



Traitement des congestions dans les réseaux de transport et dans un environnement dérégulé

Vincent Manzo

► To cite this version:

Vincent Manzo. Traitement des congestions dans les réseaux de transport et dans un environnement dérégulé. Sciences de l'ingénieur [physics]. Institut National Polytechnique de Grenoble - INPG, 2004. Français. NNT: . tel-00408307

HAL Id: tel-00408307

<https://theses.hal.science/tel-00408307>

Submitted on 30 Jul 2009

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

N° attribué par la bibliothèque

□□□□□□□□□□

T H E S E

pour obtenir le grade de

DOCTEUR DE L'INPG

Spécialité : « **Génie Electrique** »

préparée au Laboratoire d'Electrotechnique de Grenoble
dans le cadre de l'Ecole Doctorale « **Electronique, Electrotechnique, Automatique et Traitement du signal** »

présentée et soutenue publiquement

par

Vincent MANZO
Ingénieur ENSIEG

le 22 Octobre 2004

**Traitement des congestions dans les réseaux de transport et dans un
environnement dérégulé**

Directeur de thèse :

Nouredine HADJ SAID

JURY

M.	D. ROYE	, Président
M.	M. EREMIA	, Rapporteur
M.	B. MULTON	, Rapporteur
M.	N. HADJ SAID	, Directeur de thèse

AVANT-PROPOS

Ce travail de thèse a été effectué au Laboratoire d'Electrotechnique de Grenoble (LEG) dans l'équipe SYREL (Systèmes&Réseaux électriques). Je tiens à remercier chaleureusement :

Monsieur Daniel ROYE, Professeur à l'Ecole Nationale Supérieure d'Ingénieurs Electriciens de Grenoble, pour m'avoir fait l'honneur d'être Président de mon jury de thèse,

Monsieur Mircea EREMIA, Professeur à l'Université Polytechnique de Bucarest et Monsieur Bernard MULTON, Professeur à l'ENS Bretagne, pour avoir accepté d'être rapporteurs pour ma thèse et pour avoir consacré du temps à la lecture de cette thèse,

Monsieur Adi MANESCU, Enseignant-Chercheur à l'Université de Craiova, pour m'avoir encadré durant ces trois années de thèse, pour ses grandes compétences, et pour m'avoir prodigué des conseils et un soutien oh combien utile,

Monsieur Nouredine HADJ SAID, Professeur à l'Ecole Nationale Supérieure d'Ingénieurs Electriciens de Grenoble, qui m'a permis de mener cette thèse à bien. Il m'a accordé sa confiance en me laissant une grande autonomie dans la conduite de mes travaux, et j'ai pu découvrir au fil de ces trois années sa personnalité faite d'un optimisme réconfortant mais aussi d'une grande rigueur, ce qui m'a bien été utile à des moments particulièrement décisifs,

Monsieur Seddik BACHA, Professeur à l'Ecole Nationale Supérieure d'Ingénieurs Electriciens de Grenoble, qui m'a accueilli en DEA et qui m'a donné l'occasion de faire mes premiers pas dans le monde de la recherche,

Monsieur Jean-Pierre ROGNON, Ancien Directeur du Laboratoire d'Electrotechnique de Grenoble, pour m'avoir accueilli dans son laboratoire.

Je tiens aussi à remercier :

Tous les doctorants du LEG avec qui j'ai passé trois années dans une ambiance de travail chaleureuse et très détendue. Je tiens d'ailleurs à faire un salut spécial aux communautés roumaines, bulgares et vietnamiennes très présentes au sein de notre laboratoire,

Tous les permanents du LEG, pour leur grande sympathie et pour les bons moments passés en leur compagnie ,

Ianko VALERO, Ion EXTEBERRIA et Stefan STERPU qui m'ont permis de garder un excellent souvenir du voyage en Roumanie (rappelez vous !) ; c'est une expérience que je n'oublierais jamais,

Tous ceux qui ont contribué directement ou indirectement à la réussite de ce travail.

Je tiens enfin aussi à remercier toute ma famille d'ici et tous mes amis (allons-y : Sly, Mag, Francky, Lyo, Marie, Isa, Gaëlito, Pau, Marion, Jérôme, Camille, Cédric le tout fou, Juju, Nounours et tous ceux que j'aurais malencontreusement oubliés....)

Je tiens aussi à saluer toute ma famille en Italie ; je ne vais pas faire d'énumérations (ce serait trop long !), mais je pense à elle très souvent, et je ne les oublies pas, comme eux ils ne m'oublient pas malgré la distance !

Enfin, At Last But Not The Least, merci à toi Gwendoline, pour tout ce que tu as fait et continues de faire pour moi, et pour avoir donné un nouveau sens à ma vie depuis presque deux ans de cela à présent.

RESUME

TRAITEMENT DES CONGESTIONS DANS LES RESEAUX DE TRANSPORT ET DANS UN ENVIRONNEMENT DEREGULE

La restructuration du secteur de l'électricité occasionne des transferts de puissance importants guidés par une logique essentiellement économique, engendrant à leur tour de nouvelles contraintes sur les réseaux de transport appelées congestions. Une congestion dénote l'incapacité du réseau de transport à conduire les programmes du marché de l'énergie, et le traitement des congestions est le procédé assurant que les réseaux sont exploités dans leurs limites de sécurité imposées. Le contexte actuel nécessite donc de définir une méthodologie de traitement des congestions fiable, optimale du point de vue économique, et qui donne de bonnes incitations sur le long terme en vue de réduire les contraintes et de favoriser le développement du réseau. A cet effet, nous proposons dans le cadre de cette Thèse un modèle de traitement des congestions qui définit celui-ci comme un service système séparé du marché de l'énergie. Le traitement repose sur l'usage d'offres d'ajustements venant des producteurs sur une base volontaire, et dont le coût total est minimisé via un algorithme d'optimisation. Ensuite, ce coût est redistribué aux usagers du réseau suivant de nouvelles stratégies d'allocation basées sur la traçabilité de l'énergie. Enfin, cette méthodologie a été adaptée en vue de répondre au problème de la coordination supranationale du traitement des congestions. Les résultats ont montré que cette coordination facilite le traitement de contraintes difficiles, et permet d'espérer des réductions appréciables de coût de congestion, tout en assurant la confidentialité de données économiques sensibles. Les essais ont notamment été effectués sur le réseau RTS 96 comportant 72 nœuds.

ABSTRACT

CONGESTION MANAGEMENT IN POWER SYSTEMS AND IN A DEREGULATED MARKET

The restructuring of the electricity sector leads to huge power transfers guided by economical considerations essentially, creating at their turn new constraints on Power Systems called congestions. A congestion is the incapability for a power system to carry the market preferred schedules, and congestion management is the way to insure that power systems are operated within their security constraints. The current paradigm makes it necessary to define a congestion management method which is suitable, optimal from the economic perspective and which gives the right signals for constraints reduction and network improvement. For that purpose, we propose in this Phd a congestion management model which defines this latter as an ancillary service separated from energy market. The constraints are dealt thanks to adjustment bids made by producers on a voluntary basis. The overall cost of congestion is minimised through an optimisation algorithm. Then, the congestion cost is re-allocated among network users through new allocation strategies based on electricity tracing. Finally, this methodology has been adapted in order to allow supranational coordination of congestion management. Results show that coordination makes easier the management of difficult constraints, and allow hoping notable congestion cost reduction while insuring confidentiality of sensible economic datas. Tests have been carried on the RTS 96 system consisting in 72 nodes.

TABLES DES MATIERES

GLOSSAIRE	3
INTRODUCTION GENERALE	7
CHAPITRE I : LE PROBLEME DES CONGESTIONS DANS LES RESEAUX DE TRANSPORT D'ELECTRICITE A ACCES OUVERTS	9
I.1) LA DEREGULATION DU SECTEUR DE L'ELECTRICITE.....	9
I.1.1) MODELE POOL	11
I.1.2) MODELE BILATERAL	12
I.2) LE PHENOMENE DE CONGESTION DANS LES RESEAUX DE TRANSPORT.....	13
I.2.2) LIMITES DE TRANSIT IMPOSEES AUX OUVRAGES DU RESEAU DE TRANSPORT	13
I.2.1) LES MARCHES LIBERALISES FACE AUX LIMITES DU RESEAU	14
I.2.3) LE PROBLEME DES FLUX PARALLELES.....	16
I.3) LE CAS FRANÇAIS : SOLUTION DU RTE.....	20
I.4) CONCLUSION PRELIMINAIRE SUR LE PHENOMENE DE CONGESTION.....	22
CHAPITRE II : REVUE DES PRINCIPALES METHODES DE TRAITEMENT DES CONGESTIONS	24
II.1) INTRODUCTION	24
II.2) MODELES DE TRAITEMENTS DES CONGESTIONS APPLIQUES ACTUELLEMENT	24
II.2.1) LA REGIONALISATION DU MARCHÉ (OU <i>MARKET SPLITTING</i>) NORDIQUE	25
II.2.2) LA SOLUTION CALIFORNIENNE	27
II.2.3) LES COUPURES DE TRANSACTIONS : LA SOLUTION DU NERC AMERICAIN	30
II.2.4) UN OUTIL GENERALISE D'OPTIMISATION DE LA PRODUCTION : L'OPF (<i>OPTIMAL POWER FLOW</i>)	33
II.2.5) TRAITEMENT DES CONGESTIONS PAR AJUSTEMENTS DE PRODUCTION : MODELE DU <i>BUY BACK</i>	38
II.2.5.1) <i>Démarche générale</i>	39
II.2.5.2) <i>Formulation mathématique</i>	41
II.2.5.3) <i>Offres d'ajustements des producteurs</i>	42
II.3) CONCLUSION	43
CHAPITRE III : CHOIX D'UN MODELE DE TRAITEMENT DES CONGESTIONS.....	43
III.2) ETUDE D'UN RESEAU 9 NŒUDS	45
III.2.1) PRESENTATION DU RESEAU 9 NŒUDS-DONNEES DU MARCHÉ	45
III.2.2) CAS NON CONTRAINT	46
III.2.3) TRAITEMENT DES CONGESTIONS PAR AJUSTEMENTS DE PRODUCTION (BUY BACK)	47
III.2.3.1) <i>Etablissement des plans préférés du marché et vérification de ces plans</i>	47
III.2.3.2) <i>Minimisation du coût des ajustements et résolution de la contrainte sur la ligne 1-6</i>	48
III.2.4) TRAITEMENT DES CONGESTIONS PAR COUPURES DE TRANSACTIONS	52
III.3.4.1) <i>Premier cas</i>	52
III.3.4.2) <i>Deuxième cas</i>	54
III.2.5) TRAITEMENT DES CONGESTIONS PAR L'USAGE DE L'OPF ET DES PRIX NODAUX	56

III.3.6) SENSIBILITE DU COUT DE CONGESTION DANS LES MODELES DU BUY BACK ET DES PRIX NODAUX	61
III.3.6.1) Sensibilité par rapport aux offres d'entrées	61
III.3.6.2) Sensibilité par rapport à la contrainte imposée	63
III.3.6) VALIDATION DE L'ORO BASE SUR LE MODELE DC	64
III.3.6.1) Validation sur le réseau 9 noeuds	64
III.3.6.2) Validation sur le réseau New England IEEE 39 nœuds	66
III.3.7) SYNTHESE DES RESULTATS OBTENUS PAR LE BUY BACK, COUPURES DE TRANSACTIONS, ET PAR L'USAGE DE L'OPF	68
 CHAPITRE IV : STRATEGIES D'ALLOCATION DES COUTS DE CONGESTION BASEES SUR LA TRAÇABILITE DE L'ENERGIE	71
 IV.1) INTRODUCTION	71
IV.2) CONTEXTE ET ENJEUX DE L'ALLOCATION DES COUTS DE CONGESTION	72
IV.3) ALLOCATIONS DES COUTS DE CONGESTION BASEES SUR LA TRAÇABILITE DE L'ENERGIE	73
IV.3.1) LA TRAÇABILITE DE L'ENERGIE	73
IV.3.2) FORMULATIONS DES ALLOCATIONS	75
IV.3.2.1) Approches statiques : allocations basées sur les contributions au transit	76
IV.3.2.2) Approche différentielle : allocations par destination des ajustements de production	78
IV.4) RESULTATS NUMERIQUES SUR LE RESEAU 9 NŒUDS	79
IV.4.1) DONNEES FOURNIES PAR LA TRAÇABILITE DE L'ENERGIE	79
IV.4.2) RESULTATS DES ALLOCATIONS SUR LE RESEAU 9 NŒUDS BASEES SUR LA TRAÇABILITE	81
IV.4.2.1) Allocation du coût de la congestion sur la ligne 1-6	81
IV.4.2.2) Cas présentant plusieurs lignes congestionnées	84
IV.4.3) AUTRES ALLOCATIONS PHYSIQUES : RESULTATS DE METHODES BASEES SUR LES FACTEURS DE DISTRIBUTION	86
IV.4.3.1) Allocation aux consommateurs par les facteurs de distribution	86
IV.4.3.2) Allocation aux transactions par facteurs de distribution	87
IV.4.3.2) Synthèse des résultats des allocations basées sur les facteurs de distribution	89
IV.4.4) EVALUATION APPROFONDIE DES ALLOCATIONS BASEES SUR LA TRAÇABILITE DE L'ENERGIE	90
IV.4.4.1) Responsabilisation des acteurs du marché par rapport aux conséquences sur le réseau des choix économiques du marché	90
IV.4.4.2) Découragement des scénarios congestionnels	92
IV.4.4.3) Intérêt au développement du réseau	92
IV.5) CONCLUSION SUR L'ALLOCATION DES COUTS DE CONGESTION BASEE SUR LA TRAÇABILITE DE L'ENERGIE	93
 CHAPITRE V : VERS UNE COORDINATION SUPRANATIONALE DU TRAITEMENT DES CONGESTIONS	95
 V.1) INTRODUCTION	95
V.2) PRINCIPE DE COORDINATION DU TRAITEMENT DES CONGESTIONS ENTRE PLUSIEURS OPERATEURS DU SYSTEME	96
V.3) OUTIL DE REDISPATCHING OPTIMISE (ORO) COORDONNE: FORMULATION	97
V.3.1) PROBLEME INITIAL (ORO GLOBAL)	97
V.3.2) PROBLEME EQUIVALENT	97
V.3.3) PROBLEME COORDONNE	99
V.4) ETUDE DU RESEAU RTS96	101
V.4.1) ETAT INITIAL DU RESEAU	101

V.4.2) CAS CONTRAINTS : PRESENTATION DES CAS DE CONGESTION ETUDIES	102
V.4.2.1) Cas 1 : congestion sur la ligne connectant le nœud 220 au nœud 223	102
V.4.2.2) Cas 2 : congestion sur la ligne 220-223 et sur la ligne 221-215	103
V.4.2.3) Cas 3 : congestion sur l'interconnexion 107-203 reliant la zone A à la zone B	104
V.4.3) PARAMETRES DES OFFRES D'AJUSTEMENTS	105
V.4.4) TRAITEMENT DES CONGESTIONS DANS LES TROIS CAS PRESENTES : COMPARAISON DE DIFFERENTS ORO (LOCAL, GLOBAL, COORDONNE)	105
V.4.5) ALLOCATION DES COUTS DE CONGESTION SUR LE RESEAU RTS 96	111
V.4.5.1) Une allocation décentralisée des coûts de congestion	111
V.4.5.2) Allocation du coût de congestion dans le cas 1	113
V.4.5.3) Allocation dans le cas 3	116
V.4.5.4) Synthèse des résultats des allocations obtenues sur le réseau RTS 96	116
V.5) CONCLUSION SUR L'APPROCHE SUPRANATIONALE DU TRAITEMENT DES CONGESTIONS	117
 CONCLUSION GENERALE	 119
 REFERENCES BIBLIOGRAPHIQUES	 121
 ANNEXE 1 : CALCUL DE REPARTITION DE CHARGE DIT A COURANT CONTINU (MODELE DC)	 127
ANNEXE 2 : THEORIE DE L'OPTIMISATION	136
ANNEXE 3 : METHODE DES IMAGES DE CHARGES BASE SUR LE MODELE DC	150
ANNEXE 4 : RESEAU RTS 96 ET RESULTATS DES COMPLEMENTAIRES SUR LES ALLOCATIONS DE COUT	165
ANNEXE 5 : LES TRANSFORMATEURS DEPHASEURS	170
ANNEXE 6 : PREUVE DE LA CONVERGENCE DE L'ALGORITHME D'OPTIMISATION DECOUPLE	174

GLOSSAIRE

Bourse de l'électricité : structure centralisée permettant la confrontation des offres de vente des producteurs et de la demande des acheteurs (distributeurs, consommateurs industriels, etc...) jusqu'à équilibre production-consommation, équilibre fixant alors un certain prix de marché, égal au prix de la dernière tranche de production appelée.

Calcul de répartition de charge : algorithme destiné à déterminer les transits de puissance dans un réseau électrique à partir des données de production et de consommation et des paramètres du réseau

Day ahead market : (définition de PJM) programmation des échanges la veille pour le lendemain et calcul des prix nodaux et des coûts de congestion prévisionnels pour chaque heure du lendemain

Fournisseur : acteur du marché vendant de l'électricité ; peut être différent du producteur

ISO : Independent System Operator. Dénomination anglophone de l'opérateur chargé d'assurer la conduite du réseau de transport. Il peut suivant les différentes organisations avoir aussi la responsabilité de gestion d'un pool (modèle de PJM)

J-1 : échéance de temps permettant de programmer des transactions la veille pour le lendemain. Cette programmation se fait pour chaque pas horaire (1 heure en général) du jour J. Il peut exister d'autres échéances de temps plus longues en J-k, comme par exemple pour la réservation de capacité pour des échanges internationaux qui peut se faire jusqu'à 6 mois à l'avance.

Load-flow : dénomination anglophone du « calcul de répartition de charge »

Locationnal Marginal Pricing (LMP) : équivalent anglophone de la « tarification nodale » ; voir prix nodal plus loin

Marché spot de l'électricité : marché dédié au négoce de l'énergie électrique permettant de confronter les offres des producteurs et les demandes des acheteurs, en vue d'effectuer des livraisons d'électricité à une certaine échéance (immédiate, H-1, J-1,...) .

Opérateur du système : entité chargée d'assurer la conduite du réseau de transport d'électricité, de veiller à sa sécurité et à son développement

Optimal Power Flow (OPF) : algorithme destiné à trouver une répartition des transits optimale par rapport à un critère donné et satisfaisant les contraintes de fonctionnement du système (limites imposées aux ouvrages, équilibre production-consommation, etc ...). Le critère retenu est très souvent un critère économique de type minimisation des coûts de production, des coûts des pertes, etc...

PJM : Grand réseau interconnecté dans le Nord-Est américain.

Pool : modèle de marché qui centralise les offres des producteurs et les demandes d'énergie et par lequel peut se faire le négoce d'électricité. Un prix de marché est alors délivré ; ce prix dépend du temps, et peut aussi suivant le modèle adopté dépendre de la localisation (prix nœaux).

Prix nodal : prix attribué à la puissance électrique consommée en un nœud donné du réseau et à un instant donné. Il se définit comme le surcoût engendré pour consommer une unité de puissance supplémentaire (coût dit *marginal*).

Trader : acteur du marché intermédiaire gérant des livraisons d'électricité avec un certain portefeuille de clients et de fournisseurs/producteurs

Transaction : contrat d'achat d'une certaine quantité d'électricité conclut entre un ou plusieurs producteur(s) et un ou plusieurs acheteur(s) sans passer par une bourse de l'électricité. Elle se caractérise par un ou plusieurs point(s) d'injection et ou plusieurs point(s) de soutirage. Elle est dite bilatérale lorsqu'elle implique uniquement un producteur et un acheteur ; elle est multilatérale lorsqu'elle implique plus de deux acteurs et peut être coordonnée par un intermédiaire (fournisseur spécialisé dans le négoce de l'électricité, trader..).

INTRODUCTION GENERALE

Les grands réseaux d'énergie électriques sont des structures vastes et complexes dont le rôle est d'acheminer l'électricité souvent sur de longues distances, depuis des centres de production jusqu'aux centres de consommation. Ces réseaux sont soumis à des règles de fonctionnement très strictes, qui obligent les exploitants des centres de conduite à faire fonctionner le réseau dans ses limites de sécurité. Plus précisément, chaque ouvrage du réseau de transport a une limite de transit à ne pas franchir. Ces limites sont fonction de la tenue thermique des ouvrages, mais aussi des limites en tenue de tension et des limites de stabilité du réseau.

Traditionnellement, le secteur électrique était détenu par un opérateur intégré qui avait le monopole sur les fonctions de production, de transport, et distribution de l'énergie électrique. Pour satisfaire la demande, il choisissait ses unités de production par ordre croissant de coût de production (on parlait alors de *liste de mérite*), tout en satisfaisant les contraintes techniques de fonctionnement du réseau. Il pouvait éventuellement compléter sa fourniture par des imports de l'extérieur, bien connus et établis. Ayant le contrôle de sa production et des échanges intervenants aux frontières de son réseau, il s'ensuit qu'il avait une bonne maîtrise des transits de puissance, les éventuelles surcharges ou aléas étant traitées en temps réel par des mesures d'urgence classiques.

