
Définition et utilisation opérationnelle du concept d'accessibilité

Dans ce chapitre, nous proposons une définition de la notion d'accessibilité, puis, à partir des entretiens réalisés avec les acteurs, nous identifions les domaines auxquels celle-ci s'applique et nous examinons le déroulement type des études qui la mettent en œuvre.

1.1 Définition de l'accessibilité

Un *déplacement* peut être défini comme un mouvement effectué dans un espace entre deux lieux – une origine et une destination – et caractérisé par un motif.

L'*accessibilité*, quant à elle, peut être définie comme l'accomplissement du franchissement spatial entre deux points répondant au motif du déplacement de la personne. On peut également dire qu'elle est la capacité d'atteindre les biens, les services et les activités désirés par un individu. Le concept d'*accessibilité* a toutefois de multiples définitions selon les contextes et les applications. Classiquement, l'*accessibilité* est définie comme la « *plus ou moins grande facilité pour atteindre les opportunités souhaitées* » (voir référence [1] dans la bibliographie en fin d'ouvrage), comme un « *potentiel d'interaction entre les aménités d'un territoire et les individus* » [2] ou comme une « *utilité retirée des avantages du système des transports et de l'usage des sols* » [3].

La notion d'*accessibilité* revêt plusieurs formes, en fonction des objectifs de développement économique ou d'aménagement du territoire auxquels elle s'applique. « *La notion d'accès aux fonctions urbaines [...] a évolué au cours des cinquante dernières années. L'accessibilité a concerné l'accès géographique ou spatial lorsqu'il s'est agi de faciliter le développement économique en dotant les villes d'équipements de transports, qu'ils soient routiers pour l'automobile ou des infrastructures de transports collectifs [...]. On a ensuite parlé d'accessibilité sociale, pas uniquement en cherchant des solutions au désenclavement de quartiers de plus en plus dispersés mais aussi en cherchant à compenser les inégalités que la naissance ou la fortune a établies entre les individus.* » [4]

Dans [5], l'hypothèse est faite qu'un territoire peut être interprété comme étant l'imbrication de trois sous-systèmes, chacun doté de sa logique de fonctionnement et de transformation et relié aux autres sous-systèmes par des relations de causalité. L'*accessibilité* se définit alors selon les interactions considérées entre ces sous-systèmes [6] :

- la localisation des agents territoriaux, résidences et opportunités (activités, biens et services), qui constitue la *composante spatiale* de l'*accessibilité* : elle renvoie aux questions d'aménagement du territoire et d'usage du sol ;
- les pratiques et rapports sociaux (programmes d'activités) ainsi que les caractéristiques socioprofessionnelles et les perceptions des individus qui constituent la *composante individuelle* de l'*accessibilité* : elle renvoie aux besoins, capacités et opportunités des individus en termes d'*accessibilité* ;
- le système de transport : il est le support des deux précédentes composantes et détermine la distribution spatiale des déplacements et l'effort que les individus fournissent pour atteindre les opportunités du territoire (temps et coûts).

L'*accessibilité* tient compte de la répartition spatiale des opportunités et du potentiel social des individus, c'est-à-dire de leurs capacités individuelles à se déplacer pour atteindre ces opportunités, en utilisant le système de transport. En complément, il convient d'intégrer une *composante temporelle* dans l'interprétation de l'*accessibilité*, puisque celle-ci est influencée par les temps d'ouverture qui régissent l'accès aux biens et aux services à différents moments de la journée, par le temps que les individus accordent à ces activités et par la qualité du système de transport en fonction des moments de la journée (période de pointe, période creuse, soirée...).

L'illustration 1 représente les interactions entre les sous-systèmes qui influencent l'accessibilité des territoires.

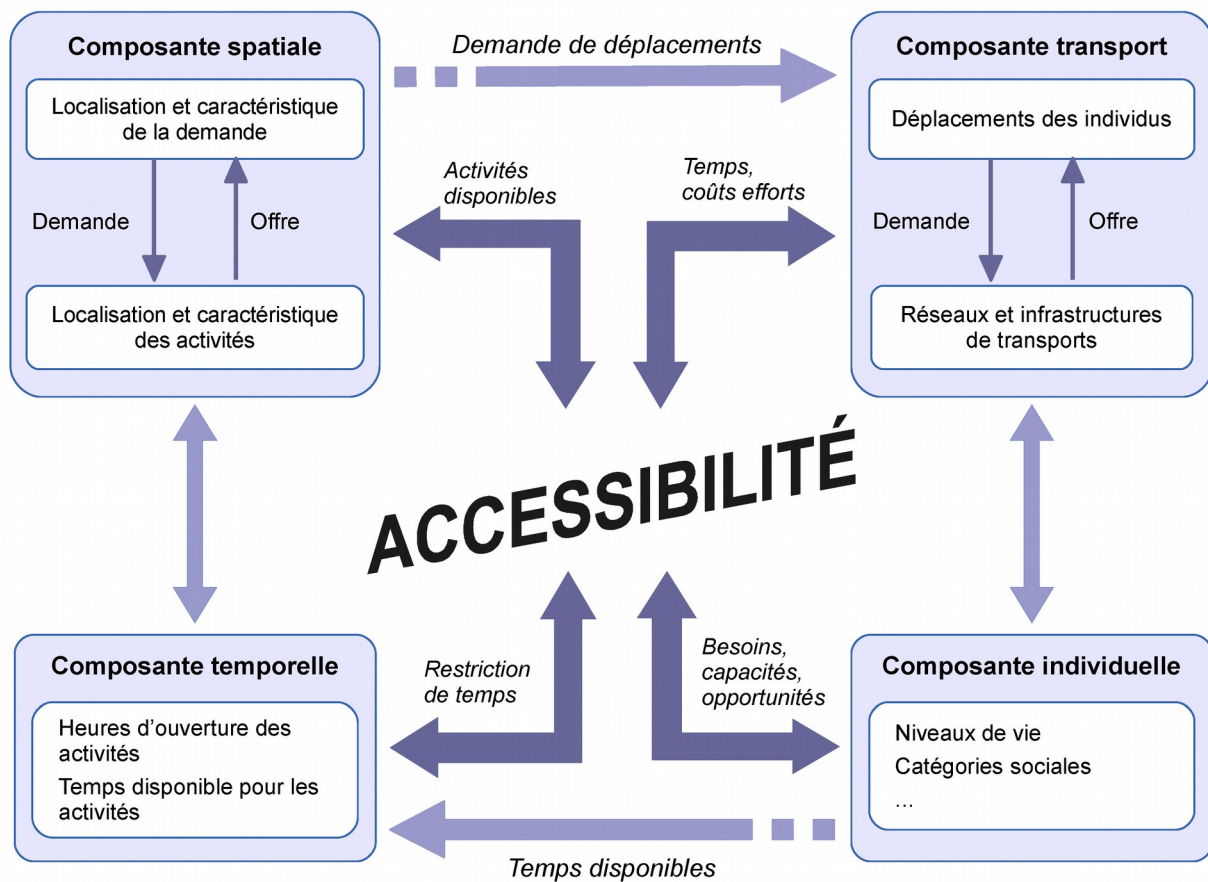


Illustration 1: Les différentes composantes de l'accessibilité

1.2 Quelles applications au concept d'accessibilité ?

Lors des entretiens, différents types de travaux allant d'études apportant une aide à la décision immédiate (par ex. étude F, présentée en annexe) à des travaux de recherche (par ex. études O, P, R et U) ont été recensés. On y retrouve les principaux domaines d'application du calcul d'accessibilité identifiés par l'action COST « *Accessibility Instruments for Planning Practice* » [7] :

1. *choix de localisation d'activités, de services ou de zones résidentielles* (par ex. études E, F, K, T ou encore [8]) ;
2. *gestion des offres de transport et évaluation des mesures à mettre en œuvre pour favoriser le report modal* (par ex. études A, B, D, E, I, L, O, P, Q et R) ;
3. *évaluation des mesures à mettre en œuvre pour tendre vers une meilleure équité sociale d'accès à la mobilité et aux opportunités offertes par la ville* (par ex. études G, H, J et N) ;
4. *stimulation de l'activité économique* (par ex. études C et M) ;
5. *réduction de l'impact environnemental des transports aux travers de leurs émissions de polluants et de la consommation d'énergie* (par ex. étude U).

Selon les études, les quatre composantes de l'accessibilité (voir. illustration 1) sont prises en compte de manière plus ou moins complète dans la situation de référence ainsi que dans les éventuels scénarios d'évolution construits.

1.2.1 Focalisation sur la composante *transport*

L'accessibilité est souvent considérée en se focalisant sur la composante *transport*, éventuellement en considérant conjointement la composante *individuelle*, mais sans prendre en compte la composante *spatiale*. On se trouve alors dans une évaluation du potentiel de déplacement offert par le réseau de transport, qui correspond à une vision simplifiée de l'accessibilité. Tous les couples origine-destination étudiés sont considérés comme représentant des potentiels d'activités équivalents, la densité des opportunités n'étant pas prise en compte. Par abus de langage, on parle malgré tout d'*accessibilité* dans ce contexte alors que le lien des transports au territoire n'est pas ou peu fait dans l'analyse.

Les applications recensées concernent l'évaluation de :

- la performance globale des réseaux pour desservir un type de lieux (par ex. étude M) ;
- la coordination entre les différentes offres de transport collectif, en vue de démontrer que des modifications d'horaires des transports collectifs ou une diminution du temps de marche en correspondance par une amélioration de l'organisation d'un pôle d'échange pourraient améliorer le service rendu aux usagers et rendre les transports collectifs plus attractifs (par ex. études E et P, ou encore [9] et [10]) ;
- la vulnérabilité du réseau en cas d'incident, ainsi que les possibilités de reroutage que celui-ci offre aux usagers (par ex. effets d'une panne du métro évalués dans l'étude O), pour aider à la décision concernant la gestion d'incidents et orienter les investissements sur les infrastructures de transport.

Ce type d'approche utilise généralement les vitesses de déplacement et les temps de parcours comme mesure de l'accessibilité. D'autres indicateurs sont parfois utilisés, quantifiant les possibilités de correspondances aux nœuds du réseau de transport collectif, par exemple l'indicateur de *cohérence de l'offre* [11], l'indicateur d'*intensité nodale* [10] ou encore l'indicateur de *contactabilité*, décrivant la possibilité de réaliser un aller-retour sur une journée pour un couple origine-destination donné (par ex. étude P).

1.2.2 Prise en compte de l'interaction entre *transports* et *aménagement du territoire*

Une vision plus complète de l'accessibilité consiste à considérer l'interaction entre les transports et les localisations des opportunités sur le territoire. Se placer dans ce contexte permet de souligner que les transports comme l'aménagement du territoire peuvent constituer des leviers d'action des politiques publiques. L'accessibilité vient alors compléter l'analyse socio-économique des projets.

C'est l'approche retenue par exemple par l'étude E, qui cherche à détecter des décalages entre les projets de localisation d'activités et le réseau de transport existant ou en projet. On retrouve également cette approche dans l'étude R, qui cherche à évaluer l'effet de différentes politiques de tarification des transports sur l'accessibilité à l'emploi dans l'agglomération lyonnaise. La veille et l'évaluation des actions sur le foncier, en fonction de ses caractéristiques d'accessibilité, entrent également dans le cadre de cette approche [8].

1.2.3 Prise en compte de l'équité sociale et territoriale

Enfin, l'accessibilité permet dans certains travaux d'évaluer les inégalités d'accès aux opportunités, à l'intérieur de la population ou entre les territoires. Ainsi, l'étude G propose « *une vision étayée et partagée entre partenaires institutionnels des enjeux et inégalités socio-urbaines (mesure de l'enclavement des quartiers) pour servir de base au Contrat Urbain de Cohésion Sociale* ». C'est aussi le cas de l'étude P proposant de « *faire avancer les connaissances sur les indicateurs d'accessibilité et sur leurs utilisations dans le cadre de l'évaluation de projets de transport, dans l'optique plus particulière d'éclairer les relations entre les transports et l'équité territoriale* ». On relève encore l'étude O, qui évalue de façon différenciée l'accessibilité des pôles d'excellence de l'agglomération lilloise selon la catégorie socioprofessionnelle, ou encore l'étude L, qui évalue la perte d'accessibilité à l'emploi pour les personnes à mobilité réduite, pour les résidents de la périphérie de la ville par rapport à ceux du centre, ainsi que les effets cumulés de ces deux situations (mobilité réduite et résidence en périphérie).

[12] et [13] soulignent également que les mesures d'accessibilité permettent de prendre en compte l'aspect social de la mobilité dans l'évaluation des projets : « *La question de l'égalité de traitement entre les territoires et les agents est souvent posée, mais est peu prise en compte par le calcul socio-économique. Un programme rentable lorsque l'on considère la société dans son ensemble n'est pas nécessairement générateur de justice, dans le cas où ce sont les usagers des régions les plus accessibles initialement qui bénéficient le plus du projet après sa mise en service. L'accessibilité peut être traitée de manière qualitative, notamment par des représentations cartographiques, en plus du calcul socio-économique. Cela peut correspondre à un traitement de l'inégalité sociale dans le cas où les zones dont la desserte est améliorée concentrent des phénomènes d'exclusion et de pauvreté. Cependant, il reste difficile de prévoir quels peuvent être les effets des infrastructures de transport sur les inégalités de répartition des richesses.* » [12]. L'accessibilité apparaît donc comme un indicateur à mettre en balance avec le calcul économique pour évaluer l'opportunité des projets de transport. Cette prise en compte d'effets non monétarisables d'un projet est un des éléments majeurs de la méthode d'évaluation des projets préconisée par le ministère du Développement durable [14].

1.3 Contexte de réalisation des études d'accessibilité en France

Les agences d'urbanisme interrogées signalent une certaine méconnaissance de la mesure d'accessibilité comme outil d'analyse de la part des maîtres d'ouvrage de projets d'aménagement et de transport. Par conséquent, il est parfois nécessaire pour elles de réaliser des productions de leur propre initiative afin de valoriser et de faire connaître ce type d'analyses, notamment en démontrant leur apport complémentaire aux modèles de déplacements statiques.

Certaines études entrent dans le cadre d'observatoires : dans ce cas, le producteur d'études réalise un travail de veille sur un sujet, met à jour son outil de calcul d'accessibilité et fournit de temps en temps une étude sur ce sujet à différents maîtres d'ouvrage. C'est souvent le cas des agences d'urbanisme sur la question de l'accessibilité aux quartiers dits de « politique de la ville ». Ce type d'organisation permet d'être réactif lorsqu'un besoin d'analyse est exprimé par un maître d'ouvrage.

Lorsqu'il y a une commande, celle-ci est parfois assez floue initialement, étant donné la complexité des questions en relation avec l'accessibilité. La méthodologie est souvent affinée au cours du travail pour correspondre au mieux aux souhaits de la maîtrise d'ouvrage et tenir compte des contraintes liées à l'étude, notamment la disponibilité des données et les capacités des outils logiciels. Le producteur de l'étude doit donc présenter au maître d'ouvrage les contraintes qui pèsent sur l'étude ainsi que lui proposer une palette d'indicateurs, de paramétrages et de modes de représentation possibles. Les choix finaux doivent être effectués en coordination avec le maître d'ouvrage, en fonction de son vécu du territoire et de son objectif : ce fut le cas par exemple dans l'étude E pour valider le choix de zones d'activités ou centres commerciaux à rayonnement métropolitain qui seraient étudiés.

L'appropriation des résultats de l'étude par le maître d'ouvrage en vue de leur valorisation nécessite la mise en œuvre d'une démarche pédagogique de la part du producteur d'études. Les entretiens ont permis de mettre en évidence la difficulté que peut représenter la vulgarisation des indicateurs autres que le temps de parcours, notamment l'indicateur gravitaire. Par ailleurs, des maîtres d'ouvrage et producteurs d'études interrogés mentionnent la réticence de certains décideurs à communiquer sur la réalité des temps de parcours en transports collectifs.

2. Méthodologies de calcul de l'accessibilité

On trouve des applications et des définitions des mesures d'accessibilité aussi bien en géographie qu'en économie. L'état de la littérature sur les mesures d'accessibilité, sur la base de [15], met en évidence les modèles utilisés dans les études théoriques ou empiriques sur les transports. Nous faisons ici un tour d'horizon des indicateurs d'accessibilité des territoires les plus classiques d'un point de vue théorique et, dans la mesure du possible, nous précisons leur finalité, leurs avantages, leurs inconvénients et leurs limites, puis nous voyons comment ils sont mis en œuvre dans les études.

2.1 Grandes familles d'indicateurs

Les indicateurs d'accessibilité sont rattachés à une définition plus ou moins précise et poussée du concept d'accessibilité. La littérature propose plusieurs classifications de ces indicateurs par niveau de complexité ou par familles, en tenant compte des contours des sous-systèmes influant sur l'accessibilité (voir chapitre 1).

[16] propose une première catégorisation qui distingue trois types d'indicateurs :

- les indicateurs dits « *réalistiques* », qui mesurent les propriétés géométriques de l'espace à travers la structure du réseau de transport (par exemple, les isochrones) ;
- les indicateurs dits « *économiques* », qui caractérisent l'accessibilité potentielle aux opportunités territoriales. Ils prennent en compte dans le calcul de l'accessibilité, plus ou moins complètement, les opportunités ainsi que les interactions spatiales (par exemple, les modèles gravitaires) ;
- les indicateurs dits « *prismes spatio-temporels* » issus de l'école de la « Time Geography », qui mesurent l'accessibilité en tenant compte des possibilités de déplacement dans le territoire dans l'espace et le temps.

Cette catégorisation prend en compte les différentes composantes de l'accessibilité (*spatiale, individuelle, transport et temporelle*), de manière partielle pour certaines de celles-ci, mais ne distingue pas les indicateurs selon leur degré de faisabilité.

