1 \otimes \mathcal{A}_2$ est alors donnée par

$$\mu_1 \otimes \mu_2(C) = \int_{\Omega_1} \mu_2(C_{\omega_1}) d\mu_1(\omega_1) = \int_{\Omega_2} \mu_1(C^{\omega_2}) d\mu_2(\omega_2). \quad (\text{A.17})$$

Nous ne détaillerons pas la preuve qui repose sur le théorème A.7 d'extension de Carathéodory, mais son principe est simple. Une fois définie sur les pavés $A_1 \times A_2$ par la relation (A.16), la mesure $\mu_1 \otimes \mu_2$ est entièrement prescrite sur $\mathcal{A}_1 \otimes \mathcal{A}_2$ car ceux-ci engendrent cette σ -algèbre. Par la proposition A.23, la projection C_{ω_1} est dans $\mathcal{A}_2$ et on peut aussi vérifier que $\omega_1 \mapsto \mu_2(C_{\omega_1})$ est mesurable dans $\mathcal{A}_1$. Ceci permet de définir la mesure

$$\forall C \in \mathcal{A}_1 \otimes \mathcal{A}_2, \quad m(C) = \int_{\Omega_1} \mu_2(C_{\omega_1}) d\mu_1(\omega_1).$$

Sur les pavés $C = A_1 \times A_2$, cette mesure coïncide avec $\mu_1 \otimes \mu_2$

$$m(C) = \mu_1(A_1) \mu_2(A_2) = \mu_1 \otimes \mu_2(C).$$

Par conséquent, m doit aussi coïncider avec $\mu_1 \otimes \mu_2$ sur tout $\mathcal{A}_1 \otimes \mathcal{A}_2$ car l'extension sur $\mathcal{A}_1 \otimes \mathcal{A}_2$ est unique par le théorème A.7. Ceci justifie la relation (A.17) qui permet de décomposer $\mu_1 \otimes \mu_2$ à partir des projections sur $\mathcal{A}_1$ ou $\mathcal{A}_2$.

A.3.2 Intégration sur les espaces produits

La mesure produit permet de définir une notion d'intégration sur les espaces produits comme dans la définition A.15.

Théorème A.25. (Fubini) Soit f une fonction mesurable dans $\mathbb{L}^1(\Omega_1 \times \Omega_2, \mathcal{A}_1 \otimes \mathcal{A}_2)$, alors ses projections sur Ω_1 et Ω_2 sont intégrables :

l'application $\omega_2 \mapsto f(\omega_1, \omega_2)$ est dans $\mathbb{L}^1(\Omega_2, \mathcal{A}_2)$ pour μ_1 presque tout ω_1 ,

l'application $\omega_1 \mapsto f(\omega_1, \omega_2)$ est dans $\mathbb{L}^1(\Omega_1, \mathcal{A}_1)$ pour μ_2 presque tout ω_2 .

De plus, l'ordre des intégrations peut être permuté

$$\begin{aligned} \int f d\mu_1 \otimes \mu_2 &= \int_{\Omega_1} \left(\int_{\Omega_2} f(\omega_1, \omega_2) d\mu_2(\omega_2) \right) d\mu_1(\omega_1) \\ &= \int_{\Omega_2} \left(\int_{\Omega_1} f(\omega_1, \omega_2) d\mu_1(\omega_1) \right) d\mu_2(\omega_2). \end{aligned} \quad (\text{A.18})$$

En général, l'inversion des intégrales (A.18) n'est vraie que si f appartient à $\mathbb{L}^1$, c'est-à-dire sous l'hypothèse $\int |f| d\mu_1 \otimes \mu_2 < \infty$. Cependant dans le cas d'une fonction f positive, l'identité (A.18) est valable sans condition : si une des intégrales est infinie les 2 autres le seront aussi.

Une démonstration du théorème A.25 de Fubini peut être trouvée page 77 du livre [28]. Le schéma de la preuve est identique à celui suivi pour construire l'intégration. Pour des fonctions étagées l'inversion des intégrales (A.18) est une conséquence de l'identité (A.17) sur des ensembles de $\mathcal{A}_1 \otimes \mathcal{A}_2$. Ensuite l'intégrale de fonctions positives s'obtient par approximation par des fonctions étagées (cf. définition A.15). Finalement, les fonctions arbitraires se décomposent sous la forme $f = f^+ - f^-$ pour se ramener au cas des fonctions positives.

Annexe B

Rappels sur la théorie des probabilités

Ce chapitre reprend quelques éléments de théorie des probabilités en insistant sur les liens avec la théorie de la mesure présentée dans l'annexe A.

B.1 Variables aléatoires

La théorie des probabilités est une vaste escroquerie qui consiste à recycler des notions de théorie de la mesure à l'aide d'une nouvelle terminologie. Ainsi un espace de probabilité $(\Omega, \mathcal{A}, \mathbb{P})$ n'est rien d'autre qu'un espace mesurable $(\Omega, \mathcal{A})$ muni d'une mesure $\mathbb{P} : \mathcal{A} \rightarrow [0, 1]$ de masse totale 1, i.e. $\mathbb{P}(\Omega) = 1$. Les variables aléatoires sont simplement des fonctions mesurables sur Ω .

Définition B.1 (Variable aléatoire). Soient $(\Omega, \mathcal{A})$ et $(E, \mathcal{E})$ deux espaces mesurables. Une variable aléatoire X à valeurs dans E est une application mesurable de Ω dans E . Si X est à valeurs dans $\mathbb{R}$, son intégrale est appelée espérance et notée

$$\mathbb{E}(X) = \int_{\Omega} X(\omega) d\mathbb{P}(\omega).$$

L'espérance s'interprète comme la valeur moyenne de la variable X . En particulier, si $X = \mathbf{1}_C$ alors la probabilité de l'évènement C vaut $\mathbb{P}(C) = \mathbb{E}(\mathbf{1}_C)$.

Ce changement de terminologie s'accompagne aussi d'un changement de point de vue qui met l'accent sur la notion de loi.

Définition B.2 (Loi d'une variable aléatoire). Soit X une variable aléatoire de $(\Omega, \mathcal{A})$ dans $(E, \mathcal{E})$. La loi $\mathbb{P}_X$ de X est la mesure image de $\mathbb{P}$ par X . Il s'agit d'une mesure de probabilité sur $(E, \mathcal{E})$ associant à tout ensemble C dans $\mathcal{E}$ la mesure

$$\mathbb{P}_X(C) = \mathbb{P}(X^{-1}(C)) = \mathbb{P}(X \in C) = \mathbb{P}(\{\omega \in \Omega, \quad X(\omega) \in C\}).$$

Pour illustrer ce concept, on considère un dé à 6 faces dont la valeur est donnée par $X \in \{1, \dots, 6\}$. On pourrait coder par une variable ω tous les paramètres associés à

la condition initiale du lancer du dé et déduire de la mécanique Newtonienne la valeur $X(\omega)$ à l'issue du lancer. Une telle approche supposerait de mesurer un nombre gigantesque de paramètres dans un espace Ω et de déterminer précisément l'application $\omega \rightarrow X(\omega)$. La notion de loi permet de s'abstraire de cette description en oubliant le détail du mouvement et en définissant uniquement une probabilité $\mathbb{P}_X$ sur $\{1, \dots, 6\}$. On ne cherchera jamais à préciser l'application $\omega \rightarrow X(\omega)$ et on se contentera de quantifier des événements simples du type $\mathbb{P}_X(\{1\}) = \mathbb{P}(X = 1)$.

Si $f : E \mapsto \mathbb{R}$ est une application mesurable et $X : \Omega \mapsto E$ une variable aléatoire, alors $f(X)$ est aussi une variable aléatoire (cf. le point (i) de la proposition A.12). Son espérance est donnée en fonction de la loi de X

$$\mathbb{E}(f(X)) = \int_E f(x) d\mathbb{P}_X(x).$$

σ -algèbre engendrée par une variable aléatoire

Une variable aléatoire X sur un espace mesurable $(\Omega, \mathcal{A})$ peut ne dépendre que d'une information partielle sur Ω . Par exemple, Ω peut coder plusieurs lancers d'un dé et $X(\omega)$ est simplement le résultat du premier lancer. Dans ce cas, seulement une partie des ensembles de $\mathcal{A}$ suffira à rendre compte de X . Plus généralement, on peut associer à X la plus petite σ -algèbre de $\mathcal{A}$ qui rende X mesurable.

Définition B.3. Soit X une variable aléatoire de $(\Omega, \mathcal{A})$ dans $(E, \mathcal{E})$. La σ -algèbre engendrée par X est définie par

$$\sigma(X) = \{A = X^{-1}(B); \quad B \in \mathcal{E}\} \subset \mathcal{A}.$$

Pour une famille de variables aléatoires $\{X_i\}_{i \in I}$ de $(\Omega, \mathcal{A})$ dans $(E_i, \mathcal{E}_i)_{i \in I}$ (avec I éventuellement dénombrable), la σ -algèbre engendrée est définie par

$$\sigma(\{X_i\}_{i \in I}) = \sigma(X_i^{-1}(B_i); \quad B_i \in \mathcal{E}_i, i \in I) \subset \mathcal{A}.$$

Le lemme suivant permet de décrire toutes les variables aléatoires mesurables par rapport à $\sigma(X)$, i.e. à l'information relative à la variable aléatoire X .

Lemme B.4. On considère deux variables aléatoires sur $(\Omega, \mathcal{A}, \mathbb{P})$: X à valeurs dans $(E, \mathcal{E})$ et Y à valeurs dans $\mathbb{R}$. Alors Y est $\sigma(X)$ -mesurable si et seulement s'il existe une fonction mesurable $f : E \mapsto \mathbb{R}$ telle que $Y = f(X)$.