Cependant, ces dernières années ont vu une profonde mutation s'opérer dans le secteur électrique, occasionnant des changements organisationnels importants. Une vague de restructuration du secteur électrique s'est rapidement propagée dans le monde entier, entraînant la séparation des activités de production, de transport et de distribution de l'électricité. Elle a eu pour conséquence de multiplier le nombre d'acteurs sur le marché de l'électricité et d'introduire la concurrence entre les fournisseurs d'énergie électrique. Le choix des fournisseurs se fait alors sur de nouvelles structures de marché dont l'opérateur du réseau de transport n'a pas le contrôle, ce qui occasionne des transferts de puissance massifs guidés uniquement par une logique économique. Ces transferts de puissance engendrent alors l'apparition de contraintes de plus en plus fréquentes sur les réseaux de transport, appelées aussi congestions. Une situation de congestion est définie lorsque le système de transport n'est plus capable de conduire les transactions du marché de l'énergie sans que la limite de transit ne soit violée sur un ou plusieurs ouvrages du réseau. A cause de l'imprévisibilité de l'évolution du marché de l'énergie, de la multiplication des transactions commerciales et de la forte interaction physique des réseaux interconnectés, il est très difficile de prévenir

l'apparition des congestions sur le réseau. De plus, la possibilité de renforcement du réseau est très limitée du fait du coût élevé des investissements, des pressions écologiques freinant les constructions de nouvelles lignes et du manque de prévision sur le long terme de l'évolution du marché. Tout ceci a obligé les autorités chargées de la sécurité du réseau à mettre en place des solutions adaptées à l'environnement dérégulé en vue de gérer les contraintes dans un cadre prévisionnel. Le *traitement des congestions* est donc le procédé par lequel on va s'assurer que le système sera conduit en temps réel dans le respect des contraintes de sécurité imposées à tous les ouvrages du réseau. Il engendre en outre un coût supplémentaire d'exploitation qu'il convient de toujours minimiser.

Ainsi, depuis le début de la restructuration du secteur électrique, de nombreuses méthodes de traitement des congestions ont émergé, souvent liées au choix organisationnel décidé pour chaque pays concerné. L'efficacité d'une méthode de traitement des congestions peut se mesurer suivant ces critères essentiels :

- Elle doit être techniquement faisable et la plus flexible d'utilisation: elle doit pouvoir s'adapter à tous les types de réseaux et de marchés, et pouvoir trouver des solutions dans le plus grands nombre de cas possibles. Elle doit pouvoir aussi être facilement coordonnable au niveau international ;
- Elle doit être transparente : on doit s'assurer qu'elle modifie le moins possible les sorties du marché de l'énergie (volumes et prix). En outre, elle doit s'appuyer sur des outils techniques incontestables et ses règles de fonctionnement doivent être suffisamment simples pour être comprises de tous les participants ;
- Elle doit être économiquement efficace : le coût d'exploitation supplémentaire engendré par les congestions doit être le plus bas possible. En outre, le traitement des congestions doit donc être tarifié aux usagers de façon judicieuse et adéquate, afin de donner de bonnes incitations aux participants au marché.

L'objet de notre travail de thèse est de définir un modèle de traitement des congestions répondant le mieux possible à tous ces critères. Le présent mémoire est structuré comme suit :

Dans le premier Chapitre, nous allons rappeler les nouveaux modes de fourniture en électricité étant apparus avec la libéralisation du marché. Nous allons ensuite expliquer le lien entre ce nouveau contexte et l'apparition des contraintes sur le système, en insistant sur le rôle des flux dits « parallèles ».

Dans le second Chapitre, nous allons décrire les principales méthodes de traitement des congestions ayant été appliquées jusqu'à maintenant, en les explicitant sur le plan théorique.

Dans le troisième Chapitre, nous allons présenter une étude de cas sur un réseau 9 nœuds, où nous comparons la méthode du « buy-back » aux méthodes dites de « coupures de transactions » et des « prix nodaux ». Nous ferons en outre une analyse poussée sur le plan de l'efficacité technique et économique de chacun des modèles, et concluons sur l'intérêt de la méthode du buy-back.

Le quatrième Chapitre sera consacré à l'allocation des coûts de congestion. Après avoir brièvement rappelé les méthodes d'allocation existantes, nous présenterons de nouvelles stratégies d'allocation basées sur la traçabilité de l'énergie que nous proposons. Nous illustrerons ces allocations sur le cas du réseau 9 nœuds, en examinant très attentivement la pertinence des signaux économiques envoyés. Nous comparerons en outre les allocations proposées à des allocations basées sur les facteurs de distribution, et soulignerons d'une part la cohérence des résultats obtenus, et d'autre l'apport des nouvelles allocations basées sur la traçabilité. Enfin, nous discuterons de façon plus approfondie sur la pertinence de chaque version d'allocation proposée.

Enfin, le cinquième Chapitre sera consacré à la coordination du traitement des congestions au niveau supranational. Nous présenterons une méthode permettant de découpler le traitement des congestions entre plusieurs opérateurs, tout en coordonnant leurs actions afin d'atteindre l'optimum global du système interconnecté. Nous décrirons ensuite les résultats de nos travaux ayant porté sur l'étude du réseau IEEE RTS 96, en montrant notamment les bienfaits de la coordination du traitement des congestions sur le coût global de congestion et sur la faisabilité du traitement. Les allocations basées sur la traçabilité présentées au quatrième Chapitre seront aussi appliquées sur les cas étudiés sur le réseau RTS 96, ce qui donnera lieu de nouveau à une analyse approfondie des allocations présentées. Cette analyse nous permettra de mieux valider les conclusions du quatrième chapitre, voire de les nuancer dans certains cas.

Enfin, nous clôturerons cette Thèse par une conclusion générale dans laquelle nous allons mettre en avant l'apport général délivré par nos travaux, ainsi que les limitations de la méthode de traitement des congestions présentée. Nous présenterons aussi les perspectives qui pourront faire suite à ces travaux.

CHAPITRE I

Le problème des congestions dans les réseaux de transport d'électricité à accès ouverts

I.1) La dérégulation du secteur de l'électricité

Traditionnellement, le secteur de l'électricité est détenu par un seul opérateur historique, qui gère à la fois la production de l'énergie, son transport et sa distribution vers ses clients. C'était une situation dite de « monopole », où les clients, hormis quelques gros consommateurs industriels ou ceux raccordés à de rares distributeurs indépendants, n'ont pas le choix de leur fournisseur.

Une première expérience significative de libéralisation du secteur de l'électricité a été effectuée en Amérique du Sud dans les années 80, suivi par le Royaume-Uni au début des années 90. A partir de là, ce mouvement de libéralisation et de restructuration du secteur, qui a progressivement mis fin à son ancienne structure verticalement intégrée s'est progressivement propagé aux Etats-Unis, en Nouvelle-Zélande, en Scandinavie et finalement dans le reste de l'Europe. Cette libéralisation se traduit pour les consommateurs par la possibilité de choisir un fournisseur autre que le fournisseur historique duquel ils étaient « captifs ».

La restructuration du secteur, a entraîné la séparation des activités de production, de transport et de distribution de l'énergie électrique (Fig. I.1 et Fig. I.2). Les activités de production et de distribution (partie commerciale et services aux consommateurs) sont devenues dans la majeure partie des cas privées, avec l'éclatement de l'ancien opérateur intégré en plusieurs compagnies électriques de production et de distribution en concurrence.

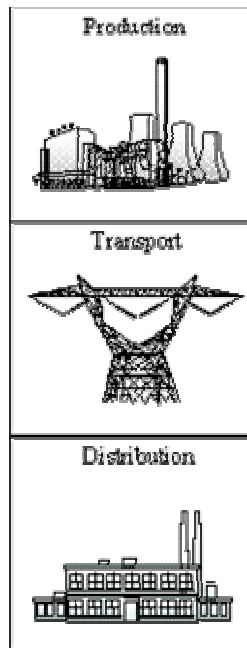


Figure I.1 : ancienne structure verticalement intégrée du secteur de l'électricité

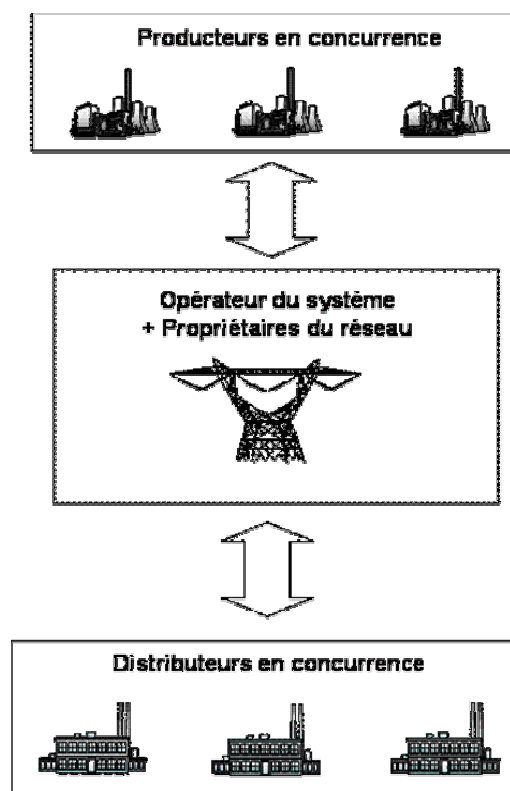


Figure I.2 : séparation des activités de production, transport et distribution

La restructuration de l'activité de transport a donné naissance à plusieurs nouvelles formes d'organisation qui diffèrent suivant les pays et les états concernés. En France, le réseau de

transport reste un bien d'utilité publique, et sa gestion est confiée à un opérateur du système indépendant appelé aussi Gestionnaire du Réseau de Transport (GRT) à but non lucratif. Son rôle est de garantir un accès au réseau non discriminatoire et de veiller à son entretien et à son développement. Par contre aux Etats-Unis, plusieurs formes d'organisation ont vu le jour [HOG1 99]. Parmi celles-ci, on peut citer les :

- TRANSCO (*Transmission Company*) : ce sont des sociétés qui conjuguent les fonctions de propriétés du réseau et de sa gestion opérationnelle, et qui peuvent être à but lucratif
- ISO (*Independent System Operator*) : On le trouve notamment en Californie ou dans l'état de New York. Cet opérateur du système ne possède pas le réseau de transport, et a uniquement pour mission de gérer le réseau et d'assurer sa sécurité
- ISO mixte : l'opérateur du système cumule le rôle de gestion et de supervision du réseau, et celui de gestion d'un marché de l'énergie, mais ne possède pas nécessairement le réseau. C'est la situation de l'opérateur de PJM (*Pennsylvania-New Jersey-Maryland*) au nord-est des Etats-Unis.

La libéralisation du secteur a aussi entraîné l'émergence de nouvelles structures de marché de l'électricité, dont les 2 plus répandues sont le modèle *pool*, qui a la forme d'une bourse centralisée, et le modèle *bilatéral*, où un producteur et un consommateur concluent un contrat pour une certaine fourniture en énergie à un prix négocié librement entre eux. Ces deux modes de fourniture peuvent d'ailleurs très bien coexister au sein d'une même région. Dorénavant, nous désignerons les gros consommateurs industriels raccordés au réseau de transport ainsi que les distributeurs indépendants par le terme général de « consommateurs ».

I.1.1) Modèle pool

Dans le modèle *pool*, le négoce d'énergie est gérée de façon centralisée par un opérateur de bourse qui collecte les offres des producteurs et les demandes des consommateurs jusqu'à obtenir l'équilibre production-consommation. Les producteurs spécifient pour chaque tranche de puissance proposée un prix de vente laissé à leur choix. Les consommateurs quant à eux précisent des commandes fermes d'achat, et éventuellement un prix au-delà duquel ils préfèrent retirer leur demande de la bourse. Il peut cependant exister des modèles de bourse dans lesquels les consommateurs peuvent varier leur demande en fonction du prix auquel ils auront à payer leur fourniture ; on parle dans ces cas-là d'*élasticité* de la demande. L'opérateur de la bourse classe alors les offres des producteurs de la moins chère vers la plus

chère, et les demandes des consommateurs du plus offrant vers le moins offrant. Ce processus d'agrégation peut être mis sous forme de courbes d'offres de production et de demande tel que le montre la Figure I.3. L'intersection des deux courbes nous donne le point d'équilibre production-consommation (donc le volume total d'énergie contracté à la bourse pour la tranche horaire donnée), ainsi que le prix auquel a été fixé l'énergie contractée. Ce prix correspond au prix de la dernière tranche (dite tranche marginale) prise en compte (*Market Clearing Price-MCP*).

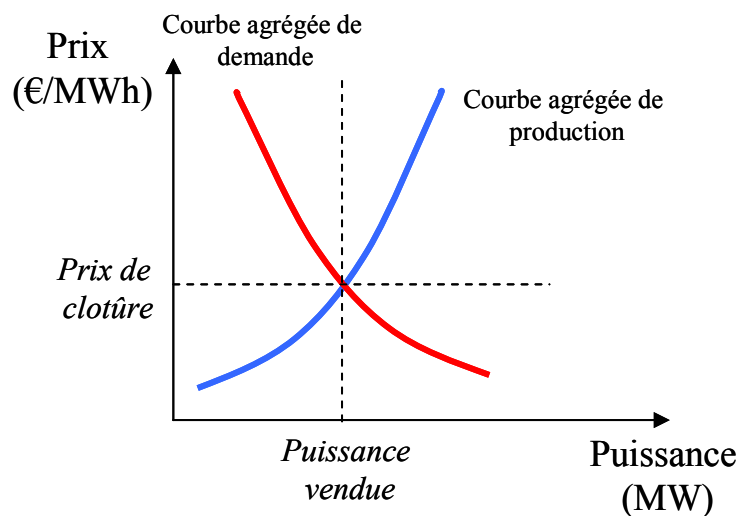


Figure I.3 : principe de fonctionnement d'un marché pool

Parmi les pays qui ont choisi le modèle pool figure le Royaume-Uni, qui a imposé au début de la dérégulation de son secteur de l'électricité une bourse de l'énergie unique et obligatoire pour tous les participants. Les bourses de l'électricité fonctionnent en général la veille pour le lendemain : cette échéance de temps est dite « J-1 » ou *day ahead*. La fourniture d'électricité est alors négociée pour chaque tranche horaire (1 heure ou ½ heure) du lendemain. C'est sur ce principe que fonctionne par exemple la bourse française Powernext. Certaines bourses, telles que CalPX en Californie, négocient même des produits à fournir l'heure suivante ; on parle alors d'une échéance « H-1 » (ou *hour ahead*).

I.1.2) Modèle bilatéral

Dans le modèle bilatéral, le consommateur contracte directement avec un fournisseur de son choix pour assurer sa fourniture en énergie. Ils se mettent aussi d'accord sur le prix de vente de l'énergie contractée. On parle alors ici de *transaction bilatérale*. Dans certaines régions du globe, ce mode peut être le principal moyen de fourniture en électricité, comme

c'est le cas en Scandinavie. D'autres marchés, comme PJM possède un marché spot centralisé, et laisse la possibilité à un certain nombre d'acteurs de se fournir par contrats bilatéraux. Enfin, en Espagne, ce mode fourniture semble être plutôt marginalisée en comparaison d'une bourse de l'énergie quasi obligatoire.

Le modèle bilatéral peut être étendu à plus de un producteur ou consommateur ; On peut alors parler dans ces cas-là de *transaction multilatérale*. Des acteurs de marché spécifiques appelés traders peuvent mettre en relation plusieurs fournisseurs et plusieurs consommateurs. En Californie, des opérateurs du marché spécifiques appelés *Scheduling Coordinators* (nous reviendrons sur le rôle par rapport au traitement des congestions plus loin) gèrent des groupes (ou portefeuilles) de participants au marché de l'énergie dont le bilan électrique (somme des puissances vendues et somme des puissances achetées) est nul. En outre, en France, les traders deviennent *responsables d'équilibre* lorsqu'ils prennent la responsabilité des écarts constatés durant la conduite entre leurs fournitures et leurs livraisons déclarées.

I.2) Le phénomène de congestion dans les réseaux de transport

I.2.2) Limites de transit imposées aux ouvrages du réseau de transport

Dans les réseaux de transport, des limites en terme de puissance maximale pouvant transiter sur une ligne peuvent être imposées en fonction :

- **Des limites thermiques:** pour des lignes dites « courtes » (< 80 km), ces sont surtout des limites thermiques qui sont rencontrées en premier. Le courant circulant dans les conducteurs provoque un échauffement (par effet Joule), qui, en cas de forte surcharge, peut détériorer les conducteurs
- **Des limites de tenue en tension :** les limites en tenue de tension sont plus contraignantes pour les lignes de longueur « moyenne » (entre 80 et 250 km) que les limites thermiques. Plus la puissance active circulant dans ces lignes est importante, plus on observe un phénomène de chute de tension du à l'impédance de la ligne. Dans les cas les plus critiques, cela peut provoquer un effondrement de tension en bout de ligne qui une fois entamé nécessite des délestages au niveau de la charge. Ces effondrements peuvent aussi mener à la perte de l'ensemble du réseau (blackout).

- **Des limites de stabilité de synchronisme** : ces contraintes apparaissent pour les lignes longues (>250 km). Des perturbations sur le réseau (perte d'un générateur, défaut...) peuvent occasionner des oscillations entre deux centres de production relié par une ligne longue. Si ces oscillations ne sont pas amorties, elles peuvent mener jusqu'au déclenchement de la ligne.

La Figure I.4 nous donne les limites de transit habituelles imposées aux lignes en fonction du niveau de tension (tensions US) [ABB 01] et de leur longueur :

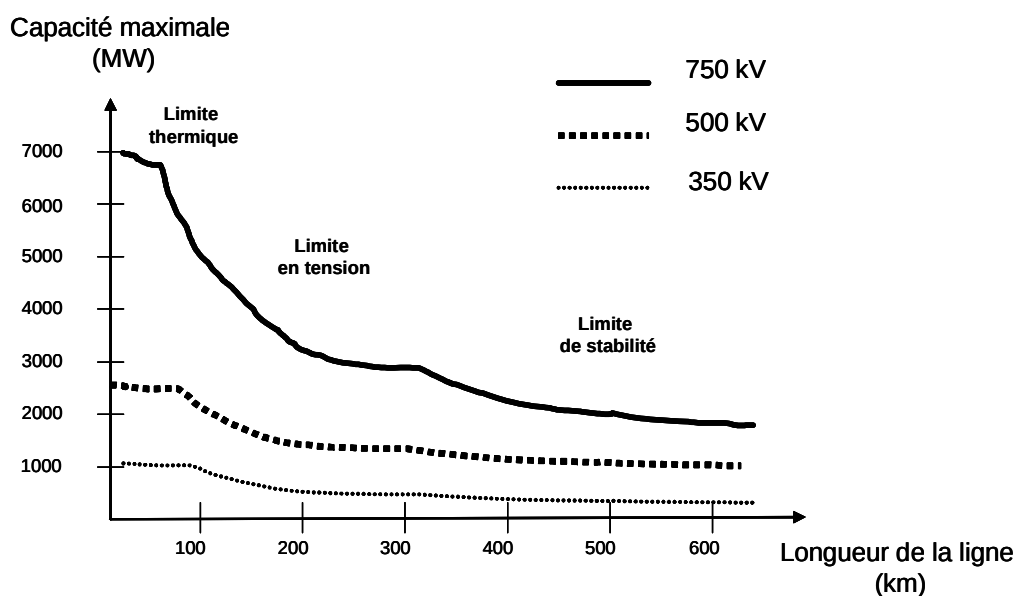


Figure I.4 : limites thermiques, de tension et de stabilité de synchronisme des lignes de transport en fonction du niveau de tension et de leur longueur [ABB 01]

I.2.1) Les marchés libéralisés face aux limites du réseau

Le développement des grands réseaux de transport a été jusqu'à récemment assuré par des monopoles nationaux qui adaptaient leur parc de production et le renforcement de leur réseau à leurs prévisions de consommation à long terme. Les réseaux conçus dans une logique « monopoliste », étaient donc bien adaptés au marché intérieur et aux imports/exports connus et établis. Cependant, la libéralisation du secteur de l'électricité a entraîné une internationalisation des échanges, et l'entrée de nouveaux arrivants sur le marché. Cela a pour effet de changer de façon conséquente la répartition des transits de puissance sur le réseau, de les rendre plus imprévisibles, et finalement de pousser toujours au plus près de ses limites un

réseau qui n'a pas encore été adapté à ce changement. On peut alors se retrouver dans une situation où les exigences du marché, qui voudraient que les réseaux fonctionnent « comme des plaques de cuivre », se heurtent aux réalités physiques du fonctionnement des réseaux.

Pour analyser l'interaction des marchés libéralisé avec le fonctionnement des réseaux, prenons un cas de figure très basique où l'on a deux zones A et B reliées par une interconnexion AB (Fig. I.4). Supposons que la zone B dispose d'une énergie produite plus coûteuse que celle de la zone A, et qu'elle importe 500 MW de la zone A. La capacité de l'interconnexion AB conçue pour un tel import/export est de 750 MW. Les deux zones consomment chacune 1000 MW. Cette situation peut être assimilée à celle pouvant prévaloir avant la libéralisation du marché.

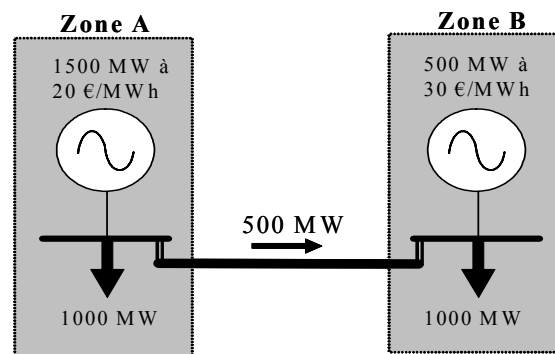


Figure I.4 : cas de deux zones avec import/export de 500 MW

Supposons à présent qu'après la libéralisation, certains des consommateurs de la zone B ayant désormais le choix de leur fournisseur préfèrent se fournir en zone A où l'énergie est moins chère. Cela va se traduire par un changement de production qui va modifier les transits et notamment le transit sur l'interconnexion AB qui va augmenter. Si la capacité de l'interconnexion est seulement de 750 MW, cela signifie que le changement de production maximum permis qui respecte les contraintes du réseau serait de 250 MW (fig I.5). Si, dans ces conditions limites, d'autres consommateurs de la zone B veulent encore changer de fournisseur, on aurait dépassé la capacité maximale de l'interconnexion AB. On serait alors dans une situation de **congestion**, qui obligerait les opérateurs du système des deux zones à prendre des mesures correctives sans quoi, les conditions de sécurité sur l'interconnexion seraient violées.

