

MARCEL KADOSCH

FRANÇOIS GIRAUD

**Limites d'emploi imposées par la contrainte
de sélectivité aux systèmes de transport
guidé en site propre**

*Revue française d'automatique, d'informatique et de recherche
opérationnelle. Recherche opérationnelle*, tome 8, n° V2 (1974),
p. 63-73.

http://www.numdam.org/item?id=RO_1974__8_2_63_0

© AFCET, 1974, tous droits réservés.

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

NUMDAM

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

LIMITES D'EMPLOI IMPOSEES PAR LA CONTRAINTE DE SELECTIVITE AUX SYSTEMES DE TRANSPORT GUIDE EN SITE PROPRE ⁽¹⁾

par Marcel KADOSCH et François GIRAUD ⁽²⁾

Sommaire. — Pour assurer un service sélectif entre n stations (c'est-à-dire sans arrêt aux stations intermédiaires), avec un temps d'attente en station ne dépassant pas une limite supportable t_a , il faut mettre en circulation un trafic minimal: $n(n-1)/t_a$ véhicules/unité de temps. Si tous ces véhicules empruntent la même voie, formant un circuit fermé sur lequel les stations sont placées en dérivation, et si les véhicules successifs sont séparés par un intervalle de temps Δ de sécurité, le temps d'attente est $t_a \sim 0,5 n(n-1)\Delta$ (effet de queue non compris).

NOTION DE SELECTIVITE

Si l'on considère un réseau de transport qui dessert n stations, on dit que le service est *absolument sélectif* si un passager qui se présente à une station quelconque peut se rendre à n'importe quelle autre station sans arrêt intermédiaire.

A l'inverse, un moyen de transport, « omnibus » s'arrêtant à toutes les stations intermédiaires du réseau, comme le métro parisien par exemple, est *absolument non sélectif*. Une sélectivité relative peut être définie comme une fonction décroissante du nombre d'arrêts intermédiaires. Par exemple si un véhicule effectuant une rotation complète traverse n stations et s'arrête dans s stations, le rapport $\frac{n-s}{n-2}$, définit un degré de sélectivité variant de 1 pour la sélectivité absolue, à zéro pour le cas d'un service omnibus.

(1) Cette publication reflète une partie des travaux menés par les auteurs dans le cadre d'un contrat d'étude de l'Institut de Recherche des Transports.

(2) Société Cytec. 78 Coignières.

Dans ce qui suit, par souci de simplicité, nous considérerons uniquement la sélectivité absolue, ou service direct de la station d'origine à celle de destination.

Le service rendu par une automobile privée ou par un taxi offre cette qualité, quand les conditions suivantes sont remplies :

- les passagers du véhicule ont la même destination;
- le véhicule n'est pas obligé de s'arrêter pour laisser passer d'autres véhicules;
- le véhicule peut dépasser un véhicule qui s'arrête devant lui.

Un système de transport guidé en site propre, susceptible d'assurer un service sélectif entre n stations, doit satisfaire à des conditions analogues, qui déterminent son infrastructure et sa gestion :

Première solution : on relie toutes les stations deux à deux par des voies doubles sans intersection : comme il faut $\frac{n(n-1)}{2}$ tronçons de voie double, le coût et l'emprise au sol d'un tel réseau sont prohibitifs dans presque tous les cas (fig. 1 a).

Deuxième solution : un circuit relie toutes les stations (fig. 1 b), celles-ci étant placées en dérivation sur la voie principale : on peut se rendre à n'importe quelle station sans arrêt intermédiaire, puisqu'on peut dépasser les véhicules qui s'arrêtent aux stations intermédiaires en empruntant les dérivations.

Ce réseau est le plus économique puisqu'il ne nécessite que n tronçons de voie simple. En contrepartie les temps de trajet sont augmentés très sensiblement. On peut les réduire :

en reliant tout ou partie des stations par plusieurs circuits indépendants, entre lesquels on partage les $n(n-1)$ couples origine-destination à desservir sélectivement, en sous-ensembles de n' ($n' - 1$) couples ($n' < n$) (fig. 1 c);

en shuntant le circuit, de manière à constituer un réseau composé de plusieurs mailles : mais tout shunt introduit au moins une intersection de deux voies convergentes, où la priorité d'un des flux de véhicules sur l'autre se traduira en général par un arrêt d'attente pour insertion (fig. 1 d).