Une autre approche, mise en avant par l'agence d'urbanisme de la région stéphanoise [17] conduit à considérer trois niveaux d'indicateurs d'accessibilité :

- *mesures de temps de parcours* depuis ou vers un lieu particulier en décrivant un niveau d'offre de transport (par exemple, les isochrones) : ce premier type d'indicateur se rapproche des mesures « réalistiques » en ne considérant que la composante transport de l'accessibilité ;
- *croisements entre les informations de temps de parcours et des données socio-économiques du territoire* (par ex. pourcentage de population ou nombre d'emplois à moins de 30 minutes d'une gare) : ces indicateurs visent à décrire sommairement le territoire, mais se confondent partiellement avec les mesures « réalistiques » (notamment, par les isochrones) et se rapprochent également des mesures « économiques » au sens où ils intègrent à minima la complexité des interactions entre opportunités et réseaux de transports (isochrones avec superposition d'informations, modèle gravitaire simple). Ils prennent en compte la composante transport et partiellement la composante spatiale de l'accessibilité ;
- *croisement des données d'offre et de demande de transport* « issu d'une analyse de l'utilité des modes qui s'appuie sur le postulat que, pour un motif donné, le choix de destination des usagers est conditionné par le niveau d'offre de transport. » [17]. Par exemple, on cherchera à mesurer l'utilité des transports collectifs pour le motif de déplacements domicile-travail sur l'ensemble d'un territoire. Ces indicateurs, proches des mesures « économiques », s'appuient sur la théorie de maximisation de l'utilité [18] et ont donc tendance à considérer partiellement la composante *sociale* de l'accessibilité.

On peut trouver bien d'autres classifications dans la littérature. Nous choisissons ici d'utiliser une classification en trois familles, proches des deux précédentes, se rapprochant des définitions retenues dans le chapitre précédent, en tentant de classer les indicateurs par finalité et par ordre croissant de complexité.

- *Mesures d'accessibilité basées sur une composante spatiale fixe*
Ces indicateurs prennent en compte principalement la composante *transport* et de manière croissante la composante *spatiale* et leurs interactions. Ils se basent sur le système de transport, sur la localisation fixe

des individus et des opportunités au sein du territoire. Les indicateurs les plus complets prennent en compte la résistance des déplacements ainsi que les interactions spatiales des concurrences entre individus et opportunités.

- *Mesures d'accessibilité basées sur l'individu et le temps, dites « prismes spatio-temporels »*
Ces indicateurs prennent en compte les composantes *transport* et *spatiale* (sans les interactions liées aux concurrences) ainsi que la composante *temporelle* et partiellement *sociale*. Ils se basent sur l'interprétation de l'accessibilité en fonction des contraintes d'espace-temps liées aux déplacements et aux opportunités utilisées par les individus (programmes d'activités quotidiens).
- *Mesures d'accessibilité basées sur le maximum d'utilité*
Ces indicateurs prennent en compte les composantes *transport*, *spatiale* et *individuelle* dans la détermination de l'accessibilité, celle-ci étant interprétée comme la maximisation de l'utilité individuelle parmi l'ensemble des destinations permettant d'accéder à une opportunité d'un type donné. On trouve ici les indicateurs de la théorie des choix discrets [18]. Les indicateurs les plus complets (et donc les plus complexes à mettre en œuvre) vont jusqu'à prendre en compte la composante *temporelle* en couplant les indicateurs d'utilité avec des « prismes spatio-temporels ».

2.1.1 Mesures basées sur une composante spatiale fixe

Ces indicateurs d'accessibilité considèrent la localisation des individus, des activités, des biens et des services ainsi que le système de transports. On recense dans cette catégorie les indicateurs suivants :

- isochrone ;
- modèle gravitaire simple ;
- modèle gravitaire avec effets de concurrence liés à la demande ;
- modèle gravitaire avec effets de concurrence liés à la demande et à l'offre (facteurs d'équilibres inverses).

Isochrone

Courbe isochrone

L'isochrone désigne une courbe géolocalisée, délimitant un territoire où chaque point est accessible depuis une origine fixée (ou ayant accès à une destination fixée), avec un coût de déplacement inférieur à une valeur x . Ce coût peut correspondre à une distance, un temps ou un coût généralisé (il s'agit alors plutôt d'une courbe « isovaleur »). L'illustration 2 donne un exemple de courbe isochrone pour deux cas monomodaux (voiture et transports en commun) et pour un cas intermodal (voiture + transports en commun).

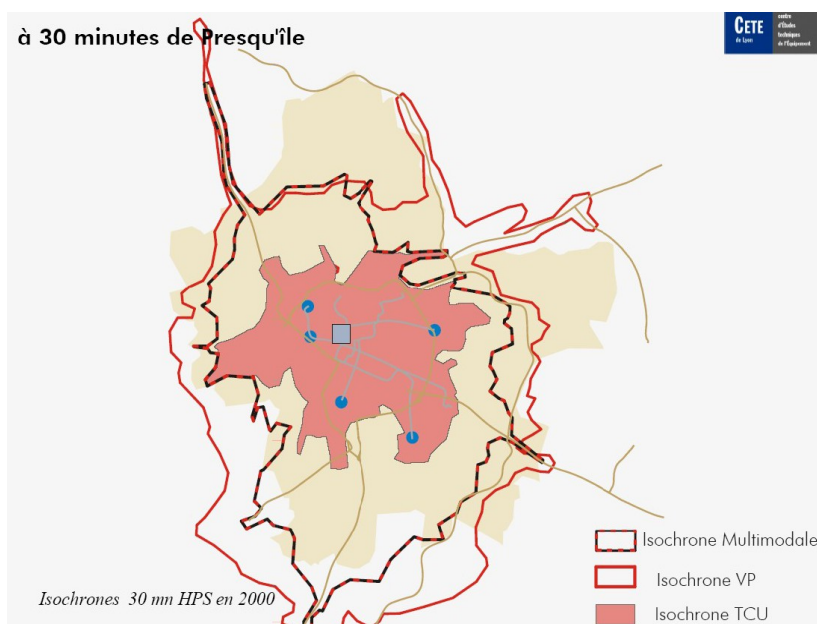


Illustration 2 : Courbes isochrones à 30 minutes depuis le centre de Lyon (unipolaires) en heure de pointe du soir en 2000 en voiture, en transports en commun et en combinant les deux (isochrone multimodale), source [68]

L'illustration 3 présente des variations de courbes isochrones, pour différents scénarios d'offre de transport.

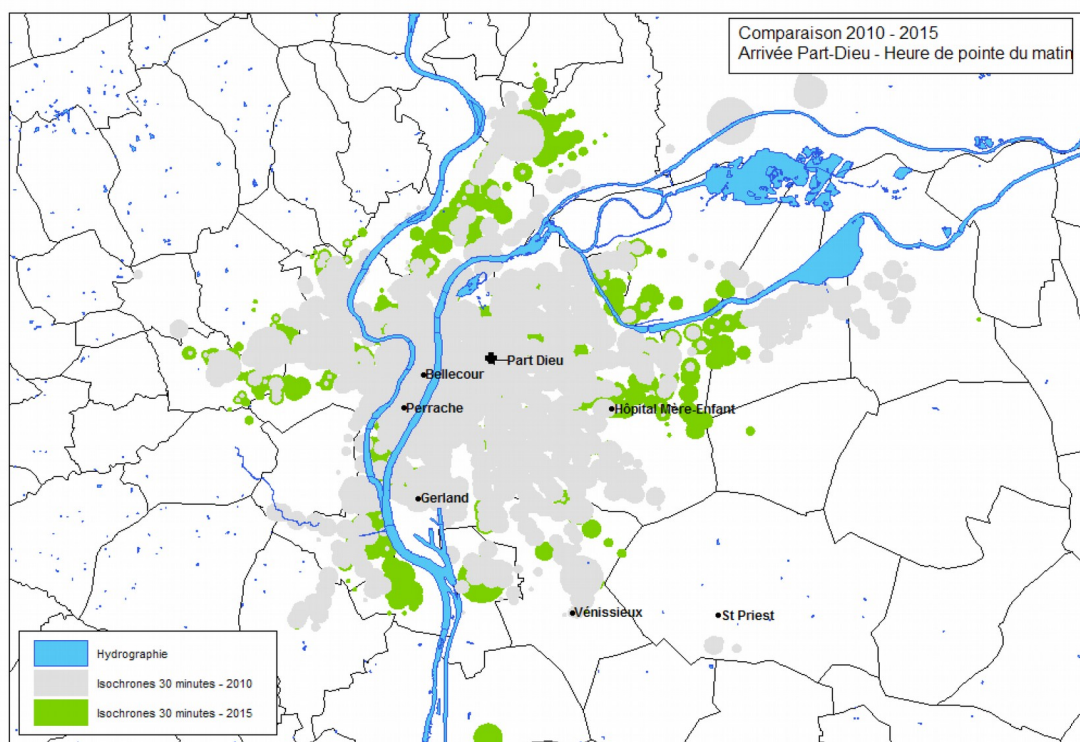


Illustration 3 : Isochrones vers le quartier de la Part-Dieu à Lyon, en transport collectif, pour deux scénarios d'offre de transport, en 2010 et 2015, source [60]

On peut également construire une courbe isochrone multipolaire, c'est-à-dire, par exemple, en reprenant l'illustration 3, représenter la zone accédant en moins de 30 minutes à au moins une des deux gares principales lyonnaises : Lyon Part-Dieu et Lyon Perrache. Cette zone correspond alors à la réunion des zones situées à l'intérieur des courbes isochrones associées à chaque gare.

Indicateur dérivé de la courbe isochrone

La courbe isochrone ne constitue donc pas à proprement parler un indicateur numérique d'accessibilité, toutefois il est possible d'en construire un en dénombrant les opportunités, par exemple les emplois, situées à l'intérieur de cette courbe.

L'indicateur « isochrone » est alors défini par :
$$A_i = \sum_{j/d_{ij} \leq X} O_j$$

avec :

- i , la zone d'origine ;
- j , la zone de destination du déplacement réalisé par un individu depuis la zone i ;
- X , le coût maximum de déplacement acceptable ;
- A_i , l'accessibilité des individus localisés dans la zone i ;
- O_j , le volume d'opportunités de la zone j ;
- d_{ij} , le coût (distance, temps ou coût généralisé) de déplacement entre la zone i et la zone j .

Encadré 1 : Présentation de l'indicateur « isochrone »

Ainsi, l'illustration 4 présente l'accessibilité à la population depuis différents quartiers lyonnais. Chaque point de la courbe représente en ordonnée le décompte de la population comprise dans la courbe isochrone, dont l'amplitude X est donnée en abscisse.

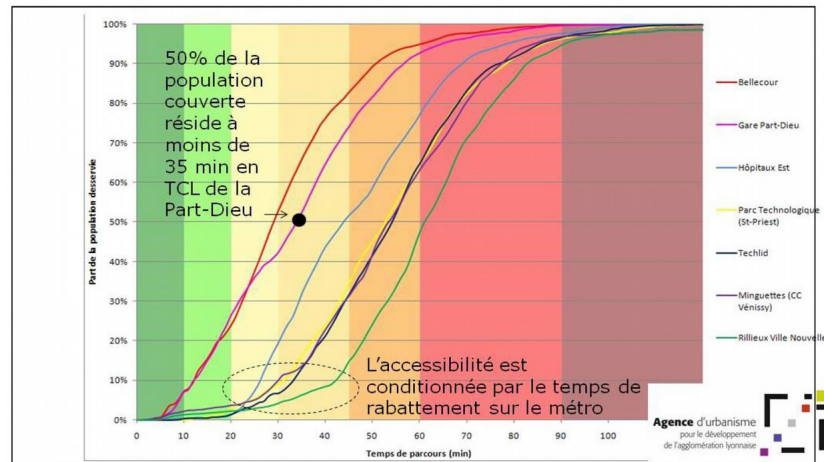


Illustration 4 : Courbe cumulée de la population accessible en fonction du temps de parcours depuis différents quartiers de Lyon, source [60]

Au lieu de fixer des paliers sur la valeur du coût maximum X , on peut en fixer sur le volume d'opportunités atteintes pour une valeur de X donnée, et cartographier par exemple les zones accédant en moins de 60 minutes à 1, 2 ou 3 et plus établissements d'enseignement supérieur, comme présenté dans l'illustration 5.

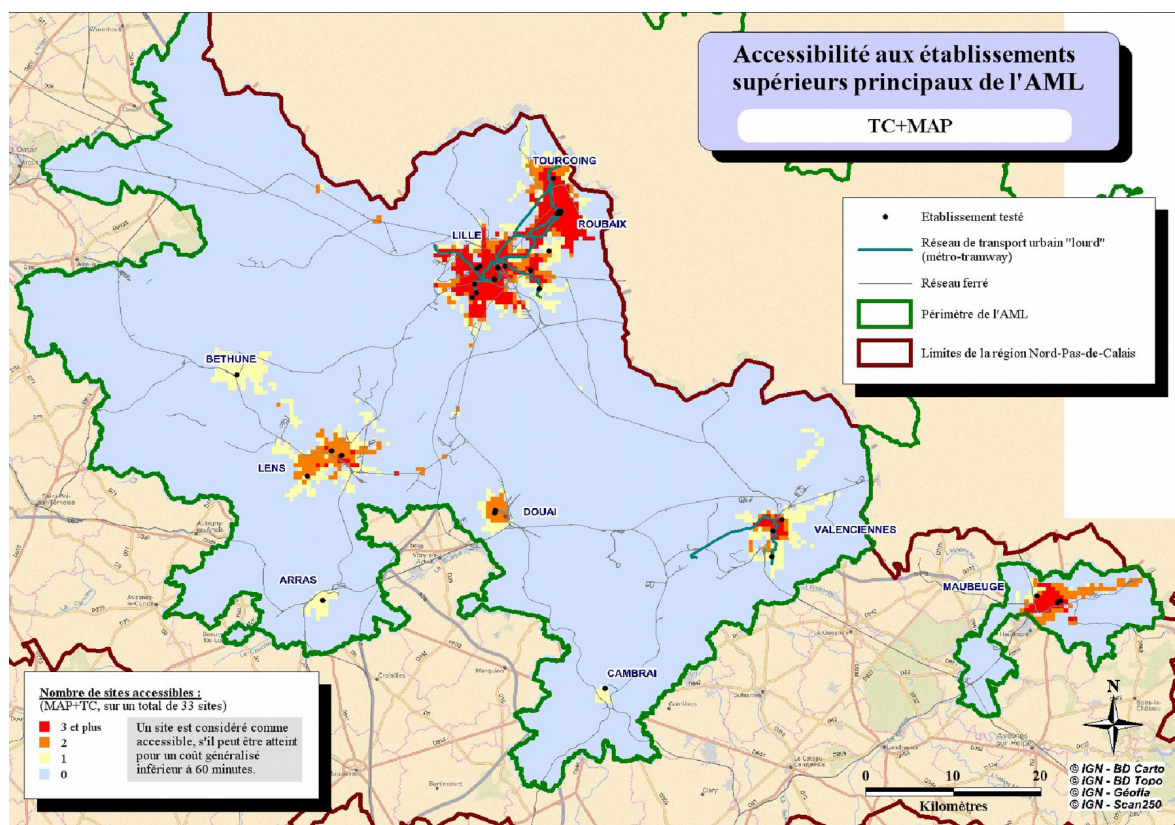


Illustration 5 : Accessibilité aux établissements d'enseignement supérieur en moins de 60 minutes de temps généralisé en transports collectifs + marche à pied dans l'aire métropolitaine lilloise, mesurée par dénombrement des établissements source [60]

Cet indicateur dérivé de l'isochrone est parfois appelé « isochrone multipolaire ». En effet, il permet de représenter au moyen d'un indicateur synthétique l'accessibilité depuis une zone donnée à plusieurs opportunités comparables en termes d'utilité pour l'utilisateur et résulte du croisement de plusieurs courbes isochrones : par exemple, une zone qui accède à deux établissements d'enseignement supérieur est une zone située à l'intérieur de deux courbes isochrones ayant chacune pour origine un des deux établissements.

Enfin, on peut envisager de croiser l'information sur les opportunités à la destination et la demande à l'origine. Après avoir calculé l'indicateur « isochrone » pour l'ensemble du territoire, il faut alors comparer ses valeurs au niveau de demande de la zone à laquelle ils sont associés, par exemple la population dans le cas de l'accès aux établissements d'enseignement supérieur, comme présenté dans l'illustration 6. Ceci permet d'identifier l'adéquation (ou la non-adéquation) entre l'accessibilité aux équipements depuis la zone i (l'offre) et la population de cette zone i (la demande potentielle).

