

Chaînes de Markov

Recherche opérationnelle
Dimitri Watel - ENSIIE

2024

Dans ce chapitre, on va voir comment il est possible de modéliser certains comportements aléatoires spécifiques sous forme de chaînes de Markov, c'est à dire des graphes augmentés avec des probabilités. Les propriétés de ces graphes vont ensuite permettre de déduire des informations sur le système. Les applications des chaînes sont généralement des recherches de pannes dans un système, l'estimation de la position d'une population dans un ensemble d'états, et, comme nous le verrons dans le chapitre associé, les modélisations de files d'attente. A partir des informations sur ces systèmes, on peut en déduire une action adaptée qui ne dépendra pas des aspects probabilistes du systèmes.

1 Définitions

1.1 Processus stochastique et chaîne de Markov

Définition 1. Une variable aléatoire X est une fonction qui prend des valeurs dans un ensemble S d'états. Un processus stochastique est une famille de variables aléatoires $(X_t)_{t \in T}$ où $T \subset \mathbb{R}^+$.

Par exemple, vous pouvez considérer la météo comme un phénomène aléatoire. Au jour t , X_t sera la variable aléatoire représentant la météo de ce jour, il y a aura une certaine probabilité de voir cette variable se réaliser sous forme de beau temps, de pluie, de temps maussade, On considère la suite des variables comme un processus stochastique.

Certains processus ont des propriétés particulières. La première est celle d'un processus Markovien, c'est un processus sans mémoire, où connaître le dernier état du processus à l'instant t est suffisant pour en déduire les probabilités de réalisation à l'instant suivant. Avant d'aller dans la définition, voici trois exemples.

Supposons que vous ayez deux sacs. Le premier contient respectivement de n_1 boules blanches et m_1 noires. Le second contient respectivement de n_2 boules blanches et m_2 noires. Au premier tour, on pioche dans le sac 1. A chaque tour suivant, si la dernière boule piochée est blanche, on pioche dans le sac 1, sinon on pioche dans le sac 2. Une fois une boule piochée, on remet la boule dans son sac. Dans ce cas, on a bien un processus Markovien.

Seule la dernière boule est importante. Si c'est blanc, on a une probabilité $\frac{n_1}{n_1+m_1}$ de piocher blanc et $\frac{m_1}{n_1+m_1}$ de piocher noir. Si c'est noir, on a une probabilité $\frac{n_2}{n_2+m_2}$ de piocher blanc et $\frac{m_2}{n_2+m_2}$ de piocher noir. Savoir si les boules d'avant étaient blanches ou noires ne nous donne aucune information supplémentaire.

Supposons que vous ayez un seul sac rempli de n boules blanches et m noires. Vous piochez dedans puis remettez la boule dans le sac avant de recommencer. Dans ce cas, chaque pioche est indépendante, vous avez la même probabilité $\frac{n}{n+m}$ de piocher blanc et $\frac{m}{n+m}$ de piocher noir. Cet exemple est une version extrême d'un processus Markovien. Il n'est même pas nécessaire de connaître la dernière boule pour déterminer la probabilité de la pioche suivante.

Supposons enfin que vous ayez un seul sac rempli de n boules blanches et m noires. Vous piochez dedans mais ne remettez pas la boule dans le sac avant de recommencer. Dans ce cas, chaque pioche dépend de toutes les précédentes. Sans connaissance de ces précédentes pioches, vous ne savez pas combien de boules il reste dans le sac et donc vous ne pouvez pas déduire les probabilités de pioche au tour suivant.

Définition 2. Un processus $(X_t)_{t \in T}$ est Markovien si et seulement si *le futur dépend uniquement du présent* :

$$\forall t_1 < t_2 < \dots < t_n < t_{n+1} \in T, \forall A \subset S$$

$$P(X_{t_{n+1}} \in A | X_{t_1} X_{t_2}, \dots, X_{t_n}) = P(X_{t_{n+1}} \in A | X_{t_n})$$

Un processus Markovien signifie que la fonction $P(X_{t_{n+1}} \in A | X_{t_1} X_{t_2}, \dots, X_{t_n})$ ne dépend pas de l'état des variables $X_{t_1} X_{t_2}, \dots, X_{t_{n-1}}$. C'est une fonction qui ne dépend que de l'état de X_{t_n} . C'est un processus sans mémoire. Il regarde dans quel état il est à un instant donné, et cela définit la variable aléatoire de l'état suivant. Attention, connaître la réalisation de X_{t_n} ne suffit pas à connaître celle de $X_{t_{n+1}}$ mais suffit à connaître les probabilités de réalisation de cette variable.