Démonstration. Si $Y = f(X)$ alors Y est $\sigma(X)$ -mesurable d'après (i) de la Proposition A.12.

Réciproquement supposons que Y soit $\sigma(X)$ -mesurable. Si $Y = \sum_{i=1}^n a_i \mathbf{1}_{A_i}$ est une fonction étagée, alors les ensembles A_i sont de la forme $X^{-1}(B_i)$ car ils appartiennent à $\sigma(X)$. On en déduit que $Y = f(X)$ où f est la fonction $\mathcal{E}$ -mesurable définie par

$$\forall x \in E, \quad f(x) = \sum_{i=1}^n a_i \mathbf{1}_{B_i}(x).$$

En général, on peut approcher la variable aléatoire Y par une suite de variables $\{Y_n\}$ étagées et $\sigma(X)$ -mesurables. Pour cela, il suffit de décomposer $Y = Y^+ - Y^-$ et d'appliquer la proposition A.11. L'étape précédente permet de construire $\{f_n\}$, une suite de fonctions $\mathcal{E}$ -mesurables telles que $Y_n = f_n(X)$. Pour tout ω , on a alors la convergence vers une fonction f telle que $f(X(\omega)) = \lim_{n \rightarrow \infty} f_n(X(\omega)) = Y(\omega)$. Ceci conclut le Lemme B.4. $\square$

B.2 Intégration de variables aléatoires et espaces $\mathbb{L}^p$

Les résultats établis dans l'annexe A sur l'intégration des fonctions mesurables se transposent naturellement aux variables aléatoires.

B.2.1 Théorèmes de convergence

Les théorèmes fondamentaux de convergence sont rappelés ci-dessous dans le formalisme probabiliste.

Soit $\{X_n\}_{n \geq 1}$ une suite de variables aléatoires à valeurs dans $\mathbb{R}$ qui converge vers X presque sûrement.

Théorème de convergence monotone :

Si X_n converge de façon croissante vers X et $X_n \geq 0$ alors

$$\lim_{n \rightarrow \infty} \mathbb{E}(X_n) = \mathbb{E}(X). \quad (\text{B.1})$$

Théorème de convergence dominée :

Supposons qu'il existe une variable aléatoire Y intégrable, i.e. $\mathbb{E}(|Y|) < \infty$, telle que $|X_n| \leq Y$ pour tout n , on a alors

$$\lim_{n \rightarrow \infty} \mathbb{E}(|X_n - X|) = 0. \quad (\text{B.2})$$

Lemme de Fatou :

Si $X_n \geq 0$ alors

$$\liminf_{n \rightarrow \infty} \mathbb{E}(X_n) \geq \mathbb{E}\left(\liminf_{n \rightarrow \infty} X_n\right) = \mathbb{E}(X). \quad (\text{B.3})$$

Dérivation et espérance :

Soit X une variable aléatoire à valeurs dans $\mathbb{R}$ et f une fonction telle que

$$\forall u \in \mathcal{U} \subset \mathbb{R}, \quad \mathbb{E}(|f(u, X)|) < \infty.$$

S'il existe une fonction g vérifiant

$$\forall (u, x) \in \mathcal{U} \times \mathbb{R}, \quad \left| \partial_u f(u, x) \right| \leq g(x) \quad \text{avec} \quad \mathbb{E}(g(X)) < \infty,$$

alors la fonction $u \rightarrow \mathbb{E}(f(u, X))$ est différentiable et

$$\forall u \in \mathcal{U}, \quad \partial_u \mathbb{E}(f(u, X)) = \mathbb{E}(\partial_u f(u, X)). \quad (\text{B.4})$$

B.2.2 Inégalités classiques

Cette section rappelle quelques inégalités importantes en théorie de l'intégration.

Proposition B.5 (Inégalité de Hölder). Soient X, Y deux variables aléatoires à valeurs dans $\mathbb{R}$ telles que

$$\mathbb{E}(|X|^p) < \infty \quad \text{et} \quad \mathbb{E}(|Y|^q) < \infty \quad \text{avec} \quad p > 1 \quad \text{et} \quad \frac{1}{p} + \frac{1}{q} = 1.$$

L'inégalité de Hölder s'écrit

$$\mathbb{E}(|XY|) \leq \mathbb{E}(|X|^p)^{1/p} \mathbb{E}(|Y|^q)^{1/q}. \quad (\text{B.5})$$

Démonstration. On vérifie facilement que

$$\forall u, v \geq 0, \quad u^{1/p} v^{1-1/p} \leq \frac{1}{p}u + \left(1 - \frac{1}{p}\right)v.$$

En appliquant cette inégalité à

$$U = \frac{|X(\omega)|^p}{\mathbb{E}(|X|^p)} \quad \text{et} \quad V = \frac{|Y(\omega)|^q}{\mathbb{E}(|Y|^q)}$$

puis en prenant l'espérance, on en déduit que

$$\frac{1}{\mathbb{E}(|X|^p)^{1/p}} \frac{1}{\mathbb{E}(|Y|^q)^{1/q}} \mathbb{E}(|X|^{p/p} |Y|^{q(1-1/p)}) \leq \frac{1}{p} + \left(1 - \frac{1}{p}\right) = 1,$$

car $\mathbb{E}(U) = \mathbb{E}(V) = 1$. Comme $q(1 - \frac{1}{p}) = 1$, on retrouve l'inégalité (B.5). $\square$

Proposition B.6 (Inégalité de Minkowski). Soient X, Y deux variables aléatoires à valeurs dans $\mathbb{R}$ telles que

$$\mathbb{E}(|X|^p) < \infty \quad \text{et} \quad \mathbb{E}(|Y|^p) < \infty \quad \text{avec} \quad p > 1.$$

L'inégalité de Minkowski s'écrit

$$\mathbb{E}(|X + Y|^p)^{1/p} \leq \mathbb{E}(|X|^p)^{1/p} + \mathbb{E}(|Y|^p)^{1/p}. \quad (\text{B.6})$$

Démonstration. En utilisant l'inégalité

$$|X + Y|^p \leq |X| |X + Y|^{p-1} + |Y| |X + Y|^{p-1},$$

on en déduit

$$\mathbb{E}(|X + Y|^p) \leq \mathbb{E}(|X| |X + Y|^{p-1}) + \mathbb{E}(|Y| |X + Y|^{p-1}).$$

On applique ensuite à chaque terme l'inégalité de Hölder (B.5) avec $q = \frac{p}{p-1}$

$$\begin{aligned} \mathbb{E}(|X + Y|^p) &\leq \mathbb{E}(|X|^p)^{1/p} \mathbb{E}(|X + Y|^{q(p-1)})^{1/q} + \mathbb{E}(|Y|^p)^{1/p} \mathbb{E}(|X + Y|^{q(p-1)})^{1/q} \\ &\leq \left(\mathbb{E}(|X|^p)^{1/p} + \mathbb{E}(|Y|^p)^{1/p} \right) \mathbb{E}(|X + Y|^p)^{1/q}. \end{aligned}$$

Ceci prouve l'inégalité de Minkowski (B.6). $\square$

Proposition B.7 (Inégalité de Markov). *Soit X une variable aléatoire à valeurs dans $\mathbb{R}$. Si $f : \mathbb{R} \mapsto \mathbb{R}_+$ est une fonction borélienne croissante et positive, alors l'inégalité de Markov s'écrit pour tout $c \in \mathbb{R}$*

$$f(c) \mathbb{P}(X \geq c) \leq \mathbb{E}(f(X)). \quad (\text{B.7})$$

En particulier, on en déduit que pour $c > 0$

$$\mathbb{P}(|X| \geq c) \leq \frac{1}{c} \mathbb{E}(|X|) \quad \text{et} \quad \mathbb{P}(|X| \geq c) \leq \frac{1}{c^2} \mathbb{E}(X^2). \quad (\text{B.8})$$

Démonstration. Comme f est positive et croissante, l'espérance de $f(X)$ est bornée inférieurement par

$$\mathbb{E}(f(X)) \geq \mathbb{E}(f(X) \mathbf{1}_{\{X \geq c\}}) \geq f(c) \mathbb{P}(X \geq c).$$

Ceci conclut la preuve de l'inégalité de Markov (B.7).

Les inégalités (B.8) sont une conséquence de (B.7) appliquée à $|X|$ avec

$$f(x) = |x| \quad \text{et} \quad f(x) = x^2.$$

□

Proposition B.8 (Inégalité de Jensen). *Soient une variable aléatoire X dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ et une fonction convexe $g : \mathbb{R}^d \mapsto \mathbb{R} \cup \{\infty\}$ telle que $\mathbb{E}(|g(X)|) < \infty$. L'inégalité de Jensen s'écrit*

$$\mathbb{E}(g(X)) \geq g(\mathbb{E}(X)). \quad (\text{B.9})$$

Démonstration. ¹ Étant donné $M > 0$, la fonction $g_M(x) = \sup\{g(x), -M\}$ est aussi convexe. Nous allons d'abord montrer l'inégalité de Jensen pour g_M

$$\mathbb{E}(g_M(X)) \geq g_M(\mathbb{E}(X)). \quad (\text{B.10})$$

L'inégalité (B.9) se déduit ensuite par passage à la limite $M \rightarrow \infty$, en utilisant le théorème de convergence dominée car $|g_M| \leq |g|$ et $|g(X)|$ est intégrable.