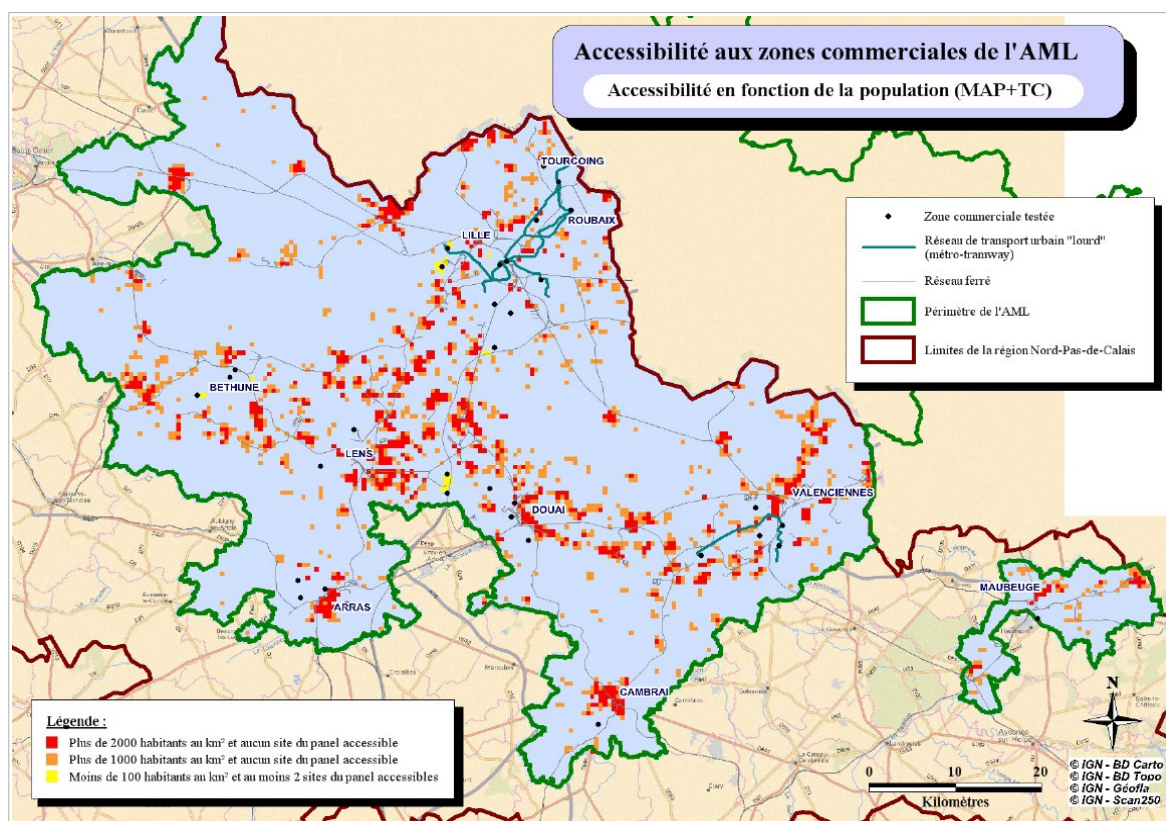


Illustration 6 : Identification des zones pour lesquelles il existe un décalage entre densité de population (demande d'opportunités) et zones commerciales accessibles en moins de 60 minutes de temps généralisé (offre d'opportunités) : ces zones présentent soit une forte densité de population et pas d'opportunité accessible, soit une faible densité de population et plusieurs opportunités accessibles, source [61]

Par rapport à la courbe isochrone, l'indicateur qui en dérive permet de représenter de manière synthétique un niveau d'accessibilité du territoire à un certain type d'opportunités et agrège donc l'information qui serait contenue sur de multiples cartes si on utilisait une courbe isochrone (une carte par opportunité).

Utilisation dans les études

Les mesures d'accessibilité par des courbes isochrones sont très largement utilisées dans la planification (mesure des évolutions d'accessibilité à un type d'opportunité depuis un territoire en situation actuelle et en prospective), en tant qu'outil d'aide à la décision (maximisation des satisfactions par gains d'accessibilité) et dans les études à vocation géographique. Ce succès s'explique par une relative facilité de mise en œuvre (les isochrones nécessitent peu de données en comparaison avec les autres indicateurs), d'interprétation et de communication.

L'isochrone est utilisée la plupart du temps dans un contexte monomodal (par ex. études A1 et A2), et éventuellement pour comparer l'accessibilité procurée par différents modes (par ex. étude A3). On ne rencontre que très exceptionnellement des analyses intermodales (voir illustration 2). Les coûts considérés sont le plus souvent des temps de parcours ou des temps généralisés (voir définitions au paragraphe 2.2.2) et n'intègrent donc pas d'autres dimensions du coût généralisé.

La courbe isochrone n'est pas systématiquement croisée avec des données d'opportunités. Certaines études considèrent uniquement le réseau de transport (par ex. étude A1), valorisant les gains de coût du déplacement et se rapportant donc majoritairement à la composante *transport* de l'accessibilité (voir illustration 3). On compare alors des variations de surface de territoire accessible entre deux scénarios. D'autres études considèrent, conjointement au réseau de transport, la population ou les opportunités accessibles à l'intérieur de l'isochrone (par ex. études A2 et B), valorisant la densité des opportunités ou de la population (voire des deux) conjointement aux gains sur le coût du déplacement et considérant ainsi l'accessibilité de façon plus complète (voir illustrations 5 et 6).

La plupart des isochrones sont basées sur le temps de parcours. Le niveau de service des transports collectifs n'est donc pas complètement pris en compte, notamment parce que le nombre de correspondances nécessaires et les fréquences de passage ne sont pas utilisées pour construire l'indicateur. Pour traiter ce manque, on recense dans l'étude O une approche multicritère de l'accessibilité : les critères *temps*, *intensité* (fréquences) et *pénibilité* (nombre de correspondances) sont évalués et des classes de qualité de desserte sont définies, l'appartenance des zones à ces différentes classes étant ensuite cartographiée. Cette approche permet d'expliquer les différences d'accessibilité des zones de façon distincte par les trois indicateurs, contrairement à un temps généralisé qui agrégerait ces trois dimensions dans un seul indicateur.

Avantages et inconvénients

Malgré sa facilité de mise en œuvre et de communication, l'isochrone ne rend compte qu'imparfaitement des interactions entre les composantes *transport* et *spatiale* de l'accessibilité et ne permet pas une prise en compte poussée de la composante *individuelle* :

- elle ne rend pas compte de la différence d'effort à fournir pour atteindre l'une ou l'autre des opportunités situées à l'intérieur de la courbe isochrone ;
- la dimension temporelle de l'accessibilité ne peut être que partiellement prise en compte (au travers des variations de qualité du système de transport) ;
- les caractéristiques individuelles ne peuvent être prises en compte qu'au travers des propriétés, pour une population donnée (par ex. personnes à mobilité réduite), d'accessibilité physique et de coûts de déplacement des réseaux de transport ;
- les valeurs bornant les différentes zones isochrones sont évaluées de façon arbitraire et, en modifiant ce paramètre, on peut influencer considérablement l'allure de la carte d'accessibilité. « Cela crée une différence artificielle entre les occasions situées à 399 mètres (valorisées) et celles situées à 401 mètres (qui n'ont aucune valeur) » [15].

Modèle gravitaire simple

Présentation de l'indicateur

C'est au milieu du xx^e siècle que le concept de potentialité est introduit pour décrire l'accessibilité reposant sur un modèle gravitaire – au sens physique du terme [2]. Cet indicateur représente le volume potentiel d'opportunités qu'on peut atteindre au sein d'un territoire, pondéré par une fonction de résistance (ou d'impédance) au déplacement entre une zone d'origine i et une zone de destination j . Cette fonction de résistance traduit l'effort que doit fournir l'individu en se déplaçant pour atteindre une opportunité dont il a besoin. [2] présente cet indicateur comme étant le produit d'une fonction d'attraction (les opportunités des zones de destinations) et d'une fonction de résistance (le coût du déplacement).

L'indicateur gravitaire est défini par la formule suivante : $A_i = \sum_j O_j \times F(d_{ij})$

avec :

- i , la zone d'origine ;
- j , la zone de destination du déplacement réalisé par un individu depuis la zone i ;
- A_i , l'accessibilité des individus localisés dans la zone i ;
- O_j , le volume d'opportunités de la zone j ;
- d_{ij} , le coût (distance, temps ou coût généralisé) de déplacement entre la zone i et la zone j .

Encadré 2 : Présentation de l'indicateur gravitaire

$F(d_{ij})$ correspond à la fonction de résistance (ou d'impédance) liée au système de transport, qui est définie de façon générique par la formule $F(d_{ij}) = e^{-\alpha d_{ij}} \times d_{ij}^{-\beta}$.

Le plus souvent, les indicateurs sont construits en supposant $\beta=0$ et la fonction de résistance est alors de forme Logit $F(d_{ij}) = e^{-\alpha d_{ij}}$, comme le proposait Hansen [2].

Le modèle de Kirchhoff suppose quant à lui $\alpha=0$, et donc $F(d_{ij}) = d_{ij}^{-\beta}$ [19].

Notons qu'on peut aussi trouver une variante dans la formulation de la fonction de résistance appelée fonction Box-Cox : $F(d_{ij}) = e^{-\alpha d_{ij}^{\lambda} - 1/\lambda}$ [19].

Pour éclairer l'interprétation des indicateurs gravitaires par rapport aux isochrones, supposons que, depuis un lieu de résidence donné, on évalue l'accessibilité aux commerces alimentaires. Nous obtenons une isochrone à 30 minutes au sein de laquelle nous comptons l'ensemble des commerces alimentaires présents, peu importe qu'ils soient proches ou lointains du lieu de résidence, pourvu que le temps d'accès soit inférieur à 30 minutes. Il est évidemment plus aisé d'atteindre le commerce alimentaire localisé à 5 minutes du lieu de résidence que celui localisé à 25 minutes, qui sont tous deux dans l'isochrone à 30 minutes. Dans l'indicateur gravitaire, le temps de parcours pour atteindre l'activité, moindre dans le premier que dans le second cas, se traduit en résistance, cette résistance augmentant de plus en plus vite lorsque le temps de parcours augmente (fonction non linéaire du temps de parcours), à l'instar du modèle gravitaire en physique.

Ainsi, deux zones ayant accès chacune à exactement 10 commerces en moins de 30 minutes – et donc une même accessibilité au sens de l'isochrone – n'auront pas la même accessibilité gravitaire aux commerces alimentaires.

Encadré 3 : Différence entre indicateur gravitaire et isochrone au travers d'un exemple

On remarque que, lorsque la fonction de résistance comporte une composante « puissance » ($\beta \neq 0$), le coût intrazonal d_{ii} (coût d'un déplacement ayant son origine et sa destination dans la zone i) ne peut pas être considéré comme nul (pour éviter un terme nul au dénominateur). En revanche, dans une formulation de type Logit ($\beta=0$), un coût intrazonal nul revient à compter pour 1 chaque opportunité de la zone.

[20] a cherché à définir des coûts intrazonaux non nuls, en utilisant la formule suivante :

$$d_{ii} = \sqrt{\frac{\sqrt{S_i}}{\pi}} \sqrt[3]{d_{i1} \cdot d_{i2} \cdot d_{i3}}$$

où S_i est la surface de la zone i ,

d_{ii} est la distance intrazonale de la zone i ,

d_{i1}, d_{i2}, d_{i3} correspondent aux distances entre la zone i et les 3 zones les plus proches.

Utilisation dans les études

Comme l'isochrone, l'indicateur gravitaire est généralement utilisé dans un contexte monomodal, pour étudier des variations d'accessibilité entre différents scénarios d'offre de transport, d'aménagement du territoire ou entre différents modes. Sans unité, l'accessibilité gravitaire peut être formulée en référence au point ayant l'accessibilité la plus élevée (voir illustration 7, référence de valeur 100 à la gare de la Part-Dieu) ou encore donnée en valeur brute (voir illustration 8, donnant, sans transformation, l'indicateur d'accessibilité gravitaire aux actifs de Montréal, utile par exemple à une entreprise qui s'implante pour évaluer son potentiel de salariés accessibles).

Accessibilité aux emplois en transport en commun en heure de pointe sur la zone Lyon - La Verpillière

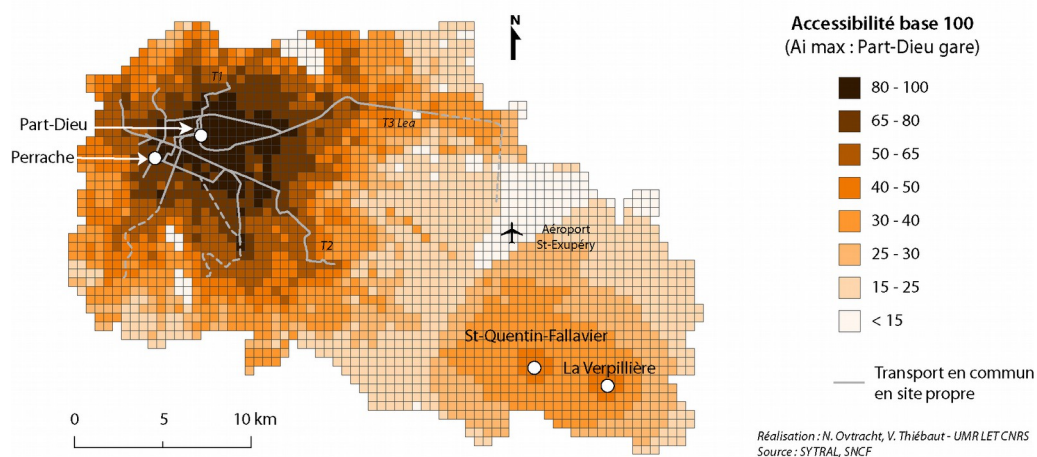


Illustration 7 : Accessibilité gravitaire aux emplois de la zone Lyon - La Verpillière en transport en commun en heure de pointe, source [19]

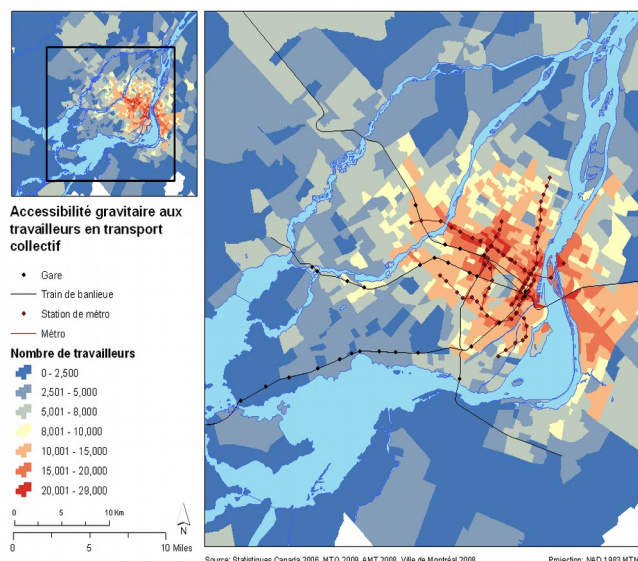


Illustration 8 : Accessibilité gravitaire en transports collectifs aux actifs de Montréal, source [15]

L'indicateur gravitaire reste relativement peu utilisé dans les études (par ex. études A3, K, R), notamment en raison de la difficulté de calibrage du paramètre α qui doit correspondre aux comportements réels des usagers. Il doit être calé à partir de données de mobilité sur la zone d'étude. Plusieurs producteurs d'études interrogés mentionnent leur souhait de disposer de valeurs types calées sur différentes agglomérations pour ce paramètre. En revanche, lorsqu'un modèle de déplacements est disponible, le paramètre α est connu, généralement avec une distinction par motif.

L'indicateur gravitaire est plus complexe à interpréter que l'isochrone puisqu'il donne une valeur d'indice d'interaction potentielle qui n'a pas de sens physique direct : un moyen de rendre plus intelligibles les résultats est d'évaluer des variations d'accessibilité gravitaire par rapport à un scénario, un mode ou un territoire de référence. Ce type d'analyse rend bien compte des forces et des faiblesses comparées d'accessibilité entre deux situations.

Avantages et inconvénients

La mesure de l'accessibilité par le modèle gravitaire simple possède plusieurs avantages :

- elle introduit une résistance au déplacement traduite par une fonction non linéaire de l'éloignement entre le lieu d'origine et le lieu de destination du déplacement, rejoignant ainsi un raisonnement intuitif fait par les individus par rapport à leur environnement géographique ;
- elle tient partiellement compte des interactions entre la composante *transport* et la composante *spatiale* de l'accessibilité (voir illustration 1), notamment en supprimant l'isotropie supposée quant à la localisation des opportunités au regard du lieu d'origine ou de destination considéré.

Malgré cela, les limites et les inconvénients du modèle gravitaire sont nombreux :

- le résultat de la mesure dépend du coût généralisé et du coefficient α choisi, dont le paramétrage est déterminant et délicat à réaliser (voir paragraphe 2.2.3 pour le paramétrage) ;
- si on considère un zonage pour le calcul des temps de déplacements (cas des modèles de déplacements), les opportunités à l'intérieur d'une même zone sont pondérées de la même manière : il faudra donc veiller à utiliser un zonage fin pour limiter ce biais ;
- la fonction de résistance ne prend pas en compte les variations de l'effort ressenti pouvant provenir d'autres facteurs que celui lié à l'éloignement des opportunités et qui relèvent davantage de la composante sociale de l'accessibilité – ici tous les individus sont considérés comme identiques, avec les mêmes fonctionnements, états et capacités [21] ;
- l'indicateur ne prend pas en compte l'effet des concurrences entre les opportunités et entre les besoins des individus au sein du territoire (relation interne à la composante *spatiale* dans l'illustration 1) : au sens économique du terme, cela reviendrait à introduire la relation entre l'offre et la demande d'opportunités. Si le territoire n'est pas à l'équilibre, des effets de concurrence peuvent apparaître, soit sur l'offre (forte par rapport à la demande), soit sur la demande (forte par rapport à l'offre). De même, l'indicateur ne permet pas de prendre en compte le caractère non isotrope de la demande ainsi que les contraintes de capacité de l'offre [22].

Modèle gravitaire avec effets de concurrence liés à la demande

Présentation de l'indicateur

Une des manières d'introduire les effets de concurrence liés à la demande dans les modèles gravitaires est proposée par [22]. Elle se traduit par la prise en compte du potentiel de la demande (c'est-à-dire le nombre d'individus cherchant un type donné d'opportunité) dans le calcul de l'indicateur d'accessibilité gravitaire.

L'offre d'opportunités dans la zone j est divisée par la demande pour cette même zone j . Cela revient donc, selon la formule suivante, à évaluer, pour chacune des zones j , le ratio des offres par rapport à la demande en exerçant une concurrence entre les présents et les entrants sur un territoire qui veulent un type donné d'opportunité.