La seconde propriété importante des chaînes de Markov est l'homogénéité. Cela signifie que les probabilités ne varient pas au cours du temps. Les 2 premiers exemples avec les sacs de boules sont homogènes. Si, par contre, à chaque tour, vous rajoutez une boule noire et une boule blanche dans les sacs, les probabilités vont varier au cours du temps (jusqu'à atteindre environ 0.5

dans ce cas). Ce processus n'est donc pas homogène. Considérons un autre exemple. Si on a une voiture qui tombe régulièrement en panne. On modélise le fait de tomber en panne avec un processus: cette voiture, quand elle est en panne une semaine, a une probabilité 0.5 de retomber en panne la semaine suivante. Si elle n'est pas en panne, alors elle a probabilité 0.01 de tomber en panne la semaine suivante. Ce processus est un processus Markovien (seul l'état de la voiture la semaine précédente est important) et il est homogène (les probabilités ne dépendent pas de la semaine où on regarde si elle va tomber en panne). Supposons maintenant que le modèle soit le suivant: quand elle est en panne une semaine t , a une probabilité $1 - 0.5e^{-t}$ de retomber en panne la semaine suivante. Si elle n'est pas en panne la semaine t , alors elle a probabilité 0.01 de tomber en panne la semaine suivante. Ainsi plus la voiture tombe en panne tard, plus elle a de chance de rester en panne longtemps. Ce processus n'est pas homogène. Par contre il est toujours Markovien.

Définition 3. Un processus $(X_t)_{t \in T}$ est *homogène* si et seulement si les *probabilités conditionnelles ne dépendent pas du temps* :

$$\forall t, t' \in T, s > 0 \setminus t + s, t' + s \in T, A \subset S$$

$$P(X_{t+s} \in A | X_t) = P(X_{t'+s} \in A | X_{t'})$$

On a maintenant tous les éléments pour définir un chaîne de Markov.

Définition 4. Une **chaîne de Markov** est un processus stochastique $(X_t)_{t \in T}$ Markovien et homogène.

Dans la suite, on considèrera un cas simple de chaîne de Markov où le temps T est dénombrable et où l'espace des états S est fini et discret. Mais il faut bien noter que les définitions de ce cours se généralisent au cas où le temps est continu et au cas où l'espace des états est continu.

1.2 Caractérisation par un graphe

Compte tenu de l'homogénéité et du fait qu'une chaîne de Markov est un processus Markovien, on observe que, connaissant la réalisation de la variable X_t , on peut en déduire les probabilités de réalisation de la variable X_{t+1} . Cette probabilité ne dépend pas de l'itération t .

Définition 5. On appelle p_{ij} la probabilité que le système passe de l'état i à l'état j en un pas de temps. Cette probabilité ne dépend pas du moment $t \in T$ où le système est dans l'état i .

La valeur p_{ij} est la *probabilité de transition* de l'état i à j . La matrice P des coefficients p_{ij} est appelée *matrice de transition*.

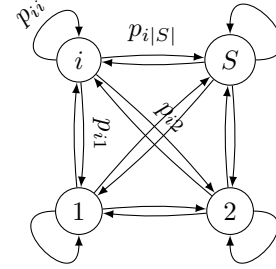
$$\forall i, j \in S, \exists p_{ij} \quad \setminus \quad \forall t \in \mathbb{N} \quad P(X_{t+1} = j | X_t = i) = p_{ij}$$

$$P = \begin{pmatrix} 1 & 2 & \cdots & j & \cdots & |S| \\ p_{11} & p_{12} & \cdots & p_{1j} & \cdots & p_{1|S|} \\ p_{21} & p_{22} & \cdots & p_{2j} & \cdots & p_{2|S|} \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ p_{i1} & p_{i2} & \cdots & p_{ij} & \cdots & p_{i|S|} \\ \vdots & \vdots & \ddots & \vdots & \ddots & \vdots \\ p_{|S|1} & p_{|S|2} & \cdots & p_{|S|j} & \cdots & p_{|S||S|} \end{pmatrix} \begin{matrix} 1 \\ 2 \\ \\ i \\ \\ |S| \end{matrix}$$