On considère une suite $\{X_i\}_{i \geq 1}$ de variables aléatoires indépendantes et identiquement distribuées de même loi que X . On obtient par la convexité de g_M que

$$\mathbb{E} \left(g_M \left(\frac{\sum_{i=1}^n X_i}{n} \right) \right) \leq \mathbb{E} \left(\frac{1}{n} \sum_{i=1}^n g_M(X_i) \right) = \mathbb{E}(g_M(X)).$$

La loi des grands nombres implique la convergence presque sûre

$$g_M(\mathbb{E}(X)) = \lim_{n \rightarrow \infty} g_M \left(\frac{\sum_{i=1}^n X_i}{n} \right).$$

Comme g_M est bornée inférieurement, le lemme de Fatou permet de conclure

$$g_M(\mathbb{E}(X)) = \mathbb{E} \left(\liminf_{n \rightarrow \infty} g_M \left(\frac{\sum_{i=1}^n X_i}{n} \right) \right) \leq \liminf_{n \rightarrow \infty} \mathbb{E} \left(g_M \left(\frac{\sum_{i=1}^n X_i}{n} \right) \right) \leq \mathbb{E}(g_M(X)).$$

□

1. Cette démonstration, due D. Chafaï, repose sur des arguments probabilistes. Une preuve plus classique se trouve page 61 de [28].

B.2.3 Espaces $\mathbb{L}^p$

La notion d'espace $\mathbb{L}^1$ introduite page 217 peut être généralisée. Pour tout p dans $[1, +\infty[$, on notera $\mathbb{L}^p = \mathbb{L}^p(\mathcal{A}, \mathbb{P})$ l'ensemble des variables aléatoires X sur l'espace de probabilité $(\Omega, \mathcal{A}, \mathbb{P})$ telles que

$$\|X\|_p = (\mathbb{E}(|X|^p))^{1/p} < \infty.$$

Pour $p = \infty$, l'espace $\mathbb{L}^\infty$ est l'ensemble des variables aléatoires X vérifiant

$$\|X\|_\infty = \inf \{C \in \mathbb{R}^+; |X| \leq C \text{ presque sûrement}\} < \infty.$$

Comme dans le cas de l'espace $\mathbb{L}^1$, on identifiera les variables aléatoires de $\mathbb{L}^p$ qui coïncident presque sûrement.

L'inégalité de Minkowski (B.6) n'est rien d'autre que l'inégalité triangulaire dans $\mathbb{L}^p$

$$\|X + Y\|_p \leq \|X\|_p + \|Y\|_p.$$

Les espaces $\mathbb{L}^p$ sont décroissants, c'est-à-dire que $\mathbb{L}^p \subset \mathbb{L}^r$ pour $1 \leq r \leq p \leq \infty$. En effet, en appliquant l'inégalité de Jensen (B.9), on montre que pour toute variable aléatoire X dans $\mathbb{L}^p$

$$\|X\|_r \leq \|X\|_p \quad \text{avec} \quad 1 \leq r \leq p \leq \infty. \quad (\text{B.11})$$

L'espace $\mathbb{L}^2$ joue un rôle privilégié car l'application $(X, Y) \mapsto \mathbb{E}[XY]$ définit un produit scalaire sur $\mathbb{L}^2$ et confère ainsi à $\mathbb{L}^2$ une structure hilbertienne. Dans ce cas particulier, l'inégalité de Hölder (B.5) correspond à l'inégalité de Schwarz. Ainsi pour toutes variables aléatoires X, Y dans $\mathbb{L}^2$, on a

$$|\mathbb{E}(XY)| \leq \|X\|_2 \|Y\|_2. \quad (\text{B.12})$$

Le théorème suivant montre que les espaces $\mathbb{L}^p$ sont complets. Ceci sera essentiel pour établir la convergence de suites de variables aléatoires.

Théorème B.9. *Pour $p \geq 1$, l'espace $\mathbb{L}^p$ muni de la norme $\|\cdot\|_p$ est un espace de Banach. Plus précisément, si $\{X_n\}_{n \geq 1}$ est une suite de Cauchy dans $\mathbb{L}^p$, i.e.*

$$\|X_n - X_m\|_p \longrightarrow 0 \quad \text{quand} \quad n, m \rightarrow \infty,$$

alors cette suite converge vers une variable aléatoire X appartenant à $\mathbb{L}^p$

$$\lim_{n \rightarrow \infty} \|X_n - X\|_p = 0.$$

Démonstration. Pour montrer que l'espace $\mathbb{L}^p$ est complet, nous considérons une suite de Cauchy $\{X_n\}_{n \geq 1}$ et nous allons montrer qu'elle converge. Il existe une suite d'entiers k_n strictement croissante telle que

$$\|X_{k_{n+1}} - X_{k_n}\|_p \leq \frac{1}{2^n}. \quad (\text{B.13})$$

On en déduit, par l'inégalité (B.11), que

$$\mathbb{E}(|X_{k_{n+1}} - X_{k_n}|) \leq \|X_{k_{n+1}} - X_{k_n}\|_p \leq \frac{1}{2^n},$$

et par conséquent

$$\mathbb{E}\left(\sum_{n \geq 1} |X_{k_{n+1}} - X_{k_n}|\right) < \infty.$$

Ainsi, la série $\sum_{n \geq 1} |X_{k_{n+1}} - X_{k_n}|$ est finie presque sûrement et on peut définir la variable aléatoire

$$X = X_{k_1} + \sum_{n=1}^{\infty} (X_{k_{n+1}} - X_{k_n})$$

qui appartient à $\mathbb{L}^1$. Comme X_{k_ℓ} peut se réécrire sous la forme d'une série télescopique

$$X_{k_\ell} = X_{k_1} + \sum_{n=1}^{\ell-1} (X_{k_{n+1}} - X_{k_n}),$$

la suite extraite $\{X_{k_\ell}\}_{\ell \geq 1}$ converge presque sûrement (et dans $\mathbb{L}^1$) vers X . L'inégalité (B.13) permet de vérifier que X appartient à $\mathbb{L}^p$ et que la convergence est aussi valable dans $\mathbb{L}^p$. $\square$

B.3 Convergences

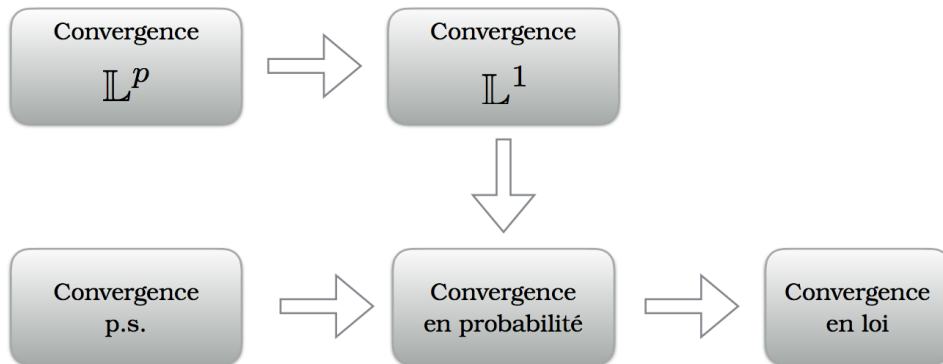
Les variables aléatoires étant des fonctions mesurables les notions de convergence presque sûre et de convergence dans $\mathbb{L}^p$ s'appliquent aussi. Soient $\{X_n\}_{n \geq 1}$ et X des variables aléatoires à valeurs dans $\mathbb{R}^d$ définies sur l'espace de probabilité $(\Omega, \mathcal{A}, \mathbb{P})$.

La suite $\{X_n\}_{n \geq 1}$ converge presque sûrement vers X si

$$\mathbb{P}(\{\omega \in \Omega; \ X(\omega) = \lim_{n \rightarrow \infty} X_n(\omega)\}) = 1.$$

Pour $p \geq 1$, la suite $\{X_n\}_{n \geq 1}$ converge vers X dans $\mathbb{L}^p$ si

$$\lim_{n \rightarrow \infty} \mathbb{E}(|X_n - X|^p) = 0.$$



Nous allons aussi définir deux autres types de convergence plus faibles qui seront plus faciles à démontrer. Le schéma ci-dessus résume les liens entre les différentes notions de convergence évoquées dans cette section.

B.3.1 Convergence en probabilité

Dans cette section, on considère $\{X_n\}_{n \geq 1}$ et X des variables aléatoires définies sur le même espace de probabilité $(\Omega, \mathcal{A}, \mathbb{P})$.

Définition B.10 (Convergence en probabilité). *On dit que la suite $\{X_n\}_{n \geq 1}$ converge en probabilité vers X si*

$$\lim_{n \rightarrow \infty} \mathbb{P}(|X_n - X| \geq \varepsilon) = 0 \quad \text{pour tout } \varepsilon > 0.$$

Le lemme suivant montre que la convergence en probabilité est moins forte que les convergences presque sûre et dans $\mathbb{L}^p$.

Lemme B.11.

1. La convergence presque sûre implique la convergence en probabilité.
2. La convergence dans $\mathbb{L}^p$ (avec $p \geq 1$) implique la convergence en probabilité.

Démonstration. Pour la première assertion, il suffit de remarquer que

$$\mathbb{P}(|X_n - X| \geq \varepsilon) = \mathbb{E}(\mathbf{1}_{\{|X_n - X| \geq \varepsilon\}})$$

et que $\mathbf{1}_{\{|X_n - X| \geq \varepsilon\}}$ converge presque sûrement vers 0 car $X_n - X$ converge presque sûrement vers 0. On en déduit la convergence en probabilité de la suite $\{X_n\}_{n \geq 1}$ vers X en utilisant le théorème de convergence dominée (B.2).

La seconde assertion est une conséquence immédiate de l'inégalité de Markov (B.7) appliquée avec la fonction $f(x) = |x|^p$:

$$\mathbb{P}(|X_n - X| \geq \varepsilon) \leq \frac{1}{\varepsilon^p} \mathbb{E}(|X_n - X|^p).$$

Ainsi la convergence dans $\mathbb{L}^p$ implique la convergence en probabilité. □

L'hypothèse d'uniforme intégrabilité, définie ci-dessous, va permettre de renforcer la convergence en probabilité pour retrouver une convergence dans $\mathbb{L}^1$.