Cet indicateur se définit donc ainsi :

$$A_i = \sum_j \frac{O_j}{D_j} \times F(d_{ij})$$

avec

- i la zone d'origine ;
- j la zone de destination du déplacement réalisé par un individu depuis la zone i ;
- A_i l'accessibilité des individus localisés dans la zone i ;
- O_j le volume d'opportunités de la zone j ;
- D_j le volume de la demande (individus voulant bénéficier des opportunités de la zone j) ;
- $F(d_{ij})$ la fonction de résistance (ou d'impédance) du système de transport ;
- d_{ij} le temps (distance ou coût) de déplacement entre la zone i et la zone j .

Encadré 4 : Construction de l'indicateur gravitaire avec effets de concurrence liés à la demande

L'illustration 9 montre un exemple à Montréal [15] d'accessibilité à l'emploi en voiture particulière en tenant compte de la concurrence entre les travailleurs. La formule précédente s'interprète ainsi :

O_j désigne le nombre d'emplois en zone j et D_j le nombre d'actifs en concurrence pour cette zone j ,

avec $D_j = \sum_i actifs_i \times F(d_{ij})$

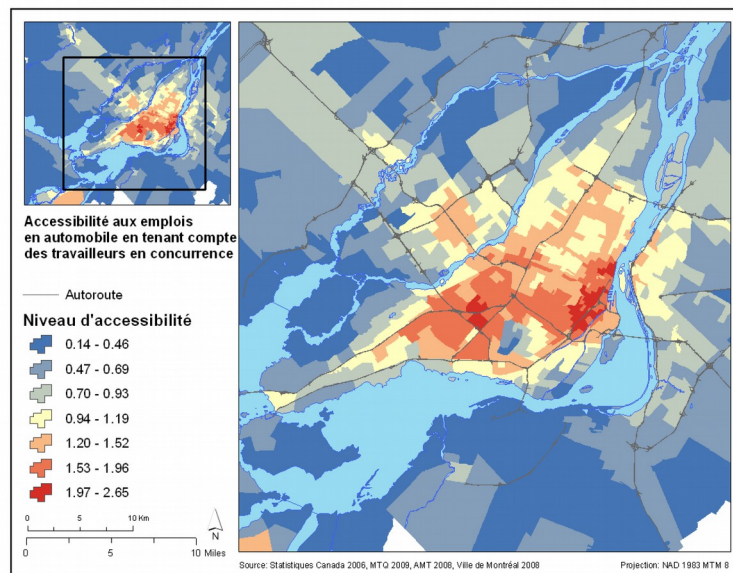


Illustration 9 : Accessibilité concurrentielle (sur la demande) en automobile aux emplois, source [15]

Avantages et inconvénients

Cet indicateur est adapté pour mettre en évidence des potentialités d'accès à l'emploi, car il tient compte de la concurrence entre les demandeurs. Il n'est que partiellement satisfaisant puisqu'il ne prend pas en compte la concurrence liée à l'offre entre l'ensemble des zones de destination. Enfin, l'inconvénient majeur est la difficulté de son interprétation, et donc de sa valorisation et de sa communication.

Modèle gravitaire avec effets de concurrence liés à la demande et à l'offre : facteurs d'équilibre inverses

Présentation de l'indicateur

En plus des effets de concurrence liés à la demande, cet indicateur cherche à prendre en compte la concurrence liée à l'offre. Il vise donc à être plus proche des flux de déplacements observés. Les flux T_{ij} issus d'une zone i ne dépendront pas uniquement du potentiel d'opportunités accessibles en j , mais ils se répartiront parmi l'ensemble des zones j en concurrence. Pour prendre en compte cette double approche des concurrences liées à l'offre et à la demande, et donc pleinement tenir compte des interactions spatiales entre les individus et les opportunités, [23] et [24] formalisent un modèle d'interaction, appelé « facteurs d'équilibre inverses » qui se base sur les principes de maximisation entropique sous contraintes. Ainsi, cela revient à évaluer, de manière itérative, un potentiel d'opportunités et un potentiel de demande pour chaque zone en fonction des flux entrant et sortant. La mise en œuvre de cet indicateur est présentée ci-dessous :

Soit $T_{ij} = a_i b_j O_j P_i F(d_{ij})$

avec :

- T_{ij} le nombre de déplacements entre les zones d'origine i et de destination j ;
- a_i et b_j des paramètres représentant selon la configuration urbaine des coefficients de rareté ou d'abondance, appelés « facteurs inverses d'équilibre » [24] ;
- P_i , le nombre d'individus résidant dans la zone d'origine i ;
- O_j , le nombre de ressources (ex. les emplois) dans la zone de destination j ;
- $F(d_{ij})$ la fonction de résistance (ou d'impédance) du système de transport ;
- d_{ij} , le temps (distance ou coût) de déplacement entre la zone i et la zone j .

En tenant compte de la concurrence liée à la demande ($\sum_i T_{ij} = O_j$) et de la concurrence liée à l'offre

($\sum_j T_{ij} = P_i$), on arrive aux deux équations suivantes :

$$\frac{1}{a_i} = \sum_j (b_j O_j F(d_{ij})) \quad \text{et} \quad \frac{1}{b_j} = \sum_i (a_i P_i F(d_{ij}))$$

La valeur a_i permet d'égaliser le nombre de déplacements partant de la zone d'origine i au nombre d'habitants de cette zone i . La valeur b_j permet d'égaliser le nombre de déplacements arrivant dans la zone j au nombre de ressources présentes dans cette zone j .

Les facteurs a_i et b_j sont mutuellement dépendants. Ils doivent donc être estimés de manière itérative, en imposant une valeur initiale à b_j (cette valeur initiale est souvent prise égale à 1).

Le facteur $1/a_i$, représente la concurrence des destinations possibles perçues par un individu localisé en i . Ce facteur conduit à une mesure de l'accessibilité. Si on atteint un état d'équilibre après plusieurs itérations, a_i peut être multiplié par la valeur moyenne (sur l'ensemble des zones) de b_j .

Le volume d'opportunités de l'ensemble d'un territoire qui peut être atteint depuis une zone i , est alors exprimé :

$$A_i \approx \sum_j O_j F(d_{ij}) \approx \frac{1}{a_i \times \bar{b}}$$

où $\bar{b} = \frac{1}{N} \sum_j b_j$ et N est le nombre de zones de l'espace urbain.

*Encadré 5 : Modèle gravitaire – facteurs d'équilibre inverses,
sources : d'après [23], [24], [13], cités dans [25]*

L'illustration 10 montre un exemple à Montréal d'accessibilité équilibrée en voiture particulière à l'emploi. Elle est à comparer à l'illustration 9, prenant partiellement en compte les effets de concurrence.

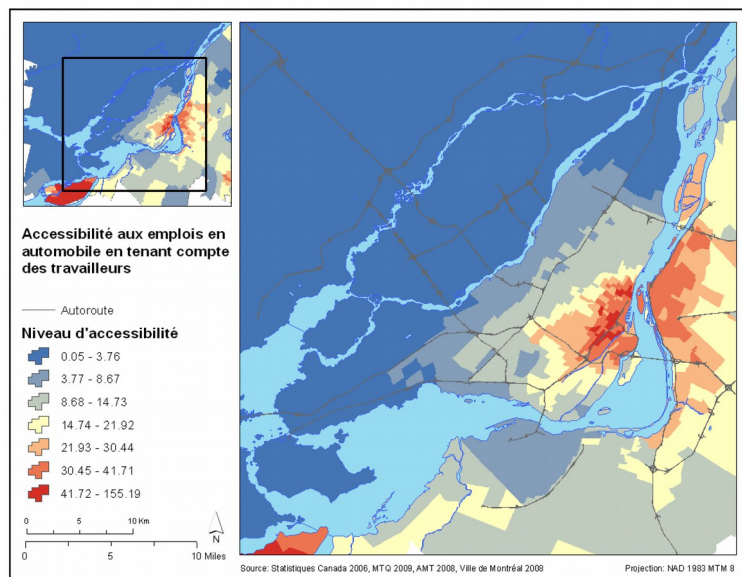


Illustration 10 : Accessibilité équilibrée (concurrence sur l'offre et sur la demande) en automobile à tous les emplois, source [15]

Avantages et inconvénients

Cet indicateur est adapté lorsqu'on souhaite mesurer l'attractivité effective d'un territoire (flux depuis et vers une zone), car il tient compte de la concurrence des offres situées dans les autres zones. Il reste empreint d'un inconvénient majeur lié au fait qu'il résulte d'un processus itératif rendant les résultats difficilement interprétables et communicables et le rendant relativement difficile à calculer. Pour mieux en cerner le fonctionnement et gagner en lisibilité, il convient de séparer les effets de certains facteurs (évolutions de l'usage des sols, des transports...), toutes autres choses étant égales par ailleurs.

Synthèse

Sans chercher à être exhaustifs, nous avons présenté les principaux modèles d'indicateurs d'accessibilité basés sur une composante spatiale fixe. Le tableau ci-dessous propose une synthèse des principaux enseignements à retenir sur ces modèles.

Indicateurs	Prise en compte des composantes				Difficulté d'interprétation, de communication	Limites et inconvénients
	Transport	Spatiale	Sociale	Temporelle		
<i>Isochrone</i>	Oui	Partielle (localisation isotrope des opportunités)	Non	Partielle (prise en compte possible des variations de niveaux de service des transports)	Facile	Pas de prise en compte de la résistance au déplacement Pas de prise en compte des caractéristiques individuelles et temporelles
<i>Modèle gravitaire simple</i>	Oui + résistance au déplacement	Partielle (isotropie partielle)	Non	Partielle (prise en compte possible des variations de niveaux de service des transports)	Assez difficile	Pas de prise en compte des caractéristiques individuelles et temporelles Pas de prise en compte des effets de concurrence
<i>Modèle gravitaire avec effets de concurrence liés à la demande</i>	Oui + résistance au déplacement	Partielle (isotropie partielle + concurrence liée à la demande)	Non	Partielle (prise en compte possible des variations de niveaux de service des transports)	Assez difficile	Pas de prise en compte des caractéristiques individuelles et temporelles Prise en compte partielle des effets de concurrence
<i>Modèle gravitaire avec effets de concurrence liés à la demande et à l'offre</i>	Oui + résistance au déplacement	Oui (effets de concurrence totaux, relation offre-demande)	Non	Partielle (prise en compte possible des variations de niveaux de service des transports)	Difficile	Complexe à calculer Pas de prise en compte des caractéristiques individuelles et temporelles

Tableau 1 : Synthèse des indicateurs basés sur une composante spatiale fixe

2.1.2 Mesures basées sur l'individu et le temps

Comme nous l'avons vu, les indicateurs basés sur une composante spatiale fixe prennent peu en compte la composante *individuelle* et la composante *temporelle*. De plus, ils ne mesurent l'accessibilité que depuis un lieu de résidence ou un lieu d'opportunité, sans tenir compte des boucles de déplacements (et des modes de transports associés) des individus tout le long de la journée. Ces indicateurs sont dits « statiques » [26]. Ils ont tendance à surévaluer l'accessibilité réellement perçue par les individus au sein des territoires, en tenant compte de localisations inaccessibles (contraintes spatio-temporelles) et en ne considérant pas la disponibilité horaire des opportunités.

Prismes spatio-temporels

Présentation des indicateurs

Les mesures d'accessibilité qui tentent de dépasser le cadre « fixe » des précédents indicateurs s'appuient sur la géographie des espaces-temps ([27], [28]). Elles sont développées seulement à la fin des années 1990 ([29], [30]). Ces modèles d'accessibilité intègrent les deux composantes *individuelle* et *temporelle* et prennent en compte les dynamiques de la composante *spatiale* et les interactions des opportunités avec les individus dans le temps. Elles se placent du point de vue de l'individu, à l'inverse des approches précédentes qui moyennent des contraintes et des comportements dans la construction des indicateurs.

Ces modèles dits « prismes spatio-temporels » évaluent alors les volumes d'opportunités potentiellement accessibles par les individus depuis un lieu donné du territoire, à un moment donné de la journée. Ces mesures d'accessibilité varient donc sous les contraintes temporelles et spatiales des individus et des opportunités au cours de la journée. Sur la base de cette observation, [27] et [28] définissent l'*espace-temps contraint* et les *prismes spatio-temporels*.

L'« espace-temps contraint » fait référence à l'effet réducteur des caractéristiques temporelles et spatiales dans le choix des activités. Trois types de contraintes peuvent intervenir :

- les contraintes de « capacité individuelle » [21] : ensemble des états contraignant une personne pour atteindre une destination (par exemple, ne pas avoir de moyens de transports disponibles) ;
- les contraintes de « couplage » : durée pendant laquelle un individu n'est pas disponible du fait de sa présence en un lieu donné pour bénéficier d'une opportunité dont il a besoin (par exemple, la durée de travail) ;
- les contraintes d'« autorité » : règlements et restrictions s'appliquant aux opportunités (horaires d'ouverture et de fermeture, disponibilité en soirée, week-end, saisonnière...).

Un « prisme spatio-temporel » identifie alors l'espace des activités réalisables en respectant un ensemble de contraintes d'espace-temps (contraintes liées aux programmes d'activités, à la localisation géographique des individus et des opportunités ainsi qu'aux limites mises sur les déplacements dans l'espace et dans le temps). Le prisme spatio-temporel représente l'espace potentiel d'action – le volume d'opportunités maximal potentiellement accessible sous les contraintes précédentes ([31], [29]) – d'une personne localisée à un endroit donné, à un moment de la journée, en fonction de son programme d'activités.

Utilisation dans les études

L'indicateur de « *contactabilité* », notamment utilisé dans l'étude P, est un cas simplifié de prisme spatio-temporel : il mesure le potentiel de contact en face-à-face sur une journée entre personnes situées dans des villes différentes. Ici les opportunités sont donc des villes ou des agglomérations urbaines. La *contactabilité* mesure pour une ville donnée les capacités du réseau de transport à permettre une liaison aller-retour avec d'autres villes en une journée tout en autorisant un temps suffisant pour accomplir des activités professionnelles sur place. Aucun autre exemple d'indicateur spatio-temporel n'a été recensé dans les études examinées.

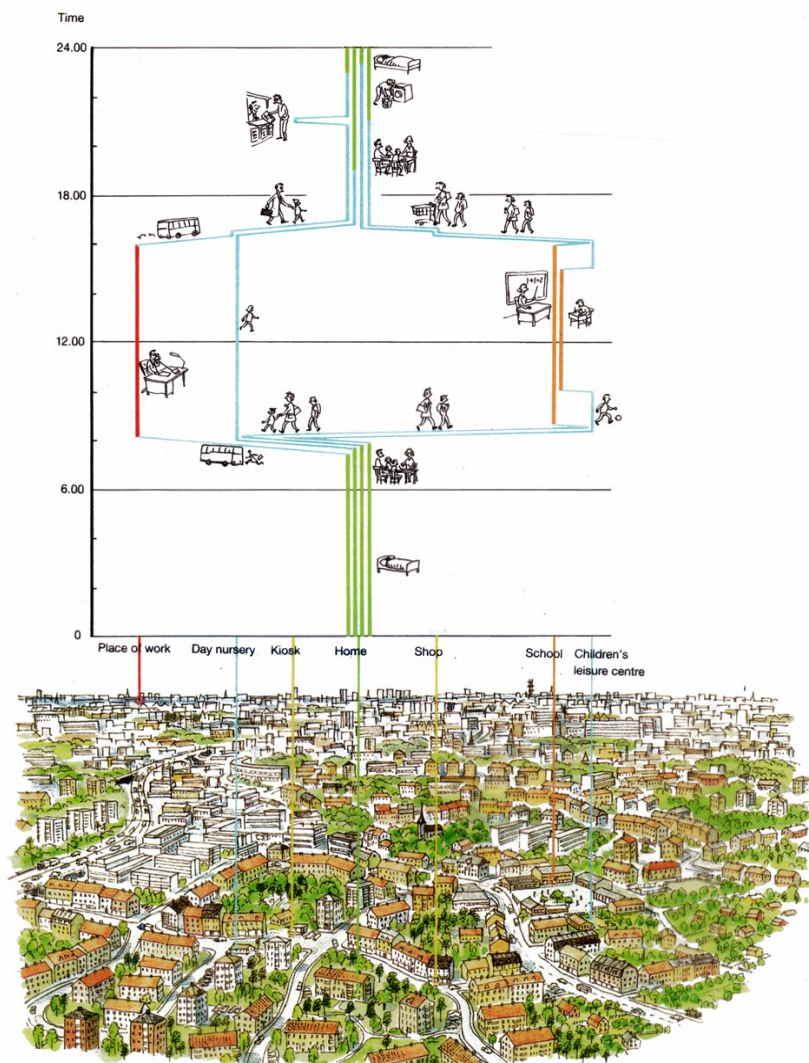


Illustration 11 : Exemple de programmes d'activité, source [69]

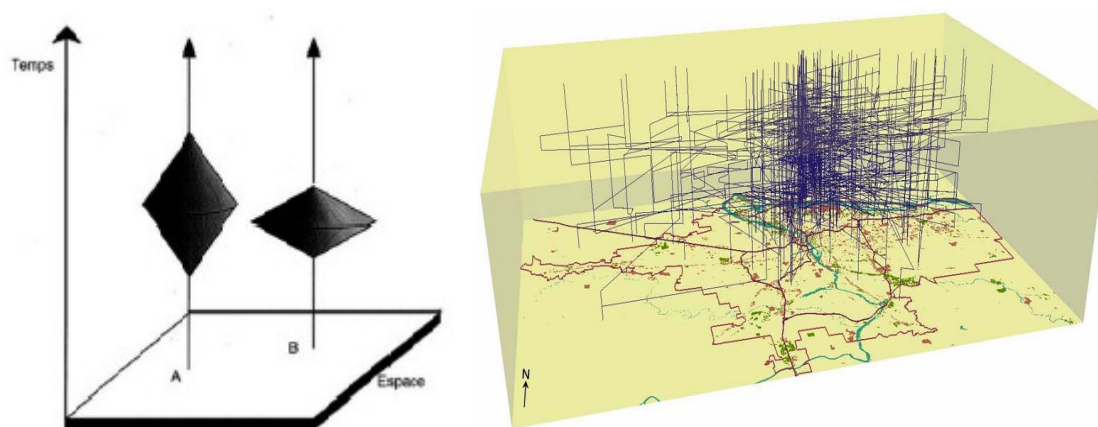


Illustration 12 : Exemples de représentations des prismes spatio-temporels sur une journée ; à gauche représentation idéalisée, source [70], à droite représentation réaliste, source [69]

Avantages et inconvénients

Ce type d'indicateurs a l'avantage de prendre en compte les composantes *temporelle* et *spatiale* (opportunités et individus sous contraintes) de l'accessibilité – il est donc adapté pour évaluer les niveaux d'accessibilité géographique – tout en intégrant aussi la composante *individuelle*. Autrement dit, l'avantage de ce modèle est de considérer l'accessibilité dans l'ensemble de ses composantes et dans leurs interactions.