Il y a également deux variantes possibles pour la gestion des véhicules :

- 1) le véhicule en partance est préaffecté, par exemple par un contrôle centralisé;
- 2) le véhicule est banalisé : le premier voyageur qui monte décide de sa destination.

Dans ce qui suit, à titre d'illustration, nous considérerons uniquement le cas d'un circuit reliant n stations, desservies par des véhicules préaffectés.

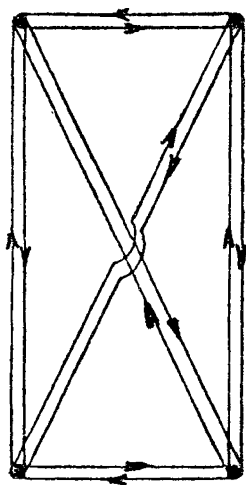


Figure 1 a

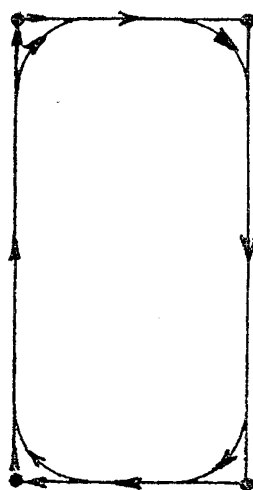


Figure 1 b

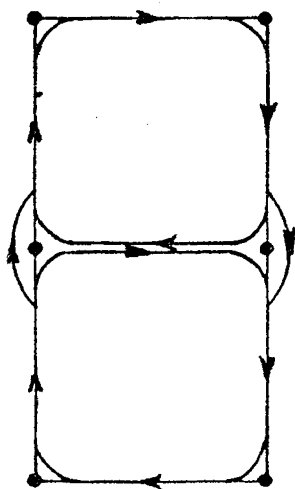


Figure 1 c

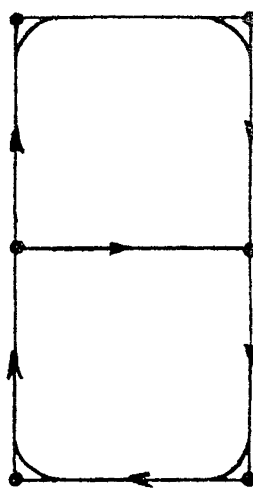


Figure 1 d

POSITION DU PROBLEME

Un système de transport sélectif en site propre ne présente pas grand intérêt si les passagers doivent attendre pendant un nombre excessif de minutes le véhicule qui doit les emmener à leur destination ou si les stations sont trop éloignées l'une de l'autre.

Notre propos est de montrer les limites d'emplois d'un système sélectif auquel on imposerait en outre les contraintes de service suivantes :

— un temps maximal d'attente dans les stations, t_a , supportable par l'ensemble des usagers, que nous appellerons « patience ». Par exemple : $t_a = 3$ minutes;

— éventuellement : une distance maximale, d_M , entre stations, s'il y a lieu de limiter la distance de marche à pied pour accéder au système. Par exemple : $d_M = 300$ à 500 mètres.

L'origine de ces limites d'emploi est la suivante :

si n stations sont ouvertes au public, il peut se présenter à n'importe quel moment des passagers pour $n(n - 1)$ couples origine-destination. Si les véhicules nécessairement différents, assurant ce service sont astreints à emprunter la même voie principale, sur laquelle ils doivent se suivre à des intervalles de temps Δ supérieurs à un minimum Δ_m de sécurité, pour éviter les risques de collision, le temps d'attente dans une station A d'un véhicule affecté à une destination B sera proportionnel à $n(n - 1)\Delta_m$: en le limitant à t_a , on ne peut servir sélectivement qu'un nombre maximal n_M de stations, couvrant une distance maximale $n_M d_M$. Que les véhicules affectés à chaque couple origine-destination soient indépendants, groupés en rames, ou portés par une voie mobile (câble, convoyeur), l'intervalle Δ_m est, dans tous les cas, au minimum celui qui est nécessaire à chaque véhicule pour entrer en station ou en sortir en empruntant une voie de dérivation, le risque de collision avec les autres véhicules étant limité à un maximum acceptable par la communauté des usagers [1].