Définition B.12 (Uniforme intégrabilité). *Une suite $\{X_n\}_{n \geq 1}$ de variables aléatoires est dite uniformément intégrable si*

$$\lim_{c \rightarrow \infty} \sup_{n \geq 1} \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| \geq c\}}) = 0.$$

Le résultat suivant permet d'énoncer une forme de réciproque au Lemme B.11.

Théorème B.13. *Soient $\{X_n\}_{n \geq 1}$ et X des variables aléatoires dans $\mathbb{L}^1$. Alors la suite $\{X_n\}_{n \geq 1}$ converge vers X dans $\mathbb{L}^1$ si et seulement si les 2 conditions suivantes sont vérifiées*

- (1) la suite $\{X_n\}_{n \geq 1}$ converge vers X en probabilité ;
- (2) la suite $\{X_n\}_{n \geq 1}$ est uniformément intégrable.

Démonstration. Supposons d'abord que les conditions (1) et (2) sont satisfaites et montrons qu'elles impliquent la convergence dans $\mathbb{L}^1$. Étant donné $c > 0$, on considère la fonction

$$\forall x \in \mathbb{R}, \quad \varphi_c(x) := -c\mathbf{1}_{\{x < -c\}} + x\mathbf{1}_{\{-c \leq x \leq c\}} + c\mathbf{1}_{\{x > c\}}.$$

Le point de départ consiste à décomposer $\|X_n - X\|_1$

$$\|X_n - X\|_1 \leq \|X_n - \varphi_c(X_n)\|_1 + \|\varphi_c(X_n) - \varphi_c(X)\|_1 + \|\varphi_c(X) - X\|_1. \quad (\text{B.14})$$

Nous allons ensuite montrer que chaque terme tend vers 0 quand n , puis c tendent vers l'infini.

Comme $\varphi_c(x)$ est une fonction lipschitzienne et bornée, la suite $\{\varphi_c(X_n)\}_{n \geq 1}$ converge dans $\mathbb{L}^1$ vers $\varphi_c(X)$. En effet pour tout $\varepsilon > 0$, on peut écrire

$$\|\varphi_c(X_n) - \varphi_c(X)\|_1 = \mathbb{E}(|\varphi_c(X_n) - \varphi_c(X)|) \leq \varepsilon + c\mathbb{E}(\mathbf{1}_{\{|X_n - X| \geq \varepsilon\}}).$$

Il suffit ensuite de laisser tendre n vers l'infini puis ε vers 0, pour déduire la convergence

$$\lim_{n \rightarrow \infty} \|\varphi_c(X_n) - \varphi_c(X)\|_1 = 0.$$

Il ne reste plus qu'à montrer que les deux termes restants dans (B.14) tendent vers 0 quand c tend vers l'infini (uniformément en n). On remarque que

$$\|\varphi_c(X_n) - X_n\|_1 = \mathbb{E}(|\varphi_c(X_n) - X_n|) \leq \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| > c\}}). \quad (\text{B.15})$$

Par conséquent, l'uniforme intégrabilité implique que

$$\lim_{c \rightarrow \infty} \sup_{n \geq 1} \|\varphi_c(X_n) - X_n\|_1 = 0.$$

Comme X est intégrable, l'inégalité (B.15) permet aussi de déduire que

$$\lim_{c \rightarrow \infty} \|\varphi_c(X) - X\|_1 = 0.$$

En utilisant les résultats précédents, on conclut que $\{X_n\}_{n \geq 1}$ converge vers X dans $\mathbb{L}^1$ car chaque terme de (B.14) tend vers 0.

Réciproquement, supposons que $\{X_n\}_{n \geq 1}$ converge vers X dans $\mathbb{L}^1$, alors la convergence en probabilité (1) est une conséquence du lemme B.11. Il suffit donc de montrer l'uniforme intégrabilité (2). Pour tout $\varepsilon > 0$, la convergence $\mathbb{L}^1$ implique l'existence d'un entier N tel que

$$\mathbb{E}(|X_n - X|) < \varepsilon \quad \text{pour tout } n > N. \quad (\text{B.16})$$

Par ailleurs, d'après le lemme A.20, il existe $\delta > 0$ tel que pour tout ensemble mesurable A vérifiant $\mathbb{P}(A) < \delta$ alors

$$\sup_{n \leq N} \mathbb{E}(|X_n| \mathbf{1}_A) < \varepsilon \quad \text{et} \quad \mathbb{E}(|X| \mathbf{1}_A) < \varepsilon. \quad (\text{B.17})$$

En appliquant l'inégalité de Markov aux ensembles $A_n := \{|X_n| > c\}$, on constate qu'il existe c assez grand tel que

$$\sup_{n \geq 1} \mathbb{P}(A_n) \leq \frac{1}{c} \sup_{n \geq 1} \mathbb{E}(|X_n|) < \delta. \quad (\text{B.18})$$

Par les deux inégalités précédentes (B.17) et (B.18), on obtient

$$\begin{aligned} \sup_{n \geq 1} \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| > c\}}) &= \max \left\{ \sup_{n \leq N} \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| > c\}}), \sup_{n > N} \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| > c\}}) \right\} \\ &\leq \max \left\{ \varepsilon, \sup_{n > N} \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| > c\}}) \right\} \\ &\leq \max \left\{ \varepsilon, \sup_{n > N} \mathbb{E}(|X| \mathbf{1}_{\{|X_n| > c\}}) + \mathbb{E}(|X - X_n| \mathbf{1}_{\{|X_n| > c\}}) \right\} \\ &\leq \max \left\{ \varepsilon, \sup_{n > N} \mathbb{E}(|X| \mathbf{1}_{A_n}) + \mathbb{E}(|X - X_n|) \right\}. \end{aligned}$$

Il suffit alors d'utiliser à nouveau (B.16) et (B.17), pour en déduire que

$$\sup_{n \geq 1} \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| > c\}}) < 2\varepsilon.$$

Cette inégalité reste vraie pour tout $\varepsilon > 0$ en choisissant c suffisamment grand. Ceci implique l'uniforme intégrabilité de la suite $\{X_n\}_{n \geq 1}$. $\square$

L'uniforme intégrabilité est une condition plus forte qu'une borne uniforme dans $\mathbb{L}^1$

$$\sup_{n \geq 1} \mathbb{E}(|X_n|) < \infty.$$

Cependant, dès que $p > 1$, une borne uniforme dans $\mathbb{L}^p$

$$\sup_{n \geq 1} \mathbb{E}(|X_n|^p) < \infty \quad (\text{B.19})$$

suffit à impliquer l'uniforme intégrabilité. En effet, en appliquant successivement l'inégalité de Hölder, puis l'inégalité de Markov, on obtient

$$\begin{aligned} \mathbb{E}(|X_n| \mathbf{1}_{\{|X_n| \geq c\}}) &\leq \mathbb{E}(|X_n|^p)^{1/p} \mathbb{P}(|X_n| \geq c)^{1-1/p} \\ &\leq \frac{1}{c^{1-1/p}} \mathbb{E}(|X_n|^p)^{1/p} \mathbb{E}(|X_n|)^{1-1/p}. \end{aligned}$$

Ceci montre que l'uniforme intégrabilité de la suite $\{X_n\}_{n \geq 1}$ est une conséquence de la borne uniforme (B.19) dans $\mathbb{L}^p$ (et donc aussi dans $\mathbb{L}^1$).

B.3.2 Convergence en loi

Définition B.14 (Convergence en loi). Soient X et $\{X_n\}_{n \geq 1}$ des variables aléatoires à valeurs dans $(E, \mathcal{E})$. On dit que $\{X_n\}_{n \geq 1}$ converge en loi vers X si

$$\lim_{n \rightarrow \infty} \mathbb{E}(f(X_n)) = \mathbb{E}(f(X)) \quad \text{pour toute fonction } f \text{ continue bornée sur } E. \quad (\text{B.20})$$

La convergence en loi de $\{X_n\}_{n \geq 1}$ vers X est équivalente à la convergence des lois $\{\mathbb{P}_{X_n}\}_{n \geq 1}$ vers la loi $\mathbb{P}_X$. Ainsi (B.20) se réécrit

$$\lim_{n \rightarrow \infty} \int_{\Omega} f(X_n(\omega)) d\mathbb{P}_{X_n}(\omega) = \int_{\Omega} f(X(\omega)) d\mathbb{P}_X(\omega) \quad \text{pour } f \text{ continue bornée sur } E.$$

Contrairement aux précédentes notions de convergence, la convergence en loi ne permet pas de comparer des réalisations des variables aléatoires $X_n(\omega)$ et de la limite $X(\omega)$. Ainsi si X est une variable gaussienne centrée $\mathcal{N}(0, 1)$, alors la suite $X_n = X$ converge en loi vers X , mais aussi vers $-X$. La convergence en loi ne nécessite pas que les variables X_n et X soient définies sur le même espace de probabilité $(\Omega, \mathcal{A}, \mathbb{P})$.

Pour établir la convergence en loi, on peut se contenter d'analyser la convergence (B.20) sur une classe réduite de fonctions continues. Dans cette perspective, nous allons définir la *fonction caractéristique* d'une variable aléatoire $X = (X_1, \dots, X_d)$ à valeurs dans $\mathbb{R}^d$ par

$$\forall \lambda \in \mathbb{R}^d, \quad \Phi_X(\lambda) = \mathbb{E}(\exp(i \lambda \cdot X)) \quad \text{avec} \quad \lambda \cdot X = \sum_{\ell=1}^d \lambda_{\ell} X_{\ell}, \quad (\text{B.21})$$

où i représente le nombre complexe $i^2 = -1$. La fonction caractéristique est l'analogue de la transformée de Fourier de la loi $\mathbb{P}_X$ et elle détermine entièrement la loi de X . On peut montrer que si deux variables aléatoires ont la même fonction caractéristique, alors elles ont la même loi. Le théorème suivant joue un rôle clef pour établir la convergence en loi.