Les limites de ce type d'indicateurs sont les suivantes :

- les mesures d'accessibilité liées à l'espace-temps sont difficiles à mettre en œuvre et à rendre opérationnelles, du fait d'un besoin important de données individuelles. Ces données peuvent provenir, entre autres, des enquêtes ménages-déplacements quand celles-ci sont disponibles sur le territoire d'étude. Toutefois, la totalité des données nécessaires peut être difficile à obtenir ;
- les mesures sont, par conséquent, faites sur un nombre restreint d'individus. Elles ne prétendent pas à l'exhaustivité des potentiels d'accessibilité observés ;
- bien que les représentations soient relativement aisées à comprendre et à communiquer, elles deviennent inintelligibles lorsqu'un grand nombre d'estimations est fait (voir illustrations précédentes) ;
- même si cet indicateur de prisme spatio-temporel a pour vocation d'avoir une approche exhaustive de la l'accessibilité, il connaît deux limites dans la composante « spatiale » de l'accessibilité qui sont partiellement levées dans les indicateurs considérant une composante spatiale fixe : la résistance aux déplacements et la concurrence liée à l'offre et la demande d'opportunités ne sont pas apparentes.

Ces indicateurs issus de la géographie des espaces-temps sont relativement récents. Ils sont toujours à l'étape de construction et de perfectionnement théorique et opérationnel.

2.1.3 Mesures basées sur le maximum d'utilité

Les modèles d'accessibilité basés sur la théorie des choix discrets ou le maximum d'utilité apparaissent au milieu des années 1970 et montrent comment un choix est déterminé parmi un ensemble d'alternatives qui satisfont toutes essentiellement les mêmes besoins [18] [32].

Maximum d'utilité

Présentation de l'indicateur

On désigne par « coût généralisé du déplacement » le critère à minimiser dans le problème de calcul d'itinéraire sous-jacent au calcul d'accessibilité. Il peut s'agir :

- d'un simple temps de parcours ;
- d'un temps généralisé, dans lequel certaines composantes du temps de parcours ont été pondérées pour représenter les différences de pénibilité entre les différentes étapes du déplacement ;
- ou encore d'une combinaison d'un temps généralisé et d'autres indicateurs (coût financier, coût environnemental du déplacement...) dans une somme pondérée.

[18] postule que la mesure de la contrainte d'accessibilité individuelle exprimée en termes de coût généralisé ne prend en compte qu'un aspect limité des déplacements : la contrainte qu'ils représentent. Son modèle d'accessibilité consiste donc à évaluer l'utilité retirée par un individu de la possibilité de se rendre dans différents lieux d'un territoire afin d'accéder aux activités dont il a besoin. Pour cela, on suppose que les individus associent une utilité cardinale à chaque choix de destination auquel ils sont confrontés et qu'ils choisissent celle qui leur procure une utilité maximale.

L'indicateur est représenté par la formule suivante dans laquelle U_i représente l'utilité ou la satisfaction de l'individu situé dans la zone i par rapport à l'ensemble des choix possibles des opportunités et est assimilée à une mesure de l'accessibilité :

$$U_i = \lambda \times \log \left(\sum_j O_j F(d_{ij}) \right)$$

où :

- λ est une constante d'ajustement ;
- O_j est le volume d'opportunités de la zone de destination j que désirent atteindre les individus. Cet élément d'« attraction » représente l'objectif pour lequel l'individu souhaite accéder à la zone de destination de son déplacement ;
- $F(d_{ij})$ est une fonction du coût généralisé d_{ij} du déplacement entre les zones i et j , dite « fonction de résistance ». On prendra généralement, comme pour le modèle gravitaire : $F(d_{ij}) = e^{-\alpha d_{ij}}$.

Encadré 6 : Définition de l'indicateur du maximum d'utilité

Partant des propositions faites par Koenig, Ben Akiva et Lerman [32] définissent alors l'accessibilité comme étant l'espérance de l'utilité maximale que retire l'individu de son choix parmi un certain nombre d'alternatives possibles pour atteindre les activités dont il a besoin : $A_i = E \left(\max_j (U_{ij}) \right)$.

Notons que, contrairement aux modèles gravitaires, cet indicateur introduit les préférences individuelles de déplacement (composante *individuelle* de l'accessibilité). Toutefois, il ne diffère pas fondamentalement du modèle gravitaire de Hansen [2]. Neuburger [33] fut le premier à établir la relation entre les indicateurs de Hansen et la mesure pour les individus des avantages retirés des localisations des opportunités offertes par le territoire. Ces mesures basées sur la théorie des choix discrets sont cohérentes avec la théorie économique du choix du consommateur (calcul de surplus), tout comme les mesures gravitaires.

Utilisation dans les études

L'agence d'urbanisme de la région stéphanoise réalise ce type d'analyses en interne et réfléchit à utiliser plus largement un tel indicateur dans de futures études ou publications (voir illustration Erreur : source de la référence non trouvée).

Par ailleurs, certaines études d'évaluation des effets d'un projet de transport sont réalisées selon la méthodologie proposée par Jean Poulit, décrite dans l'annexe 2 de l'instruction cadre de 2004 [34]. Basée sur le principe du maximum d'utilité, cette méthode fait débat en raison de ses simplifications calculatoires et de la transposition des gains d'utilité en gains de PIB pour le territoire¹.

1 - Un groupe de réflexion sur la « valorisation de l'accessibilité aux territoires » a été créé en 2012, à l'initiative de Jean Poulit. Ce groupe a approfondi la notion d'accessibilité aux territoires et sa prise en compte dans l'évaluation des effets d'un projet de transport. Pour proposer des améliorations à la méthode décrite dans [34], ce groupe a ainsi publié un rapport en 2014 : <http://www.cdu.urbanisme.developpement-durable.gouv.fr>

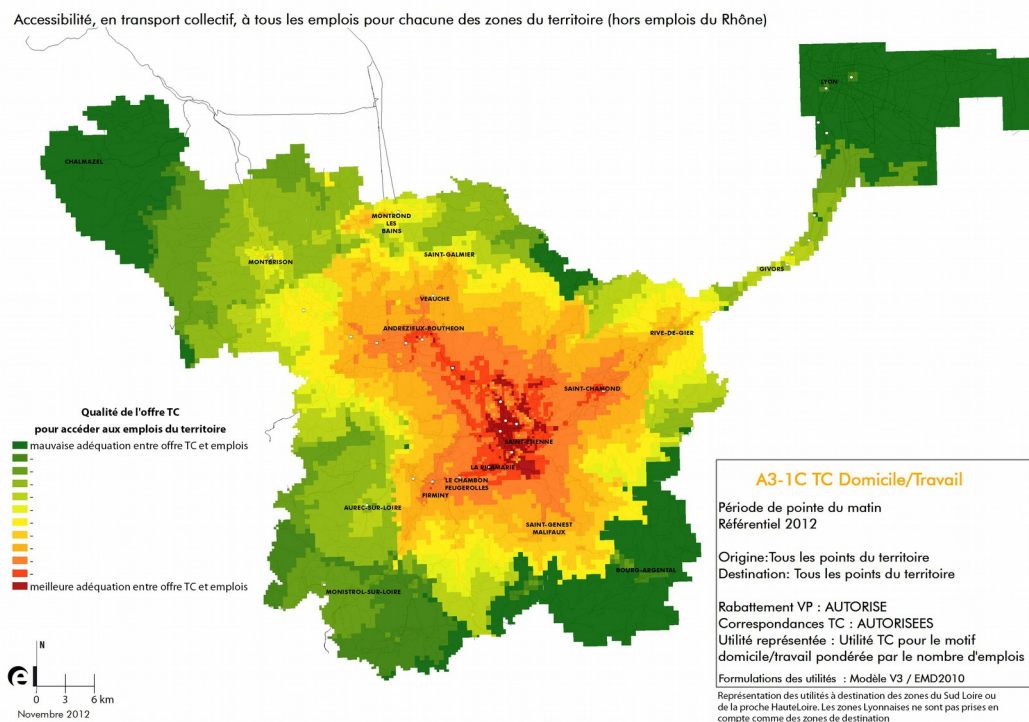


Illustration 13 : Indicateur du maximum d'utilité pour l'accessibilité aux emplois en transports collectifs
Source : « epures », agence d'urbanisme de la région stéphanoise, novembre 2012

Avantages et inconvénients

Cet indicateur basé sur le maximum d'utilité présente l'avantage d'être fondé sur un cadre théorique (celui de la théorie des choix discrets) permettant d'introduire les choix individuels et l'attractivité des destinations, mais également de produire un résultat convertible en valeur monétaire. Il reste toutefois très difficile à interpréter, à vulgariser et à communiquer. Il nécessite des traitements lourds et de nombreuses bases de données pour pouvoir être mesuré. Enfin, les résultats sont difficilement comparables d'un territoire à un autre, l'utilité individuelle étant bien plus subjective que la théorie ne le mentionne.

Enfin, on peut envisager de coupler les programmes d'activités, voire les prismes spatio-temporels avec la maximisation de l'utilité. Le but est alors de prendre en compte l'enchaînement quotidien des déplacements et des activités accomplis par les individus pour choisir les opportunités à utiliser selon le principe de maximisation de l'utilité.

Les indicateurs qui en découlent sont toutefois très complexes à mesurer et très peu utilisés en pratique : données nécessaires nombreuses et difficiles à collecter, résultats complexes à interpréter et à vulgariser. Nous n'avons pas recensé d'études mettant en œuvre ce type d'indicateurs.

Synthèse

Nous avons présenté les principaux indicateurs basés sur la maximisation de l'utilité individuelle. Le tableau ci-dessous dresse une synthèse des principaux enseignements à retenir sur ces modèles.

Indicateurs	Prise en compte des composantes				Difficulté d'interprétation, de communication	Limites et inconvénients
	Transport	Spatiale	Sociale	Temporelle		
Maximum d'utilité	Oui + résistance au déplacement	Oui	Oui (satisfaction individuelle pouvant varier selon les niveaux de vie, etc.)	Non	Difficile	Nombreuses bases de données nécessaires Complexité des calculs
Couplage du maximum d'utilité avec les programmes d'activités ou avec les prismes spatio-temporels				Oui	Très difficile	

Tableau 2 : Synthèse des indicateurs basés sur la maximisation de l'utilité individuelle

2.2 Mise en œuvre des indicateurs

Après avoir présenté les indicateurs d'accessibilité, leurs avantages et leurs inconvénients, il est utile de s'intéresser de plus près à leur mise en œuvre pratique. Dans ce document figurent notamment les valeurs de paramètres des méthodes présentées au paragraphe 2.1 et observées dans les études recensées. Ces informations n'ont pas valeur de recommandations mais permettent aux chargés d'études n'ayant pas de référence disponible sur leur territoire pour un paramètre de connaître l'ordre de grandeur des valeurs communément utilisées. Ce paragraphe ne couvre pas toutefois l'ensemble des paramètres qui peuvent être utilisés dans les méthodes présentées au paragraphe 2.1, certains n'étant pas ou peu utilisés dans les études recensées. Pour la mise en œuvre de ces indicateurs, une prise de connaissance complémentaire de la littérature technique et scientifique est nécessaire.

2.2.1 Représentation des réseaux de transport et des temps de parcours

Les réseaux de transport sont classiquement représentés par des graphes. Ce sont des structures de données composées de nœuds, représentant les points de bifurcation dans le réseau, et d'arcs, représentant les infrastructures ou les services de transport reliant les nœuds. La théorie des graphes propose de nombreuses méthodes pour trouver un chemin optimal dans un graphe, les critères d'optimisation pouvant varier : minimisation du temps de parcours, de la distance, du nombre de correspondances, d'un coût généralisé agrégeant plusieurs indicateurs, etc. La modélisation choisie pour le réseau de transport conditionne les solutions retournées par les algorithmes d'optimisation d'itinéraires.

Représentation du réseau routier

Dans le calcul d'accessibilité, le réseau routier sert de support aux déplacements individuels : marche, vélo ou voiture. Il faut prêter une attention particulière à sa finesse, aux contraintes de circulation qui seront prises en compte ainsi qu'à la manière de définir les temps de déplacement qui sont associés aux arcs routiers.

Finesse du réseau

La finesse (ou *capillarité*) du réseau routier est fonction :

- des données disponibles ;
- de l'utilisation éventuelle d'un modèle de déplacements – dans ce cas, la finesse du réseau est définie en fonction du zonage du modèle, lui-même étant contraint par le découpage géographique des données socio-économiques et des données mesurant la demande de déplacement ;
- de la puissance de calcul des machines et des logiciels utilisés, limitant la taille du graphe manipulable.

[35] a montré, à partir d'un exemple, qu'utiliser un réseau routier d'une très grande finesse par rapport à l'échelle d'analyse des résultats apporte un gain de précision peu significatif dans les mesures d'accessibilité.

La finesse du réseau peut en revanche être cruciale lorsqu'on s'intéresse à l'accessibilité en modes actifs. La non-prise en compte de certaines voiries accessibles uniquement aux modes actifs peut en effet conduire à sous-estimer l'accessibilité de certaines zones (difficulté notamment mentionnée pour l'étude E).

Contraintes de circulation

La prise en compte des contraintes de circulation sur le réseau routier (sens uniques, mouvements tournants interdits) permet d'améliorer le réalisme des itinéraires calculés. Ce besoin est d'autant plus fort que le niveau géographique auquel les résultats sont analysés est fin.

Comme pour la définition du réseau en lui-même, les contraintes de circulation spécifiques aux modes actifs sont généralement mal recensées dans les sources de données disponibles. Leur prise en compte peut pourtant modifier significativement l'accessibilité de certaines zones (prise en compte de contre-sens cyclables par exemple). De plus, nous verrons par la suite que les logiciels ne permettent pas tous une complète prise en compte des mouvements tournants interdits dans le calcul d'itinéraires.

Temps théoriques (en situation fluide) en voiture

Selon les études, les estimations de temps de parcours des véhicules particuliers tiennent compte ou non de la charge des réseaux. Si la congestion n'est pas prise en compte, on peut utiliser une classification des infrastructures basée sur la vitesse moyenne observée à vide ou sur la vitesse limite légale (par ex. études J et S). [35] montre qu'il y a un véritable impact des sources de données routières utilisées sur les estimations de temps de parcours obtenues (comparaison des données Multinet® et BD Carto®), en raison des différences observées dans la géométrie des réseaux et donc dans le calcul des longueurs de tronçons, et dans les classes de vitesses à vide. En plus des classes de voiries, des paramètres comme le fait de se situer en agglomération ou en dehors peuvent être utilisés pour affiner la définition du temps de parcours à vide (par ex. étude U).

Temps en charge en voiture

L'estimation de temps de parcours tenant compte de la charge des réseaux est réalisée au travers d'une des deux approches suivantes :

- l'utilisation de données historiques de suivi GPS fournissant des temps de parcours mesurés à différentes heures de la journée (l'étude M utilise par exemple de telles données pour calculer le temps de parcours des 70 derniers kilomètres de trajet d'accès à une station de sports d'hiver, le temps de parcours du reste du trajet étant estimé sur la base d'un réseau à vide) : une telle approche reste encore peu répandue aujourd'hui, d'une part parce que la couverture du réseau par les données est parfois trop parcellaire et d'autre part parce que le prix d'achat de ces données peut constituer un obstacle, notamment lorsqu'on travaille à l'échelle de larges territoires et de réseaux très maillés ;
- l'utilisation des résultats d'affectation d'un modèle de déplacements (par ex. études B2., I et R) : c'est généralement le besoin d'avoir des temps de parcours en charge qui justifie l'utilisation d'un modèle de déplacement pour le calcul de l'accessibilité (on devra toutefois veiller à ce que le modèle ait bien été calé par rapport à des temps de parcours observés et pas seulement sur des comptages de véhicules).

Temps de recherche de stationnement

Malgré son importance dans le temps de parcours total en voiture, en particulier sur des trajets urbains, le temps de recherche de stationnement n'est quasiment jamais pris en compte. Lorsque le calcul d'accessibilité est réalisé à partir d'un modèle de déplacements, il faut distinguer la prise en compte du stationnement par le modèle (pour la distribution et le choix de mode), de sa prise en compte pour le calcul d'accessibilité en lui-même. Par exemple, dans le modèle détenu par l'agence d'urbanisme de la région grenobloise, un temps de recherche de stationnement est pris en compte pour le choix modal mais pas pour le calcul d'accessibilité. Le modèle détenu par l'agence d'urbanisme de la région stéphanoise prend en compte le temps de recherche de stationnement dans le calcul d'accessibilité au travers d'un coefficient de pénalité distinct selon les trois types de zones : hypercentre, zones bleues et autres.