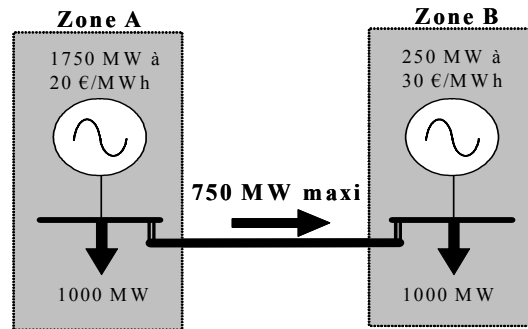


Figure I.5 : capacité maximale d'échange entre les deux zones atteinte dans le sens A->B

Les interfaces saturées des réseaux de transport interconnectés sont souvent appelés *goulots d'étranglements* et sont généralement dues à des faiblesses d'interconnexions entre réseaux maillés. Ceci étant dit, le phénomène de congestion n'est pas particulier aux interconnexions, mais à tout ouvrage du réseau fortement chargé, tel que lignes, transformateurs, etc... Toutefois, lorsque nous parlerons de congestions par la suite, nous nous référerons essentiellement aux congestions sur les lignes et interconnexions du réseau de transport. D'autre part, nous considérerons les congestions comme un problème de transit de puissance active, et les modes d'action que nous examinerons par la suite seront exclusivement liés à la puissance active.

I.2.3) Le problème des flux parallèles

Lorsque deux participants au marché veulent conclure une transaction bilatérale internationale, il est très courant dans les marchés libéralisés de procéder à des réservations auprès des opérateurs du système concernés de capacités disponibles sur les interconnexions. Pour ces réservations, on définit un « chemin contractuel » de la transaction du point source au point de soutirage. Le cheminement choisi a le plus souvent un caractère purement administratif et sert à régler les réservations de capacité sur les réseaux. Cependant, le cheminement de l'électricité obéit à des lois physiques bien précises et dépend fortement des caractéristiques des réseaux, et non pas d'un chemin contractuel purement arbitraire. Tout flux imprévu attribuable à une transaction et ne circulant pas sur le chemin contractuel est qualifié de *flux parallèle*.

Pour analyser le phénomène des flux parallèles, considérons l'exemple de la Figure I.7. Le producteur situé en zone A conclut une transaction bilatérale avec le consommateur en zone B et ils réservent auprès du gestionnaire de l'interconnexion la capacité sur la ligne 1 qui est choisie comme chemin conventionnel de la transaction.

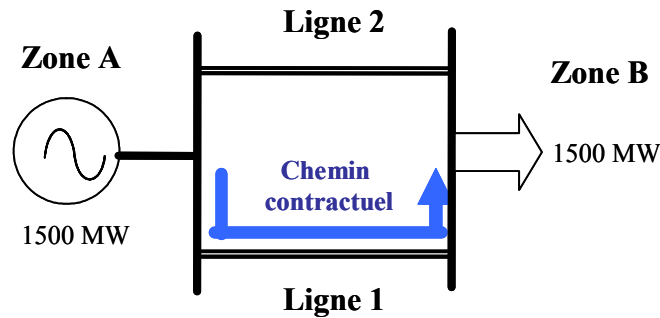


Figure I.7 : transaction bilatérale et réservation suivant le chemin contractuel choisi

Cependant, la manière dont se répartissent les flux physiques dans le réseau est tout autre. L'électricité obéit en effet à des lois physiques bien précises, connues sous le nom de lois de Kirchhoff, qui font que les flux se répartissent principalement suivant les impédances des lignes composant le réseau. Dans notre exemple, nous supposons que la ligne 1 a une impédance de 1 unité réduite¹ (u.r.) et que l'impédance de la ligne 2 a une valeur de 2 u.r.. Il en résulte que 2/3 seulement de la puissance contractée passe par la ligne 1 et que 1/3 passe par la ligne 2. (fig.I.8)

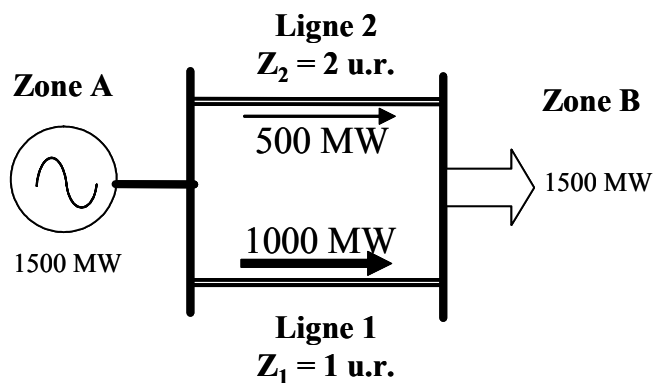


Figure I.8 : flux physiques résultants de la transaction bilatérale

Le phénomène des flux parallèles est fondamental dans les réseaux maillés, et plus le maillage et l'interdépendance des réseaux interconnectés seront forts, plus les flux parallèles seront importants. Ceci explique entre autre pourquoi un ensemble maillé peut difficilement être géré « par morceaux » et que les réseaux, malgré la restructuration du secteur de l'électricité, restent des monopoles naturels. Un exemple « grandeur nature » connu nous

¹ En ce qui concerne l'impédance d'une ligne, l'unité réduite représente le rapport entre la valeur réelle de l'impédance (en Ohm) et la valeur de l'impédance de base choisie pour le réseau étudié (en Ohm). D'autres grandeurs électriques (tension, courant, puissances ...) peuvent aussi être exprimés en unités réduites.

montre l'impact que peut avoir une transaction entre la Belgique et l'Italie sur l'ensemble des réseaux interconnectés (Fig. I.9). Dans cet exemple plus complexe que le précédent, les flux résultants ne doivent pas être interprétés en terme de chemins réels empruntés par la transaction. En effet, il est évident que les 100 MW produits en Belgique ne vont pas tous être consommés en Italie, vu que l'électricité ne suit pas le chemin dit contractuel. Cependant, on peut dire qu'une augmentation de production de 100 MW en Belgique et une augmentation de la consommation de 100 MW en Italie provoque des changements dans les transits de puissance de l'ensemble du réseau interconnecté. Les flèches jaunes sur la figure I.9 indiquent dans quel sens se font ces changements et suivant quelle proportion :

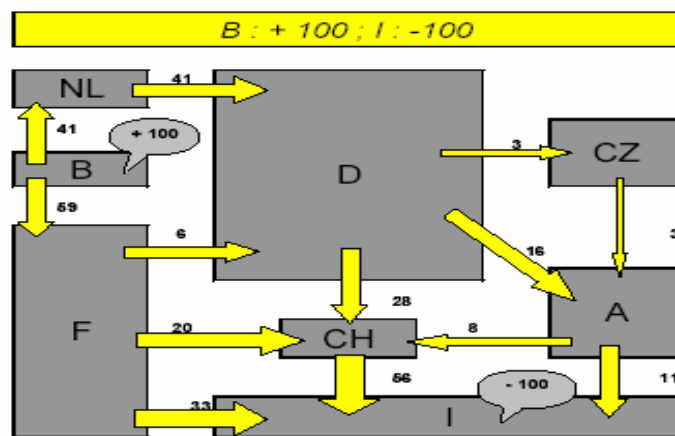


Figure I.9 : impact² d'une transaction de 100 MW entre la Belgique et l'Italie sur un ensemble des réseaux interconnectés [ETSO 00]

Cet exemple montre clairement le danger d'estimer les flux sur des hypothèses non électrotechniques, qui conduiraient à avoir une image faussée de la répartition des transits de puissance, ce qui peut mener à des situations de congestion inattendues. Aux Etats-Unis, au tout début de la dérégulation des marchés de l'électricité, le concept du chemin contractuel était fréquemment utilisé par les utilisateurs souhaitant programmer des transactions interrégionales. Les valeurs des capacités disponibles étant postées publiquement sur un site Web (appelé OASIS [CHR 00]), ils pouvaient réserver les capacités dont ils avaient besoin suivant le cheminement conventionnel de leur transaction. Cette approche a connu très rapidement des problèmes techniques après son implémentation. Des congestions de plus en plus fréquentes et imprévues ont rapidement affecté l'ensemble des régions interconnectées. Cela a obligé les opérateurs du système à annuler à la dernière minute certaines transactions

² Cette analyse se fonde sur un modèle linéarisé de calcul de répartition de charge qui est une approximation acceptable du mode de fonctionnement réel. Ce modèle est le modèle de calcul de répartition de charge dit « à courant continu » (voir Ch.2)

qui avaient été réservées et à prendre des mesures d'urgence en temps réel, ce qui démontrait que la sécurité du réseau n'était pas réellement assurée.

Des problèmes du même type existent en Europe lorsqu'il s'agit de calculer les NTC (*Net Transfer Capacity*) entre chaque zone [ETS0 01]. Une NTC d'une zone A vers une zone B se définit comme la valeur de la puissance totale transmissible de la zone A vers la zone B, moins une certaine marge de sécurité. Elle est calculée en déplaçant de façon itérative une quantité de production de la zone B vers la zone A (Fig. I.10). L'algorithme s'arrête lorsque la première contrainte est rencontrée. Les NTC sont utilisées pour permettre aux participants du marché à programmer leurs transactions internationales.

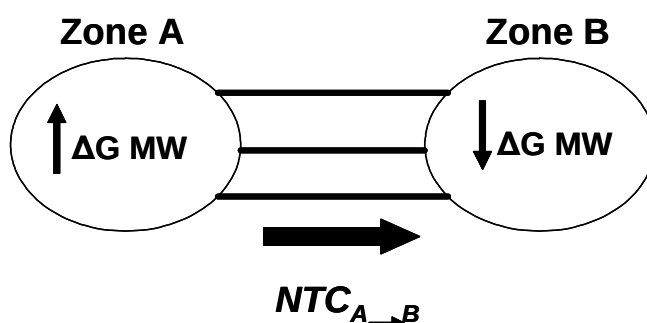


Figure I.10 : Calcul de la NTC de la zone A vers la zone B

La première difficulté de l'approche décrite ci-dessus est que le choix des productions à diminuer et à augmenter peut influencer de façon conséquente la valeur de la NTC trouvée. Une deuxième réside dans l'architecture complexe des grands réseaux interconnectés, qui peuvent comporter bien plus que deux zones. Pour pouvoir calculer de façon assez précise une NTC d'une zone A vers une zone B, on devrait faire des hypothèses valides sur les transactions se déployant sur l'ensemble du système interconnecté. Par exemple, si nous considérons l'exemple de la Figure I.9, un calcul de NTC de l'Allemagne (D) vers l'Autriche (A) devrait tenir compte du flux parallèle engendré par la transaction allant de la Belgique à l'Italie pour être suffisamment précis. Or, face à la taille du système interconnecté et à l'imprévisibilité de l'évolution du marché de l'énergie, il est très difficile d'avoir une représentation précise de l'état de charge du réseau longtemps à l'avance. Les flux parallèles rendent en outre les valeurs de NTC interdépendantes, ce qui rend la publication de scénarios réalistes pratiquement impossible, le nombre de combinaisons à analyser étant trop grand.

Ainsi, comme le reconnaît l'ETSO³ (*European Transmission System Operators*), les valeurs de NTC publiées sont très indicatives et ne sont pas fiables si le scénario réel dévie du scénario hypothétique envisagé. Les NTC ne peuvent donc pas prévenir efficacement l'apparition de congestions sur le système.

Ainsi, la forte présence de flux parallèles dans les réseaux maillés rend le concept de chemin contractuel peu réaliste, et rend très difficile le calcul de capacités d'échanges entre pays. Chaque opérateur du système doit donc procéder à une vérification des plans de production/consommation durant la préparation à la conduite en vue de détecter d'éventuelles congestions. Si des congestions sont détectées, alors des mesures de traitement efficaces devront être prises.

I.3) Le cas français : solution du RTE [RTE1 02]

Jusqu'à présent, il y a relativement peu de cas de congestion observés à l'intérieur du territoire français, à l'exception de la région niçoise qui peut souffrir de difficultés d'approvisionnement en cas d'indisponibilité de l'unique ligne 400 kV chargée de l'alimenter.

Les problèmes de congestion se situent plutôt au niveau des interconnexions frontalières qui sont souvent saturées. Comme il existe une demande de plus en plus forte de la part d'acteurs de marché pour pouvoir accéder à ces interconnexions, le gestionnaire du réseau de transport RTE a récemment proposé une solution pour gérer les problèmes de congestion aux frontières.

Habituellement, il alloue les capacités disponibles sur les interconnexions frontalières en se basant sur un calcul prévisionnel de NTC. Si toute la capacité disponible a été réservée, RTE peut proposer des capacités supplémentaires en cherchant à soulager la contrainte qui limite les échanges transfrontaliers. Ainsi, deux jours avant la conduite effective (J-2) RTE essaye de déterminer le réaménagement des programmes de production le moins coûteux possible permettant de soulager la contrainte. Les capacités supplémentaires dégagées sur les interconnexions influençant la contrainte sont alors calculées. RTE peut alors proposer aux clients intéressés les capacités supplémentaires dégagées, ceux-ci se partageant le surcoût de congestion associé.

Le surcoût de congestion obtenu en J-2 est d'abord réparti équitablement entre toutes les interconnexions considérées. Il est ensuite divisé par 24h et par capacité supplémentaire dégagée pour obtenir un coût en €/MWh.

³ L'ETSO est une association à but non lucratif réunissant les différents Opérateurs de système européens. Elle a pour but de favoriser un développement du marché de l'énergie en accord avec la sécurité des réseaux de transport

Finalement, pour allouer le coût de congestion aux transactions commerciales, RTE empile les transactions prévues sur chaque interconnexion par ordre de priorité jusqu'à obtenir la NTC de cette interconnexion. Ces transactions ne sont pas soumises à un coût de congestion. Au-delà de cette valeur, tous les contrats empilés voulant bénéficier de la capacité supplémentaire dégagée doivent se partager au prorata du volume transité le coût de congestion associé à l'interconnexion.

Ce calcul se faisant pour chaque pas horaire du jour J, chaque transaction bénéficiant des capacités supplémentaires dégagées se voit attribuée un coût de congestion global pour la journée J.

Un exemple fourni par RTE permet d'illustrer la méthodologie proposée [RTE1 02]. Cet exemple est celui de la Figure I.11 :

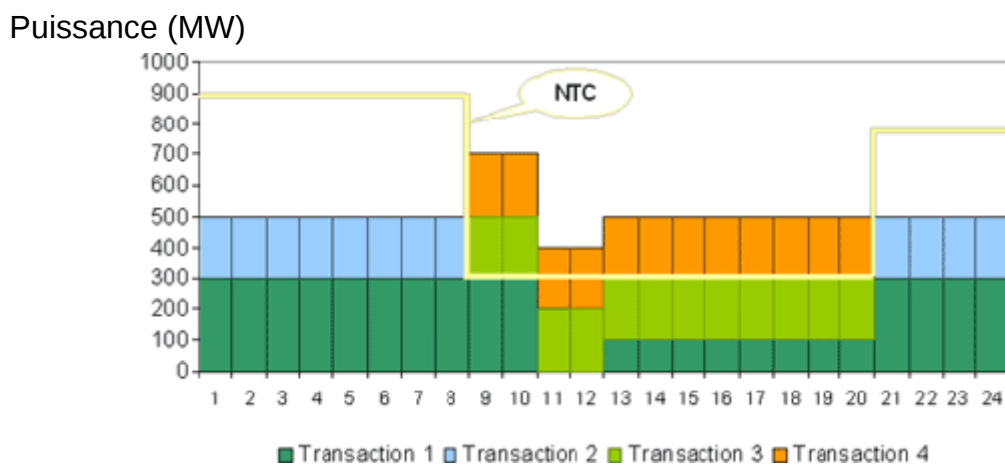


Figure I.11 : exemple illustratif de la méthode proposée par RTE

Les transactions 1 et 2 ne se voient pas affecter de coût de congestion puisque leur puissance cumulée ne dépasse jamais la NTC sur la journée. La transaction 3 est affectée d'un coût de congestion sur 2 points horaires. La transaction 4 est affectée d'un coût de congestion sur la période 8h-20h.

Par ailleurs, si le coût de congestion pour la neuvième heure est de 2 €/MWh, les transactions 3 et 4 ont chacune un coût de congestion de 400 €/h, ayant toutes deux un volume de 200 MW. De même, si le coût de congestion de la quinzième heure est de 1.5 €/MWh, la transaction 4 supporte seule le coût de congestion facturé pour cette heure-ci, qui est de 300 €/h.

La solution mise en place en place par RTE a l'avantage d'être relativement simple et transparente. Toutefois, elle a l'inconvénient de favoriser les premiers arrivés au détriment des derniers : c'est le principe du « premier arrivé, premier servi » qui peut être contesté par certains acteurs du marché, et qui ne peut être justifié sur aucune base objective. En outre, le partage au prorata du coût de congestion entre plusieurs transactions, bien que simple, ne donne pas nécessairement les bons signaux économiques car il ne reflète pas correctement la contribution physique réelle de chaque transaction sur la congestion. Nous développerons ce point de façon plus détaillée dans la suite de cette Thèse.

I.4) Conclusion préliminaire sur le phénomène de congestion

Dans ce chapitre, nous avons défini le phénomène de congestion et dégagé ses principales causes d'occurrence qui sont :

- Des limitations à imposer de façon nécessaire aux lignes et interconnexions du réseau de transport qui sont fonction du niveau de tension et de la longueur de l'ouvrage.
- Des réseaux bien adaptés à l'ancien modèle monopoliste, mais qui n'ont pas encore eu le temps d'évoluer dans un contexte libéralisé, et qui sont exploités toujours au plus près de leurs limites. La libéralisation du secteur de l'électricité, en entraînant une plus grande volatilité du marché de l'énergie, bouleverse la répartition des transits et provoque de plus en plus de flux parallèles non prévus à l'origine, et dont la cause peut devenir difficile à cerner.

De fait, le **traitement des congestions** est le procédé grâce auquel on s'assure que le système est conduit dans le respect des limites de transits imposées.

Nous devons d'ores et déjà faire la distinction entre le phénomène de congestion et ce qu'on appelle communément « surcharges » sur le réseau:

- Une surcharge se produit lorsque le transit en temps réel d'une ligne met en péril la sécurité de cette ligne. Ces surcharges surviennent de façon inattendue lors de la conduite du réseau en temps réel et sont dues principalement à des aléas survenant sur le réseau (perte d'une ligne, d'une centrale, ...). Elles ne sont pas spécialement liées à la restructuration du secteur de l'électricité et sont résolues par des mesures d'urgence classiques (protections, reconfiguration du réseau, etc...) sans optimisation économique.

- Le phénomène de congestion se rapporte plus à une incapacité du réseau à conduire tels quels les plans établis par le marché de l'énergie. C'est un problème technique, mais dont la cause est économique. La pression du marché génère des contraintes chroniques sur le système qui doivent être traitées dans un cadre prévisionnel. Les méthodes de traitement de congestions cherchent pour la plupart à optimiser les coûts induits par ces contraintes

D'autre part, il existe une règle importante qui stipule que le réseau doit être sécurisé non seulement en présence de tous ses éléments, mais aussi si un élément vient à être perdu (perte d'un générateur, perte d'un transformateur, perte d'une ligne, etc...). Cette règle est la règle dite du « N-1 ». Si donc aucune contrainte n'est violée avec un réseau complet, mais que la perte d'un de ses éléments menait à la violation d'une contrainte, on peut considérer que le réseau est congestionné [RTE2 02]. Toutefois, les mêmes principes peuvent être appliquées pour le traitement des congestions en « N » que pour le traitement des congestions en « N-1 ». Ainsi, nous n'explicitons pas systématiquement la contrainte du N-1, tout en sachant qu'elle doit être prise en compte sur un cas pratique⁴.

Dans le contexte actuel de libéralisation du marché de l'énergie, on peut s'attendre à un phénomène de congestion de plus en plus récurrent sur les réseaux de transport d'électricité. Le traitement des congestions devient alors un sujet d'étude de grande importance, dont la sécurité du réseau est l'enjeu principal. Nous avons récemment vu les conséquences sociales et économiques que pouvaient avoir un réseau insécurisé, lors du blackout nord américain d'août 2003 ou du blackout italien survenu quelque temps après. Bien que ces incidents de grande ampleur ne soient pas dus directement à des problèmes de congestions, nous pouvons néanmoins penser que des congestions non traitées peuvent avoir des conséquences tout aussi sévères.

⁴ le blackout italien du 28 septembre 2003 a été causé par la perte d'une interconnexion reliant la Suisse à l'Italie à un moment où cette règle du N-1 n'était pas assurée ; la conséquence a été un effet de « cascade » qui a conduit à les pertes de l'ensemble des interconnexions alimentant l'Italie

CHAPITRE II

Revue des principales méthodes de traitement des congestions

II.1) Introduction

Dans ce chapitre, nous allons décrire les principales méthodes de traitement des congestions qui ont vu le jour depuis l'avènement de la dérégulation. Nous ne discuterons pas ici des techniques traditionnelles utilisées en temps réel pour éliminer les surcharges sur le réseau, telles que les reconfigurations de réseau ou délestages de charge. Ces méthodologies ont en effet été largement utilisées sous les monopoles régulés, et n'ont pas de lien direct avec la dérégulation. De même, les solutions consistant au renforcement et au développement du réseau, plutôt qu'à la gestion des contraintes, sont des solutions de long terme qui ne seront pas évoquées pour l'instant.