Théorème B.15 (Théorème de Lévy). Une suite de variables aléatoires $\{X_n\}_{n \geq 1}$ à valeurs dans $\mathbb{R}^d$ converge en loi vers X si et seulement si

$$\forall \lambda \in \mathbb{R}^d, \quad \lim_{n \rightarrow \infty} \Phi_{X_n}(\lambda) = \Phi_X(\lambda).$$

La démonstration de ce théorème étant longue et délicate, nous l'omettons et renvoyons à [28] pour une preuve détaillée.

Les précédentes notions de convergence impliquent la convergence en loi.

Lemme B.16. La convergence en probabilité implique la convergence en loi.

Démonstration. Soit $\{X_n\}_{n \geq 1}$ une suite de variables aléatoires convergeant en probabilité vers X . Étant donné $\delta > 0$, on peut borner l'écart entre les fonctions caractéristiques

$$\begin{aligned} |\Phi_{X_n}(\lambda) - \Phi_X(\lambda)| &\leq \mathbb{E}(|\exp(i \lambda \cdot (X_n - X)) - 1|) \\ &\leq \mathbb{E}(|\exp(i \lambda \cdot (X_n - X)) - 1| \mathbf{1}_{\{|X_n - X| \leq \delta\}}) + 2\mathbb{P}(|X_n - X| > \delta). \end{aligned}$$

Pour λ fixé et $\varepsilon > 0$, le paramètre δ peut être choisi suffisamment petit pour obtenir une majoration uniforme en n

$$\mathbb{E} \left(\left| \exp(i\lambda \cdot (X_n - X)) - 1 \right| \mathbf{1}_{\{|X_n - X| \leq \delta\}} \right) \leq \varepsilon.$$

Le paramètre δ étant fixé, la convergence en probabilité implique que pour tout n suffisamment grand

$$\mathbb{P}(|X_n - X| > \delta) \leq \varepsilon.$$

Ceci démontre que $\Phi_{X_n}(\lambda)$ converge vers $\Phi_X(\lambda)$ quand n tend vers l'infini. La convergence en loi se déduit alors du théorème de Lévy B.15. $\square$

B.4 Indépendance

B.4.1 Variables aléatoires indépendantes

Deux évènements A, B sont indépendants si

$$\mathbb{P}(A \cap B) = \mathbb{P}(A)\mathbb{P}(B).$$

Cette identité implique que la réalisation de l'évènement B n'influe pas sur la probabilité de l'évènement A .

L'indépendance se généralise au cas de plusieurs évènements.

Définition B.17 (Évènements indépendants). *Les évènements $A_1, \dots, A_n$ sont indépendants si*

$$\mathbb{P}(A_{i_1} \cap A_{i_2} \cap \dots \cap A_{i_k}) = \mathbb{P}(A_{i_1})\mathbb{P}(A_{i_2}) \dots \mathbb{P}(A_{i_k}),$$

pour tout $k \leq n$ et tout sous-ensemble $\{i_1, \dots, i_k\}$ de $\{1, \dots, n\}$.

Il est important de noter que la relation

$$\mathbb{P}(A_1 \cap A_2 \cap \dots \cap A_n) = \mathbb{P}(A_1)\mathbb{P}(A_2) \dots \mathbb{P}(A_n)$$

ne suffit pas à assurer l'indépendance des n évènements.

On peut aussi définir une notion d'indépendance pour des σ -algèbres.

Définition B.18 (σ -algèbres indépendantes). *Les σ -algèbres $\mathcal{A}_1, \dots, \mathcal{A}_n$ incluses dans $\mathcal{A}$ sont dites indépendantes si toute famille d'ensembles $A_1 \in \mathcal{A}_1, \dots, A_n \in \mathcal{A}_n$ vérifie la relation*

$$\mathbb{P}(A_1 \cap A_2 \cap \dots \cap A_n) = \mathbb{P}(A_1)\mathbb{P}(A_2) \dots \mathbb{P}(A_n). \quad (\text{B.22})$$

La définition B.17 est équivalente à l'indépendance des σ -algèbres

$$\sigma(A_k) = \{\Omega, \emptyset, A_k, A_k^c\}, \quad k \leq n,$$

engendrées par les évènements $A_1, \dots, A_n$.

Cette définition permet d'étendre la notion d'indépendance aux variables aléatoires.

Définition B.19 (Variables indépendantes). *Les variables aléatoires $X_1, \dots, X_n$ à valeurs dans $(E_1, \mathcal{E}_1), \dots, (E_n, \mathcal{E}_n)$ sont indépendantes si les σ -algèbres correspondantes $\sigma(X_1), \dots, \sigma(X_n)$ sont indépendantes.*

Ceci est équivalent à dire que pour toute famille d'ensembles $(A_1, \dots, A_n)$ de $\mathcal{E}_1 \times \dots \times \mathcal{E}_n$ l'identité suivante est satisfaite

$$\mathbb{P}(\{X_1 \in A_1\} \cap \dots \cap \{X_n \in A_n\}) = \mathbb{P}(X_1 \in A_1) \dots \mathbb{P}(X_n \in A_n). \quad (\text{B.23})$$

L'indépendance des variables aléatoires est liée aux propriétés des espaces produits.

Proposition B.20. *L'indépendance des variables aléatoires $X_1, \dots, X_n$ à valeurs dans $(E_1, \mathcal{E}_1), \dots, (E_n, \mathcal{E}_n)$ est équivalente à chacune des assertions suivantes :*

(i) *La loi du n -uplet $\{X_1, \dots, X_n\}$ est une mesure produit*

$$\mathbb{P}_{(X_1, \dots, X_n)} = \mathbb{P}_{X_1} \otimes \dots \otimes \mathbb{P}_{X_n}.$$

(ii) *Pour toutes fonctions $f_i : E_i \mapsto \mathbb{R}$ mesurables bornées*

$$\mathbb{E} \left(\prod_{k=1}^n f_k(X_k) \right) = \prod_{k=1}^n \mathbb{E}(f_k(X_k)).$$

Démonstration. Pour toute famille d'ensembles $(A_1, \dots, A_n)$ de $\mathcal{E}_1 \times \dots \times \mathcal{E}_n$, on peut écrire

$$\mathbb{P}_{(X_1, \dots, X_n)}(A_1 \times A_2 \times \dots \times A_n) = \mathbb{P}(\{X_1 \in A_1\} \cap \{X_2 \in A_2\} \cap \dots \cap \{X_n \in A_n\})$$

et

$$\mathbb{P}_{X_1} \otimes \dots \otimes \mathbb{P}_{X_n}(A_1 \times A_2 \times \dots \times A_n) = \mathbb{P}(X_1 \in A_1) \mathbb{P}(X_2 \in A_2) \dots \mathbb{P}(X_n \in A_n).$$

Ainsi la propriété d'indépendance (B.23) est équivalente à l'égalité des deux mesures $\mathbb{P}_{(X_1, \dots, X_n)}$ et $\mathbb{P}_{X_1} \otimes \dots \otimes \mathbb{P}_{X_n}$ sur les pavés, i.e. les ensembles de la forme $A_1 \times A_2 \times \dots \times A_n$. Par le théorème A.24, ceci est équivalent à l'égalité des deux mesures sur la σ -algèbre produit $\mathcal{E}_1 \otimes \dots \otimes \mathcal{E}_n$.

Le théorème de Fubini A.25 permet d'établir l'équivalence entre les assertions (i) et (ii). $\square$

B.4.2 Comportement asymptotique

Nous rappelons maintenant les deux théorèmes fondamentaux qui caractérisent le comportement asymptotique de suites de variables aléatoires indépendantes et identiquement distribuées.

Théorème B.21 (Loi forte des grands nombres). *Soit $\{X_n\}_{n \geq 1}$ une suite de variables aléatoires indépendantes, identiquement distribuées et intégrables, i.e. $\mathbb{E}(|X_1|) < \infty$. Alors*

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i = \mathbb{E}(X_1) \quad \text{presque sûrement.}$$

Une preuve de ce théorème, reposant sur la convergence des martingales, est faite théorème 11.4.

Dans le cas de variables aléatoires positives, la loi des grands nombres reste valable même si l'espérance est infinie.

Corollaire B.22. Soit $\{X_n\}_{n \geq 1}$ une suite de variables aléatoires positives, indépendantes et identiquement distribuées. Alors

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i = \mathbb{E}(X_1) \quad \text{presque sûrement.}$$

Démonstration. Il suffit simplement de traiter le cas $\mathbb{E}(X_1) = \infty$. Étant donné $K > 0$, nous allons appliquer la loi des grands nombres du théorème B.21 à la suite de variables tronquées $\{\inf\{X_1, K\}\}_{n \geq 1}$

$$\liminf_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n X_i \geq \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{i=1}^n (\inf\{X_i, K\}) = \mathbb{E}(\inf\{X_1, K\}) \quad \text{presque sûrement.}$$

Le théorème de convergence monotone implique que $\mathbb{E}(\inf\{X_1, K\})$ tend vers $\mathbb{E}(X_1) = \infty$ quand K tend vers l'infini. Ceci prouve le corollaire. $\square$

Le théorème central limite permet de quantifier les fluctuations de la moyenne empirique.