CONDITION D'UN SERVICE SELECTIF A ATTENTE LIMITEE

A chacune des n stations du réseau de numéro i , un passager qui se présente et veut aller directement de i vers n'importe laquelle des $(n - 1)$ autres stations, sans arrêt intermédiaire, prend place dans l'une de $(n - 1)$ files d'attente existant à la station i , pendant un temps T_q , et finit par trouver un véhicule et une place dans ce véhicule pour la station j , où il désire aller : il attend encore pendant un temps T_r , que ce véhicule charge d'autres voyageurs et ferme ses portes, soit au total $T_a = T_q + T_r$.

Si l'espérance mathématique $E(T_a)$ est trop grande, c'est-à-dire si après avoir essayé souvent le système, le passager trouve qu'en moyenne il attend plus qu'une limite supportable t_a , il abandonnera le système.

Nous cherchons la limite de validité de ce service, quand on impose au système la contrainte :

$$E(T_a) \leq t_a, \text{ patience de l'utilisateur (en secondes),}$$

qui définit une qualité de service à respecter.

Nous définissons encore deux matrices origine-destination $[Y]$ et $[\Theta]$ représentant respectivement :

$y(i, j)$ (passagers/heure) : demande de transport à la station i pour la station j ;

$\theta(i, j)$ (véhicules/heure) : trafic de véhicules allant de la station i à la station j .

On suppose que les véhicules ont une capacité de K places.

Plaçons nous à une station i (d'origine) et observons les passagers et véhicules allant vers une station j (de destination). A l'instant où les portes d'un véhicule reliant ces deux stations se ferment, on se trouve devant la situation suivante :

- V personnes sont montées dans le véhicule : $V \leq K$;
- W personnes désirant aller de i en j restent sur le quai.

$X = V + W$ est le nombre total de personnes, d'espérance $E(X)$ qui à l'instant considéré souhaitent aller de i en j ; si ces personnes séjournent en moyenne dans la station pendant le temps $E(T_a)$, le rapport $E(X)/E(T_a)$ est par définition le taux moyen d'arrivée des passagers, donc :

$$E(X) = \frac{y(i, j)}{3\,600} E(T_a).$$

Le coefficient de remplissage $g(i, j)$ des véhicules considérés satisfait de son côté aux équations :

$$\begin{aligned} y(i, j) &= g(i, j) K \theta(i, j) && \text{(passagers/heure)} \\ E(V) &= g(i, j) K && \text{(passagers/véhicule),} \end{aligned}$$

d'où l'on déduit :

$$\frac{y(i, j)}{3\,600} = \frac{E(X)}{E(T_a)} = \frac{E(V)}{3\,600} \theta(i, j),$$

pourvu que $g(i, j)$ ne soit pas nul, ce qui implique qu'on soit astreint à servir la liaison (i, j) pendant la période considérée.

Comme : $E(X) = E(V) + E(W) \geq E(V)$,

il est nécessaire que :

$$\frac{E(X)}{E(V)} = E(T_a) \frac{\theta(i, j)}{3\,600} \geq 1.$$

La contrainte de service $E(T_a) \leq t_a$ impose donc la mise en circulation d'un trafic minimal :

$$\theta(i, j) \geq 3\,600/t_a \text{ véhicules/heure,}$$

sur la liaison (i, j) .

Si le service dans la station i est sélectif pour les $(n - 1)$ destinations j possibles, on trouve que le trafic de véhicules D_i traversant la station est tel que :

$$\sum_j \theta(i, j) = D_i \geq 3\,600(n - 1)/t_a.$$

Pour que le service soit sélectif dans toutes les stations, il faut mettre en circulation un trafic total Θ tel que :

$$\Theta = \sum_i D_i \geq 3\,600n(n - 1)/t_a \text{ véhicules/heure.}$$

Cette condition ne dépend que du nombre n de stations quel que soit le réseau de voies qui les relie.