On peut déjà noter une propriété intéressante de la matrice P . Elle est dite *stochastique* car la somme des éléments d'une ligne vaut 1.

$$\sum_{j=1}^{|S|} p_{ij} = 1$$

Définition 6. Le graphe $G = (V, A)$ d'une chaîne de Markov est un **graphe orienté** où chaque nœud est un état de S (donc $V = S$), et chaque arc reliant i à j est valué avec p_{ij} . On ne met pas un arc si $p_{ij} = 0$



Il est possible de déduire de P des probabilités de transition plus complexe.

Théorème 1.1. Soit $k \in \mathbb{N}$ alors $P(X_{t+k} = j | X_t = i) = P_{ij}^k$.

Remarque 1. Il peut sembler étrange de considérer cette probabilité sachant que, dans les exemples de processus Markovien de la section précédente, on a beaucoup insisté sur le fait que seule la dernière réalisation, X_{t+k-1} , était nécessaire pour connaître la probabilité de réalisation de X_{t+k} . Mais c'est parce que nos exemples sont imprécis. Si on regarde bien la définition de processus Markovien, on ne dit pas que seule la réalisation de l'itération précédente permet de déduire des probabilités, on dit que si on a connaissance de plusieurs réalisations, alors seule la dernière est utile. Dit autrement, si on connaît la réalisation de X_t et celle de X_{t+k-1} alors la connaissance de X_t est superflue. Mais si on ne connaît pas X_{t+k-1} on pourra quand même tirer des informations (moins précises) de X_t .

Proof. Démontrons ce théorème par récurrence sur k . Notons $R_{kij} = P(X_{t+k} = j | X_t = i)$. Montrons que $R_k = P^k$ pour tout k .

Si $k = 0$ alors la probabilité d'être dans l'état j à l'itération t sachant qu'on est dans l'état i est 1 si $i = j$ et 0 sinon. Donc $R_0 = I = P^0$. Donc la propriété est bien prouvée pour $k = 0$.

Supposons maintenant la propriété prouvée pour une certaine valeur $k - 1$. D'après la formule des probabilités totales:

$$\begin{aligned} R_{kij} &= \sum_{s \in S} P(X_{t+k} = j | X_{t+k-1} = s) \cdot R_{k-1, is} \\ &= \sum_{s \in S} P_{sj} \cdot R_{k-1, is} \end{aligned}$$

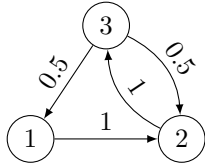
Par hypothèse de récurrence

$$\begin{aligned} R_{kij} &= \sum_{s \in S} P_{sj} \cdot P_{is}^{k-1} \\ &= P_{ij}^k \end{aligned} \quad \square$$

1.3 Vecteur de distribution de probabilité des états

On souhaite étudier $q_i(t) = P(X_t = i)$. On note $Q(t)$ le vecteur des probabilités $q_i(t)$. On peut remarquer qu'il n'est pas possible de déduire $Q(t)$ à partir de P uniquement. Une raison simple est que Q dépend de t alors P est constante. Cependant, connaître $Q(t)$ permet de connaître $Q(t+k)$ pour tout $k \in \mathbb{N}$.