Certains travaux ont cherché à estimer ce temps à partir de caractéristiques mesurables du quartier [36] : taux d'occupation et type de stationnement (payant, zone bleue...). Toutefois, les données de taux d'occupation par période de la journée manquent généralement sur les zones étudiées, ce qui limite les possibilités d'application d'une telle méthode.

Vitesses en modes actifs

La vitesse de déplacement des modes actifs est généralement définie de façon statique et homogène sur l'ensemble du réseau. Le tableau 3 résume les valeurs classiquement utilisées pour les vitesses de déplacement en modes actifs. Dans les études examinées, nous n'avons pas observé de pénalisation de la vitesse en fonction du dénivelé. Dans certaines zones, le dénivelé pourrait pourtant être à la fois un facteur de pénalisation du temps de parcours et une contrainte sur l'itinéraire interdisant certains tronçons fortement dénivelés, notamment à vélo ou pour des personnes en situation de handicap.

Les vitesses et les durées maximales de déplacement acceptables relevés pour les modes actifs sont résumées dans le tableau 3. On remarque qu'aucune valeur n'a été recensée pour la durée ou la distance maximales de déplacement à vélo et que les temps maximaux de déplacement à pied sont très variables d'une étude à l'autre.

<i>Vitesse moyenne de marche</i>	Pour un déplacement quelconque : 3 à 4,8 km/h. Pour un déplacement bien connu : 5 km/h [8]. Pour un accès à la voiture : 4 km/h [19] (intégrant le temps de prise du véhicule) En situation de difficulté de mobilité : 3 km/h pour les personnes âgées, 3,6 km/h pour les personnes en fauteuil, 2,7 km/h pour les personnes ayant un handicap lié à la marche (étude N).
<i>Vitesse moyenne de déplacement à vélo</i>	Pour un vélo classique : 12 à 16 km/h. Pour un vélo électrique : 20 km/h (étude D).
<i>Temps maximal de marche</i>	<i>Sur l'ensemble du trajet :</i> – 30 min (étude O), – 240 min (dont 20 min pour la marche en début d'itinéraire et 20 min pour la marche en fin d'itinéraire) [37]. <i>En rabattement et diffusion avec les transports collectifs :</i> – 8 min pour un arrêt de bus et 12 min pour une gare ferroviaire [38], – 10 à 15 min en rabattement vers les TC, aucune limite en diffusion (étude F), – 20 min en rabattement vers les TC et 20 min en diffusion (étude B), – 10 min en rabattement vers les TC et 10 min en diffusion (étude H).

Tableau 3 : Paramètres du temps de parcours des modes actifs

Représentation du réseau de transport collectif

Les temps de parcours des transports collectifs peuvent être pris en compte au moyen d'une représentation des services en :

- horaires théoriques (études A1, A2, A3, E, G, I, N, O, P, S, U) ;
- fréquences et temps de parcours moyens par tronçon (études B, D, F, H, K et L).

La prise en compte des horaires théoriques assure une analyse plus fine des potentiels individuels de déplacement. Elle permet d'adopter le point de vue de l'individu et non plus un point de vue moyennant des réalités parfois très diverses. Elle permet notamment d'étudier les performances du réseau en termes de correspondances (par ex. les études O et P). Les indicateurs obtenus pour différentes heures de départ possibles pourront être agrégés a posteriori (voir paragraphe 2.2.4).

En revanche, l'utilisation des fréquences permet une compacité de représentation (moins de données à stocker en mémoire) et une économie de temps de calcul. Elle présente également l'avantage de nécessiter moins de données d'entrée. L'utilisation des fréquences et des temps de parcours moyens permettent également de tenir compte des tranches horaires (les études B et H définissent ces indicateurs pour la période de pointe du matin, pour la période creuse et pour la période pointe du soir).

Le choix d'une solution de représentation (horaires ou fréquences) doit donc être effectué en fonction des objectifs assignés à l'étude, des données disponibles et des capacités informatiques de la machine et du logiciel utilisés.

Lorsque le calcul d'accessibilité est effectué au travers d'un modèle de déplacements, la représentation des temps de parcours est conditionnée par les choix effectués dans le modèle, une représentation en fréquences étant généralement utilisée pour les services très fréquents (urbains) et une représentation par les horaires théoriques étant utilisée pour des modes moins fréquents (interurbains généralement).

Le temps nécessaire pour effectuer une correspondance entre points d'arrêt proches peut être estimé de manière homogène pour tous les couples d'arrêts, ou bien plus finement, à partir d'une mesure sur le terrain (par ex. étude F) ou encore d'une modélisation des cheminements piétons dans les stations (pas d'exemple recensé). Les principaux paramétrages du temps de parcours en transports collectifs sont résumés dans le tableau 4.

<i>Temps d'attente maximum en correspondance</i>	60 min sur l'ensemble du trajet (étude O, [37]) En interurbain (étude M) : 2 heures pour les correspondances avion → bus ou train → bus, 1 heure pour les correspondances bus → bus
<i>Nombre maximum de correspondances</i>	Défini explicitement : 4 correspondances (étude A), 5 correspondances [37] ou limité implicitement par le temps maximal d'attente en correspondance
<i>Méthode d'estimation du temps d'attente</i>	À partir des horaires théoriques (études A1, A2, A3, E, G, I, N, O, P, S et U). ou Demi-temps inter-véhiculaire (études B, D, F, H, K et L). Parfois pénalisé par un temps additionnel destiné à représenter le manque de fiabilité des horaires : 0,75 min pour les trains et 2 min pour les bus dans [38]. Parfois plafonné pour l'attente à l'origine afin de prendre en compte la connaissance des horaires par les usagers (15 min pour les TCU et 20 min pour le ferroviaire dans les études B et H ; 10 min pour les TCU dans l'étude F).
<i>Temps de précaution en correspondance ou à l'origine (temps minimal d'attente pour éviter de rater une correspondance)</i>	1 à 3 min pour les transports collectifs urbains et le train 45 min pour l'avion (étude P) En interurbain (étude M) : 1 heure pour la correspondance avion → bus, 30 minutes pour la correspondance train → bus, 15 min pour la correspondance bus → bus

Tableau 4 : Paramètres du temps de parcours en transports collectifs

Représentation de la multimodalité

Dans ce document, la *multimodalité* désigne l'usage potentiel de plusieurs modes de transport (individuels et collectifs) pour réaliser un déplacement. L'*intermodalité* désigne, quant à elle, l'utilisation effective au cours d'un déplacement de plusieurs modes de transport successifs, avec au moins un mode de transport collectif et un mode individuel (ex : un déplacement VP + TER + marche est considéré comme intermodal).

Les déplacements intermodaux sont globalement marginaux (de l'ordre de 2 % des déplacements). Ils peuvent néanmoins être non négligeables sur certains itinéraires, en particulier ceux partant du périurbain et allant vers les cœurs d'agglomérations, avec des phénomènes de rabattement sur des gares ou stations de transport collectif en limite des zones denses. Ils sont également beaucoup plus importants parmi ceux ayant une extrémité au domicile, pour des raisons de disponibilité du véhicule. De nombreux producteurs d'études mentionnent l'intérêt qu'ils portent à la question de la multimodalité et la difficulté qu'ils ont à la prendre en compte avec les outils et les méthodes actuellement disponibles, notamment pour représenter les rabattements sur les parcs-relais. Nous allons voir comment cette question est traitée dans les études.

Représentations simplifiées de l'accessibilité en transports collectifs

Dans certains cas, le réseau de transport collectif est représenté isolément, sans lien avec un réseau routier (par ex. études A1, A2, F). Le calcul d'itinéraire est alors réalisé d'arrêt à arrêt et des « halos » concentriques sont définis autour des arrêts pour représenter les rabattements sur le réseau qui peuvent être effectués avec différents modes (marche, vélo, voiture...). La portée du rabattement est mesurée à vol d'oiseau à partir d'une vitesse moyenne et par conséquent ne tient pas compte de la forme du réseau viaire et de ses caractéristiques de circulation (voir illustration 14). Les paramètres de rabattement rencontrés sont repris dans le tableau 5.

Un coefficient multiplicateur de la distance à vol d'oiseau est parfois introduit pour représenter la différence avec la distance réellement parcourue. L'étude R utilise par exemple un coefficient de $\sqrt{2}$. [39] utilise un coefficient variable en fonction de l'éloignement entre l'origine et la destination. En effet, plus l'origine et la destination sont éloignées, plus la distance en utilisant le réseau s'approche de la distance à vol d'oiseau, donc plus le coefficient correctif doit être faible. Les valeurs du coefficient varient entre 1,1 et 1,4 et la distance est définie par la formule suivante :

$$D_r = D_{vo} \times (1,1 + 0,3 \times e^{-D_{vo}/20}) \quad \text{pour } D_{vo} \leq 20 \text{ km} \quad \text{et}$$

$$D_r = D_{vo} \times 1,1 \quad \text{pour } D_{vo} > 20 \text{ km} ,$$

où D_r représente la distance estimée, nommée *distance rectilinéaire pondérée* et D_{vo} représente la distance à vol d'oiseau.

Vitesses de rabattement en voiture	10 km/h [35] (en zone urbaine) 5 km/h [40] (en zone urbaine) 50 km/h [19] (vers une gare TER)
Vitesses de rabattement à pied	3,5 km/h [40] 3,5 km/h [19]
Distance maximale de rabattement en voiture	Entre 5 et 20 km selon les quartiers et les études (étude R)
Distance maximale de rabattement à pied	1 km sur une gare TER ou une station de métro (étude R) 500 m sur un arrêt de bus (étude R)

Tableau 5 : Exemples de paramétrages du rabattement sur un réseau à vol d'oiseau

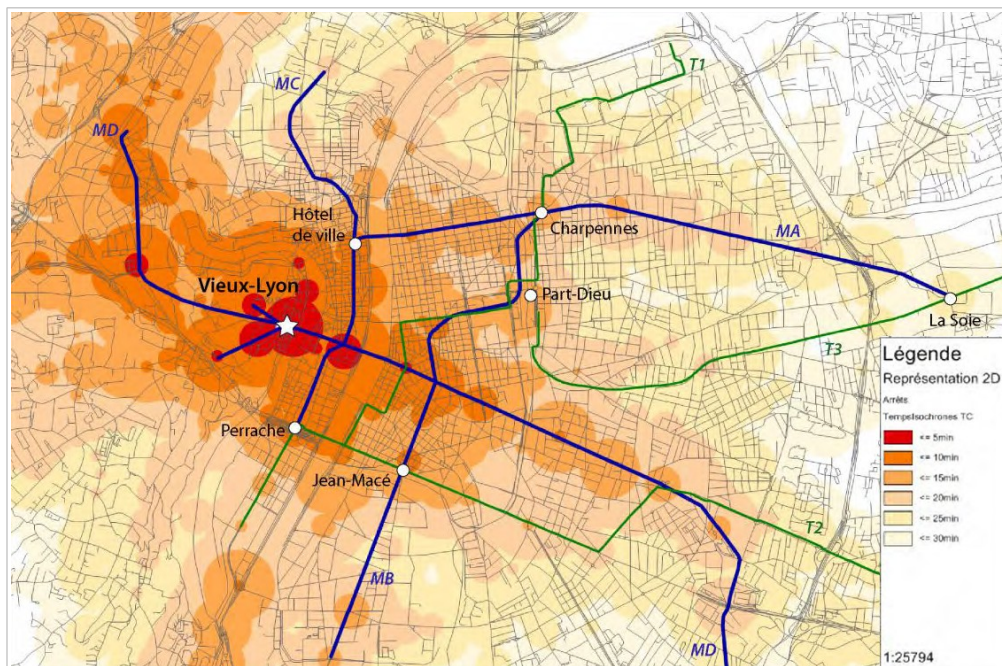


Illustration 14 : Isochrone de 30 minutes en heure de pointe en transports collectifs à partir de l'arrêt Vieux-Lyon, agglomération lyonnaise, source [19]

Parfois, c'est uniquement l'accessibilité aux arrêts de transport collectif qui est mesurée, sans préjuger de la qualité de service offerte par ces arrêts aux usagers. On estime donc quelle partie du territoire a accès à un arrêt de métro, de tramway... en moins de x minutes (par ex. étude T), soit en utilisant des « halos » de rabattement, soit en calculant l'accessibilité à partir du réseau routier à proximité de l'arrêt.

Dans les modèles de déplacements actuels, l'intermodalité n'est, à notre connaissance, pas complètement représentée non plus, ceci pour deux raisons [41] :

- les affectations pour les transports collectifs et pour le mode routier motorisé sont effectuées séparément. Les modes individuels (marche, vélo, voiture) sont généralement représentés par un arc connecteur, qui définit a priori la liaison entre une zone d'origine et un arrêt de transport collectif. Par conséquent, les points d'intermodalité possibles au départ d'une zone d'origine sont présélectionnés (définition d'un bassin d'attraction pour chaque arrêt de transport collectif), de même que les combinaisons modales. On ne tient pas toujours compte du niveau de service sur le trajet de rabattement, lorsque celui-ci est effectué en voiture ;
- les rabattements depuis différents lieux d'une même zone sont représentés à partir du centroïde de zone, donc moyennés en temps de parcours. Les vitesses moyennes de rabattement utilisées classiquement sont listées dans le tableau 5.

Le projet MOSART, calculant des accessibilités à partir d'un modèle de déplacements, a contourné la seconde difficulté en utilisant pour origines et destinations des déplacements des microzones (carroyage de 250 à 500 mètres de côté), subdivisant ainsi les zones utilisées par le modèle de demande de déplacements [19]. Un réseau routier extrêmement fin, tendant vers l'exhaustivité, permet ensuite de calculer avec précision les rabattements depuis les centroïdes des microzones vers le réseau routier principal modélisé et vers les arrêts de

transport collectif. Les indicateurs d'accessibilité peuvent finalement être calculés et cartographiés à l'échelle des microzones, quel que soit le mode de rabattement utilisé (voir illustration 15 avec un rabattement piéton), les temps de parcours pour le mode voiture étant estimés sur la base d'une situation fluide, puisque le réseau routier fin ne fait pas partie du modèle et ne bénéficie donc pas d'une estimation du temps de parcours en charge.

Accessibilité aux emplois en transport en commun en heure de pointe sur la zone Lyon - La Verpillière

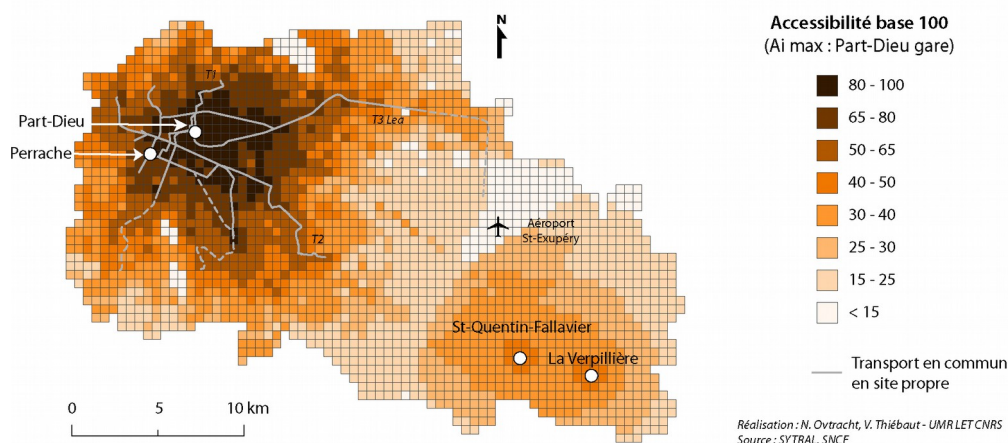


Illustration 15 : Représentation de l'accessibilité à l'échelle des microzones, les temps de parcours étant calculés grâce à un modèle statique de déplacements, source [19]

Représentations « connectées » des réseaux multimodaux

Si on peut légitimement considérer que l'intermodalité a un impact mineur sur un modèle de déplacements d'agglomération, celle-ci peut toutefois représenter un potentiel de déplacement relativement important à une échelle plus vaste (régionale par exemple) ou sur des itinéraires particuliers (déplacements radiaux notamment) et contribuer significativement à l'accessibilité de certains territoires ou équipements. On peut alors calculer des itinéraires multimodaux sur un réseau intégrant les modes individuels et les modes collectifs, et non sur deux réseaux distincts. Cette représentation de la multimodalité laisse le calculateur d'itinéraire déterminer la combinaison modale et les lieux d'intermodalité de l'itinéraire, qui ne sont pas connus a priori, au contraire des précédentes approches.