CAS DU CIRCUIT HAMILTONIEN

Suivant la terminologie adoptée en théorie des graphes, on appelle circuit hamiltonien un circuit reliant n stations (mises en dérivation pour assurer la sélectivité) et passant par chacune une fois et une seule.

Commençons par rappeler quelques propriétés d'un tel circuit (fig. 2) :

a) Le débit Φ véhicules/heure passant en tout point de la voie principale est constant et se partage à chaque station en :

$$\begin{aligned} A_i &= D_i \text{ véhicules/heure empruntant la voie de dérivation;} \\ (\Phi - A_i) &\text{ véhicules/heure empruntant la voie de transit;} \end{aligned}$$

b) Le débit de ligne est égal à :

$$\Phi = \sum_{i > j} \theta(i, j) \text{ véhicules/heure} = 3\,600/\Delta \leq \Phi_M = 3\,600/\Delta_m,$$

en appelant Δ (secondes) l'intervalle de temps séparant deux véhicules, qui ne peut en moyenne être supérieur à un minimum Δ_m déterminé par une condition de sécurité;

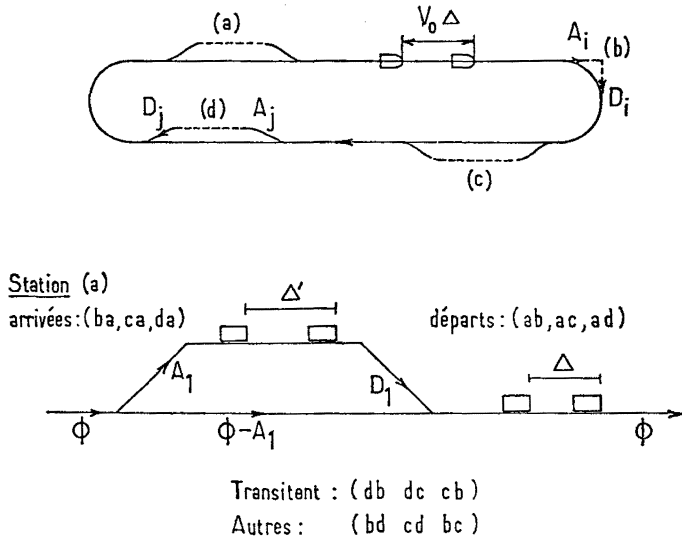


Figure 2

Exemple de circuit hamiltonien avec stations placées en dérivation sur la voie principale : détail d'une station

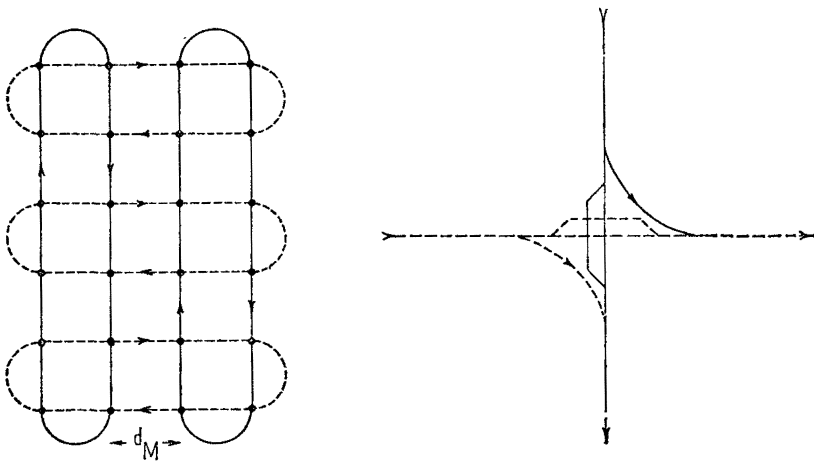


Figure 3

Réseau de Manhattan. Voies et station type

c) Si D est la longueur totale du circuit, et D_v le trajet moyen effectué par un véhicule à chaque voyage, le trafic Θ et le flux Φ vérifient la relation :

$$D\Phi = \Theta D_v \text{ (véhicules kilomètres/heure)}$$

ou

$$\Theta = \frac{D}{D_v} \Phi = \frac{D}{D_v} \frac{3\,600}{\Delta_m}$$

L'inéquation ci-dessus s'écrit alors :

$$t_a \geq \frac{D_v}{D} n(n-1)\Delta_m$$

Le cas le plus favorable (mais aussi le moins probable) est celui où presque tous les passagers ne veulent parcourir que la distance minimum $D_v = \frac{D}{n}$.