Considérons par exemple la chaîne suivante:



Sans information supplémentaire, il n'est pas possible de savoir quelle est la probabilité de réalisation de $X(t)$. Supposons qu'on connaisse, par un moyen détourné, cette probabilité quand $t = 0$. On sait qu'au démarrage de la chaîne, on a une probabilité $q_1(0) = 0.25$ d'être dans l'état 1, $q_2(0) = 0.5$ d'être dans l'état 2 et $q_3(0) = 0.25$ d'être dans l'état 3. Une autre manière de voir ces probabilités est de supposer qu'on étudie une population avec de nombreux individus, par exemple 1000000 personnes. Il y aurait alors au début 250000 personnes dans l'état 1, 250000 personnes dans l'état 3 et 500000 personnes dans l'état 2. L'obtention des probabilités $Q(0)$ se fait alors juste avec un comptage statistique, la réalisation de $X(0)$ n'est alors que le tirage au sort d'une de ces personnes. Après une itération de la chaîne, toute personne qui était dans l'état 1 va, avec une probabilité de 1, dans l'état 2. De même toute personne dans l'état 2 se déplace

en 3. Enfin, toute personne en 3 a une chance sur deux de se déplacer dans l'état 1 ou dans l'état 2. On aurait donc à l'issue de ce déplacement, 125000 personnes dans l'état 1, 375000 dans l'état 2 et 500000 dans l'état 3, soit $Q(1) = (0.125, 0.375, 0.5)$. On peut enfin noter la propriété suivante: si on multiplie $Q(0)$ par P , on obtient $Q(1)$.

$$(0.25 \quad 0.5 \quad 0.25) \cdot \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 0.5 & 0.5 & 0 \end{pmatrix} = (0.125, 0.375, 0.5)$$

Ainsi la connaissance de $Q(0)$ permet de déduire $Q(1)$ à l'aide de P . Cette propriété peut se généraliser et se démontrer.

Théorème 1.2. Soit $k \in \mathbb{N}$ alors $Q(t+k) = Q(t) \cdot P^k$.

Proof. La preuve est très similaire à celle du théorème 1.1. Elle se démontre par récurrence et utilise la formule des probabilités totales. $\square$

Attention: pour calculer $Q(t+k)$, il faut multiplier $Q(t)$ par P^k et non l'inverse.

2 Distribution stationnaire et distribution limite

Ainsi, en poursuivant les calculs de l'exemple précédent, on peut obtenir $Q(t)$ pour tout t à partir de $Q(0)$. Cette propriété est intéressante mais nous oblige à connaître $Q(0)$ pour avoir des informations. On peut cependant remarquer un détail étrange. Si on calcule $Q(1000)$, on obtient

$$Q(0) \cdot P^{1000} = (0.2 \quad 0.4 \quad 0.4)$$

Supposons maintenant qu'on parte avec un autre vecteur initial, disons $Q'(0) = (0.33, 0.33, 0.33)$. Si on calcule $Q'(1000)$, on obtient

$$Q'(0) \cdot P^{1000} = (0.2 \quad 0.4 \quad 0.4)$$

On peut recommencer avec d'autres vecteurs initiaux, on tombera toujours sur le même vecteur. Pour être exact, il ne s'agit pas du même vecteur mais de vecteurs très proches au point qu'une précision très élevée soit nécessaire pour les différencier. Mais il semblerait que, quel que soit le vecteur de probabilité initial, la distribution converge toujours vers le même vecteur. Cette propriété est intéressante car il n'est plus nécessaire de connaître le vecteur initial de distribution de la population, il suffit d'exécuter suffisamment d'itérations pour transformer toute distribution en sa distribution limite. Dans la suite de cette section, on va caractériser les chaînes qui ont une telle distribution limite, et on va décrire une méthode pour la calculer. Une chose

intéressante à constater est que la caractérisation passe purement par une caractérisation du graphe de la chaîne, sans tenir compte des valeurs de probabilités sur les arcs.

2.1 Classification des états

Les concepts utilisés dans les chaînes de Markov qui sont décrits ici sont les mêmes que ceux utilisés dans la théorie des graphes, même si le vocabulaire diffère. On considère dans la suite une chaîne dont le graphe est G .