Dans un réseau multimodal, la finesse de représentation du réseau routier conditionne celle du réseau de transport collectif. En effet, un réseau routier grossier ne permet pas de positionner et de différencier des arrêts de transport collectif proches géographiquement. Si la combinaison modale de l'itinéraire final n'est pas connue a priori, certaines combinaisons modales irréalistes doivent être éliminées (par ex. un trajet ayant pour chaîne modale : voiture particulière + train + voiture particulière). Les logiciels permettant ce type d'optimisation multimodale définissent généralement a priori les combinaisons modales réalistes (par exemple Musliw, qui le définit au travers de la méthode de construction du réseau). La définition de contraintes supplémentaires sur les itinéraires permettra alors d'augmenter le réalisme des solutions générées et de limiter les temps de calcul (lorsque le logiciel le permet) :

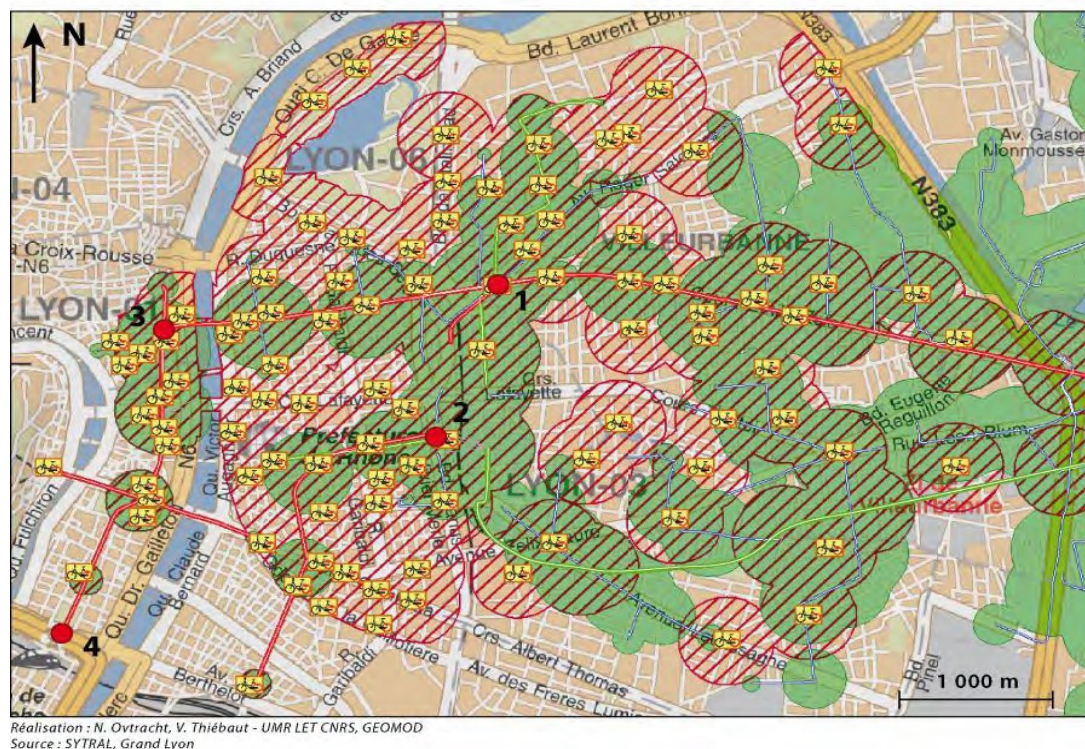
- limiter le temps ou la distance totale parcourue avec un mode ;
- limiter le temps ou la distance parcourue avec un mode pour chaque composante du trajet multimodal (par exemple, pas plus de x minutes de marche à l'origine et à la destination) ;
- limiter le nombre de correspondances...

Prise en compte des modes partagés

Certains travaux prennent également en compte les systèmes de véhicules partagés dans les calculs d'accessibilité (voir illustration 16). Du point de vue de la codification du réseau, cela nécessite simplement de localiser par rapport au réseau routier les points de location et de retour des véhicules, au même titre que les points d'arrêt du réseau de transport collectif. Pour la prise en compte de l'intermodalité, au travers des rabattements piétons sur les stations de location et des rabattements vélo sur les arrêts de transport collectif

(ce second rabattement nécessitant la présence d'une station de location à proximité immédiate de l'arrêt), trois solutions similaires à celles exposées précédemment sont envisageables :

- calcul de l'itinéraire uniquement sur le réseau de transport collectif et représentation simplifiée des rabattements par des halos, la distance étant calculée à vol d'oiseau ;
- intégration du réseau routier pour calculer les rabattements vélos et piétons sur les arrêts de transport collectif, l'itinéraire en transport collectif étant calculé séparément ;
- utilisation d'un réseau multimodal « connecté » entre les transports collectifs et la route, sans connaissance a priori des stations de location initiale et finale utilisées.



Accessibilité :

-  Métro
-  Tramway
-  Bus
-  Vélo V

Principales stations TC atteintes :

- 1 - Charpennes : MA - MB - T1
- 2 - Part-Dieu : MB - T1 - T3 Lea/Lesly
- 3 - Hôtel de Ville : MA - MC
- 4 - Perrache : MA

-  Marche à pied 300 m autour des stations TC atteintes
-  Marche à pied 300 m autour des stations Vélo V

Illustration 16 : Isochrones de 40 minutes en transports collectifs et avec possibilité d'une diffusion en vélo libre-service, depuis la zone industrielle de Meyzieu (agglomération lyonnaise), source [19]

Traitement des chaînes d'activités

La notion d'intermodalité amène à considérer les déplacements non plus isolément mais comme des chaînes, puisque l'utilisation des modes faite sur un déplacement pourra avoir un impact sur les combinaisons modales pour les autres déplacements de la boucle. Ce problème s'illustre particulièrement lorsque l'on utilise des mesures de l'accessibilité des individus au travers de programmes d'activités (voir paragraphe 2.1.2). Toutefois, il est très peu traité dans la littérature et, à notre connaissance, les outils logiciels actuels ne permettent pas sa prise en compte.

2.2.2 Définition du coût généralisé du déplacement

Nous allons voir comment sont définis et paramétrés classiquement le temps généralisé et les autres composantes du coût généralisé (voir définition p. 24).

Construction du temps généralisé

Le temps de parcours correspond au temps réellement expérimenté par l'utilisateur pour se déplacer. Certaines études utilisent directement le temps de parcours pour évaluer l'accessibilité (par ex. les études L et P). D'autres utilisent une notion de *temps généralisé*, au travers de laquelle on cherche à quantifier le ressenti de l'utilisateur vis-à-vis de ce temps, en pondérant certaines dimensions du trajet pour représenter leur pénibilité (par ex. étude E). Des valeurs types pour chacun des paramètres utilisés pour définir le temps généralisé sont présentées dans le tableau 6, sur la base des valeurs observées dans les études et dans la bibliographie. Selon les objectifs de l'étude, on choisira de prendre en compte ou non chacun de ces paramètres.

1. Paramètres du temps généralisé pour les déplacements en transports collectifs	
<i>Pénalité de temps généralisé sur les correspondances (en plus du temps réellement passé à attendre)</i>	Entre 5 et 15 min [42] par correspondance 10 min (étude E) par correspondance
<i>Pondération du temps d'attente dans le temps généralisé</i>	* 1 dans les cas où la pénalité est prise en compte par un temps ajouté ([42] ou étude E) * 1,5 ([12], p. 149) * 2 ([43], cité dans l'étude R) * 5 (étude O)
<i>Pondération du temps passé dans un véhicule en congestion</i>	* 1,5 ([44], page 193) Formule de pondération selon le taux de charge du véhicule en nombre de personnes debout par m ² ([12], p. 154)
<i>Valorisation de l'irrégularité du temps de parcours</i>	Fonction de la probabilité de retard, de l'ampleur du retard et du motif ([12], p. 158-159)
2. Paramètres du temps généralisé pour les déplacements en modes actifs	
<i>Pondération du temps de marche dans le temps généralisé par rapport au temps en transports collectifs</i>	* 1 (étude U) * 1,5 (étude E, étude O), pour tenir compte du contexte peu favorable à la marche en zone périurbaine * 2 ([12], p. 149) en pré et postacheminement ou en correspondance * 2 ([43], utilisé dans l'étude R)
<i>Pondération du temps de déplacement à vélo dans le temps généralisé par rapport au temps en transports collectifs</i>	* 1,5 (étude E), pour tenir compte du contexte peu favorable au vélo en zone périurbaine
3. Paramètres du temps généralisé pour les déplacements en voiture particulière	
<i>Valorisation de l'irrégularité du temps de parcours</i>	Fonction du 90 ^e percentile et de la médiane des temps de parcours observés ([12], p. 157)

Tableau 6 : Paramètres du temps généralisé

Certains producteurs d'études mentionnent leur souhait de prendre en compte la congestion dans les transports collectifs pour calculer l'accessibilité. En effet, la notion de temps d'accès est une mesure insuffisante de l'accessibilité pouvant conduire à encourager l'urbanisation le long de lignes de transport collectif déjà saturées. Le rapport sur l'évaluation socio-économique des investissements publics [12] propose une approche pour prendre en compte l'inconfort lié à la saturation en mesurant le nombre de personnes debout par m² dans les véhicules. Toutefois, cet aspect n'est généralement pas considéré dans les études d'accessibilité, même dans le cas où on dispose d'un modèle de déplacement, du fait du manque de données sources.

Construction du coût généralisé

Les études Q et S minimisent les distances parcourues en plus du temps de parcours. On peut imaginer la minimisation d'autres coûts : le coût financier du déplacement, le volume de CO₂ généré... Ces différents critères d'optimisation peuvent être agrégés, sous la forme d'une somme pondérée, au sein d'une fonction de coût généralisé du déplacement.

Coût financier

Pour la voiture, le coût financier peut être pris en compte au travers d'un coût marginal (uniquement le coût direct du déplacement) ou au travers d'un coût total (incluant les coûts d'amortissement du véhicule, d'assurance et d'entretien). Les valeurs typiques utilisées dans les études examinées sont présentées dans le tableau 7. Le budget de l'automobiliste, publié annuellement par l'Automobile club association [45], donne une estimation de ces coûts.

Pour les transports collectifs urbains, certaines études prennent en compte le coût du ticket à l'unité, d'autres réalisent une moyenne des coûts entre abonnés et non abonnés, en pondérant par le pourcentage d'usagers abonnés sur le réseau lorsque cette donnée est disponible. Pour les usagers non abonnés du train, on utilise généralement un coût kilométrique par tranche de distances auquel s'ajoute une constante, qui permet de reconstituer approximativement le coût réel du billet SNCF pour les trajets en TER et en train Intercités :

$$\text{Prix} = a + b \times d$$

où d est la distance, b est le prix kilométrique, a est une constante, et où a et b ont des valeurs définies par tranche de valeurs de d , les valeurs étant disponibles sur le site de la SNCF pour les trains Intercités [46].

Pour le TGV comme pour l'avion, les variations de prix sont très fortes sur un même train (elles dépendent du moment de réservation du billet et des éventuelles cartes de réductions ou promotions). Par conséquent, une stratégie d'évaluation utilisée par certains producteurs d'études consiste à interroger un calculateur à intervalles réguliers et à agréger les coûts obtenus en un coût moyen.

Une étude commandée par la Fédération nationale des associations d'usagers des transports a permis d'estimer les coûts moyens supportés en 2011 par les usagers pour différents modes de transport [47]. D'après cette étude, le coût moyen kilométrique par voyageur serait :

- sur l'ensemble des réseaux d'Île-de-France de 0,11 € ;
- sur le réseau Transilien de 0,10 € ;
- sur les réseaux urbains de province de 0,13 € ;
- sur les réseaux des conseils généraux de 0,04 € ;
- sur les réseaux TER de 0,08 € ;
- sur le réseau Intercités de 0,09 € ;
- sur le réseau TGV de 0,11 € ;
- sur le réseau aérien traditionnel de 0,15 € ;
- sur les compagnies aériennes à bas coût de 0,06 € ;
- sur l'ensemble des compagnies aériennes de 0,12 € ;
- pour les véhicules particuliers de 0,25 € de coût total dont 0,09 € de coût marginal (supposant un taux d'occupation des véhicules de 1,39) ;
- pour les deux-roues motorisés de 0,33 € de coût total dont 0,08 € de coût marginal (supposant un taux d'occupation de 1,15) ;
- pour les deux-roues non motorisés de 0,15 € de coût total dont 0 € de coût marginal.

Cette étude propose des modulations des coûts moyens selon la distance parcourue ainsi qu'une décomposition des coûts pour certains modes de transport.

	Par véhicule	Par voyageur
<i>Coût total</i>	En 2012 : 0,69 €/km pour une voiture essence, 0,52 €/km pour une voiture diesel [45] (notre calcul indique 0,56 €/km pour un véhicule moyen) En 2011 : 0,34 €/km (recalculé à partir des chiffres figurant dans [47], et notamment le taux d'occupation moyen de 1,39 voyageur par véhicule) En 2008 : 0,49 €/km [19], calcul réalisé à partir du <i>Budget de l'automobiliste</i> 2008 [45] qui estime le coût kilométrique à 0,59 € pour une voiture essence et 0,48 € pour une voiture diesel (notre calcul indique plutôt 0,51 €/km pour un véhicule moyen)	En 2011 : 0,25 €/km [47]
<i>Coût marginal</i>	En 2012 : sur la base de [45], notre calcul indique 0,09 €/km pour un véhicule moyen En 2011 : 0,12 €/km (recalculé à partir des chiffres figurant dans [47], et notamment le taux d'occupation moyen de 1,39 voyageur par véhicule) En 2008 : 0,14 €/km [19], calcul réalisé à partir du <i>Budget de l'automobiliste</i> 2008 [45] (mais avec une définition de coût marginal différente, incluant aussi les frais de réparation. Si on enlève ces frais de réparation, on obtient un coût marginal de 0,08 €/km)	En 2011 : 0,09 €/km [47]

Tableau 7 : Exemples de coûts kilométriques de la voiture particulière

Valeur du temps

Pour agréger le temps généralisé et le coût financier, on introduit la notion de valeur du temps. Celle-ci permet de convertir le temps généralisé en une valeur financière. Des valeurs tutélaires ont été définies au niveau national ([12], p. 146-150). Elles tiennent compte du milieu (urbain ou interurbain) ainsi que :

- du motif de déplacement (professionnel, domicile-travail, autre) en milieu urbain ;
- ou de la distance, du motif et du mode en milieu interurbain.

Elles sont situées dans une fourchette de 6,8 €/heure (valeur pour motif personnel, pour des distances inférieures à 20 km) à 72,9 €/heure (valeur pour motif professionnel, avec le mode aérien).

Toutefois, la valeur du temps utilisée dans les études est généralement unique pour l'ensemble des déplacements du modèle (par ex. étude R où on utilise la valeur du temps correspondant au motif domicile-travail).

Plafonnement des coûts généralisés

Certaines études fixent un plafond de temps de parcours, de temps généralisé ou de coût généralisé au-delà duquel on ne considère plus les opportunités comme accessibles. L'étude E utilise par exemple un plafond de 60 minutes de temps généralisé. Ce plafond est à fixer en fonction des objectifs de l'étude et en particulier des motifs de déplacement visés.

2.2.3 Calibrage de la fonction de résistance

Pour les mesures basées sur le maximum d'utilité ou les différentes mesures gravitaires, une fonction de résistance de forme combinée est souvent utilisée : $F(d_{ij}) = e^{-\alpha d_{ij}} \times d_{ij}^{-\beta}$

Pour la forme Logit (avec $\beta=0$), l'annexe 13 de l'instruction de 2007 pour l'évaluation économique des projets routiers interurbains [48] proposait une valeur de $\alpha=0,0047$ pour le motif travail : la formule proposée fait décroître l'accessibilité en fonction du temps entre une zone i et une zone j avec la fonction $F(T_{ij}) = e^{-0,0047 \times d_{ij}} = e^{-0,47 \times T_{ij}}$, où T_{ij} est exprimé en heures et d_{ij} en kilomètres, en supposant une vitesse moyenne sur le réseau routier national (routes nationales et autoroutes) de 100 km/h [49]. Cette formulation se base sur des travaux d'Orus et Sanchez [50], qui ont calibré cette fonction de résistance exponentielle pour quatre types de motifs différents : travail, services, tourisme, affaires professionnelles. La valeur de 0,0047 reprise dans l'annexe 13 de l'instruction de 2007 correspond en fait au calibrage pour le motif « affaires professionnelles ».

La fonction combinée est souvent utilisée dans les modèles de déplacements, car le calibrage du modèle peut ainsi traduire une résistance à la distance qui pourra prendre une forme variable selon le motif : parfois proche d'une fonction exponentielle et parfois proche d'une fonction puissance. Dans le projet MOSART, ces différentes fonctions ont été calibrées sur l'agglomération de Lyon, pour différents motifs et différentes périodes horaires, et les paramètres retenus sont indiqués dans [19].

On peut trouver dans la littérature de nombreuses valeurs proposées ([30], [51], [52], [53]), mais il faut être attentif au type et à l'unité de la variable d'utilité retenue (temps, distance, coût...).

2.2.4 Agrégation et représentation des indicateurs d'accessibilité

Agrégation temporelle des indicateurs

Les mesures d'accessibilité avec représentation du temps de parcours des transports collectifs par les horaires peuvent être difficiles à analyser en raison de la très forte variabilité des indicateurs d'une minute à l'autre. Ainsi, comme le soulignent Richer et Palmier [16], le calcul d'accessibilité pour un départ à 9 heures peut donner un résultat très différent du même calcul pour un départ à 9 h 04. On peut également envisager des variations sur le temps de parcours routier, même si celles-ci sont a priori moindres.

Pour fournir un indicateur représentant la période dans son ensemble, il faudra donc agréger les données selon une des deux approches ci-dessous :

- agrégation des données de coût par arc, avant le calcul des itinéraires :
 - pour les transports collectifs, les temps de parcours sont alors moyennés sur la période au moyen d'une approche par fréquences et temps de parcours moyen,
 - pour le mode routier, les temps de parcours sont généralement déjà moyennés sur la période, qu'ils soient issus d'un modèle de déplacements ou d'observations GPS ;
- agrégation a posteriori des coûts d'itinéraires obtenus pour différentes heures de départ :
 - on pourra analyser les résultats en termes de moyenne, de médiane, d'écart-type, de valeur minimale ou maximale... (voir étude R qui utilise le temps de parcours minimal sur la période),
 - les analyses pourront porter sur différents éléments de coût séparément (temps de parcours, coût financier, nombre de correspondances, indicateur de pénibilité de l'itinéraire...), sur le coût généralisé dans son ensemble, ou encore sur un indicateur d'accessibilité calculé à partir du coût généralisé.

L'agrégation des indicateurs en sortie du calcul d'itinéraires, plus lourde en charge de calculs, reflète mieux la réalité vécue par l'utilisateur que l'agrégation en fréquences et temps de parcours avant le calcul d'itinéraires. Lorsque les données et les outils logiciels le permettent et que l'on souhaite adopter le point de vue de l'utilisateur,

il est donc préférable de calculer les indicateurs d'accessibilité en transports collectifs à partir des horaires sur l'ensemble d'une période, avec par exemple un départ toutes les minutes, puis d'agréger les résultats obtenus sur l'ensemble de la période.