Toute méthode de traitement des congestions inclue obligatoirement un modèle de calcul de répartition de charge. Dans ce chapitre, les méthodes de traitement des congestions présentées seront décrites à l'aide du modèle de calcul de répartition de charge dit « à courant continu » ou modèle DC. Ce modèle est décrit en détail dans l'Annexe 1, où les raisons du choix de ce modèle sont aussi explicitées.

II.2) Modèles de traitements des congestions appliqués actuellement

Depuis le début de la dérégulation, plusieurs solutions ont été proposées et appliquées pour gérer les contraintes de transit dans les réseaux de transport. Certaines de ces solutions sont plutôt à caractère général et peuvent théoriquement être appliquées sur la plupart des marchés libéralisés (coupures de transactions, outils d'optimisation de la production), d'autres peuvent être plus spécifiques à un modèle de marché, voire à une certaine configuration de réseau (modèle californien, régionalisation du marché scandinave). Dans ce chapitre, nous allons présenter les méthodes de traitement des congestions les plus connues. Il servira de base pour la discussion sur le choix d'un modèle de traitement des congestions (Chapitre III).

II.2.1) La régionalisation du marché (ou *Market Splitting*) nordique [GLA 02], [ETSO 01]

Le réseau scandinave est de configuration radiale, avec une répartition des échanges d'énergie plutôt hétérogène. En effet, les principaux centres de production sont situés au nord, tandis que la consommation est plutôt massée au sud, ce qui crée des transits importants du nord vers le sud. Parmi les pays nordiques, La Norvège a mis en place une solution originale et très spécifique pour résoudre les congestions, qui fait intervenir la bourse de l'énergie comme agent de régulation des flux. Si il n'y a pas congestion, le marché spot se clôture sur un prix de marché donné (le *Market Clearing Price*-MCP) qui est le même sur toute la zone donnée (voir Ch.I §I.1.1). Par contre s'il y a congestion, la bourse de l'énergie est séparée (*Market Splitting*) en zones délimitées par les goulots d'étranglements et affichant des prix de marché différents.

Prenons un exemple sur deux zones A et B reliées par une interconnexion AB (fig.II.5), dont le transit s'écoule de la zone A vers la zone B :

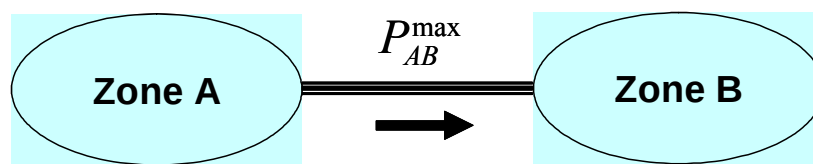


Figure II.5 : exemple à deux zones

L'interconnexion AB est limitée à $P_{AB}^{\max}$. Dans le cas sans congestion (cas non contraint), les courbes de production et de consommation sont agrégées sur l'ensemble des deux zones et à la clôture du marché nous avons :

$$P_{GA} + P_{GB} = P_{CA} + P_{CB}$$

avec P_{GA} , P_{GB} production totale en zone A et B

P_{CA} , P_{CB} consommation totale en zone A et B

Le prix de clôture est dans ce cas-là l'*Unconstrained Market Clearing Price* (UMCP, prix du marché spot sans congestion). Le flux sortant de la zone A (et donc le flux entrant dans la zone B) est égal au déséquilibre production-consommation de la zone A (respectivement de la zone B) :

$$P_{AB} = P_{GA} - P_{CA} = P_{CB} - P_{GB}$$

Ce flux apparaît aussi comme la distance entre les courbes de production et de consommation de chaque zone (fig.II.6). Si ce flux dépasse la capacité de l'interconnexion, le marché unique est séparé suivant les deux zones, et le prix dans chaque zone est réglé de manière à ce que le flux sortant de la zone A (et donc le flux entrant dans la zone B) soit égal à la valeur maximale permise :

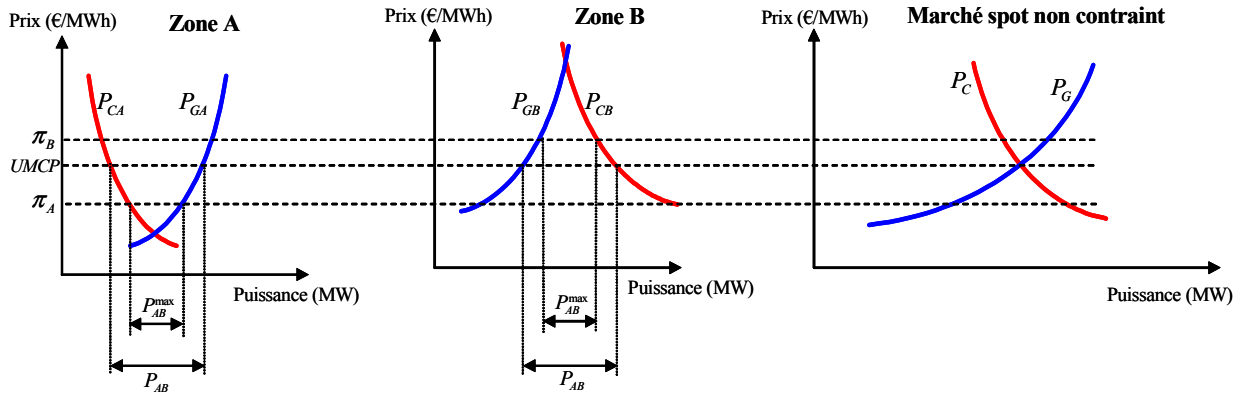


Figure II.6 : marché spot non contraint et régionalisation du marché en cas de congestion sur l'interconnexion AB [CHR 00]

Ce processus de régionalisation du marché revient à baisser la production et à augmenter la consommation⁵ dans les zones fortement exportatrice (cas de la zone A) et à faire l'inverse dans les zones fortement importatrice (cas de la zone B). Cela a pour effet de rééquilibrer la répartition de la production et de la consommation sur les deux zones. L'apparition d'une différence de prix entre la zone A et la zone B définit un prix de congestion égal à $\pi_B - \pi_A$. En ce qui concerne les transactions bilatérales entre les deux zones, bien que n'étant pas censées participer à la régionalisation du marché, elles doivent payer un coût de congestion pour tenir compte de leur participation à la congestion. On définit alors pour chaque zone une charge financière appelée *capacity fee* (CF) comme la différence de prix entre cette zone et le prix du marché sans congestion. Par exemple pour la zone A qui est la zone excédentaire, la *capacity fee* est :

$$CF_A = UMCP - \pi_A$$

et pour la zone B qui est la zone déficitaire :

$$CF_B = \pi_B - UMCP$$

⁵ si on considère la consommation élastique par rapport au prix du marché

Donc, si nous avons une transaction bilatérale de T MW entre un producteur de la zone A et un consommateur de la zone B, le producteur en question devra payer un surcoût égal à $CF_A * T$ et le consommateur un surcoût égal à $CF_B * T$. Cependant, si nous avons une transaction allant en sens inverse du flux congestionné, le producteur en zone B et le consommateur en zone A recevront au contraire une compensation, basée aussi sur la *capacity fee*. Cette propriété du modèle norvégien permet d'envoyer des signaux économiques sur le long terme pour une localisation optimale des futurs centres de production et consommation par rapport aux contraintes du réseau.

Cependant, bien qu'étant efficace sur le réseau norvégien, la régionalisation du marché souffre d'un manque d'adaptation sur les réseaux plus maillés. En effet, dans les réseaux maillés, il est très difficile de définir des zones de prix à cause des flux parallèles. Si ceux-ci deviennent importants et que le réseau est fortement maillé, traiter les congestions par régionalisation du marché devient quasi insoluble. Or, la plupart des réseaux de transport existants sont maillés voire fortement maillés, les configurations exclusivement radiales faisant plutôt figure d'exception.

Ainsi, la Norvège a mis en place une procédure de gestion des congestions efficace dans les conditions d'exploitation de son réseau, mais de portée limitée.

II.2.2) La solution californienne [CAI 99], [CAI 03]

Le système californien est caractérisé par la présence d'une bourse centrale de l'énergie (le CalPX) et d'un certain nombre d'acteurs de marché très spécifiques appelés *Scheduling Coordinators* (SC). Ces SC sont des coordinateurs possédant un portefeuille de plusieurs participants, producteurs et consommateurs, et ont avec le CalPX un rôle important dans la gestion des contraintes du système. Le traitement des congestions en Californie se fait sur une base prévisionnelle en J-1 ou H-2. Une des particularités du système californien est qu'il distingue un traitement des congestions *interzonales* et un traitement des congestions *intrazonales*. En Californie, les congestions les plus importantes surviennent en effet souvent au niveau des interconnexions qui relient les zones densément maillées. Les congestions intrazonales sont peu nombreuses et demandent des ajustements relativement mineurs comparés aux congestions interzonales. Pour le traitement des congestions interzonales, le CalPX et les différents SC transmettent à l'opérateur du système les plans de production-consommation associés à leur portefeuille de participants, ainsi que des offres d'ajustements. Si l'opérateur du système, en vérifiant la faisabilité de ces plans sur son système détecte une congestion, il va faire appel aux offres d'ajustement pour reprogrammer les plans des SC et de

CalPX, tout en maintenant le bilan de leur portefeuille constant. Ainsi, une autre particularité que l'on peut distinguer est un traitement des congestions en deux étapes successives:

- d'abord, l'établissement des plans de production-consommation des SC et de la bourse sans tenir compte des contraintes du système. Ces plans sont qualifiés de *Plans Initiaux Souhaités* (PIS)
- ensuite, la reconfiguration de ces plans par l'opérateur du système à l'aide des offres d'ajustement fournis par les SC et par la bourse CalPX, tout en maintenant leur bilan global constant pour chaque SC et pour la bourse [GRI 98].

Pour illustrer le fonctionnement du système californien, considérons l'exemple de la figure II.7 :

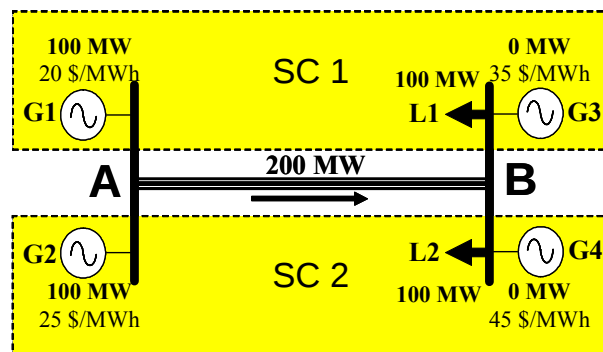


Figure II.7 : PIS des traders SC1 et SC2

Nous avons deux zones A et B couvertes par deux SC, appelé SC1 et SC2. Ces deux SC doivent fournir 100 MW de demande situés en zone B et font appel pour cela à leur ressources les moins chères situées en zone A. Le flux résultant sur l'interconnexion AB est donc de 200 MW. Toutefois, si le transit maximum de l'interconnexion est fixé à 150 MW, l'opérateur du système doit reconfigurer les PIS proposés. Pour cela, il fait donc appel aux offres d'ajustement de SC1 et SC2 :

- SC1 dispose d'une ressource à 35 \$/MWh en zone B. SC1 fait donc une offre pour l'usage de la capacité AB égal à 15 \$/MWh, différence de prix entre sa ressource en A et sa ressource en B
- SC2 dispose d'une ressource à 45 \$/MWh en zone B. SC2 fait donc une offre pour l'usage de la capacité AB égal à 20 \$/MWh

Comme le but de l'opérateur du système est de reconfigurer les PIS tout en minimisant le coût de congestion, il va donc attribuer 100 MW de capacité de transfert à SC2, et les 50 MW restant à SC1. Le prix de congestion est donc de 15 \$/MWh (fig. II.8)

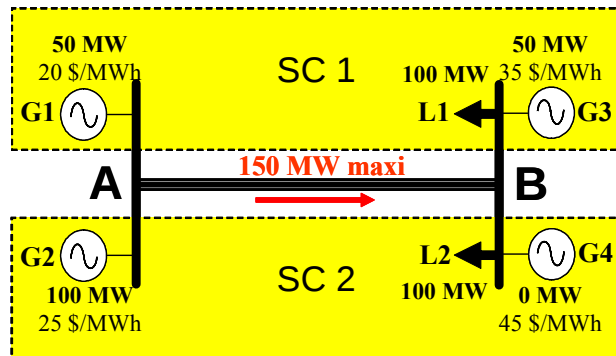


Figure II.8 : reconfiguration des PIS

La figure II.9 nous montre les différents flux financiers générés par le traitement des congestions. On peut y distinguer les paiements attribués aux producteurs après reconfiguration des PIS et le coût de congestion payé par SC1 et SC2 à l'opérateur du système, qui est finalement reversé aux propriétaires du réseau :

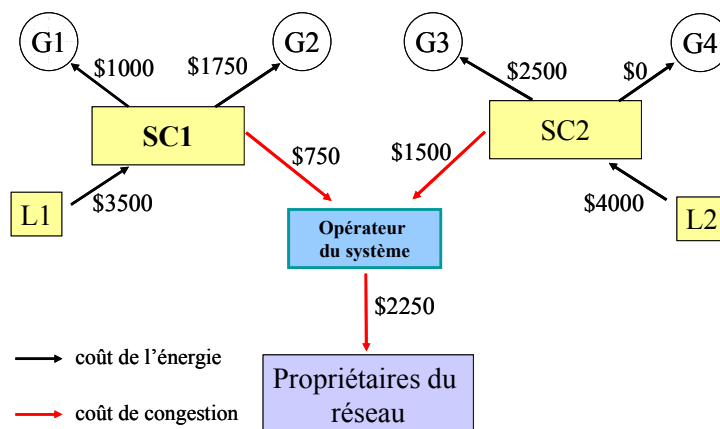


Figure II.9 : flux financiers créés par le traitement des congestions dans le modèle californien

A partir de cet exemple, nous pouvons faire plusieurs remarques :

- La présence de contraintes sur le système assure un revenu aux propriétaires du réseau. Ce revenu est d'autant plus important que le système est contraint. Ce mode

de fonctionnement risque de ne pas réellement inciter les propriétaires du réseau à investir pour renforcer le réseau.

- Par rapport aux PIS (fig. II.7), les charges L1 et L2 vont payer chacune 1500 \$/h de plus à cause de la congestion.

La congestion a donc créé pour les consommateurs un coût supplémentaire de 3000 \$/h, dont 750\$/h généré par la reconfiguration de la production et 2250 \$/h sous forme de revenus aux propriétaires du réseau.

On peut cependant noter qu'il y avait un moyen plus simple, plus efficace et moins coûteux de gérer la congestion dans ce cas-là. En effet, si on ne cherchait pas à garder le bilan électrique des SC constant, il aurait suffi de baisser de 50 MW le producteur G2 (appartenant au SC2) et d'augmenter de 50 MW le producteur G3 (appartenant au SC1). De ce fait, le coût de congestion obtenu est de :

$$50 * 35 - 50 * 25 = 500 \$/h$$

Cette option assure que le coût de congestion obtenu est celui strictement nécessaire à résoudre la congestion, sans fournir de revenus pour les propriétaires du système. Cette reconfiguration de la production est économiquement plus efficace que la précédente, dont le montant s'élevait à 750 \$/h.

Le système californien, en gardant le bilan des SC nul permet d'assurer une certaine continuité entre le marché de l'énergie et le traitement des congestions et donner un signal de prix pour l'usage d'une interface congestionnée. Toutefois, comme nous l'avons montré, cette solution peut être plus contraignante dans certains cas [KIR 98].

II.2.3) les coupures de transactions : la solution du NERC américain [CHR 00], [NER 03]

Pour pallier les défauts du concept du chemin contractuel discuté au Chapitre I, le NERC⁶ a mis en place une procédure de traitement des congestions par coupure de transactions. Cette procédure utilise un outil d'analyse basé sur le modèle DC (voir Annexe 1) appelé *Interchange Distribution Calculator* (IDC). L'IDC intègre toutes les données des grands réseaux américains et des transactions commerciales programmées et calcule l'impact de chaque transaction sur les transits à l'aide des facteurs de distribution. L'impact d'une

⁶North America Electricity Reliability Council. C'est une organisation chargée de veiller à l'application des règles de sécurité dans la conduite des réseaux de transport

transaction de volume T d'un nœud i à un nœud j sur le transit d'une ligne l se calcule de cette façon :

$$P_l^{ij} = \alpha_l^{ij} * T_{ij} = [A(l,i) - A(l,j)] * T_{ij} \quad (\text{II.11})$$

avec α_l^{ij} que nous pouvons définir comme le Facteur de Distribution Bilatéral (FDB) de la transaction T_{ij} sur la ligne l (*Power Transfer Distribution Factors-PTDF* en appellation anglophone). Cet impact est interprété comme la part du transit dans une ligne l qui est attribuable à une transaction d'un nœud i vers un nœud j . $A(l,i)$ et $A(l,j)$ sont des éléments de la matrice $\mathbf{A}$. Cette matrice est la matrice des facteurs de distribution reliant les injections nodales aux transits (pour plus de détails sur le calcul de cette matrice, voir l'Annexe 1).

Lorsqu'il y a danger potentiel de surcharge sur une ligne du réseau ou que cette ligne soit déjà chargée au-delà de sa capacité maximale, l'opérateur du système peut lancer une procédure d'urgence appelée *Transmission Loading Relief* (TLR). Cette procédure possède plusieurs niveaux suivant la gravité de la situation et son caractère d'urgence. Elle consiste à élaguer des transactions ayant une contribution physique allant dans le même sens que le transit de la ligne congestionnée. Le choix des transactions à élaguer se fait aussi suivant un ordre de priorité, les transactions qualifiées de *non fermes* étant moins prioritaires que celles qualifiées de *fermes*. Ainsi, le TLR coupe d'abord l'ensemble des transactions non fermes contribuant au transit d'une ligne congestionnée (fig. II.10). Ensuite, si le problème n'est pas résolu en coupant les transactions non fermes, le TLR accède à des niveaux supérieurs où il a recours à des reconfiguration du réseau de transport et enfin aux coupures des transactions dites fermes pour les situations les plus critiques. Le TLR est donc une procédure purement technique, sans aucune optimisation économique.

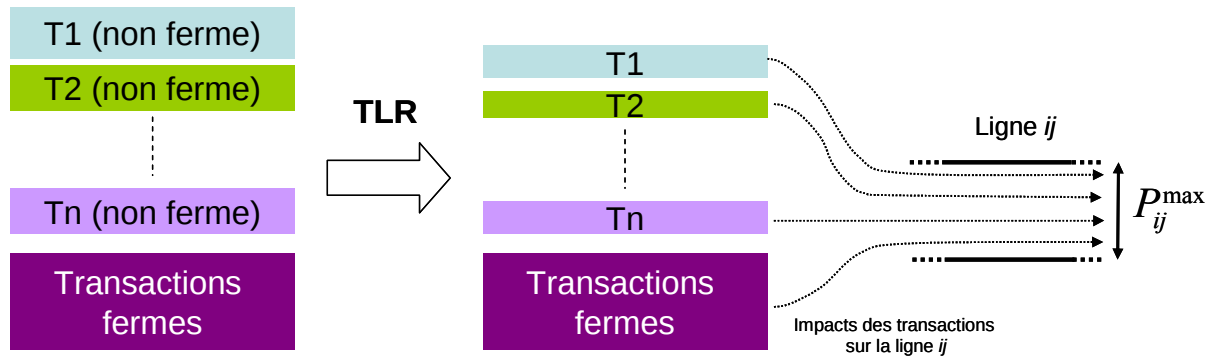


Figure II.10 : coupures de transactions non fermes suivant leur impact sur le transit d'une ligne congestionnée

Lorsque l'on a Nt transactions de même priorité sur un transit congestionné l , le TLR utilise une formule mathématique basée sur les FDB pour déterminer les quantités de transactions coupées. Ainsi, le volume coupé pour une transaction allant d'un point i à un point j et de volume T est déterminé de la façon suivante [NER 03] :

$$\Delta T_{ij} = \frac{FDB_{ij}^l * T_{ij} * \frac{FDB_{ij}^l}{\sum_{Nt} FDB_{ij}^l}}{\sum_{Nt} FDB_{ij}^l * T_{ij} * \frac{FDB_{ij}^l}{\sum_{Nt} FDB_{ij}^l}} * \frac{\Delta P^l}{FDB_{ij}^l} = \frac{FDB_{ij}^l * T_{ij}}{\sum_{Nt} FDB_{ij}^l * T_{ij} * FDB_{ij}^l} * \Delta P^l \quad (II.12)$$

avec

FDB_{ij}^l Facteur de Distribution Bilatéral représentant l'impact d'une transaction entre le nœud i et j sur une ligne l

T_{ij} Volume de la transaction effectuée entre le nœud i au nœud j

ΔP^l surcharge sur le transit de la ligne l

Nt nombre de transactions à couper

Le concept de coupure de transaction implique que la réduction touche aussi le producteur que son client, et dans les mêmes quantités. Ainsi, si un contrat bilatéral a dû être réduit de 20 MW à cause d'une congestion, cela signifie que le producteur doit réduire de 20 MW sa production et que le consommateur devrait théoriquement réduire sa consommation de la même quantité ; toutefois en pratique, pour assurer la continuité de fourniture et de service au client, on peut lui proposer soit de compléter les 20 MW manquant de fourniture auprès d'un autre fournisseur ou à la bourse (traitement prévisionnel), soit de recourir à des redispatchings qui pourront lui être facturés (traitement en temps réel).