Définition 7. Un état j est dit *accessible* depuis l'état i s'il existe un **chemin** de i vers j dans G .

$$\exists k \setminus (P^k)_{ij} > 0$$

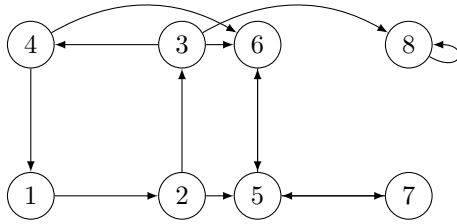
Définition 8. Deux états i et j sont dits *communiquant* si chacun est accessible depuis l'autre.

Définition 9. Une *classe d'états communiquant* est une composante fortement connexe de G , autrement dit un ensemble maximal d'états communiquant deux à deux.

Définition 10. Un état i est *transitoire* s'il existe un état j accessible depuis i tel que i n'est pas accessible depuis j .

Définition 11. Un état i est *permanent* ou *récurrent* si pour tout état j accessible depuis i , i est accessible depuis j .

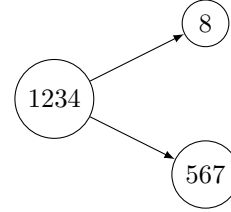
Considérons l'exemple suivant.



- Tous les états sont accessibles depuis 1
- Seuls 5, 6 et 7 sont accessibles depuis 5.
- 1 et 3 communiquent car chacun est accessible depuis l'autre.
- Cette chaîne possède trois classes d'états communiquants: 1234, 567 et 8.
- 1, 2, 3 et 4 sont transitoires. En effet, depuis n'importe lequel de ces états, on peut atteindre par exemple l'état 5. Mais depuis l'état 5 il est impossible de revenir sur ces 4 états. **Attention**, on pourrait croire que 1 n'est pas transitoire car depuis 1, on doit nécessairement rejoindre 2, depuis lequel on peut réatteindre 1. Mais la définition ne parle pas de transition directe, elle parle d'accessibilité. Depuis 1, on peut atteindre 2 puis 5. Donc il existe une possibilité que 5 soit atteint, et une fois cette possibilité réalisée, il n'est plus possible de revenir en 1.

- 5, 6, 7 sont permanents. En effet, depuis chacun de ces états, on ne peut atteindre que 5, 6 et 7. On peut donc toujours revenir à la position de départ.
- 8 est aussi permanent, car aucun n'état n'est accessible depuis 8 excepté lui-même.

On peut assez facilement distinguer les états transitoires et permanents de la manière suivante. Créer un nouveau graphe G_c . Dans ce graphe, ajouter un nœud par classe d'états communiquant de G . Ajouter dans G_c un arc entre deux nœuds u et v s'il existe dans G un arc reliant la classe correspondant à u et celle correspondant à v . Dans la chaîne précédente, cela donnerait le graphe suivant:



On nomme G_c la *contraction* de G . Ce graphe a deux propriétés intéressantes:

Théorème 2.1. Soit G_c la contraction de G alors G_c est sans circuit et un état de G est permanent si et seulement si le nœud correspondant à sa classe d'états communiquant dans G_c est un puits.

Proof. Si G_c contient un circuit entre les nœuds $u_1, u_2, \dots, u_p$ dont les p classes d'états communiquant associées sont $C_1, C_2, \dots, C_p$, alors il existe dans G un circuit D passant successivement par les nœuds de $C_1, C_2, \dots, C_p$ puis qui revient en C_1 . Soit v un nœud de $C_1 \cap D$ et w un nœud de $C_2 \cap D$, alors, puisque D est un circuit, v est accessible depuis w et inversement. Donc les deux communiquent. Puisque tous les états de C_1 communiquent avec v et tous les états de C_2 communiquent avec w alors tous les états de C_1 communiquent avec tous les états de C_2 . Ainsi $C_1 \cup C_2$ est un ensemble d'états communiquant, ce qui signifie que C_1 et C_2 n'étaient pas maximaux. Cela contredit le fait qu'il s'agissait de classes d'états communiquants. Il y a donc contradiction, donc G_c est sans circuit.