Théorème B.23 (Théorème central limite). Soit $\{X_n\}_{n \geq 1}$ une suite de variables aléatoires indépendantes, identiquement distribuées et de variance finie $\sigma = \mathbb{E}(X_1^2) - \mathbb{E}(X_1)^2 < \infty$. Après renormalisation, la moyenne converge en loi vers $\mathcal{N}(0, \sigma)$ une gaussienne centrée de variance σ

$$\sqrt{n} \left(\frac{1}{n} \sum_{i=1}^n X_i - \mathbb{E}(X_1) \right) \xrightarrow[n \rightarrow \infty]{(loi)} \mathcal{N}(0, \sigma).$$

Démonstration. Quitte à considérer les variables $Y_i = X_i - \mathbb{E}(X_1)$, on peut toujours se ramener au cas où $\mathbb{E}(X_1) = 0$. On définit $S_n := \frac{1}{\sqrt{n}} \sum_{i=1}^n X_i$. Comme les variables sont indépendantes et ont la même loi, la fonction caractéristique se factorise

$$\Phi_{S_n}(\lambda) = \prod_{i=1}^n \Phi_{\frac{X_i}{\sqrt{n}}}(\lambda) = \left(\Phi_{X_1} \left(\frac{\lambda}{\sqrt{n}} \right) \right)^n.$$

En utilisant que $\mathbb{E}(X_1) = 0$, le développement de Taylor au second ordre permet d'écrire

$$\Phi_{X_1} \left(\frac{\lambda}{\sqrt{n}} \right) = 1 - \frac{1}{2} \frac{\lambda^2}{n} \sigma + o \left(\frac{1}{n} \right).$$

Ainsi la fonction caractéristique converge quand n tend vers l'infini

$$\lim_{n \rightarrow \infty} \Phi_{S_n}(\lambda) = \exp \left(-\frac{\lambda^2}{2} \sigma \right).$$

La limite étant la fonction caractéristique de la gaussienne $\mathcal{N}(0, \sigma)$, on conclut grâce au théorème de Lévy B.15. $\square$

B.5 Espérance conditionnelle

Dans cette section, nous allons établir des résultats énoncés au chapitre 9.

B.5.1 Construction de l'espérance conditionnelle

Chapitre 9, nous avons vu que l'espérance conditionnelle est définie naturellement dans $\mathbb{L}^2$ comme une projection orthogonale. Considérons $(\Omega, \mathcal{A}, \mathbb{P})$, un espace de probabilité, et $\mathcal{F}$ une sous- σ -algèbre de $\mathcal{A}$. Par exemple $\mathcal{F} = \sigma(Y)$ si elle est définie à partir de la variable aléatoire Y et dans ce cas, le sous espace $\mathbb{L}^2(\mathcal{F}, \mathbb{P})$ des variables aléatoires de carré intégrable et mesurables par rapport à $\mathcal{F}$ s'écrit

$$\mathbb{L}^2(\mathcal{F}, \mathbb{P}) = \left\{ h(Y); \quad h \text{ fonction mesurable telle que } \mathbb{E}(h(Y)^2) < \infty \right\}.$$

L'espace $\mathbb{L}^2(\mathcal{F}, \mathbb{P})$ coïncide avec l'espace $\mathcal{H}$ défini en (9.5).

Si X est dans $\mathbb{L}^2(\mathcal{A}, \mathbb{P})$, on définit son espérance conditionnelle $\mathbb{E}(X|\mathcal{F})$ par rapport à $\mathcal{F}$ comme la projection orthogonale de X sur $\mathbb{L}^2(\mathcal{F}, \mathbb{P})$ pour le produit scalaire sur $\mathbb{L}^2(\mathcal{A}, \mathbb{P})$. Par construction $\mathbb{E}(X|\mathcal{F})$ est une variable aléatoire mesurable par rapport à $\mathcal{F}$ et définie presque sûrement. Le lemme suivant résume quelques propriétés importantes de l'espérance conditionnelle définie comme une projection orthogonale dans $\mathbb{L}^2$.

Lemme B.24. *Soit X une variable aléatoire dans $\mathbb{L}^2(\mathcal{A}, \mathbb{P})$. L'espérance conditionnelle de X par rapport à $\mathcal{F}$ satisfait la condition d'orthogonalité suivante : pour toute variable aléatoire W bornée et $\mathcal{F}$ -mesurable*

$$\mathbb{E}(XW) = \mathbb{E}(\mathbb{E}(X|\mathcal{F})W). \quad (\text{B.24})$$

Soient X_1, X_2 deux variables aléatoires dans $\mathbb{L}^2(\mathcal{A}, \mathbb{P})$. L'espérance conditionnelle admet les propriétés suivantes :

(i) *Linéarité :*

$$\mathbb{E}(X_1 + X_2 | \mathcal{F}) = \mathbb{E}(X_1 | \mathcal{F}) + \mathbb{E}(X_2 | \mathcal{F}).$$

(ii) *Monotonie : si $X_1 \geq X_2$ presque sûrement alors*

$$\mathbb{E}(X_1 | \mathcal{F}) \geq \mathbb{E}(X_2 | \mathcal{F}) \quad \text{presque sûrement.}$$

Démonstration. Comme W appartient à $\mathbb{L}^2(\mathcal{F}, \mathbb{P})$, l'identité (B.24) est une conséquence de la projection orthogonale sur $\mathbb{L}^2(\mathcal{F}, \mathbb{P})$

$$\mathbb{E}((X - \mathbb{E}(X|\mathcal{F}))W) = 0.$$

La linéarité (i) est aussi une conséquence directe de la projection orthogonale. Pour prouver la propriété (ii), il suffit de considérer l'ensemble $F = \{\omega, \mathbb{E}(X_1|\mathcal{F}) < \mathbb{E}(X_2|\mathcal{F})\}$ qui est mesurable par rapport à $\mathcal{F}$ par construction. En appliquant l'identité (B.24) à la variable aléatoire $\mathbf{1}_F$, on en déduit que

$$\mathbb{E}\left(\left[\mathbb{E}(X_1|\mathcal{F}) - \mathbb{E}(X_2|\mathcal{F})\right]\mathbf{1}_F\right) = \mathbb{E}((X_1 - X_2)\mathbf{1}_F) \geq 0,$$

car $X_1 \geq X_2$. Ceci implique que $(\mathbb{E}(X_1|\mathcal{F}) - \mathbb{E}(X_2|\mathcal{F}))\mathbf{1}_F = 0$ et conclut la preuve de la propriété (ii). $\square$

L'étape suivante consiste à étendre la notion d'espérance conditionnelle aux variables à valeurs dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$. Avant de démontrer le théorème 9.3, rappelons son énoncé.

Théorème B.25. *Soit X une variable aléatoire $\mathcal{A}$ -mesurable appartenant à $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$, i.e. telle que $\mathbb{E}(|X|) < \infty$. On considère $\mathcal{F} \subset \mathcal{A}$ une autre σ -algèbre. Il existe une unique variable aléatoire Z (définie presque sûrement) telle que*

(a) Z est $\mathcal{F}$ -mesurable,

(b) $\mathbb{E}(|Z|) < \infty$,

(c) Pour toute variable aléatoire W bornée et $\mathcal{F}$ -mesurable alors $\mathbb{E}(XW) = \mathbb{E}(ZW)$.

On définit l'espérance conditionnelle de X sachant $\mathcal{F}$ par $\mathbb{E}(X|\mathcal{F}) := Z$.

Soient X_1, X_2 deux variables aléatoires dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$. L'espérance conditionnelle admet les propriétés suivantes :

(i) Linéarité :

$$\mathbb{E}(X_1 + X_2 | \mathcal{F}) = \mathbb{E}(X_1 | \mathcal{F}) + \mathbb{E}(X_2 | \mathcal{F}).$$

(ii) Monotonie : si $X_1 \geq X_2$ presque sûrement alors

$$\mathbb{E}(X_1 | \mathcal{F}) \geq \mathbb{E}(X_2 | \mathcal{F}) \quad \text{presque sûrement.} \quad (\text{B.25})$$

Démonstration. La première étape consiste à montrer l'existence de l'espérance conditionnelle pour une variable X dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$. Comme pour la construction de l'espérance, i.e. de l'intégrale, il suffit de traiter le cas de variables aléatoires positives et d'utiliser ensuite la décomposition $X = X^+ - X^-$ pour conclure dans le cas général. La stratégie consiste à approcher $X \geq 0$ par des variables aléatoires $X_n := \inf\{X, n\}$ bornées qui appartiennent donc à $\mathbb{L}^2(\mathcal{A}, \mathbb{P})$. Ainsi par le lemme B.24, l'espérance conditionnelle $Z_n := \mathbb{E}(X_n | \mathcal{F})$ est bien définie à l'aide de la projection orthogonale et elle vérifie les conditions (a,b,c) du théorème B.25. En particulier, pour toute variable aléatoire W bornée et mesurable par rapport à $\mathcal{F}$ alors

$$\mathbb{E}(Z_n W) = \mathbb{E}(X_n W). \quad (\text{B.26})$$

La suite $\{Z_n\}_{n \geq 1}$ est croissante par la propriété (ii) du lemme B.24. Elle admet donc une limite presque sûrement

$$Z := \lim_{n \rightarrow \infty} \uparrow Z_n \in [0, +\infty].$$

Comme les variables Z_n sont mesurables par rapport à $\mathcal{F}$, leur limite Z l'est aussi (cf. proposition A.13). Les suites $\{X_n\}_{n \geq 1}$ et $\{Z_n\}_{n \geq 1}$ étant croissantes, le théorème de convergence monotone permet de déduire de (B.26) que

$$\mathbb{E}(Z W) = \mathbb{E}(X W).$$

Notons que pour appliquer le théorème de convergence monotone, il faut décomposer W sous la forme $W = W^+ - W^-$. Ainsi Z satisfait les conditions (a,c). Comme Z est positive, il suffit de poser $W = 1$ pour vérifier la condition (b).