Etant basé sur la gestion des transactions, une procédure semblable au TLR peut théoriquement être appliquée sur la plupart des marchés bilatéraux, ou contenant une forte composante bilatérale. Cependant, le TLR a dû faire face à certaines critiques ; on lui a reproché en outre de ne pas tenir compte de l'efficacité économique, du fait qu'il donne la priorité de passage à ceux qui ont réservé le plus tôt, sans vraiment donner l'occasion à chaque participant au marché de défendre son droit à l'accès au moment où une contrainte apparaît sur le système. En outre, d'un point de vue plus technique, des problèmes décisionnels peuvent apparaître lorsque l'on est en présence de deux congestions : il peut arriver qu'une transaction donnée ait un fort impact sur une congestion, mais contribue à

modérer la seconde. La couper soulagera donc la première contrainte, mais aura pour effet d'aggraver la seconde, ce qui suppose des coupures supplémentaires.

D'autres stratégies basées sur les coupures de transactions ont aussi été proposées [RAI1 01], [FAN1 99], [FAN2 99].

II.2.4) Un outil généralisé d'optimisation de la production : l'OPF (*Optimal Power Flow*) [CHR 00], [GED 99]

L'OPF est une technique utilisée depuis plus de 35 ans dans le secteur électrique. Son but est de déterminer une répartition de charge optimale du point de vue des coûts de production, tout en respectant des contraintes techniques liées au fonctionnement du réseau. Avant la dérégulation, le coût total de production à minimiser était calculé à partir du coût de chaque unité de production (pour une tranche horaire donnée) qui est fonction de la puissance de sortie de l'unité de production. Les fonctions de coût individuelles de chaque générateur étaient basées sur :

- la caractéristique d'entrées-sorties donnant pour chaque unité l'équivalence thermique de l'énergie électrique produite
- les coûts de combustible

Avec l'avènement de la dérégulation, la même technique a continué d'être appliquée, mais les courbes de coût de chaque unité de production ont été remplacées par des courbes d'*offres/prix* fournis par chaque producteur. Ces courbes d'offres intègrent les coûts fixes et les coûts variables, ainsi qu'une marge laissée au choix du producteur pour son profit personnel. Ces offres spécifient le prix fixé par chaque producteur en fonction d'une certaine quantité de puissance proposée sur le marché.

Les offres des producteurs et les demandes des consommateurs sont confrontées dans le cadre d'un marché centralisé (à l'image d'une bourse de l'électricité) appelé *marché spot*. L'OPF minimise les coûts de production tout en satisfaisant la demande, comme cela se fait dans le modèle pool classique. Toutefois, cette optimisation devra tenir compte du fonctionnement du réseau et des contraintes imposées sur les ouvrages de transport. Ces contraintes peuvent alors entraîner une différenciation géographique du prix de l'énergie. Le marché spot peut alors être vu comme une bourse de l'électricité, mais dont le prix peut varier géographiquement. L'OPF traite le marché spot et les contraintes techniques du système au sein d'un même processus d'optimisation ; ainsi, son usage implique l'existence d'un

opérateur du système mixte qui a à la fois la responsabilité de la gestion du système et celle du marché spot (à l'exemple de l'opérateur américain PJM).

L'OPF peut être formulé de cette manière à l'aide du modèle DC (Annexe1) pour un réseau à N nœuds :

Fonction Objectif de l'optimisation : Minimiser les coûts⁷ totaux de production dans le cadre du marché spot [PJM 00]

$$\text{Min } \sum_i C_i(P_{Gi}) \quad (\text{II.13})$$

Avec

$C_i(P_{Gi})$ offre du producteur donnant le prix proposé en fonction d'une quantité P_{Gi} offerte sur le marché spot.

P_{Gi} production de puissance au nœud i

Contraintes à respecter :

Equations du calcul de répartition de charge (modèle DC)

$$\mathbf{P} = \mathbf{B} * \boldsymbol{\theta} \quad (\text{II.14a})$$

$$\mathbf{P} = \mathbf{P}_G - \mathbf{P}_C \quad (\text{II.14b})$$

$$\mathbf{P}_{ij} = \mathbf{H} * \boldsymbol{\theta} \quad (\text{II.14c})$$

avec

$\mathbf{P}$ vecteur des injections de puissance nettes nodales de dimension N

$\mathbf{B}$ matrice des admittances nodales du réseau de dimension $N*N$

$\boldsymbol{\theta}$ vecteur des phases des tensions nodales de dimension N

$\mathbf{P}_G, \mathbf{P}_C$ vecteur des productions nodales et des charges nodales de dimension N

$\mathbf{P}_{ij}$ vecteur des transits (puissance active) de dimension Nb

⁷ Pour plus de commodité, nous parlerons souvent de « coûts » même en faisant référence aux offres des producteurs, tout en ayant à l'esprit que ces offres renferment les coûts de production proprement dit, ajoutés d'une marge choisie par le producteur

H matrice de dimension Nb*N reliant les phases des tensions nodales du réseau aux transits de puissance

Equilibre production-consommation

$$\sum_i P_{Gi} - \sum_j P_{Ci} = 0 \quad (\text{II.15})$$

Limites imposées aux lignes

$$-P_{ij}^{\max} \leq P_{ij} \leq P_{ij}^{\max} \quad (\text{II.16})$$

Limites imposées aux productions

$$P_G^{\min} \leq P_G \leq P_G^{\max} \quad (\text{II.17})$$

avec

$P_G^{\min}$, $P_G^{\max}$ vecteurs donnant la quantité de production minimum et maximum offerte par chaque producteur

$P_{ij}^{\max}$ vecteur donnant les limites maximales de transit imposées à chaque ligne

Les offres des producteurs doivent être linéaires, et sont mises le plus souvent sous une forme quadratique :

$$C_i(P_{Gi}) = a_i + b_i P_{Gi} + c_i P_{Gi}^2 \quad (\text{II.18})$$

avec a_i , b_i , c_i des constantes fixées par le producteur i

Dans le modèle explicité ci-dessus, nous avons considéré que la demande en chaque nœud reste constante. Cependant, des modèles plus complets peuvent tenir compte de l'élasticité de la demande. Dans ce cas, l'élasticité de la demande est modélisée sous forme de fonctions linéaires qui sont rajoutées à la fonction objectif décrite par l'expression II.13. D'autre part, ce modèle peut implicitement prendre en compte la contrainte du N-1.

La méthode n'implique pas nécessairement la modification des volumes vendus par transactions bilatérales, bien que l'on doive tenir compte de leur présence dans le modèle. Toutefois, si l'opérateur n'est pas parvenu à trouver une solution faisable par le biais de

l'OPF, il doit faire appel alors à la procédure du TLR, les transactions dites fermes ayant la préséance sur les non fermes (voir paragraphe II.2.3).

L'usage de l'OPF dans un contexte dérégulé s'accompagne d'une tarification dite *marginal* ou *nodale* [RUD 95]. Cette méthode est notamment appliquée aux Etats-Unis, mais aussi en Argentine, au Chili et en Nouvelle-Zélande [HUA 03]. Il s'agit de fixer en chaque nœud du réseau un prix auquel sera vendu ou acheté l'énergie. Les paiements sont en outre centralisés par l'opérateur du système. Ces prix nodaux sont tirés des multiplicateurs de Lagrange du problème d'optimisation. En effet, tout problème d'optimisation revient en réalité à minimiser une fonction objectif à laquelle on associe les contraintes à respecter. La fonction résultante se nomme le Lagrangien. Dans le cadre du problème décrit par les relations (II.13 à II.17), le Lagrangien s'écrit [GED 99]:

$$\Gamma = \sum_i C_i(P_{Gi}) - \lambda(\mathbf{B}^* \boldsymbol{\theta} - \mathbf{P}) - \gamma(\sum_{k \in N} P_{Gk} - \sum_{k \in N} P_{Ck}) - \boldsymbol{\mu}(\mathbf{P}_{ij}^{\max} - \mathbf{H}^* \boldsymbol{\theta}) - \mathbf{v}(\mathbf{H}^* \boldsymbol{\theta} - \mathbf{P}_{ij}^{\max}) - \boldsymbol{\alpha}(\mathbf{P}_G^{\max} - \mathbf{P}_G) - \boldsymbol{\beta}(\mathbf{P}_G - \mathbf{P}_G^{\min}) \quad (\text{II.19})$$

avec

λ vecteur des multiplicateurs associés aux équations du calcul de répartition en chaque nœud de dimension N

γ multiplicateur associé à l'équilibre production-consommation (toujours strictement positif)

$\boldsymbol{\mu}, \mathbf{v}$ vecteurs des multiplicateurs associé aux contraintes sur les transits

$\boldsymbol{\alpha}, \boldsymbol{\beta}$ vecteurs des multiplicateurs associés aux contraintes sur les unités de production

Les prix nodaux sont définis comme le coût incrémental induit pour satisfaire une unité de consommation supplémentaire en chaque nœud k [HOG 97]:

$$\rho_{(k)} = \frac{\partial \Gamma}{\partial P_{Ck}} = -\lambda_{(k)} + \gamma \quad (\text{II.20})$$

où ρ est le vecteur des prix nodaux

Dans la théorie des prix nodaux, un participant au marché spot doit acheter ou vendre son énergie au prix du nœud où il est connecté [SCH 88]. Toute transaction bilatérale (ferme ou non ferme) est soumise à un coût de congestion pour tenir compte de son influence sur les contraintes du système. Ainsi, pour une transaction bilatérale dont le point d'injection est un

nœud i et le point de soutirage un nœud j , le coût de congestion est défini comme la différence des prix nodaux entre les nœuds i et j , fois le volume de la transaction :

$$Pay_{T_{ij}} = (\rho_i - \rho_j) * T_{ij} \quad (II.21)$$

Tous les flux financiers créés par la tarification marginale sont centralisés par l'opérateur du système.

Pour illustrer l'usage de l'OPF dans un contexte dérégulé et la tarification marginale qui lui est associée, prenons un exemple de réseau à 3 nœuds (fig.II.11 et II.12) que nous avons aussi introduit en Annexe1 et supposons que l'on ait un marché spot. Dans ce marché, le producteur au nœud 1 (G1) propose 200 MW à 25 €/h et le producteur au nœud 2 (G2) propose 200 MW à 20 €/h. Au nœud 3, il y a une charge de 200 MW (C3). Supposons qu'aucune contrainte de transit ne soit imposée sur les lignes. L'OPF va simplement sélectionner la production la moins chère pour satisfaire la demande au nœud 3 (fig.II.11). En l'absence de contraintes, le prix nodal est identique en chaque nœud et égal au prix de la dernière unité appelée, c'est-à-dire 20 €/h. Ce prix peut être assimilé à l'*UMCP* que nous avons défini lorsque nous avons expliqué le principe du *Market Splitting*.

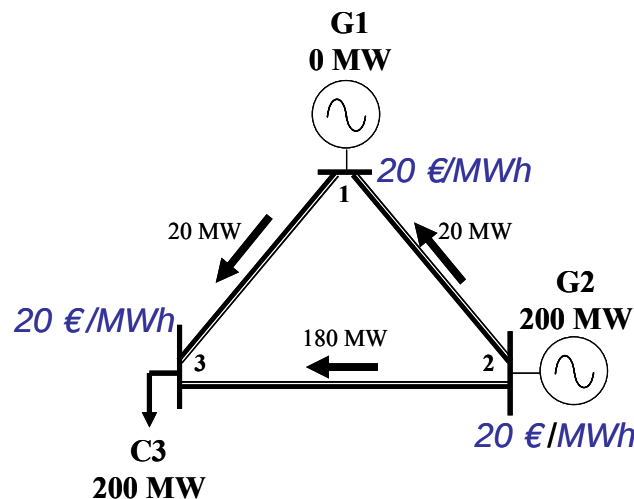


Figure II.11 : établissement des plans de production sans contrainte sur les lignes

Soit à présent un transit maximum de 150 MW imposé sur la ligne connectant le nœud 2 au nœud 3. L'OPF va déterminer une configuration de la production en tenant compte de cette contrainte, tout en minimisant le coût total de la production. La contrainte fait apparaître une différence de prix entre les nœuds (fig.II.12)

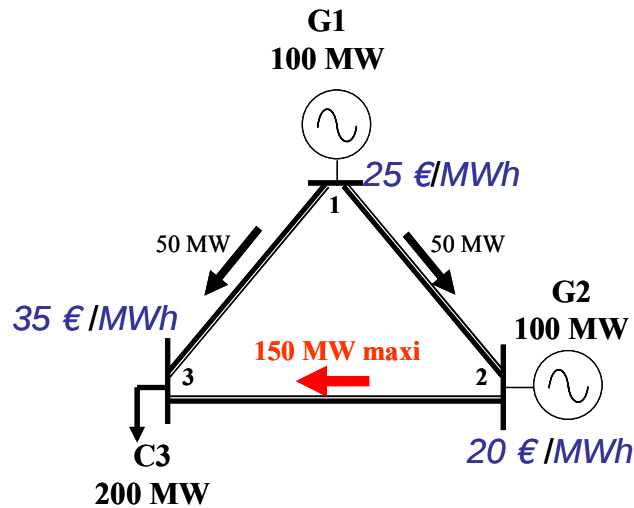


Figure II.12 : établissement des plans de production avec contrainte sur la ligne reliant le nœud 2 au nœud 3

Nous pouvons remarquer qu'il y a à présent différenciation géographique du prix du marché spot. Examinons en outre les sommes collectées du consommateur par l'opérateur du système et reversées aux producteurs : dans le premier cas sans contrainte, l'opérateur du système collecte 4000 \$/h à C3 pour les reverser intégralement à G2. Dans le deuxième cas, il collecte 7000 \$/h et reverse 2500 \$/h à G1 et 2000 \$/h à G2. En tout, il reverse donc 4500 \$/h aux producteurs, mais retient 2500 \$/h. Il y a donc apparition d'un surplus financier résultant de la contrainte imposée. De façon générale, lorsqu'il y a contrainte de transit sur un réseau, les méthodes basées sur la tarification marginale dégagent toujours un surplus financier collecté par l'opérateur du système [RAI2 01]. Ce surplus financier peut être reversé intégralement aux propriétaires du réseau, comme cela se fait par exemple dans l'état de New York aux Etats-Unis [EEI 04]. On peut formuler la même remarque que pour le modèle californien sur le danger de donner l'opportunité aux propriétaires du réseau de faire du profit sur les dysfonctionnements du réseau. Cependant, ce modèle de traitement des congestions a l'avantage de bien s'adapter à tous types de réseaux (radiaux, maillés) et de tenir compte des flux parallèles.

II.2.5) Traitement des congestions par ajustements de production : modèle du *buy back* [CHR 00], [SVE 03]

Le buy back (littéralement « rachat ») consiste pour l'opérateur du système à « racheter » des ajustements en vue de résoudre les congestions. Ces ajustements consistent à retoucher les programmes de production/consommation après clôture du marché de l'énergie.

Le modèle du buy back repose sur trois fondements :

- Il suppose l'existence d'un opérateur du système parfaitement indépendant et distinct des opérateurs du marché
- la séparation du marché de l'énergie et du traitement des congestions : à l'instar du modèle californien, les acteurs du marché établissent d'abord leurs préférences sans tenir compte des contraintes du système, et les éventuelles contraintes sont gérées par la suite. Le traitement des congestions devient un *marché* à part entière
- L'usage d'un outil d'optimisation de la production : à l'instar de l'OPF, ce modèle utilise un outil permettant de reconfigurer la production tout en tenant compte des contraintes du système et des lois physiques de l'électricité

Alors qu'en Norvège, le buy back est une solution secondaire pour traiter des contraintes secondaires, il constitue la principale méthode de traitement des congestions en Suède. Le Royaume-Uni utilise aussi une procédure similaire, où le traitement des congestions est successif à la clôture du marché et où les coûts de réajustement de la production sont intégrés au sein de l'*uplift*, qui est une charge tarifaire s'appliquant au prorata sur les usagers du réseau [MUR 98]. Le traitement des congestions devient alors un *service système* séparé du marché de l'énergie. Le buy back est le modèle de traitement des congestions qui va particulièrement nous intéresser par la suite. Nous allons donc particulièrement étoffer ce modèle, en décrivant d'abord sa démarche générale et en donnant sa formulation mathématique.

II.2.5.1) Démarche générale

Dans le marché de l'énergie, des volumes d'énergie bien déterminés sont négociés et vendus à la bourse, ou par le biais de contrats conclus entre producteurs et clients. Une fois que le marché de l'énergie est clos, les volumes échangés dans ce marché par la bourse ou par contrats bilatéraux sont fixés définitivement, tout comme les paiements qui en découlent.

Après avoir reçu toutes les données sur les volumes négociés sur le marché de l'énergie, l'opérateur du système procède alors aux premières vérifications des plans de production/consommation reçus compte tenu des différentes contraintes du système, parmi lesquelles les limites de transits sur les lignes. Si des congestions sont détectées, l'opérateur du système pourra utiliser les offres d'ajustement déposées par des fournisseurs (volontaires) pour résoudre la ou les contrainte(s) sur les transits. Une fois la congestion résolue, l'opérateur du système pourra enfin valider les plans de production/consommation définitifs pour la tranche horaire donnée du jour J. La figure II.13 nous montre le cheminement général de la méthode du buy back.

Les offres d'ajustements sont déposées par des fournisseurs volontaires et qualifiés pour le service requis. Théoriquement, aussi bien producteurs que consommateurs pourraient

participer à ce service. Toutefois, les consommateurs sont en règle générale peut sensibilisés à ce genre de service, ayant plutôt tendance à demander une quantité fixée d'énergie pour une tranche horaire donnée, avec peu de plage de variation possible. De plus, faire participer les consommateurs aux ajustements peut donner de mauvaises incitations, ce qui peut être délicat à gérer. Ils peuvent en effet être tentés de demander une quantité surévaluée de puissance au marché de l'énergie, pour provoquer volontairement des congestions et ensuite proposer un ajustement à la baisse avec une forte compensation financière. Pour toutes ces raisons, nous allons considérer dans notre modèle que le buy back ne recueille que des offres des producteurs.

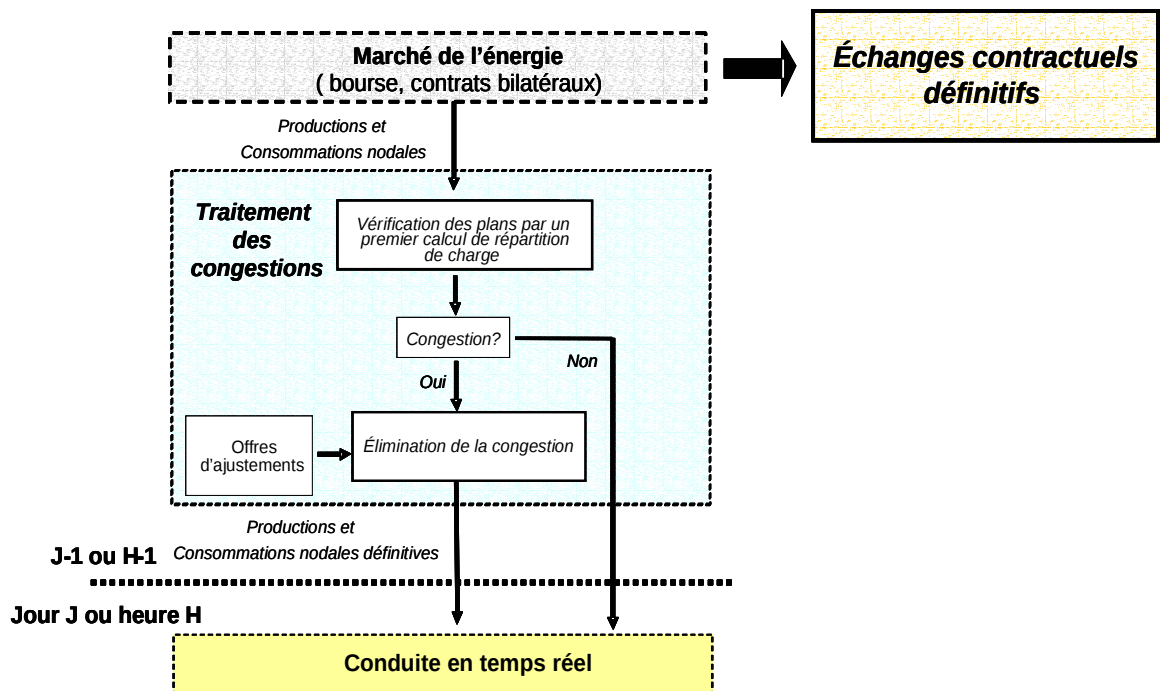


Figure II.13 : procédure générale du modèle du buy back en J-1

A ce stade, précisons un point important : les ajustements décidés serviront à trouver une nouvelle répartition de transits en vue de résoudre les congestions ; ils devront ainsi modifier des injections nodales résultant des plans du marché de l'énergie, et non les contrats décidés ni les plans de la bourse de l'énergie (les paiements du marché énergie restent immuables). Cette propriété du modèle que nous décrivons résulte du découplage qu'il fait entre le traitement des congestions et le marché de l'énergie. Par exemple, considérons un producteur ayant vendu 100 MW à 30 €/MWh, par le biais d'un contrat, et ayant été abaissé de 10 MW en vue de traiter des congestions. Ce producteur injectera ainsi 90 MW en temps réel. Toutefois, il pourra garder inchangé son contrat de 100 MW (il recevra 3000 €/h pour sa production d'énergie). En effet, son ajustement à la baisse a été décidé dans un marché à part

et complètement distinct du marché de l'énergie. Il n'y a donc pas lieu de modifier le contrat ayant été conclu auparavant.

Il peut donc y avoir décalage entre les échanges commerciaux, intouchés par le modèle de traitement des congestions ainsi introduit, et les échanges physiques. On peut noter toutefois qu'il existe, en fait, en permanence un décalage de ce type spécifique dû au fonctionnement des réseaux. En effet, comme nous l'avons déjà dit, l'acheminement de l'électricité se faisant suivant les lois de Kirchhoff, les flux électriques ne suivent pas nécessairement un chemin contractuel supposé. Notons aussi que la séparation du marché de l'énergie et du traitement des congestions est une propriété qui rend le modèle du buy back différent du *Market Splitting*, de l'usage d'un OPF centralisé ou des coupures de transactions.