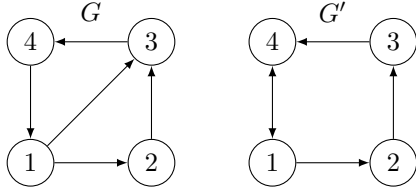
Soit un puits u de G_c et C la classe d'états communiquant associée. Soit $v \in C$. Puisqu'il n'y a pas d'arc sortant de u dans G_c alors aucun état en dehors de C n'est accessible depuis v . Donc si w est accessible depuis v , alors $w \in C$. Or w et v communiquent donc v est accessible depuis w . Donc v est permanent.

Soit un nœud u de G_c qui n'est pas un puits et C la classe d'états communiquants associée. Soit $v \in C$. Puisque u n'est pas un puits, il existe un voisin u' de u dans G_c . Soit C' la classe associée. Il existe entre C et C' un arc (w, w') . Puisque $w \in C$ alors v et w communiquent. Donc w est accessible depuis v . Puisqu'il

il y a un arc de w vers w' alors w' est accessible depuis v . Puisque G_c ne contient aucun circuit, alors il n'existe pas dans G_c de chemin de u' vers u . Donc il ne peut exister de chemin dans G de w' vers v . Donc v n'est pas accessible depuis w' , donc v est transitoire. $\square$

2.2 Période des états

La période d'un état mesure la capacité d'une chaîne à faire revenir la population de manière périodique dans un état après que ces personnes aient l'ait quitté. Considérons les deux exemples suivants.



En supposant que les probabilités soient 1 partout sauf pour les arcs sortant de l'état 1 où elles sont fixées à 0.5, et en supposant que l'état initial $Q(0)$ (à gauche) et $Q'(0)$ (à droite) vérifient $q_1(0) = q'_1(0) = 1$ et $q_i(0) = q'_i(0) = 0$ pour $i > 1$, on obtient les résultats suivants:

$$\begin{aligned} Q(1000) &= (0.28, 0.14, 0.28, 0.28) \\ Q(1001) &= (0.28, 0.14, 0.28, 0.28) \\ Q(1002) &= (0.28, 0.14, 0.28, 0.28) \\ Q(1003) &= (0.28, 0.14, 0.28, 0.28) \\ Q'(1000) &= (0.67, 0, 0.33, 0) \\ Q'(1001) &= (0, 0.33, 0, 0.67) \\ Q'(1002) &= (0.67, 0, 0.33, 0) \\ Q'(1003) &= (0, 0.33, 0, 0.67) \end{aligned}$$

On constate que, dans Q' , une régularité s'est formée avec une période 2. On pourrait considérer que Q est aussi régulier avec une période 1. On dit malgré tout que Q est apériodique alors que Q' est périodique.

Pour bien comprendre cette notion de périodicité, il faut comprendre qu'on parle ici d'une périodicité dans les probabilités et non pas dans les réalisations de ces probabilités. Ainsi, si une population se place sur l'état 1 de la chaîne dont le graphe est G' , alors au bout de 1000 itérations, on constatera que la population effectue un mouvement régulier sur la chaîne. Aux itérations paires, il y aura des gens sur les états 1 et 3 (avec environ 2 fois plus dans l'état 1). Aux itérations impaires, ces gens seront dans les états 2 et 4 (avec environ deux fois plus de gens dans l'état 4). Mais cela ne signifie pas que chaque membre de la population aura un mouvement périodique. A l'échelle de l'individu, on ne pourra pas constater quoi que ce soit. C'est seulement à l'échelle de la population qu'on voit émerger cette période.

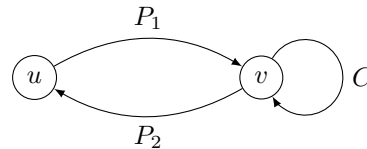
Définissons maintenant formellement cette période. On peut noter que $(P^k)_{ii} > 0$ s'il existe un circuit avec k arcs de poids non nul passant par i .

Définition 12. Soit $d_i = \text{PGCD}(k | (P^k)_{ii} > 0, k > 0)$ avec $d_i = 1$ si $(P^k)_{ii} = 0 \quad \forall k$. Un état i est dit *périodique* si $d_i > 1$, apériodique sinon. d_i est la *période* de l'état i .