Le cas le plus défavorable correspond à une demande symétrique : les trajets (i, j) et (j, i) sont également demandés et $D_v = D/2$. Dans la majeure partie des cas, D_v/D est compris entre 0,3 et 0,5 (voir appendice).

Application numérique : $t_a = 180$ secondes (3 minutes).

Si $\Delta_m = 20$ secondes : $n \leq n_M = 5$ stations.

Si $\Delta_m = 2$ secondes : $n \leq n_M = 14$ stations.

Considérons par exemple le réseau dit « de Manhattan » représenté (fig. 3) : l'aire est desservie par un certain nombre de circuits nord-sud et est-ouest formant un quadrillage de pas d_M : il y a n_1 stations sur un circuit nord-sud et n_2 stations sur un circuit est-ouest : pour aller d'une station à une autre n'appartenant pas à la même ligne, le temps d'attente total est supérieur à $\frac{\Delta_m}{2} [n_1(n_1 - 1) + n_2(n_2 - 1)]$. Si ce temps ne doit pas dépasser 3 minutes, on ne peut guère desservir une aire plus étendue que celle représentée par la figure, même en ne tenant pas compte du temps d'attente possible aux nombreuses intersections du réseau.

CONCLUSIONS ET GENERALISATIONS

Des considérations très simples nous ont permis de définir la géométrie d'un réseau de transport en site propre sur lequel un service absolument sélectif est possible. Les stations sont nécessairement placées en dérivation sur la voie principale de circulation.

Nous avons montré qu'un service absolument sélectif entre n stations requiert la mise en circulation de :

$$\Theta \geq 3\,600 \frac{n(n-1)}{t_a} \text{ véhicules/heure,}$$

si l'on veut limiter à t_a secondes le temps moyen d'attente des passagers dans la station de départ; ceci quel que soit le réseau de voies reliant les stations.

Dans la solution la plus économique, où ces véhicules empruntent une même voie formant un circuit qui relie les n stations, et sont séparés par l'intervalle de sécurité Δ_m secondes, le temps moyen d'attente ne peut être inférieur à :

$$t_a \geq \frac{D_v}{D} n(n-1)\Delta_m;$$

$\frac{D_v}{D}$, rapport du trajet moyen sur la longueur totale du circuit est compris entre 0,3 et 0,5 pour les demandes usuelles.

L'intervalle Δ_m est au minimum celui qui est nécessaire pour que les véhicules puissent entrer dans les stations ou en sortir en empruntant les voies de dérivation, le risque de collision avec les autres véhicules étant limité à un maximum acceptable par la communauté des usagers [1].

En approfondissant systématiquement l'analyse du service de transport collectif, nous avons pu calculer la moyenne et la variance du temps d'attente dans le cas le plus général, quand les arrivées de passagers et de véhicules obéissent à une loi quelconque. Il n'est pas question de développer ici ces calculs qui occuperaient une place considérable : contentons-nous de dire que les conclusions énoncées ci-dessus subsistent dans le cas le plus général, du moment que la voie est occupée à intervalles Δ par des véhicules ayant $n(n-1)$ affectations différentes dont les passagers ne disposent pas librement; si le passager peut déterminer une affectation, par exemple s'il choisit librement la destination d'un véhicule banalisé de $K=1$ place, son temps d'attente se réduit à :

$$t_a = \frac{D_v}{D} n\Delta,$$

mais un tel système a une capacité de débit de passagers forcément très limitée.

L'analyse montre que les conclusions ci-dessus ne sont mises en défaut que dans un seul cas : celui où les véhicules suivent un horaire connu des passagers. La loi d'arrivée des passagers dépend alors évidemment de leur comportement vis à vis de cet horaire.