I.2.5.2) Formulation mathématique

Le modèle du buy back utilise un outil d'optimisation qui détermine les ajustements nécessaires pour résoudre les contraintes apparaissant sur le système, tout en minimisant leur coût. On peut donc reprendre la formulation générale de l'OPF présenté au §II.2.4, tout en l'adaptant au cas du buy back. Pour marquer la différence avec l'OPF, nous pouvons appeler l'outil d'optimisation utilisé dans le cadre du buy back *Outil de Redispatching Optimisé* (ORO). Ainsi, la formulation de l'ORO est donc :

Fonction Objectif de l'optimisation : Minimiser les coûts des *ajustements* de production

$$\text{Min } \sum_i c_i^{+,-} \Delta G_i \quad (\text{II.21})$$

Contraintes à respecter :

Equations du calcul de répartition de charge

$$\Delta \mathbf{P} = \mathbf{B} * \Delta \boldsymbol{\theta} \quad (\text{II.22a})$$

$$\Delta \mathbf{P} = \Delta \mathbf{P}_G \quad (\text{II.22b})$$

$$\Delta \mathbf{P}_{ij} = \mathbf{H} * \Delta \boldsymbol{\theta} \quad (\text{II.22c})$$

Equilibre production-consommation

$$\sum_i \Delta G_i = 0 \quad (\text{II.23})$$

Limites imposées aux lignes

$$- \mathbf{P}_{ij}^{\max} \leq \mathbf{P}_{ij}^0 + \Delta \mathbf{P}_{ij} \leq \mathbf{P}_{ij}^{\max} \quad (\text{II.24})$$

Limites imposées aux productions

$$\mathbf{P}_G^{\min} \leq \mathbf{P}_G^0 + \Delta \mathbf{P}_G \leq \mathbf{P}_G^{\max} \quad (\text{II.25})$$

avec

$c_i^{+,-}$ prix proposé par le producteur i pour une offre d'ajustement à la hausse/à la baisse en fonction du volume ajusté ΔG_i

ΔG_i ajustement du producteur i déterminé par l'ORO

$\mathbf{P}_{ij}^0$ vecteur des transits initiaux avant traitement des congestions

$\mathbf{P}_G^0$ vecteur des productions initiales avant traitement des congestions

La méthode du buy back ne dégage en soit aucun surplus financier au bénéfice de l'opérateur du système. Elle crée au contraire un coût à la charge de l'opérateur du système, qui en général le récupère parmi les usagers de façon uniforme (prorata ou timbre-poste).

II.2.5.3) Offres d'ajustements des producteurs

Les offres d'ajustements déposées par les participants se définissent par :

- un *volume maximum* de l'ajustement dans les deux sens (à la hausse et à la baisse) ;
- un *prix* associé au volume ajusté.

Soit une offre d'ajustement déposée par un producteur i . Si pour le traitement des congestions il doit ajuster son injection à la hausse d'une quantité ΔG_i^+ , ce producteur sera rémunéré par l'opérateur du système d'une somme égale à $c_i^+ \Delta G_i^+$, où c_i^+ est le prix qu'il propose pour son ajustement à la hausse. Si, par contre, il doit ajuster à la baisse sa production d'une quantité ΔG_i^- , ce producteur devra rémunérer l'opérateur du système d'une somme égale à $c_i^- \Delta G_i^-$, où c_i^- est le prix qu'il propose pour son ajustement à la baisse. En formulant correctement son offre, le producteur peut s'assurer d'un gain financier en cas de congestion :

- En faisant une offre à la hausse à un prix supérieur au prix où il a vendu son énergie sur le marché ;
- En faisant une offre à la baisse à un prix inférieur du prix de vente de son énergie sur le marché.

Du fait de la séparation du traitement des congestions et du marché de l'énergie, une offre d'ajustement déposée par un producteur peut être différente de l'offre qu'il a faite sur le

marché d'énergie. En outre, il peut très bien choisir de faire une offre pour le marché de l'énergie, mais de ne pas participer au traitement des congestions et vice-versa. Enfin, il peut choisir de faire ou de ne pas faire d'offres pour le traitement des congestions indépendamment de son mode de fourniture (bourse, contrats bilatéraux, ..).

II.3) Conclusion

Dans ce chapitre, nous avons décrit sur le plan théorique les modèles de traitement des congestions les plus utilisés actuellement. La régionalisation du marché ou *Market Splitting* apparaît comme une solution bien adaptée aux réseaux à configuration radiale, mais il est reconnu que cette solution est difficilement applicable pour des réseaux maillés à cause des flux parallèles. La solution californienne avec la participation active des *Scheduling Coordinators* dans le traitement des congestions constitue une solution très particulière qui réclame une adaptation spécifique au niveau du marché de l'énergie. Le concept de coupures de transactions prévues par le TLR a la particularité d'être une solution purement technique au problème des congestions, sans aucune optimisation économique des coûts de production. Malgré ce fait qu'il lui est souvent reproché, il pourrait théoriquement être appliqué dans tout marché comportant des transactions bilatérales, sans qu'il constitue pour autant une méthode à caractère général.

De fait, l'usage de l'OPF et la méthode du buy back apparaissent comme les modèles les plus généraux de traitement des congestions. Ces deux modèles sont en outre fondamentalement différents. Alors que le marché de l'énergie et les contraintes du système sont gérés en une seule étape par l'OPF, ils sont gérés en deux étapes clairement distinctes dans le modèle du buy back. L'OPF associé à sa tarification marginale dégage un surplus financier collecté par l'opérateur centralisé, tandis que le buy back dégage plutôt un coût à sa charge, qu'il doit ensuite répercuter sur les usagers du réseau.

Dans le chapitre suivant, nous allons élaborer une analyse approfondie sur un cas précis, en vue de montrer notamment l'intérêt de la méthode du buy back par rapport aux autres modèles fondamentaux de traitement des congestions. Les critères retenus pour l'analyse seront la pertinence technique et économique des modèles étudiés et leur degré de transparence quant au fonctionnement du marché et du traitement des congestions.

CHAPITRE III

Choix d'un modèle de traitement des congestions

III.1) Introduction

Dans ce chapitre, nous allons analyser l'efficacité technique et économique des modèles de traitement de congestions qui nous ont semblé le plus pertinent d'étudier. Nous avons choisi d'approfondir l'étude du traitement des congestions par coupures de transaction, par usage de l'OPF, et par la méthode du buy back. Nous avons donc écarté les modèles utilisés en Norvège et en Californie, trop spécifiques au fonctionnement de leur marché et à leur réseau pour nous concentrer sur les méthodes théoriquement applicables sur le plus grand nombre de marchés et de réseaux possibles. Nous avons choisi d'étudier successivement les trois méthodes précitées sur un réseau maillé 9 nœuds, et à partir d'un même cas. L'analyse portera sur la pertinence technique et économique des méthodes, ainsi que sur leur degré de transparence. Nous allons notamment montrer à l'aide du cas étudié l'intérêt de la méthode du buy back par rapport à ces critères.

Dans toutes nos simulations, le modèle DC présenté en Annexe 1 a été utilisé pour le calcul de répartition de charge. Toutefois, pour vérifier la validité de ce modèle en terme de précision, nous avons aussi effectués des simulations en utilisant le modèle de calcul de répartition de charge dit « à courant alternatif » ou modèle AC. Les tests ont été effectués sur le réseau 9 nœuds et sur le réseau IEEE 39 nœuds, et une comparaison des résultats obtenus avec les deux modèles sera faite.

III.2) Etude d'un réseau 9 nœuds

III.2.1) Présentation du réseau 9 nœuds-Données du marché

Le cas que nous allons étudier se base sur un réseau simplifié qui possède neuf nœuds et onze lignes de transport. Trois de ces nœuds sont des nœuds de production et les huit nœuds restants sont des nœuds de consommation. Nous avons choisi d'étudier ce réseau car sa petite taille facilite l'analyse, mais son degré de maillage est suffisant pour avoir une bonne représentation du comportement d'un réseau de plus grande échelle. La figure III.1 nous donne la topologie de ce réseau. Les données sur les réactances des lignes sont données par le tableau III.1.

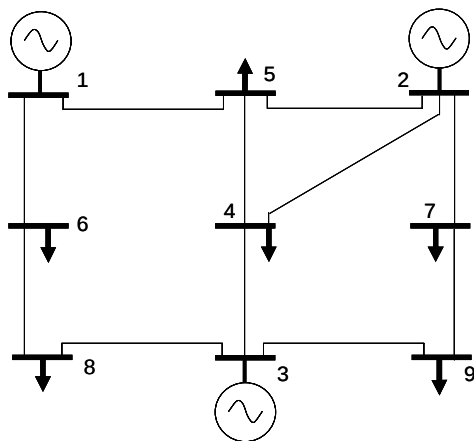


Figure III.1 : topologie du réseau 9 nœuds

TABLEAU III.1 : données des lignes

N°	Ligne		Réactance x (u.r.)	Limite de transit (MW)
	Du nœud	Au nœud		
1	1	5	0.02	600
2	5	2	0.005	300
3	1	6	0.018667	570
4	5	4	0.006667	500
5	2	4	0.007333	500
6	2	7	0.010667	500
7	6	8	0.003333	500
8	4	3	0.002	310
9	7	9	0.004667	300
10	8	3	0.004333	300
11	3	9	0.006	500

Deux marchés vont être conduits sur ce réseau pour une tranche horaire donnée:

- Un marché centralisé, que l'on peut assimiler à une bourse ou un marché spot, d'un volume total de 1170 MW
- Un marché bilatéral, dont le volume total de contrats représente 930 MW

En chaque nœud générateur on trouve deux producteurs : un producteur participant au marché spot et un producteur vendant sa production par contrats bilatéraux. De même, en chaque nœud consommateur, on a un consommateur se fournissant au marché spot (bourse de l'électricité), et un consommateur se fournissant par contrat bilatéral. Un producteur connecté au nœud i et vendant son énergie en bourse est noté Gi^B . Le producteur qui passe par des

transactions bilatérales est noté G_i^T . La même règle vaut pour les consommateurs, notés « C ».

Le tableau III.2 indique les quantités demandées par les consommateurs se fournissant en bourse, les quantités proposées par les producteurs avec le prix proposé. Les consommateurs sont considérés inélastiques au prix du marché. Le tableau III.3 indique le volume des contrats conclus, leur prix, ainsi que les participants se fournissant sur ce mode. Ici, G_3^T a choisi de ne rien vendre sur le marché de l'énergie sur la tranche horaire donnée, mais peut participer au traitement des congestions, notamment dans la méthode du buy back.

TABLEAU III.2 : offres des producteurs et demande des consommateurs au marché spot

	Participant	Quantité proposée sur le marché (MW)	Prix offert (€/MWh)
Producteurs	G_1^B	450	20
	G_2^B	450	25
	G_3^B	450	30
Consommateurs	C_4^B	150	/
	C_5^B	180	/
	C_6^B	180	/
	C_7^B	240	/
	C_8^B	120	/
	C_9^B	300	/

TABLEAU III.3 : Contrats bilatéraux

Contrat n°	Producteur	Client	Quantité contractée (MW)	Prix du contrat (€/MWh)
1	G_1^T	C_8^T	240	25
2	G_1^T	C_9^T	240	25
3	G_1^T	C_5^T	60	25
4	G_1^T	C_6^T	150	25
5	G_2^T	C_4^T	240	27

Comme l'étude est conduite pour un tranche horaire d'1 heure, les prix seront alors exprimés en €/MWh et les coûts en €/h

III.2.2) Cas non contraint

Dans le cas non contraint, on ignore les limites de transit des lignes. Dans ce cas, les plans et prix du marché peuvent être établies sans qu'il y ait traitement de congestion. Les données des contrats bilatéraux sont exactement celles fournies par le Tableau III.3 et en ce qui concerne la bourse, on va sélectionner simplement les offres des producteurs de la moins chère à la plus chère selon le principe du modèle pool exposé au Chapitre I § I.1.1. Le plan de production de la bourse trouvé est indiqué par la figure III.2 :

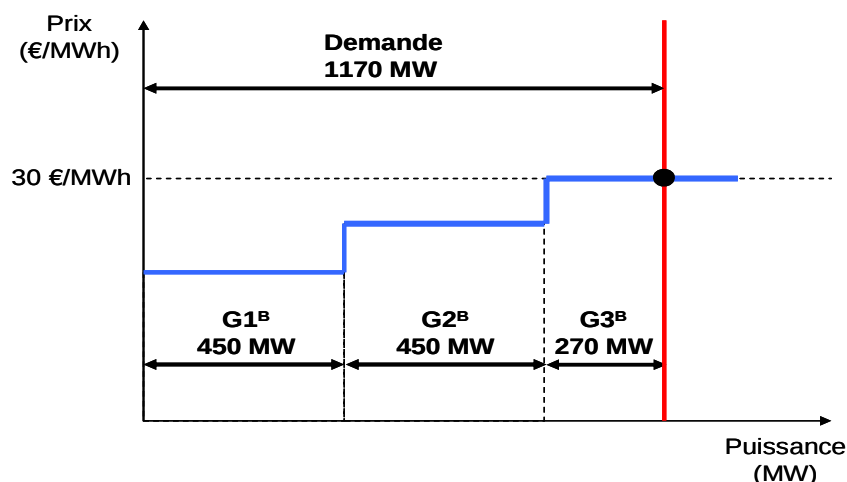


Figure III.2: plans de production de la bourse de l'énergie et prix de clôture du marché

Le marché spot a été ici clôturé sur un prix de 30 €/MWh, prix auquel se fait la vente et l'achat de l'énergie au sein de la bourse. Ce prix correspond au prix proposé par le producteur le plus cher appelé en bourse, c'est-à-dire $G3^B$. En ne considérant aucune limite sur les lignes, il n'existe aucun coût de congestion, et les plans établis à la bourse et par les transactions bilatérales peuvent être conduites pour la tranche horaire donnée du jour J sans traitement des congestions. Dans les prochains paragraphes, nous allons considérer cette fois que les limitations sur les lignes fournies par le tableau III.1 sont actives.

III.2.3) Traitement des congestions par ajustements de production (buy back)

III.2.3.1) Etablissement des plans préférés du marché et vérification de ces plans

Dans un premier temps, les données du marché de l'énergie sont établies exactement comme dans le cas non contraint sans tenir compte des limites du système, et le marché de l'énergie est clôturé, son volume et ses prix étant *définitivement fixés*.

S'ouvre à présent le traitement des congestions. Tous les plans de production/consommation doivent être vérifiés en vue de déceler d'éventuelles violations de contraintes. Dès à présent, notons *qu'il n'est pas besoin de connaître en détail* les transactions en cours. On a besoin uniquement des plans de production/consommation de chaque usager du réseau. L'opérateur du système effectue ainsi un premier calcul de répartition de charge à l'aide des données du réseau (tableau III.1) et on décèle une violation de contrainte sur la

ligne reliant le nœud 1 au nœud 6 (que nous appellerons ligne 1-6⁸ pour plus de concision). Le transit de cette ligne serait de 630.3 MW alors que sa limite est fixée à 570 MW (fig.III.3).

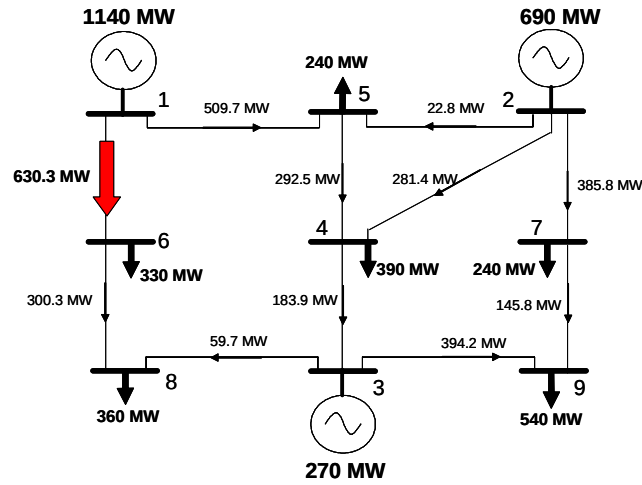


Figure III. 3: vérification des plans de production/consommation par calcul de répartition de charge et congestion sur la ligne 1-6

L'opérateur du système doit alors faire appel aux offres d'ajustement déposées par les producteurs volontaires dans le cadre du marché « traitement des congestions ».

II.2.3.2) Minimisation du coût des ajustements et résolution de la contrainte sur la ligne 1-6

On doit à présent résoudre la congestion à l'aide des offres d'ajustements déposées tout en minimisant leur coût. On va donc utiliser l'Outil de Redispatching Optimisé (ORO) décrit par les expressions II.21 à II.25 (Chapitre II). On adapte alors le problème général décrit par les expressions II.21 à II.25 (avec le nœud 1 comme nœud de référence) aux données du réseau fournies par les tableaux III.2, III.3 et III.4 et la fonction objectif de notre problème devient :

$$\text{Min} \quad C_{1B}(\Delta G1^B) * \Delta G1^B + C_{1T}(\Delta G1^T) * \Delta G1^T + C_{2B}(\Delta G2^B) * \Delta G2^B + C_{2T}(\Delta G2^T) * \Delta G2^T \\ + C_{3B}(\Delta G3^B) * \Delta G3^B + C_{3T}(\Delta G3^T) * \Delta G3^T$$

(III.26)

avec C_{1B} , C_{1T} , C_{2B} , C_{2T} , C_{3B} et C_{3T} les fonctions d'ajustement offertes par les producteurs et représentant le prix offert par les producteurs en fonction d'un certain ajustement.

⁸ L'expression générale « ligne $i-j$ » implique que le nœud i soit par convention le nœud émetteur, et que le nœud j le nœud récepteur

Etant donné le contexte dérégulé, on peut considérer que ces offres renferment les coûts fixes et variables des producteurs, ajoutés d'une certaine marge laissée aux choix du producteur. Dans notre cas, il n'est pas nécessaire de mettre en évidence les coûts de démarrage/arrêt d'une unité de production. Nous choisissons donc de modéliser ces offres d'ajustement par des fonctions affines représentant le prix offert par le producteur en fonction d'un certain ajustement à la hausse ou à la baisse. Cette représentation a aussi l'avantage de nous offrir des conditions d'optimisation plus faciles. Ainsi, le prix offert par le producteur peut s'exprimer de cette façon :

$$c_i^{+,-} = a_i + b_i \Delta G_i$$

où a_i et b_i sont des constantes fixées par le producteur. Il peut en outre fixer la constante a_i égale à son prix de vente sur le marché de l'énergie et $b_i > 0$.

Le paiement associé à un producteur est donc :

$$C_i = a_i \Delta G_i + b_i \Delta G_i^2$$

Pour résoudre le problème exprimé en (III.26), il nous faut choisir un algorithme d'optimisation. L'usage d'un algorithme déterministe s'avère ici être le choix le plus judicieux. En effet, la fonction objectif exprimée dans la relation II.26 est convexe, elle a l'avantage d'être continue et de ne présenter qu'un optimum global. Des algorithmes déterministes de type méthode de Newton ou du Point Intérieur s'adaptent très bien à la formulation de notre problème. Ici, nous avons utilisé l'algorithme d'optimisation de Quasi-Newton implanté dans l'environnement informatique Matlab. Pour plus de détails sur les méthodes d'optimisation, voir l'Annexe 2. Ici, nous avons supposé que tous les producteurs avaient fourni des offres d'ajustement. Les paramètres de ces offres sont donnés par le tableau III.4:

TABLEAU III.4 : paramètres des fonctions d'ajustements

Producteur	G1 ^B	G1 ^T	G2 ^B	G2 ^T	G3 ^B	G3 ^T
a_i (€/h)	30	25	30	30	30	30
b_i (€/MWh)	0.15	0.1	0.1	0.17	0.2	0.17
Volume à la hausse/baisse offert (MW)	+125/- 125	+125/- 125	+75/-75	+50/-50	+75/-75	+100/0

Les ajustements optimaux trouvés sont donnés par le tableau III.5. Ce dernier indique aussi les rémunérations ou paiements associés à chaque producteur et le coût total de congestion obtenu :

TABLEAU III.5 : ajustements optimaux, rémunérations/paiements des producteurs et coût de congestion

Producteur	Volume de l'ajustement (MW)	Prix de l'ajustement (€/MWh)	Rémunération ou paiement associé au producteur : (Volume ajusté) * (prix de l'ajustement) (€/h)
G1^B	-67.01	21.1	-1413.9
G1^T	-64.34	18.6	-1196.7
G2^B	+40.26	34.0	1368.8
G2^T	+24.16	34.1	825.2
G3^B	+30.42	36.1	1098.2
G3^T	+36.51	36.2	1321.7
		Coût de congestion total	2003.3 €/h

Le volume total de production « déplacée » (production déplacée du nœud 1 vers les nœuds 2 et 3) est de 131.35 MW.

Sur le plan physique, les ajustements décrits par le tableau III.5 résolve la contrainte en ramenant le transit prévisionnel de la ligne 1-6 à sa valeur maximale permise qui est de 570 MW. Tous les autres transits sont en dessous de leur limite maximale. Les figures III.4a et III.4b nous permettent de comparer les injections nodales et transits avant et après traitement des congestions. Nous rappelons aussi que ce sont les plans de production/consommation indiqués par la figure III.4b qui seront finalement conduits le jour J, si l'on ne tient pas compte d'autres modifications possibles associés à d'autres services systèmes, ainsi que des modifications possibles en temps réels.