Autrement dit, la période d'un état est le PGCD des tailles de tous les circuits passant par cet état. Dans les deux exemples précédents, dans G , l'état 1 est contenu dans deux circuits de taille 3 et 4, les circuits 1231 et 12341. Il existe d'autres circuits contenant 1 mais le PGCD de tous ces circuits ne peut être différent de 1 puisque c'est le seul diviseur commun de 3 et 4. Dans G' , il existe un infini de circuits passant par 2: les circuits 23412, 234123412, 2341412, 23414123412, ... On peut constater que les tailles de ces circuits couvrent tous les nombres pairs à partir de 4. Il n'existe pas de circuit de taille impair contenant 2. Donc la période de l'état 2 est $\text{PGCD}(4, 6, 8, \dots) = 2$. En appliquant le même procédé, on pourrait montrer que tous les états de G sont apériodiques et tous les états de G' sont périodiques de période 2.

Théorème 2.2. Deux états u et v d'une même classe d'états communiquants ont la même période.

Proof. Soit d_u la période de u , montrons que la période d_v de v est supérieure à d_u . Par symétrie, on peut en déduire qu'elles sont égales. Puisque u et v communiquent, il existe des chemins de u vers v et de v vers u . Pour chaque chemin P_1 de u vers v et chaque chemin P_2 de v vers u , on peut construire un circuit avec P_1 suivi de P_2 . Ce circuit passe par u donc $|P_1| + |P_2|$ est divisible par d_u . Soit un circuit C passant par v . Considérons le circuit constitué de P_1 puis C puis P_2 , il s'agit d'un circuit passant par u . Donc $|P_1| + |P_2| + |C|$ est divisible par d_u . On en déduit que $|C|$ est aussi divisible par d_u . Donc toutes les tailles des circuits passant par v sont divisibles par d_u donc leur PGCD vaut au moins d_u .



$\square$

2.3 Chaînes régulières

Les théorèmes démontrés ici ne sont pas tous généralisables au cas où les états sont infinis.

Définition 13. Une distribution Q est dite *stationnaire* si $Q = QP$. Une chaîne de Markov est dite *régulière* s'il existe une distribution stationnaire Q^* tel que, quel que

soit le vecteur de distribution initial $Q(0)$, $\lim_{t \rightarrow +\infty} Q(t)$ existe et vaut Q^* .

On peut noter que seules les distributions stationnaires peuvent vérifier la seconde propriété. On va montrer dans cette dernière partie qu'il est possible de caractériser les chaînes régulières.

Lemme 2.1. *Une chaîne qui contient deux classes d'états permanents ou plus n'est pas régulière.*

Proof. Supposons que la chaîne possède au moins deux classes d'états permanents C_1 et C_2 . Soit $i_1 \in C_1$ et $i_2 \in C_2$. Par définition, seuls les états de C_1 sont accessibles depuis i_1 (et communiquent avec). Donc si $Q(0)$ est le vecteur tel que $q_{i_1}(0) = 1$ et $q_j(0) = 0$ pour tout autre état j alors, pour tout $t \geq 0$, il existe $j \in C_1$ tel que $q_j(t) \neq 0$. De même seuls les états de C_2 sont accessibles et communiquent avec i_2 . Donc si $Q'(0)$ est le vecteur tel que $q'_{i_2}(0) = 1$ et $q'_j(0) = 0$ pour tout autre état j alors, pour tout $t \geq 0$ et tout $j \in C_1$, $q_j(t) = 0$. Ainsi, si ces limites existent, on a $\lim_{t \rightarrow +\infty} Q(t) \neq \lim_{t \rightarrow +\infty} Q'(t)$. La chaîne n'est donc pas régulière. $\square$

Lemme 2.2. *Une chaîne qui contient une classe d'états permanents dont tous les états sont de période $d > 1$ n'est pas régulière.*

Proof. Supposons que la chaîne possède une classe d'états permanents C dont les états sont de période d . Supposons sans perte de généralité que $1 \in C$. On note $N(i)$ les états dont la distance depuis 1 (le plus court chemin en nombre d'arcs) vaut i modulo d (il s'écrit $k \cdot d + i$ où $k \in \mathbb{N}$). Ainsi $1 \in N(0)$, tout état à distance $3 \cdot d$ est dans $N(0)$, tout état à distance $d + 2$ est dans $N(2)$ en supposant $d > 2 \dots$

Notons qu'un état de $N(i)$ ne peut être connecté avec un arc qu'à un état de $N(i + 1)$. Sinon il existerait un circuit passant par 1 dont la taille ne serait pas divisible par d . La preuve est similaire à celle du théorème 2.2.