Cette recommandation fait toutefois l'objet d'une réserve lorsque l'on souhaite projeter les réseaux en situation future. Il est en effet rare que l'on soit en mesure de construire une grille horaire détaillée pour une situation prospective et, pour rester comparable, les résultats en situation actuelle et en situation future doivent utiliser la même stratégie d'agrégation. C'est la raison pour laquelle la plupart des modèles de déplacements agrègent plutôt les temps de parcours a priori et utilisent donc une représentation en fréquence de passage et temps de parcours moyen par arc.

Si l'on utilise une grille horaire détaillée et qu'on réalise l'agrégation a posteriori, on pourra représenter la notion de *dispersion* des temps de parcours des transports collectifs par l'écart-type des temps de parcours sur l'ensemble des heures de départ ou encore par la mesure du nombre de départs possibles dans la période. En effet, si on a de faibles fréquences de desserte, on a aussi une forte variation du temps de parcours selon le moment où on décide de partir (par ex. étude O, critère d'intensité de l'offre).

Agrégation géographique des indicateurs

L'accessibilité est calculée entre nœuds du graphe représentant le réseau de transport et correspond donc au départ à une notion ponctuelle. Pour représenter sur une carte l'accessibilité depuis (ou vers) l'ensemble des points du territoire et depuis (ou vers) une zone représentant une opportunité, il est nécessaire de définir une technique de passage d'une accessibilité ponctuelle à une accessibilité surfacique. Comme pour l'agrégation temporelle, on peut choisir de réaliser l'agrégation avant ou après le calcul d'itinéraires.

Définition a priori du zonage

Pour représenter les indicateurs d'accessibilité sur un zonage prédéfini, on choisit une des approches suivantes :

- utiliser un zonage correspondant à des unités administratives (ex. : communes ou cantons) ou statistiques (par ex. : zones IRIS de l'INSEE) ou à un regroupement de celles-ci, ainsi qu'à des pôles générateurs de déplacements (zones commerciales, hôpitaux...) ;
- définir un pavage de l'espace composé de subdivisions géométriques du territoire, en vue de s'affranchir des limites administratives lors de l'analyse des résultats et de permettre de descendre à un niveau d'analyse plus fin que les découpages statistiques tels que les IRIS [54]. Le pavage pourra être de forme carrée (par ex. études B, E, F, L...) ou encore hexagonale [37].

Il faut avoir conscience que le choix d'un zonage ou d'un autre peut considérablement modifier l'allure de la carte et l'interprétation qu'on pourra faire des résultats. Ce *problème des unités spatiales modifiables*² est bien connu des géographes et a notamment été examiné par ESPON, l'Observatoire européen pour le développement et la cohésion des territoires [55].

Le zonage devra être choisi en recherchant un compromis entre une précision trop grande qui gênerait la lecture de la carte et induirait un volume important de calculs supplémentaires, et une précision trop faible qui moyennait des situations différentes et induirait des erreurs d'interprétation (notamment l'erreur écologique³). La finesse du zonage devra notamment être en lien avec la finesse du réseau routier : il est inutile d'avoir un zonage très fin si on n'a pas de réseau à l'intérieur de certaines zones et, à l'inverse, seul le réseau permettant d'aller d'une zone à une autre est utile. Un réseau trop fin par rapport au zonage introduira donc de la complexité inutile dans la codification du réseau et le calcul d'itinéraires.

On calcule alors pour chaque zone prédéfinie les itinéraires depuis ou vers cette zone en prenant comme extrémité :

- soit un nœud quelconque de la zone, central si possible (par ex. étude O prenant pour origine ou destination le nœud le plus proche du centre géométrique de la zone, pour les pôles d'excellence de Lille Métropole) ;

2 - Problème des unités spatiales modifiables ou *Modifiable Area Unit (MAU) problem* en anglais : (voir http://en.wikipedia.org/wiki/Modifiable_areal_unit_problem)

3 - Erreur écologique : notion utilisée en géographie pour désigner la transposition à une échelle fine de liens de corrélation observés à une échelle plus agrégée.

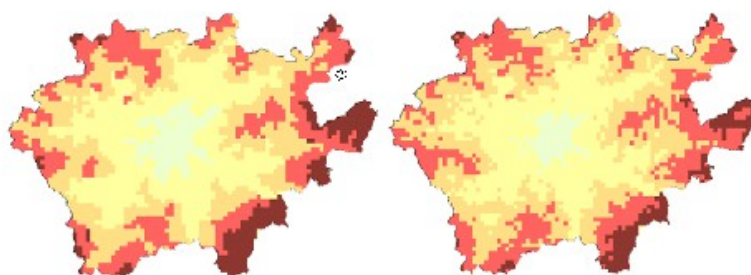
- soit plusieurs nœuds de la zone et en agrégeant a posteriori les résultats (voir [37] dans lequel les itinéraires sont calculés vers ou depuis les différents nœuds d'entrée ou de sortie des zones commerciales et les temps de parcours obtenus pour chaque zone sont moyennés, comme le montre l'illustration 17) ;
- soit en ajoutant un nœud fictif ou *centroïde* (n_0) relié à un ou plusieurs nœuds réels de la zone ($n_1, n_2, n_3, \dots$), le nœud fictif n_0 servant alors d'origine ou de destination depuis ou vers cette zone : il sera alors possible d'affecter un coût à chaque arc ((n_0, n_1) , (n_0, n_2) , (n_0, n_3) ...) qui correspond à un coût de rabattement, qui peut ensuite être estimé, par exemple en tenant compte de la densité ou de la localisation du bâti.



Illustration 17 : Utilisation des points d'entrée dans des zones commerciales comme origines ou destinations pour le calcul d'accessibilité, source [37]

Avec un coût de rabattement nul, c'est le nœud depuis ou vers lequel le coût est le plus faible qui sera retenu comme nœud extrémité de la zone. Sur un pavage carré dans lequel le centroïde d'une zone est situé en son centre géométrique, Di Salvo montre dans [35] que la prise en compte d'un temps de rabattement vers les nœuds réels du réseau peut légèrement modifier les résultats du calcul d'accessibilité (voir illustration 18). L'introduction du temps de rabattement peut en effet permettre de tenir compte de la densité du réseau routier dans le calcul, puisque naturellement le temps de rabattement moyen sera d'autant plus élevé que le réseau est peu maillé dans la zone.

Exemple calculé sur l'aire urbaine d'Agen



Interpolation au plus proche voisin

En tenant compte de la distance à celui-ci

Illustration 18 : Isochrones sans (à gauche) et avec (à droite) prise en compte des temps de rabattement des centroïdes des carreaux vers le réseau, source [35]

Élaboration d'un zonage optimal a posteriori

On peut également utiliser des techniques plus sophistiquées de « krigeage »⁴ qui permettent de définir un zonage en fonction des valeurs des indicateurs d'accessibilité calculés au niveau des nœuds du graphe : on définit la valeur de l'indicateur à attribuer à un point (x, y) de l'espace à partir de la valeur attribuée aux nœuds proches. L'illustration 19 propose un exemple de zones définies par *krigeage*. Nous verrons que peu de logiciels proposent de véritables procédures de *krigeage* à ce jour.

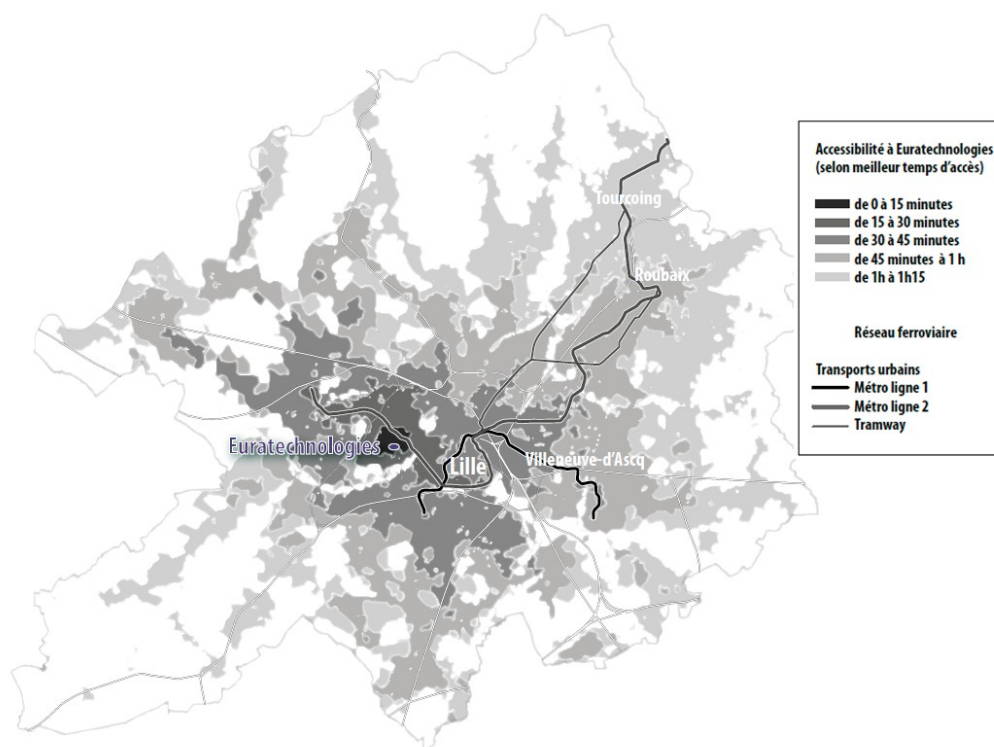


Illustration 19 : Temps d'accès au site d'Euratechnologies dans l'agglomération lilloise en transport collectif, zonage établi par « krigeage », source [16]

4 - <http://fr.wikipedia.org/wiki/Krigeage>

Autres agrégations des indicateurs

Agrégation sur les itinéraires

Dans certaines études (par ex. étude I), on considère que plusieurs itinéraires peuvent relier un nœud A à un nœud B, pouvant correspondre soit aux n chemins de coûts minimaux soit à n optimisations selon des critères différents (variation des valeurs du temps entre les usagers par exemple). On réalise alors une agrégation des indicateurs associés à tous les itinéraires considérés comme admissibles entre ces deux nœuds.

Définition de classes sur les indicateurs

La représentation cartographique nécessite de définir des classes pour les indicateurs d'accessibilité. Les seuils qui délimitent ces classes sont souvent fixés arbitrairement. Leur choix peut pourtant avoir une influence notable sur le rendu final et les possibilités d'interprétation qui en découlent. Utiliser trop de classes peut brouiller le message général de la carte. Au contraire, trop agréger les valeurs de l'indicateur final peut dissimuler des variations entre zones. Là encore, il faut trouver le bon équilibre. Certaines études utilisent des classes de largeurs variables pour mieux représenter les variations de l'indicateur (ex : des surfaces isochrones de 10 minutes pour la plage 0 à 30 minutes puis de 15 minutes pour la plage 30 à 90 minutes). On peut également décomposer les valeurs obtenues en quantiles pour définir des classes, en particulier pour les indicateurs n'ayant pas de signification physique, comme c'est notamment le cas des indicateurs gravitaires.

Agrégation de plusieurs indicateurs

Pour tenir compte de plusieurs indicateurs qualifiant l'accessibilité d'une zone, certaines études définissent des classes d'accessibilité (faible, moyenne ou forte par exemple) sur la base de seuils sur chacun des indicateurs et utilisent éventuellement une méthode de classification (voir étude O et illustration 20).

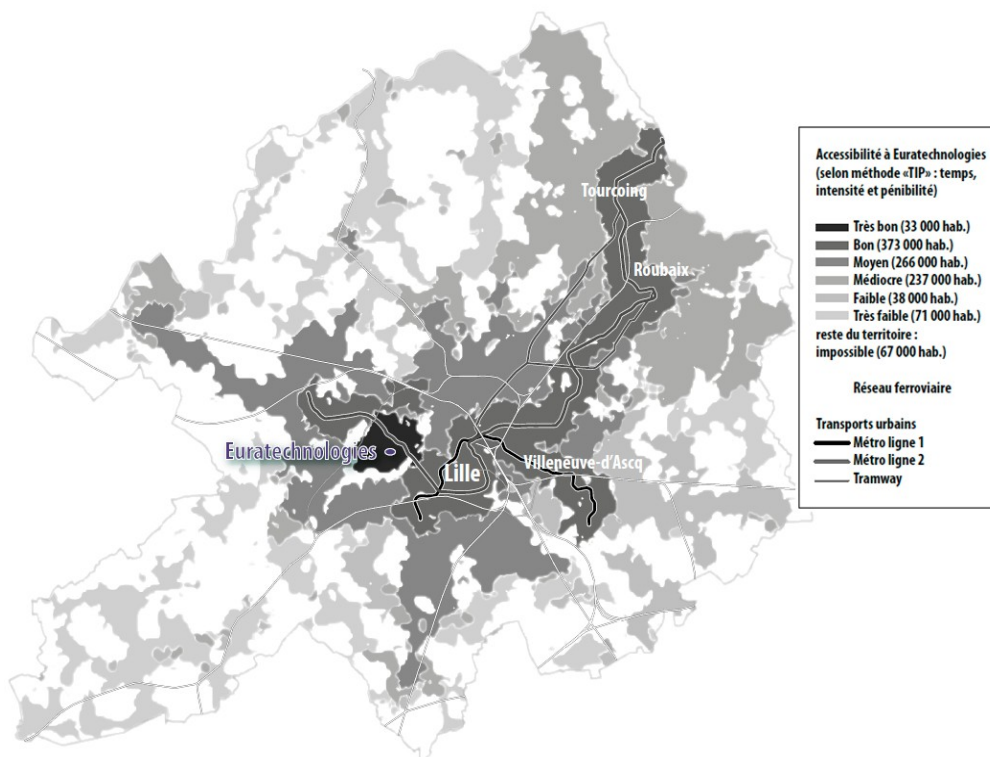


Illustration 20 : Accessibilité au site d'Euratechnologies dans l'agglomération lilloise en transport collectif, agrégée sur les indicateurs « temps », « intensité » et « pénibilité », zonage établi par « krigeage », source [16]

Scénarios

Enfin, les comparaisons de scénarios sont également une forme d'agrégation des indicateurs d'accessibilité par différence. Les éléments variant entre les scénarios pourront concerner le mode choisi, l'heure de départ, l'offre de transport, l'offre d'équipements et d'opportunités sur le territoire (voir illustration 3)...

Représentations graphiques et cartographiques de l'accessibilité

Au-delà des cartes classiques réalisant une analyse thématique sur l'indicateur d'accessibilité que l'on cherche à représenter, des représentations très diverses de l'accessibilité sont imaginables :

- *graphiques en toile d'araignée* (pour chaque lieu dont on évalue l'accessibilité, chaque rayon représente un critère, élément du coût généralisé) ;
- *histogrammes ou courbes cumulées* identifiant le nombre d'opportunités ou la population atteinte en fonction du coût du déplacement (par ex. études A1 et N).
La superposition et les écarts entre ces courbes pour plusieurs scénarios ou différents points de départ ou d'arrivée permettent de comparer les accessibilités. De plus, ces courbes dessinent des points de rupture dans l'évolution de l'accessibilité (par ex. modification de la pente de la courbe quand on entre dans le territoire couvert par le réseau de métro, voir illustration 4) ;
- *anamorphoses* (déformations de cartes pour représenter la plus ou moins grande impédance des déplacements entre zones), *cartographies en relief d'espace-temps*...

Toutefois, les chargés d'études réalisant des calculs d'accessibilité s'accordent à dire que les représentations les plus simples restent dans la majorité des cas les plus efficaces en termes de communication. Les cartes réalisant des analyses thématiques sur l'indicateur d'accessibilité, éventuellement accompagnées de courbes cumulées pour représenter un croisement de l'impédance du déplacement avec une donnée socio-économique sur le territoire, sont les représentations de l'accessibilité les plus largement utilisées.

L'allure des cartes d'accessibilité est étroitement liée aux choix qui ont été faits en matière de modélisation du réseau de transport et de définition du zonage. Par exemple, si on a fait le choix de représenter uniquement le réseau de transports collectifs (sans lien avec le réseau routier) et de calculer les rabattements à vol d'oiseau autour des stations, on aura une représentation cartographique « par halos » (par ex. illustration 14). Si au contraire on a défini un réseau routier, la cartographie de l'accessibilité pourra être ponctuelle (au niveau des nœuds du graphe) ou encore surfacique (voir illustration 18). Si aucun zonage n'a été défini, on pourra mettre en œuvre une procédure de « krigeage » (voir illustrations 19 et 20).

2.3 Synthèse

Ce chapitre 2 a montré que les mesures d'accessibilité sont nombreuses, mais que seules les plus simples sont utilisées dans des contextes d'études opérationnelles : isochrones, indicateurs gravitaires simples et, dans une moindre mesure, indicateurs basés sur le maximum d'utilité. Ces indicateurs peuvent toutefois être déclinés de multiples manières, selon la méthode utilisée pour estimer le temps de parcours, selon les modes et les combinaisons de modes pris en compte, et selon les méthodes d'agrégation temporelle et spatiale des indicateurs. Ces choix conditionnent la représentation graphique et cartographique qui pourra être faite des indicateurs. Ces dimensions méthodologiques de construction des indicateurs peuvent changer l'interprétation qui pourra être faite des résultats, et doivent donc être discutées et validées avec le commanditaire de l'étude.

Le chapitre 3 détaillera les contraintes de mise en œuvre de ces indicateurs – imposées par la disponibilité des données – ainsi que les fonctionnalités et performances des logiciels.