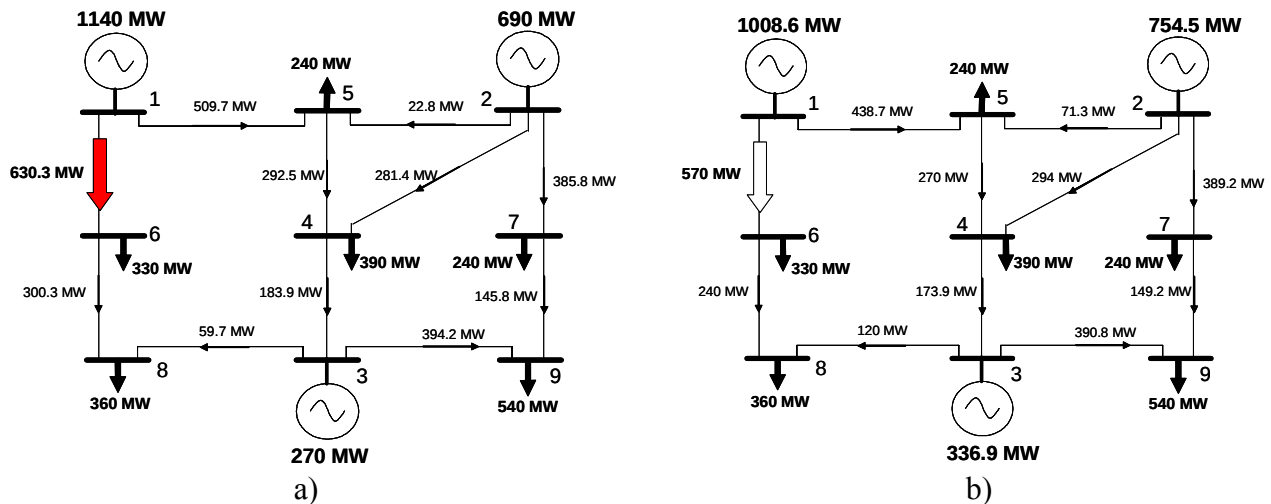


Figure III.4a et III.4b : Etat du réseau avant traitement des congestions a) et après traitement des congestions b)

Nous pouvons aussi analyser les résultats obtenus sur le plan économique à l'aide du tableau III.5 et des données du marché de l'énergie. Ici, les volumes échangés sur le marché de l'énergie et leur prix (bourse et contrats bilatéraux) et les volumes échangés dans le traitement des congestions et leur prix (ajustements de production) sont clairement distincts. Par exemple, le producteur $G1^T$ a vendu sur le marché de l'énergie 690 MW à un prix fixé de 25 €/MWh. Il recevra donc sur le marché de l'énergie une rémunération total 17250 €/h. Par contre, sur le marché « traitement des congestions », il a été ajusté à la baisse de 64.34 MW à un prix de 18.6 €/MWh, déterminé à l'aide de sa fonction d'ajustement. Il devra donc rembourser à l'opérateur du système 1196.7€/h pour cet ajustement à la baisse. Ceci dit, il garde intacts ses contrats passés avec ses clients, que ce soit en terme de volume ou de prix. En outre, ce remboursement à l'opérateur du système qui peut à première vue s'apparenter à un coût pour le producteur $G1^T$ ne l'est pas vraiment si on considère le fait qu'il garde une fourniture de 690 MW tout en injectant en temps réel 625.66 MW. Il fait donc une économie de coût de production correspondant à son volume ajusté. En formulant judicieusement son offre d'ajustement, il peut même faire en sorte de rembourser moins d'argent à l'opérateur du système qu'il en épargne en coût de production !

Le même genre d'analyse peut être conduite pour les autres producteurs. Ce modèle fait donc une distinction claire entre les transactions commerciales et les flux physiques. De par cette distinction, il peut très bien s'adapter à n'importe quelle structure de marché. Cette distinction est bénéfique sur le plan économique, car elle donne plus de transparence au marché de l'énergie, mais aussi en terme de coût de congestion. Comme nous l'avons vu lorsque nous avons décrit le modèle californien, effectuer des ajustements de production sans tenir compte des SC était plus efficace que vouloir forcer ces derniers à avoir un bilan électrique absolument équilibré (voir Chapitre II §II.2.2).

Le coût de congestion obtenu dépend bien évidemment des paramètres des fonctions d'ajustements fixés librement par les producteurs. Or ces derniers peuvent se réserver une marge de profit plus ou moins grande. Si on considère des offres comme celles utilisées ici où les marges des producteurs restent raisonnables, on peut s'attendre à un service à coût convenable. Par exemple, dans notre cas le coût du traitement des congestions représente 3.4 % du volume financier du marché de l'énergie, chiffre acceptable, mais qui peut fluctuer suivant les paramètres des offres d'ajustement proposées.

Le coût de congestion obtenu correspond en outre à un solde à la charge de l'opérateur du système. Ce coût devra donc être intégralement récupéré parmi les usagers du réseau. Nous ne traiterons pas dans ce chapitre l'allocation du coût de congestion aux usagers ; toutefois, nous lui consacrerons une analyse plus approfondie dans le Chapitre IV.

III.2.4) Traitement des congestions par coupures de transactions

Dans cette partie, nous allons illustrer comment on peut traiter les congestions sur le cas du réseau 9 nœuds en procédant par coupures de transactions. La procédure illustrée ici s'appuie sur la méthode préconisée par le NERC lorsqu'il a défini le TLR. Nous rappelons que cette procédure s'applique normalement durant les conditions d'exploitation du réseau pour traiter les contraintes de transit potentielles ou survenant en temps réel. Toutefois, nous pouvons tout aussi bien tester la méthode du TLR dans un cadre prévisionnel en appliquant les mêmes principes.

III.3.4.1) Premier cas

Ainsi, nous supposons que les données du marché de l'énergie s'établissent d'abord sans tenir compte des contraintes du système, de sorte que nous obtenons de nouveau le même cas de figure que le cas non contraint. Nous supposons aussi pour cet exemple que les transactions n°1 à 4 conclues avec le producteur G1^T ont la même priorité sur la ligne 1-6, les participants du contrat n°5 ayant réservé leur capacité plutôt sur la ligne 2-4.

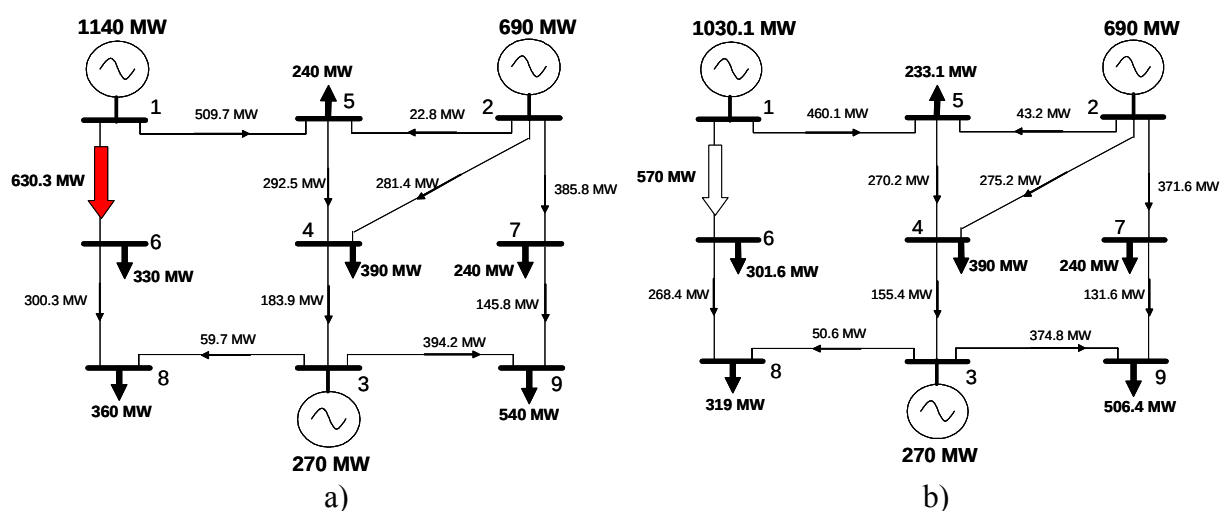
De même que précédemment, après un premier calcul de répartition de charge, nous décelons de nouveau la congestion sur la ligne 1-6. Cette fois-ci, nous allons la traiter en réduisant le volume des contrats bilatéraux qui ont un impact sur la congestion, comme le veut la procédure du TLR. En outre, la procédure du TLR précise qu'une transaction ayant un impact sur une congestion inférieur à un certain seuil (5%) ne sera pas incluse dans la stratégie de coupure. Le contrat n°5 ayant un très faible impact sur la congestion de la ligne 1-6 (<5%) ne sera donc pas coupé. De plus, comme la procédure ne s'applique théoriquement

pas aux structures boursières, nous laisserons aussi les plans de la bourse de l'énergie tels quels.

En appliquant la relation (II.12) au cas de la congestion sur la ligne 1-6, nous obtenons les coupures de transactions indiquées par le tableau III.6. Etant donné la complexité de la formule de coupure, nous avons détaillé le calcul des coupures à effectuer. Seuls les contrats n°1 à n°4 sont concernés par le traitement de la congestion de la ligne 1-6. Les résultats sur les transits et les injections nodales sont donnés par les figures III.5a et III.5b.

TABLEAU III.6 : résultats des coupures de transactions en vue de résoudre la congestion sur la ligne 1-6

	1	2	3	4	5	6	7	8
Contrat n°	Volume initial (MW)	FDB (Facteur de Distribution Bilatéral)	Impact sur la ligne (2)*(1)	FDB pondéré (2)/(2TOT)	Réduction normalisée (3)*(4)	Réduction ramenée à la surcharge (5)*(surcharge)/(5TOT)	Volume de coupure de la transaction (6)/(2)	Nouveau volume de la transaction (MW) (1)-(7)
1	240	0.576	138.24	0.278	38.43	23.62	41	199
2	240	0.474	113.76	0.228	25.94	15.94	33.63	206.37
3	60	0.385	23.1	0.186	4.30	2.64	6.86	53.14
4	150	0.64	96	0.308	29.57	18.17	28.39	121.61
TOTAUX	690	2.075	371.1		98.24	63.37	109.88	



Figures III.5a et III.5b : Etat du réseau avant traitement des congestions a) et après traitement des congestions b)

Ces résultats nous amène à formuler plusieurs observations.

Contrairement à la méthode du buy back, toute réduction de transaction a un impact direct sur le volume du contrat : par exemple, le contrat n°1 ayant été coupé de 41 MW, le consommateur C8^T doit chercher ses MW manquants auprès d'un autre fournisseur, ou auprès de la bourse, tout en veillant avec l'aide de l'opérateur du système à ne pas créer de nouveau des congestions sur le réseau.

Le modèle du TLR tel qu'il est appliqué jusqu'à maintenant ne fait pas intervenir la bourse de l'énergie. Ce fait est peu gênant pour les marchés exclusivement bilatéraux ou à forte composante bilatérale, mais il le devient si on est en présence d'une bourse de l'énergie ayant un volume relativement important. En effet, dans ce cas-là, on ne tiendrait pas compte de l'influence de la bourse de l'énergie sur les contraintes du système. Des solutions ont toutefois été proposées dans le but de mieux coordonner les différents marchés [GAL2 02].

De plus, le résultat final dépend de l'ordre de priorité des transactions. Si une transaction à forte priorité fait partie de celles qui ont l'impact le plus fort sur une congestion, il y a un risque de devoir étendre les coupures à un grand nombre de transactions ayant un couplage moins fort avec la congestion, ce qui peut générer des volumes de coupures trop importants par rapport à ce qui serait requis.

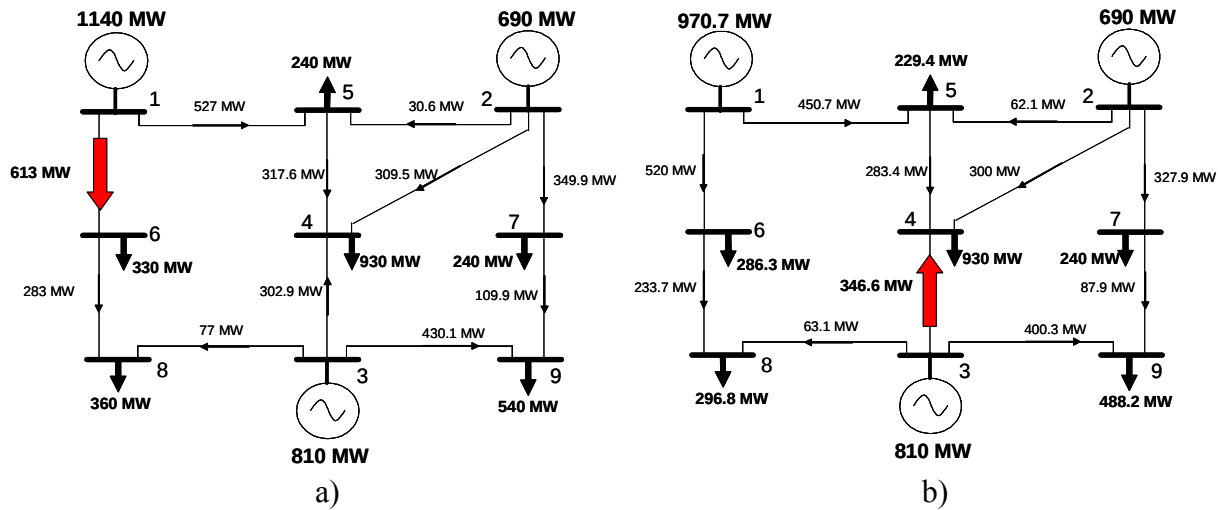
Enfin, on peut se retrouver devant des problèmes décisionnels si le délestage d'une transaction provoque une surcharge ailleurs sur le réseau. Pour illustrer ce fait, nous allons étudier dans le paragraphe suivant un deuxième cas sur le réseau 9 nœuds.

III.3.4.2) Deuxième cas

Rajoutons une sixième transaction de 540 MW entre le nœud producteur 3 et le nœud consommateur 4 et rabaissons la limite de la ligne 1-6 à 520 MW. Avec les autres données du marché de l'énergie telles qu'elles sont définies dans le cas non contraint, nous obtenons toujours une congestion sur la ligne 1-6, dont le transit prévisionnel est cette fois-ci de 613 MW. Nous allons de nouveau traiter cette congestion en utilisant la procédure du TLR, en observant cette fois-ci que le traitement de la congestion sur la ligne 1-6 fait apparaître une surcharge sur la ligne 4-3. Les résultats sont reportés par le tableau III.7 et par les figures III.6a et III.6b :

TABEAU III.7 : résultats des coupures de transactions en vue de résoudre la congestion sur la ligne 1-6 (deuxième cas)

	1	2	3	4	5	6	7	8
N° transaction	Volume initial (MW)	FDB	Impact sur la ligne (2)*(1)	FDB pondéré (2)/(2TOT)	Réduction normalisée (3)*(4)	Réduction ramenée à la surcharge (5)*(surcharge)/(5TOT)	Coupure de la transaction (6)/(2)	Nouveau volume de la transaction (MW) (1)-(7)
1	240	0.576	138.24	0.278	38.43	36.38	63.16	176.84
2	240	0.474	113.76	0.228	25.94	24.56	51.81	188.19
3	60	0.385	23.1	0.186	4.30	4.07	10.57	49.43
4	150	0.64	96	0.308	29.57	28	43.75	106.25
TOTAUX	690	2.075	371.1		98.24		169.29	



Figures III.6a et III.6b : Etat du réseau avant traitement des congestions a) et après traitement des congestions b) avec une transaction supplémentaire entre le nœud 3 et 4

Nous pouvons constater que le traitement de la congestion sur la ligne 1-6 par coupure de transactions a contribué à augmenter le transit sur la ligne 3-4 l'amenant à une valeur de 346.6 MW, ce qui est au-dessus de sa capacité maximale (310 MW). Ce phénomène est dû au fait que les transactions coupées pour soulager la contrainte sur la ligne 1-6 ont un impact en sens inverse du transit sur la ligne 3-4. Par exemple, le FDB de la transaction n°1 sur le transit de la ligne 4-3 est de 0.354. Cela signifie que la transaction n°1 surcharge un transit allant du nœud 4 au nœud 3, mais soulage un transit allant en sens inverse (ce qui est le cas, le flux étant dirigé du nœud 3 au nœud 4). L'impact de l'ensemble des coupures de transaction sur le transit de la ligne 4-3 peut aussi se calculer de cette façon :

$$\Delta P_{43} = \sum_{i=1}^4 FDB_{Ti}^{4-3} * \Delta T_i$$

avec FDB_{Ti}^{4-3} étant le Facteur de Distribution Bilatéral (tel qu'il est défini au Chapitre II §II.2.3) d'une transaction Ti sur la ligne 4-3.

Pour éliminer la congestion sur la ligne 4-3, il faudrait couper à son tour la transaction rajoutée entre le nœud 3 et le nœud 4 à hauteur de 30 MW, alors qu'elle n'était à l'origine pas responsable du problème sur la ligne 1-6. Cette dernière transaction avait même tendance à soulager légèrement la contrainte sur la ligne 1-6 (son FDB sur la ligne 1-6 est de -0.032). Ainsi, une stratégie de coupures de transaction telle qu'elle est appliquée dans le TLR peut mener suivant les cas à des décisions peu efficace sur le plan technique [CAD 01].

Une nouvelle fois, nous pouvons montrer que l'on pouvait traiter la congestion de façon bien plus simple et efficace sur le plan technique en ayant plutôt recours à des ajustements de production, sans faire référence aux transactions conclues. Avec les mêmes offres d'ajustement que celles utilisées au paragraphe I.2.3.1, on pouvait traiter le cas de la figure III.6a :

TABLEAU III.8 : ajustements de production

	$\Delta G1^B$	$\Delta G1^T$	$\Delta G2^B$	$\Delta G2^T$	$\Delta G3^B$	$\Delta G3^T$
Volume de l'ajustement (MW)	-98.16	-105.88	+69.87	+41.92	+41.92	+50.31

Le coût de congestion est relativement élevé dans ce cas (4490 €/h). Toutefois, nous devrions noter que ce deuxième cas est plus contraignant que le premier, et que là où les coupures de transactions résolvait une congestion tout en créant une surcharge ailleurs, la méthode du buy back ramène tous les transits dans leur limite fixée tout en laissant intacts les transactions commerciales.

Ce que l'on peut donc souligner pour la méthode de coupure de transaction préconisée par le TLR, ce sont les difficultés techniques pouvant surgir en procédant par coupures de transaction, le risque accru d'instabilité des contrats bilatéraux qui peut amener à une volatilité plus grande du marché de l'énergie.

III.2.5) Traitement des congestions par l'usage de l'OPF et des prix nodaux

Dans cette section, nous allons illustrer le traitement des congestions sur notre réseau 9 nœuds, en utilisant cette fois l'OPF et la méthode des prix nodaux présentés au Chapitre II §

II.2.4. Dans les deux méthodes étudiées précédemment, on établissait d'abord les plans du marché de l'énergie sans tenir compte des contraintes sur le système, et ensuite on résolvait les congestions dans un deuxième temps. Avec l'usage de l'OPF, on procède en une seule étape : on recueille les offres des participants au marché spot et les données des transactions bilatérales (Tableau III.3) et l'on va configurer le plan de production du marché spot de manière à ce que les contraintes sur le système soient respectées et que le coût total de production (du marché spot) soit minimisé. On va donc intégrer dans la fonction objectif donnée par l'expression (II.13) les offres des producteurs participants au marché spot. Le prix offert par chaque producteur sur le marché spot est donné par le Tableau III.2. La fonction objectif du problème d'optimisation devient donc :

$$\text{Min } 20 * P_{G1^B} + 25 * P_{G2^B} + 30 * P_{G3^B} \quad (\text{II.27})$$

Nous avons aussi transposé le problème d'optimisation décrit par les expressions II.14a à II.17 à l'aide des données du réseau 9 nœuds. Les puissances produites par les producteurs dispatchés peuvent varier de 0 MW à la puissance maximum offerte sur le marché spot, c'est-à-dire 450 MW (Tableau III.2). Nous avons considéré toutes les transactions bilatérales comme étant fermes ; elles sont donc intégrées au modèle mais sont traitées comme des constantes.