Afin de démontrer l'unicité de l'espérance conditionnelle, supposons qu'il existe deux variables aléatoires Z et $\tilde{Z}$ vérifiant les conditions (a,b,c). Comme Z et $\tilde{Z}$ sont mesurables par rapport à $\mathcal{F}$, la variable aléatoire $W = \mathbf{1}_{\{Z - \tilde{Z} \geq 0\}}$ est aussi mesurable pour $\mathcal{F}$. La propriété (c) permet d'affirmer que

$$\mathbb{E}((Z - \tilde{Z}) W) = 0.$$

Ceci implique que $(Z - \tilde{Z})^+ = 0$ presque sûrement. Il suffit de considérer $W = \mathbf{1}_{\{Z - \tilde{Z} \leq 0\}}$ pour obtenir que $(Z - \tilde{Z})^- = 0$ presque sûrement. Ainsi les variables aléatoires Z et $\tilde{Z}$ coïncident presque sûrement et l'espérance conditionnelle est définie de façon unique par $\mathbb{E}(X | \mathcal{F}) := Z$.

Les propriétés de linéarité (i) et de monotonie (ii) du lemme B.24 s'étendent immédiatement en tronquant les variables aléatoires de $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ et en appliquant la méthode précédente. $\square$

B.5.2 Propriétés de l'espérance conditionnelle

Commençons par démontrer les théorèmes de convergence, énoncés proposition 9.5, pour l'espérance conditionnelle. La proposition 9.5 est rappelée ci-dessous.

Proposition B.26. Soient $\{X_n\}_{n \geq 1}$ et X des variables aléatoires dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$. L'espérance conditionnelle satisfait les résultats de convergence suivants :

(i) (Convergence monotone)

Si $X_n \geq 0$ converge presque sûrement vers X en croissant, alors

$$\lim_{n \rightarrow \infty} \uparrow \mathbb{E}(X_n | \mathcal{F}) = \mathbb{E}(X | \mathcal{F}).$$

(ii) (Lemme de Fatou)

Si $X_n \geq 0$, alors

$$\mathbb{E} \left(\liminf_{n \rightarrow \infty} X_n | \mathcal{F} \right) \leq \liminf_{n \rightarrow \infty} \mathbb{E}(X_n | \mathcal{F}).$$

(iii) (Convergence dominée)

S'il existe Y dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ tel que $\sup_n |X_n| \leq Y$ et $\{X_n\}_{n \geq 1}$ converge presque sûrement vers X alors

$$\lim_{n \rightarrow \infty} \mathbb{E}(X_n | \mathcal{F}) = \mathbb{E}(X | \mathcal{F}).$$

Démonstration.

Assertion (i). Comme la suite $\{X_n\}_{n \geq 1}$ est croissante, la propriété de monotonie (B.25) implique que la suite $Z_n := \mathbb{E}(X_n | \mathcal{F})$ est aussi croissante. Par construction la variable aléatoire limite $Z = \lim_n \uparrow Z_n = \sup_n Z_n$ est aussi positive et mesurable par rapport à $\mathcal{F}$. De plus Z est intégrable, car le lemme de Fatou permet de déduire que

$$\mathbb{E}(Z) = \mathbb{E} \left(\liminf_{n \rightarrow \infty} Z_n \right) \leq \liminf_{n \rightarrow \infty} \mathbb{E}(Z_n) = \liminf_{n \rightarrow \infty} \mathbb{E}(X_n) \leq \mathbb{E}(X) < \infty.$$

Ainsi Z vérifie les propriétés (a,b) du théorème B.25 et il ne reste plus qu'à démontrer la propriété (c). Pour toute variable aléatoire W bornée et $\mathcal{F}$ -mesurable, on sait que $\mathbb{E}(X_n W) = \mathbb{E}(Z_n W)$. Comme $\sup_n |X_n| \leq X$ et $\sup_n |Z_n| \leq Z$, le théorème de convergence monotone permet de passer à la limite $\mathbb{E}(X W) = \mathbb{E}(Z W)$ et de montrer que Z satisfait la propriété (c). L'espérance conditionnelle étant unique, on en déduit que $Z = \mathbb{E}(X | \mathcal{F})$ presque sûrement. Ceci prouve l'assertion (i).

Assertion (ii). La propriété de monotonie (B.25) implique que

$$\inf_{k \geq n} \mathbb{E}(X_k | \mathcal{F}) \geq \mathbb{E} \left(\inf_{k \geq n} X_k | \mathcal{F} \right) \quad \text{pour tout } n \geq 1. \quad (\text{B.27})$$

On conclut l'assertion (ii) en utilisant le résultat de convergence monotone (i).

Assertion (iii). La preuve est identique à la démonstration du théorème A.21. On définit $Y_n := |X_n - X|$ et on remarque que $\tilde{Y} := \sup_n Y_n$ appartient à $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ grâce à la borne uniforme sur la suite $\{X_n\}_{n \geq 1}$. Le lemme de Fatou (ii) appliqué à la variable $\tilde{Y} - Y_n$ implique que

$$\mathbb{E}(\tilde{Y}) = \mathbb{E} \left(\liminf_{n \rightarrow \infty} (\tilde{Y} - Y_n) \mid \mathcal{F} \right) \leq \liminf_{n \rightarrow \infty} \mathbb{E}(\tilde{Y} - Y_n \mid \mathcal{F}) = \mathbb{E}(\tilde{Y}) - \limsup_{n \rightarrow \infty} \mathbb{E}(Y_n \mid \mathcal{F}).$$

Ceci prouve l'assertion (iii). $\square$

Les résultats de la proposition 9.8 sont rappelés et démontrés dans les propositions suivantes.

Proposition B.27 (Projections itérées). *Soit X une variable aléatoire dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ et $\mathcal{F}, \mathcal{G}$ des σ -algèbres incluses dans $\mathcal{A}$. Si $\mathcal{F} \subset \mathcal{G}$ alors*

$$\mathbb{E}(\mathbb{E}(X \mid \mathcal{G}) \mid \mathcal{F}) = \mathbb{E}(X \mid \mathcal{F}). \quad (\text{B.28})$$

Démonstration. Comme $\mathcal{F}$ contient moins d'information que $\mathcal{G}$, on observe que $\mathbb{L}^2(\mathcal{F}, \mathbb{P}) \subset \mathbb{L}^2(\mathcal{G}, \mathbb{P})$. Par conséquent pour des variables X dans $\mathbb{L}^2(\mathcal{A}, \mathbb{P})$, la relation (B.28) est une conséquence du théorème des projections itérées en algèbre linéaire. En effet, pour toute variable aléatoire W dans $\mathbb{L}^2(\mathcal{F}, \mathbb{P})$, on a bien

$$\mathbb{E} \left(\mathbb{E}(\mathbb{E}(X \mid \mathcal{G}) \mid \mathcal{F}) W \right) = \mathbb{E}(\mathbb{E}(X \mid \mathcal{G}) W) = \mathbb{E}(XW)$$

car W appartient aussi à $\mathbb{L}^2(\mathcal{G}, \mathbb{P})$.

Pour étendre la relation (B.28) aux variables aléatoires dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$, il suffit de tronquer les variables comme dans la preuve du théorème B.25. $\square$

Proposition B.28. *On considère $\mathcal{F}$ une σ -algèbre incluse dans $\mathcal{A}$. Soient X une variable aléatoire dans $\mathbb{L}^1(\mathcal{A})$ et Y une variable aléatoire $\mathcal{F}$ -mesurable telle que $\mathbb{E}(|XY|) < \infty$. Alors*

$$\mathbb{E}(XY \mid \mathcal{F}) = Y \mathbb{E}(X \mid \mathcal{F}). \quad (\text{B.29})$$

Démonstration. Supposons d'abord que Y soit une variable aléatoire bornée. Alors pour toute variable aléatoire W bornée et mesurable par rapport à $\mathcal{F}$, on peut écrire grâce à la définition de $\mathbb{E}(X \mid \mathcal{F})$

$$\mathbb{E}(Y \mathbb{E}(X \mid \mathcal{F}) W) = \mathbb{E}(\mathbb{E}(X \mid \mathcal{F}) YW) = \mathbb{E}(X YW) = \mathbb{E}(\mathbb{E}(XY \mid \mathcal{F}) W),$$

où nous avons utilisé que YW est aussi mesurable par rapport à $\mathcal{F}$. On en déduit (B.29) par l'unicité de l'espérance conditionnelle.

Le théorème de convergence dominée permet ensuite de généraliser cette relation aux variables aléatoires Y telles que $\mathbb{E}(|XY|) < \infty$. $\square$

Proposition B.29. Soient X une variable aléatoire dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ et $\mathcal{F}, \mathcal{G}$ des σ -algèbres incluses dans $\mathcal{A}$ telles que $\mathcal{G}$ est indépendante de $\sigma(\sigma(X), \mathcal{F})$. Alors

$$\mathbb{E}(X|\sigma(\mathcal{F}, \mathcal{G})) = \mathbb{E}(X|\mathcal{F}).$$

Démonstration. Il suffit de considérer une variable positive X dans $\mathbb{L}^1(\mathcal{A}, \mathbb{P})$ et de montrer que pour tout événement A dans $\sigma(\mathcal{F}, \mathcal{G})$, la relation suivante est satisfaite

$$\mathbb{E}(X\mathbf{1}_A) = \mathbb{E}(\mathbb{E}(X|\mathcal{F})\mathbf{1}_A). \quad (\text{B.30})$$

Pour des événements de la forme $A = F \cap G$ avec $F \in \mathcal{F}$ et $G \in \mathcal{G}$, nous pouvons écrire, en utilisant deux fois l'indépendance entre $\sigma(\sigma(X), \mathcal{F})$ et $\mathcal{G}$,

$$\mathbb{E}(\mathbb{E}(X|\mathcal{F})\mathbf{1}_{F \cap G}) = \mathbb{E}(\mathbb{E}(X|\mathcal{F})\mathbf{1}_F) \mathbb{E}(\mathbf{1}_G) = \mathbb{E}(X\mathbf{1}_F) \mathbb{E}(\mathbf{1}_G) = \mathbb{E}(X\mathbf{1}_{F \cap G}).$$