Considérons maintenant $Q(0)$ le vecteur tel que $q_1(0) = 1$ et $q_j(0) = 0$ pour tout $j \neq 1$. Puisque C est une classe d'états permanents, les seuls états qui communiquent avec 1 sont ceux de C . Donc, pour tout t , $\sum_{j \in C} q_j(t) = 1$. Enfin, on peut remarquer que si $t \equiv i[d]$ alors seuls les états de $N(i)$ peuvent vérifier $q_i(t) > 0$ et donc $\sum_{j \in N(i)} q_j(t) = 1$. Ainsi, par exemple aux itérations $d, 2d, 3d, \dots, kd$, toute la population est concentrée dans $N(0)$. Et aux itérations $d + 1, 2d + 1, \dots, kd + 1$, toute la population est concentrée dans $N(1)$. Puisque $N(0) \neq N(1)$, il n'est donc pas possible que $\lim_{t \rightarrow +\infty} Q(t)$ existe. Donc la chaîne n'est pas régulière. $\square$

Démontrer que les propriétés des deux lemmes précédents sont également nécessaires n'est pas très compliqué mais nécessiterait d'introduire de nombreuses notions de plus pour être bien rédigé. On donne ici un

résultat intermédiaire bien utile et la référence vers la preuve complète.

Lemme 2.3. *Si Q est une distribution stationnaire alors, pour tout état transitoire i , $q_i = 0$.*

Remarque 2. Une manière de comprendre ce lemme est de remarquer que, à chaque étape, il existe une probabilité non nulle de se déplacer de l'ensemble des états transitoires vers un état permanent. Donc, la population sur les états transitoires va se vider petit à petit. Il faut démontrer que cette décroissance converge bien vers 0.

Proof. Soit T l'ensemble des états transitoires et R l'ensemble des états permanents. Découpons Q et la matrice de transition P en blocs.

$$Q = (Q_T \quad Q_R) \quad P = \begin{pmatrix} P_T & P_{TR} \\ 0 & P_R \end{pmatrix}$$

Par définition d'état transitoire, la probabilité d'aller de R vers T est nulle. Puisque $Q = QP$ alors $Q_T = Q_T P_T$. Notons que $I - P_T$ est une matrice carré inversible (son inverse est égale à $\sum_{n \in \mathbb{N}} P_T^n$). Donc $Q_T = 0$. $\square$

Ainsi il n'est pas nécessaire de s'intéresser aux états transitoires.

Théorème 2.3. *Si une chaîne possède une unique classe d'états permanents apériodiques alors elle est régulière.*

Proof. La preuve peut être trouvée par exemple dans le théorème 4.9 de David A Levin and Yuval Peres. *Markov chains and mixing times*. Vol. 107. American Mathematical Soc., 2017. URL: <https://pages.uoregon.edu/dlevin/MARKOV/>. Jeter avant un oeil à la section 1.5 peut être utile. $\square$

2.4 Calculer Q^*

Comment calculer Q^* connaissant la régularité de la chaîne ? Il existe deux méthodes. La première déduite de la preuve du théorème 2.3 consiste à calculer la limite de P^t . On pourra alors remarquer que cette matrice converge vers la matrice où toutes les lignes sont égales à Q^* . La seconde méthode consiste à se rappeler que Q^* est stationnaire. Elle vérifie donc $Q^* = Q^*P$. Un tel système a une infinité de solutions mais, en ajoutant la contrainte que $\sum_i q_i^* = 1$, on obtient une unique solution.

Théorème 2.4. *Si Q^* existe alors il est le vecteur solution du système*
$$\begin{cases} Q^* &= Q^*P \\ \sum_i q_i^* &= 1 \end{cases}$$