L'analyse montre aussi qu'un service avec un ou plusieurs arrêts à des stations intermédiaires est de beaucoup supérieur à un service entièrement

sélectif, car la réduction du temps d'attente est considérablement plus grande que le temps perdu aux arrêts intermédiaires. Si nous reprenons la définition du degré de sélectivité $\frac{n-s}{n-2}$ donnée à l'introduction, en fonction du nombre s d'arrêts du véhicule par rotation sur le circuit ($s = 2$ pour la sélectivité absolue), la limite du temps d'attente devient :

$$t_a \geq \frac{D_v}{D} \frac{2n(n-1)}{s(s-1)} \Delta_m.$$

APPENDICE

CAS D'UNE VOIE DOUBLE ENTRE DEUX TERMINAUX

Un autre cas particulier de circuit hamiltonien est constitué par une voie double reliant deux stations terminales A et B , le long de laquelle on place en

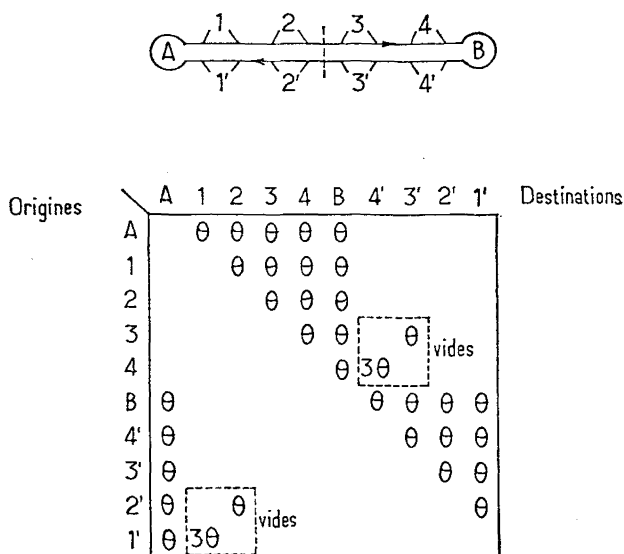


Figure 4

Voie double entre deux terminaux
Matrice de trafic des véhicules

dérivation des stations en n_1 points intermédiaires, avec un quai desservi dans le sens AB et un autre desservi dans le sens BA : chaque quai compte pour une station (fig. 4), mais il n'y a que $N = n_1 + 2$ points desservis.

Si $n_1 = 2n'$ est pair il y a au total $n = 4n' + 2$ « stations ».

Si $n_1 = 2n' - 1$ est impair il y a au total $n = 4n'$ « stations ».

Les passagers qui se présentent prennent bien entendu le quai qui leur permet d'aller à leur destination sans revenir en arrière en transitant par A ou B . Même si la demande est uniforme sur les $N = (n_1 + 2)$ destinations possibles, la matrice de trafic se présente sous la forme indiquée fig. 4 avec θ véhicules/heure par couple origine-destination. Appelons :

Θ_1 , le trafic de véhicules nécessaire pour acheminer les passagers (dans le cas d'une demande uniforme entre les N destinations);

Θ_2 , le trafic supplémentaire de véhicules vides nécessaire pour qu'il n'y ait pas accumulation des véhicules dans une station ($A_i = D_i$).

Si $n = 4n' + 2$:

$$\begin{aligned}\Theta_1 &= (2n' + 1)(2n' + 2)\theta; & \Phi &= (n' + 1)^2\theta; \\ \Theta_2 &= 2n'^2\theta; & 2 \leq \Theta/\Phi &= D/D_v \leq 6.\end{aligned}$$

Si $n = 4n'$:

$$\begin{aligned}\Theta &= 6n'^2\theta; \\ \Phi &= n'(n' + 1)\theta; & \Theta/\Phi &= D/D_v = 6n'/(n' + 1).\end{aligned}$$

Dans le 2^e cas :

$$t_a > \frac{N(N+1)}{6} \Delta.$$

BIBLIOGRAPHIE

- [1] M. KADOSCH, *Determining factors for the minimum headway in a transportation system*, 1st National Conference on Personal Rapid Transit. University of Minnesota. Minneapolis, November 1971.