Ainsi les mesures $A \mapsto \mathbb{E}(X\mathbf{1}_A)$ et $A \mapsto \mathbb{E}(\mathbb{E}(X|\mathcal{F})\mathbf{1}_A)$ sur $\sigma(\mathcal{F}, \mathcal{G})$ coïncident sur la famille d'ensembles

$$\mathcal{I} = \{F \cap G, \quad F \in \mathcal{F}, G \in \mathcal{G}\}.$$

Comme $\mathcal{I}$ est stable par intersection et engendre $\sigma(\mathcal{F}, \mathcal{G})$, le théorème A.8 assure que les 2 mesures sont égales pour tout A dans $\sigma(\mathcal{F}, \mathcal{G})$. Ceci prouve la relation (B.30) et donc la proposition B.29. $\square$

Bibliographie

- [1] R. ALBERT et A.-L. BARABÁSI – « Statistical mechanics of complex networks », *Reviews of modern physics* **74** (2002), no. 1, p. 47–97.
- [2] M. BENAÏM et N. EL KAROUI – *Promenade aléatoire : Chaînes de markov et simulations ; martingales et stratégies*, Les Éditions de l'École Polytechnique, 2004.
- [3] B. BERCU et D. CHAFAÏ – *Modélisation stochastique et simulation-cours et applications*, Dunod, 2007.
- [4] S. P. BOYD et L. VANDENBERGHE – *Convex optimization*, Cambridge university press, 2004.
- [5] P. BRÉMAUD – *An introduction to probabilistic modeling*, Undergraduate Texts in Mathematics, Springer-Verlag, New York, 1994.
- [6] X. CHEN – « Limit theorems for functionals of ergodic Markov chains with general state space », *Mem. Amer. Math. Soc.* **139** (1999), no. 664, p. xiv+203.
- [7] M. DE LARA et L. DOYEN – *Sustainable management of natural resources : mathematical models and methods*, Springer, 2008.
- [8] J.-F. DELMAS et B. JOURDAIN – *Modèles aléatoires*, Mathématiques & Applications (Berlin), vol. 57, Springer-Verlag, Berlin, 2006.
- [9] M. DUFLO – *Algorithmes stochastiques*, Springer Berlin, 1996.
- [10] R. DURRETT – *Probability : theory and examples*, Cambridge University Press, Cambridge, 2010.
- [11] — , *Random graph dynamics*, Cambridge Series in Statistical and Probabilistic Mathematics, vol. 20, Cambridge University Press, Cambridge, 2010.
- [12] A. GEORGES et M. MÉZARD – *Physique statistique*, Cours de l'École Polytechnique, 2013.
- [13] C. GRAHAM – *Chaînes de markov : cours, exercices et corrigés détaillés*, Dunod, 2008.
- [14] G. R. GRIMMETT et D. R. STIRZAKER – *Probability and random processes*, Oxford University Press, New York, 2001.
- [15] W. K. HASTINGS – « Monte carlo sampling methods using markov chains and their applications », *Biometrika* **57** (1970), no. 1, p. 97–109.
- [16] Y. HEURTEAUX – « An introduction to mandelbrot cascades », *New Trends in Applied Harmonic Analysis*, Springer, 2016, p. 67–105.
- [17] D. LAMBERTON et B. LAPEYRE – *Introduction au calcul stochastique appliqué à la finance*, second éd., Ellipses, Édition Marketing, Paris, 1997.

- [18] J. F. LE GALL – *Intégration, probabilités et processus aléatoires*, disponible sur la page web de l’auteur, 2006.
- [19] Y. LECUN, L. BOTTOU, Y. BENGIO et P. HAFFNER – « Gradient-based learning applied to document recognition », *Proceedings of the IEEE* **86** (1998), no. 11, p. 2278–2324.
- [20] N. METROPOLIS, A. W. ROSENBLUTH, M. N. ROSENBLUTH, A. H. TELLER et E. TELLER – « Equation of state calculations by fast computing machines », *The journal of chemical physics* **21** (1953), p. 1087.
- [21] S. MÉLÉARD – *Modèles aléatoires en écologie et évolution*, Cours de l’École Polytechnique, 2013.
- [22] J. NEVEU – *Martingales à temps discret*, Masson et Cie, éditeurs, Paris, 1972.
- [23] J. R. NORRIS – *Markov chains*, Cambridge Series in Statistical and Probabilistic Mathematics, vol. 2, Cambridge University Press, Cambridge, 1998.
- [24] J. G. PROPP et D. B. WILSON – « Exact sampling with coupled markov chains and applications to statistical mechanics », *Random structures and Algorithms* **9** (1996), no. 1-2, p. 223–252.
- [25] N. TOUZI – *Chaînes de markov et martingales en temps discret*, Cours de l’École Polytechnique, 2012.
- [26] — , *Calcul stochastique en finance*, Cours de l’École Polytechnique, 2015.
- [27] W. WERNER – *Percolation et modèle d’Ising*, Cours Spécialisés, vol. 16, Société Mathématique de France, Paris, 2009.
- [28] D. WILLIAMS – *Probability with martingales*, Cambridge Mathematical Textbooks, Cambridge University Press, Cambridge, 1991.
- [29] G. WINKLER – *Image analysis, random fields and markov chain monte carlo methods : a mathematical introduction*, vol. 27, Springer Verlag, 2003.

Index

- Absorbant, état, 42
- Algèbre
 - σ -algèbre, 127, 209
 - σ -algèbre borélienne, 127, 210
 - σ -algèbre engendrée, 224
 - σ -algèbre produit, 220
- Algorithme
 - de Metropolis-Hastings, 100
 - de Propp-Wilson, 103
 - de recuit simulé, 106
 - de Robbins-Monro, 192
 - du gradient stochastique, 186
 - PageRank, 77
 - stochastique, 99
- Apériodique, chaîne, 79, 82
- Cascades de Mandelbrot, 175
- Chaîne de Markov
 - convergence, 73, 80, 87
 - définition, 17
- Chapman-Kolmogorov, 21
- Condition de Doeblin, 87
- Contrôle stochastique, 204, 206
- Convergence
 - dans $\mathbb{L}^p$, 228
 - dominée, théorème, 218, 225
 - en loi, 233
 - en probabilité, 230
 - monotone, théorème, 215, 225
- Couplage
 - définition, 84
 - par le passé, 103
- Crochet d'une martingale, 155
- Dirichlet, problème, 32
- Distance en variation, 83
- Doob
 - inégalités maximales, 151
 - théorème d'arrêt, 140
- Ehrenfest, modèle, 49
- Equation de la chaleur, 27
- Ergodique, théorème, 73, 76
- Espérance conditionnelle
 - convergence dominée, 131, 239
 - convergence monotone, 131, 239
 - définition, 125
 - Fatou, 131, 239
 - inégalité de Hölder, 131
 - inégalité de Jensen, 131
 - projections itérées, 240
- Espace
 - $\mathbb{L}^1$, 217
 - $\mathbb{L}^p$, 228
 - de probabilité, 127, 132
 - mesurable, 210
- Fatou, lemme, 216, 225
- Fermée, classe, 42
- Filtration, 132
- Fonction
 - étagée, 213
 - borélienne, 213
 - caractéristique, 233
 - harmonique, 43
 - intégrable, 216, 217
 - mesurable, 212
- Galton-Watson, arbre, 64, 195
- Graphes aléatoires
 - d'Erdős-Rényi, 68
 - de Barabási-Albert, 181
- Inégalité
 - de Hölder, 226
 - de Hoeffding, 171
 - de Jensen, 136, 227
 - de Markov, 227
 - de Minkowski, 226

- de Schwarz, 228
- Indépendance
 - σ -algèbres, 234
 - événements, 234
 - variables aléatoires, 235
- Irréductible, chaîne, 42
- Ising, modèle, 101, 108
- Limite
 - inférieure, 214
 - supérieure, 214
- Loi forte des grands nombres, 156, 235
- Martingale
 - convergence $\mathbb{L}^2$, 154
 - convergence presque sûre, 157
 - définition, 135
 - fermée, 163
- Matrice de transition, 17
- Mesure
 - σ -additive, 210
 - σ -finie, 210
 - de Gibbs, 98
 - de Lebesgue, 212
 - extension de Carathéodory, 212
 - invariante, 37, 60
 - stationnaire, 37
- Modèle
 - de Cox, Ross, Rubinstein, 148
- Monte-Carlo, méthode, 34
- Option
 - européenne, 147
- Percolation, 111
- Probabilité conditionnelle, 125, 130
- Processus
 - adapté, 133
 - aléatoire, 17, 132
 - arrêté, 139
 - prévisible, 133, 138
- Processus de branchement, 64
- Programmation dynamique, 206, 207
- Propriété de Markov, 17
- Propriété de Markov forte, 26
- Récurrence aléatoire, 18
- Récurrent
 - état, 53, 55
 - nul, 53
 - positif, 53, 60
- Réseau de neurones, 189
- Réversibilité, 47
- Ruine du joueur, 30
- Segmentation, 108
- Snell, enveloppe, 198
- Sous-martingale, 135
- Stratégie
 - d'arbitrage, 145
 - de gestion, 143
- Surmartingale, 135
- Temps d'arrêt
 - définition, 25, 133
 - optimal, 197
- Théorème H, 48
- Théorème central limite
 - martingales, 167
 - chaînes de Markov, 173
 - variables indépendantes, 236
- Transitoire, état, 53, 55
- Uniforme intégrabilité, 165, 230
- Urne de Pólya, 180
- Variable aléatoire
 - définition, 223
 - loi, 223
- Voyageur de commerce, problème, 107
- Wright-Fisher, modèle, 161