Si on tient compte des limites de transit des lignes, on obtient une configuration de la production du marché spot différente de celle du cas non contraint, avec une différenciation géographique des prix nodaux (ou prix marginaux) tirés à partir de l'expression (II.20). Cette différenciation est due à la contrainte présente sur la ligne 1-6 dont le transit butte à sa limite maximale admissible. Les résultats sont présentés par le Tableau III.10. Dans le Tableau III.9, nous avons redonné les résultats du cas non contraint pour mieux pouvoir comparer les deux états. Le Tableau III.11 nous donne le coût de congestion auquel sont soumises les transactions bilatérales suivant l'expression (II.21)

TABLEAU III.9 : cas non contraint

noeud	Production du marché spot (MW)	Consommation du marché spot (MW)	Prix nodal (€/MWh)	Rémunération des producteurs (€/h)	Paiement des consommateurs (€/h)
1	450	/	30	13500	/
2	450	/	30	13500	/
3	270	/	30	8100	/
4	/	150	30	/	4500
5	/	180	30	/	5400
6	/	180	30	/	5400
7	/	240	30	/	7200
8	/	120	30	/	3600
9	/	300	30	/	9000
Totaux				35100	35100

TABLEAU III.10 : cas contraint

noeud	Production du marché spot (MW)	Consommation du marché spot (MW)	Prix nodal (€/MWh)	Rémunération des producteurs (€/h)	Paiement des consommateurs (€/h)
1	327.50	/	20	6550	/
2	450	/	28.6	12870	/
3	392.50	/	30	11775	/
4	/	150	29.3	/	4395
5	/	180	27.8	/	5004
6	/	180	33	/	5940
7	/	240	29.3	/	7032
8	/	120	31.7	/	3804
9	/	300	29.6	/	8880
Totaux				31195	35055

TABLEAU III.11 : coût de congestion associé aux transactions bilatérales

Contrat n°	Quantité contractée (MW)	Différence de prix nodal (€/MWh)	Coût de congestion (€/h)
1	240	11.7	2808
2	240	9.6	2304
3	60	7.8	468
4	150	13	1950
5	240	0.7	168
Total			7698

A partir des résultats présentés ci-dessus, nous pouvons faire plusieurs remarques :

- Alors que le paiement total des consommateurs achetant au marché spot varie peu (35055 €/h contre 35100 €/h dans le cas non contraint), l'opérateur du système bénéficie d'une production moins chère (31195 €/h contre 35100 €/h dans le cas non contraint). En outre, le producteur G1^B enregistre des revenus beaucoup moins importants.
- Les paiements liés aux transactions bilatérales sont relativement importants (7698 €/h), et plus l'impact de la transaction (mesuré à l'aide de son FDB) est grand, plus la différence des prix nodaux à laquelle elle est sujette est grande. L'importance des coûts de congestion reçus par les transactions bilatérales est connu comme pouvant constituer un frein à ce genre d'arrangement sous une tarification nodale [GAL1 02]

Dans ce modèle, l'opérateur mixte du système et du marché spot collecte un surplus financier de 11558 €/h. L'ensemble des flux financiers est présenté par la Figure III.7 :

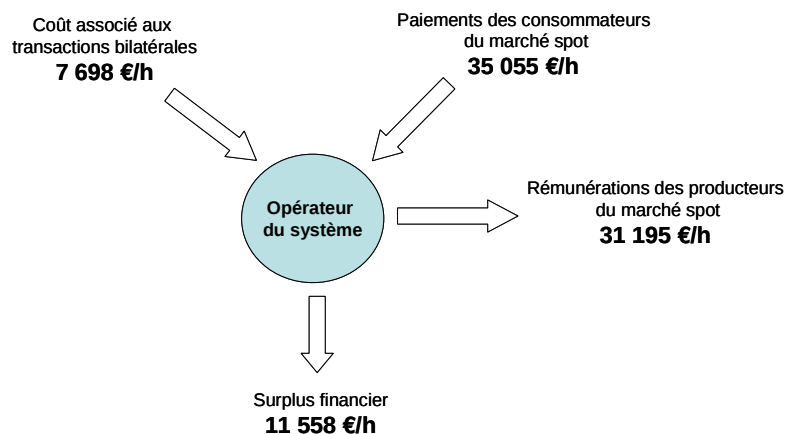


Figure III.7 : flux financiers centralisés par l'opérateur du système avec l'usage de l'OPF et des prix nodaux

[RAI2 01] contient une étude détaillée des surplus financiers amassés dans des marchés américains tels que PJM, CAISO et New York ISO. Le surplus financier dégagé par l'OPF et sa tarification marginale est souvent justifié par le fait qu'il permet de financer en partie les investissements réseau nécessaire. Toutefois, des critiques ont vu le jour, relevant le fait que ce surplus peut ne pas donner de bonnes incitations à ceux qui en bénéficient (qui sont souvent les propriétaires du réseau) [PER 95]. Alors que la méthode du buy back représente un coût raisonnable pour l'opérateur du système qui devra le récupérer ensuite de façon adéquate et sans faire de marge, la méthode basée sur l'OPF et la tarification marginale

permet clairement à l'opérateur du système de récolter des surplus pouvant être reversés au profit des propriétaires du réseau. Les coûts importants associés notamment aux transactions bilatérales suscitent chez les participants qui ont recours à ce mode de fourniture un besoin de protection contre les fluctuations souvent imprévisibles des prix nœdaux. Ainsi, des droits financiers de transport (*Financial* ou *Fix Transmission Rights*- FTR [ALO 99], [PJM 00]) peuvent être octroyé aux participants impliqués dans une transaction bilatérale ferme. Ces droits s'acquièrent par des enchères et vont donc aux plus offrants. Ils donnent aux titulaires la possibilité d'être totalement remboursés du coût de congestion associé à leur transaction. Le désavantage est qu'ils ne seraient alors plus attentifs au signal économique donné par la différence des prix nœdaux.

D'autre part, l'autre critique qui peut être adressée à ce modèle et son manque de transparence. En effet, ce modèle, en établissant les plans du marché spot en même temps qu'il gère les contraintes du système, il ne permet pas de différencier clairement dans les prix nœdaux la part qui résulte du marché de l'énergie et la part qui résulte du traitement des congestions. En outre, dans l'exemple présenté dans cette section, les participants du marché spot n'auraient pas pu savoir, sauf calcul spécifique [CHE 02], que le prix du marché sans congestion aurait été de 30 €/MWh. Ils ne peuvent donc pas clairement déterminer quelle est la part du coût à imputer *uniquement* aux congestions.

Toutefois, nous devrions souligner qu'une propriété largement reconnue et utile des prix nœdaux est l'envoi de signaux économiques représentant l'influence de chaque participant sur les contraintes du système. Par exemple, si on observe la répartition des prix nœdaux du Tableau III.10, les consommateurs connectés aux nœuds 4,5 et 7, situés en aval de la congestion, sont soumis à un prix inférieur aux consommateurs situés aux nœuds 6 et 8 juste dans le prolongement de la congestion. Or, par un raisonnement purement qualitatif, il est clair que les consommateurs des nœuds 4,5 et 7 ne contribuent physiquement pas ou très peu à la surcharge sur la ligne 1-6. L'observation inverse peut être faite pour les consommateurs situés aux nœuds 6 et 8, soutirant une bonne partie de leur puissance par la ligne congestionnée. L'envoi de signaux économiques par tarification marginale a aussi pour but d'aider à un placement optimal de nouveaux producteurs ou consommateurs vis-à-vis des contraintes du système. D'autre part, on peut observer que les différences de prix nœdaux associées aux transactions bilatérales sont proportionnelles à leur FDB par rapport à la ligne 1-6 (Fig. III.7). La même observation a d'ailleurs été faite en [CAD 01].

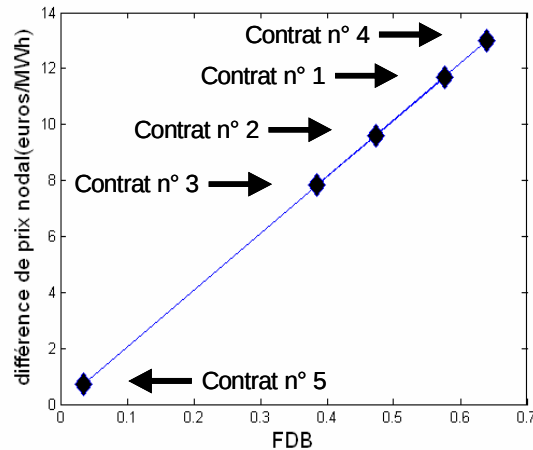


Figure III.7 : différence de prix nodal de chaque transaction en fonction de son Facteur de Distribution Bilatéral par rapport à la ligne 1-6

Nous pouvons toutefois souligner qu'il est tout à fait possible d'envoyer des signaux économiques clairs à partir de la méthode du buy back, par une allocation judicieuse du coût de congestion. Nous garderions en plus le bénéfice d'une plus grande transparence que celle apportée par l'usage de l'OPF, en séparant clairement les sorties du marché de l'énergie à celles du traitement des congestions, et en s'assurant que l'opérateur du système ne fait pas de profit sur le traitement des congestions. La somme récoltée par l'opérateur du système et résultant de l'allocation du coût de congestion serait alors strictement égale au coût du service.

III.3.6) Sensibilité du coût de congestion dans les modèles du buy back et des prix nodaux

III.3.6.1) Sensibilité par rapport aux offres d'entrées

Lorsque nous avons explicité le modèle du buy back précédemment, nous avons modélisé les offres d'ajustement des producteurs par des fonctions affines, qui ont été caractérisées par les paramètres a_i (valeur du prix à l'origine) et b_i (coefficient directeur de la fonction). Il serait intéressant d'illustrer les conséquences d'un écartement des prix des offres d'ajustements sur le coût de congestion. Ceci peut être réalisé en augmentant le coefficient directeur des offres d'ajustements. Ce faisant, le revenu des producteurs ajustés à la hausse augmente, tandis que le paiement des producteurs ajustés à la baisse diminue. Nous avons relevé le coût de congestion pour plusieurs cas de figure, et dans chacun des cas, les coefficients directeurs des offres ont été augmentés dans les mêmes proportions. Nous avons constaté que le coût de la congestion sur la ligne 1-6 augmentait linéairement, en passant de

2003.3 €/h à 6151 €/h par exemple quand tous les coefficients b_i étaient augmentés d'un rapport 3.5 :

TABLEAU III.12 : variation du coût de congestion en fonction du paramètre b_i des offres d'ajustements

Rapport b_i / b_i^0	1	1.5	2.5	3	3.5
Coût de congestion (€/h)	2003	2836	4496	5324	6151

Nous pouvons faire le même genre d'analyse pour la méthode basée sur l'OPF et les prix nodaux. Nous pouvons créer plusieurs cas de figure où l'on a un écart plus ou moins grand des offres des producteurs au marché spot :

Cas 1 (cas étudié précédemment)

Prix de $G1^B$: 20 €/MWh Prix de $G2^B$: 25 €/MWh Prix de $G3^B$: 30 €/MWh
 Somme récoltée par l'opérateur du Système : 11558 €/h

Cas 2

Prix de $G1^B$: 15 €/MWh Prix de $G2^B$: 25 €/MWh Prix de $G3^B$: 35 €/MWh
 Somme récoltée par l'opérateur du Système : 23137 €/h

Cas 3

Prix de $G1^B$: 10 €/MWh Prix de $G2^B$: 25 €/MWh Prix de $G3^B$: 40 €/MWh
 Somme récoltée par l'opérateur du Système : 34706 €/h

Nous avons constaté que plus l'écart des prix des offres des producteurs au marché spot était important, plus la différence géographique des prix nodaux était accentuée.

Ainsi, l'écartement des prix des offres d'entrées a un impact similaire sur le coût de congestion obtenu dans le modèle du buy back et sur le surplus obtenu avec l'usage de l'OPF

III.3.6.2) Sensibilité par rapport à la contrainte imposée

Nous pouvons aussi analyser l'évolution du coût de congestion dans le modèle du buy back et dans le modèle des prix nodaux en fonction de la sévérité de la contrainte imposée. La Figure III.8 illustre l'évolution du coût de congestion obtenu par ajustements de production sur la ligne 1-6 en fonction de la contrainte de transit imposée :

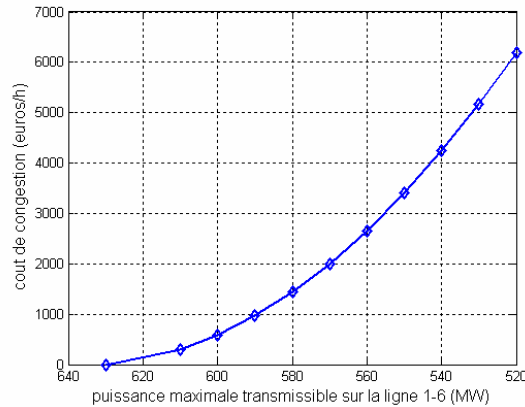


Figure III.8 : évolution du coût de congestion sur la ligne 1-6 dans le modèle du buy back en fonction de la limite de transit imposée

De la même façon, on peut représenter l'évolution de la somme récoltée par l'opérateur du Système dans le modèle des prix nodaux en fonction de la contrainte imposée sur la ligne 1-6. Voici l'évolution que nous obtenons :

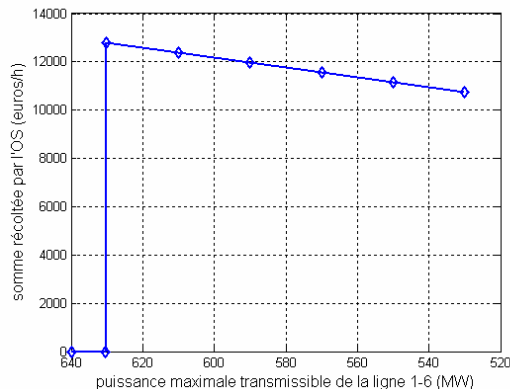


Figure III.9 : évolution du surplus financier récolté par l'opérateur du Système en fonction de la contrainte imposée sur la ligne 1-6

L'évolution du surplus financier collecté par l'opérateur du système dans le modèle de l'OPF peut aisément s'expliquer. Pour une limite fixée au-dessus de 630 MW sur la ligne 1-6, le système devient non contraint, et la solution de l'OPF devient la même que celle du cas non contraint : tous les prix nodaux sont alors égaux à 30 €/h qui est le prix du dernier producteur appelé dans le cadre du marché spot (en l'occurrence G3^B), et le surplus collecté par l'opérateur du système est nul. Dès que la limite passe en dessous de 630 MW, l'algorithme

butte sur la limite de transit de la ligne 1-6 et le système devient contraint, les prix nœdaux prenant les valeurs indiquées par le Tableau III.10. Tant que l'on ne rencontre pas une nouvelle contrainte sur le système (limite supérieure/inférieure de production, limite de transit sur une autre ligne, etc...) ces valeurs restent constantes, ce qui signifie que les paiements des consommateurs au marché spot et les coûts associés aux transactions bilatérales restent aussi constants. Seul change le coût total de production du marché spot (qui est la fonction objectif optimisée par l'OPF), qui augmente progressivement lorsque l'on baisse la limite sur la ligne 1-6. Ainsi, le surplus financier qui est le solde entre les différents paiements récoltés par l'opérateur du Système et les rémunérations des producteurs du marché spot (voir Fig.III.7) diminue progressivement au fur et à mesure que l'on baisse la limite sur la ligne 1-6.

Ainsi, l'évolution du coût de congestion par buy back en fonction de la valeur de la contrainte imposée sur la ligne 1-6 apparaît comme relativement stable. Par contre, le surplus financier dégagé dans le modèle des prix nœdaux semble varier brusquement dès qu'une contrainte apparaît sur le système. Ceci peut poser problème lorsque l'on se trouve très près d'une limite, car une petite congestion peut vite générer des surplus financiers importants.

III.3.6) Validation de l'ORO basé sur le modèle DC

III.3.6.1) Validation sur le réseau 9 noeuds

Pour vérifier que les hypothèses émises dans le modèle DC sont valides, nous avons examiné l'état des transits lorsque l'on applique les ajustements déterminés par l'ORO (Tableau III.5), ces transits étant calculés dans le modèle dit « à courant alternatif » ou modèle AC (voir Annexe 1). Nous avons introduit une résistance pour chaque ligne, en prenant un rapport résistance/réactance de 1/5. Nous avons aussi introduit des charges inductives sur les nœuds consommateurs, en respectant un rapport puissance réactive/puissance active de 0.4 maximum. La figure III.8 nous permet de comparer les transits obtenus dans le modèle DC (qui sont ceux de la figure III.4b) avec les transits obtenus par le modèle AC.

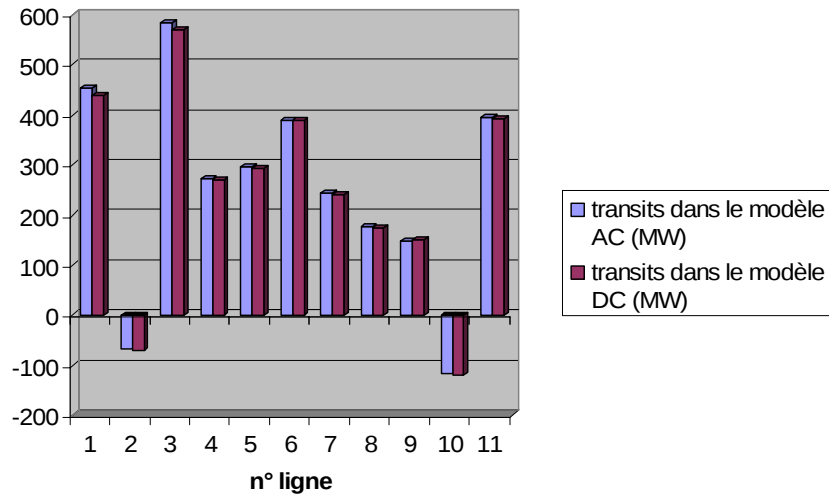


Figure III.8 : transits actifs obtenus après ajustement de la production par l'ORO : comparaison entre les transits calculés par le modèle DC et ceux calculés par le modèle AC

Sur les transits résultant du traitement des congestions par l'ORO, nous pouvons constater qu'il y a très peu de différence entre le modèle AC et le modèle DC. Par exemple, le transit de la ligne 1-6 affiche dans le modèle AC 585 MW en début de ligne et 573 MW en fin de ligne (pertes évacuées). Sur l'ensemble des lignes du réseau, nous obtenons une différence égale ou inférieure à 3% entre les transits du modèle DC et ceux du modèle AC. Cette différence est acceptable car les limites de transit sur les lignes sont généralement fixées avec une marge de sécurité plus grande ; on peut donc valider le modèle DC comme étant suffisamment fiable pour notre étude.

Nous avons aussi testé un ORO dans le modèle AC avec les mêmes hypothèses que précédemment concernant le choix des valeurs de résistances des lignes et des puissances réactives consommées. Nous avons repris la même procédure générale que celle décrite par les expressions II.21 à II.25, en remplaçant les équations du modèles DC par l'expression des transits dans le modèle AC. Ainsi, l'ORO dans le modèle AC s'écrit de la façon suivante :

Fonction Objectif de l'optimisation : Minimiser les coûts des ajustements de production

$$\text{Min } \sum_i c_i^{+,-} \Delta G_i \quad (\text{II.27})$$

Contraintes de fonctionnement du réseau :

Equations du calcul de répartition de charge

$$\mathbf{P} = \mathbf{P}^0 + \Delta \mathbf{P}_G$$

avec $\mathbf{P}^0$ étant le vecteur des injections nodales nettes initiales (à la sortie du marché)

$$P_{ij} = \frac{|V_i|^2 * \cos \varphi}{|Z_{ij}|} - \frac{|V_i||V_j| * \cos(\theta_i - \theta_j + \varphi)}{|Z_{ij}|} \quad \text{pour chaque transit du réseau}$$

avec

$|Z_{ij}|$ module de l'impédance complexe de la ligne ij

$|V_i|$ module de la tension au nœud i

φ phase de l'impédance de la ligne définie par $\tan \varphi = R / X$

Equilibre des ajustements

$$\sum_i \Delta G_i = 0 \quad (\text{II.29})$$

Limites imposées aux lignes

$$-P_{ij}^{\max} \leq P_{ij} \leq P_{ij}^{\max} \quad (\text{II.30})$$

Limites imposées aux productions

$$P_G^{\min} \leq P_G^0 + \Delta P_G \leq P_G^{\max} \quad (\text{II.31})$$

Nous obtenons des ajustements qui sont très proches de ceux déterminés par l'ORO dans le modèle DC (Tableau II.5) :

TABLEAU II.12 : ajustements déterminés avec un ORO écrit dans le modèle AC

	$\Delta G1^B$	$\Delta G1^T$	$\Delta G2^B$	$\Delta G2^T$	$\Delta G3^B$	$\Delta G3^T$
Volume de l'ajustement (MW) ORO Modèle DC	-67.01	-64.34	+40.26	+24.16	+30.42	+36.51
Volume de l'ajustement (MW) ORO Modèle AC	-70.03	-68.37	+42.11	+25.27	+32.28	+38.74

III.3.6.2) Validation sur le réseau New England IEEE 39 nœuds [PEE 01]

Pour renforcer la validité des hypothèses émises dans le modèle DC, nous avons aussi effectué des tests sur le réseau New England IEEE 39 nœuds, qui est une simplification du réseau de New England dans le nord-est américain. Il comporte 39 nœuds, dont 10 nœuds de production (numérotés de 30 à 39) et 46 lignes (fig. III.9).

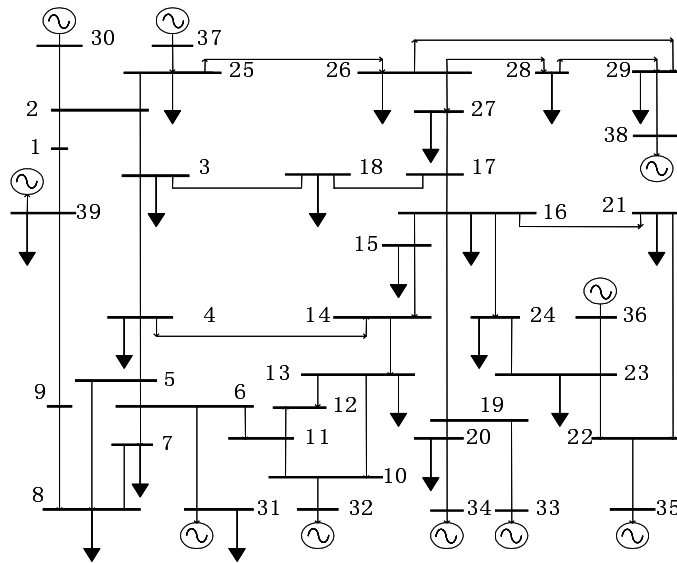


Figure III.9 : réseau IEEE 39 nœuds

Nous avons simulé un cas de congestion sur la ligne reliant le nœud 21 au nœud 22 dont le transit était de 608 MW et que nous avons contraint à 550 MW maximum. Nous avons observé comme pour le réseau 9 nœuds que l'usage de l'ORO dans le modèle DC nous donne un état du réseau très proche de celui calculé avec le modèle AC plus complet. En outre, nous avons de nouveau observé que les ajustements déterminés par un ORO dans le modèle DC et un ORO dans le modèle AC sont très comparables (Fig.III.10).

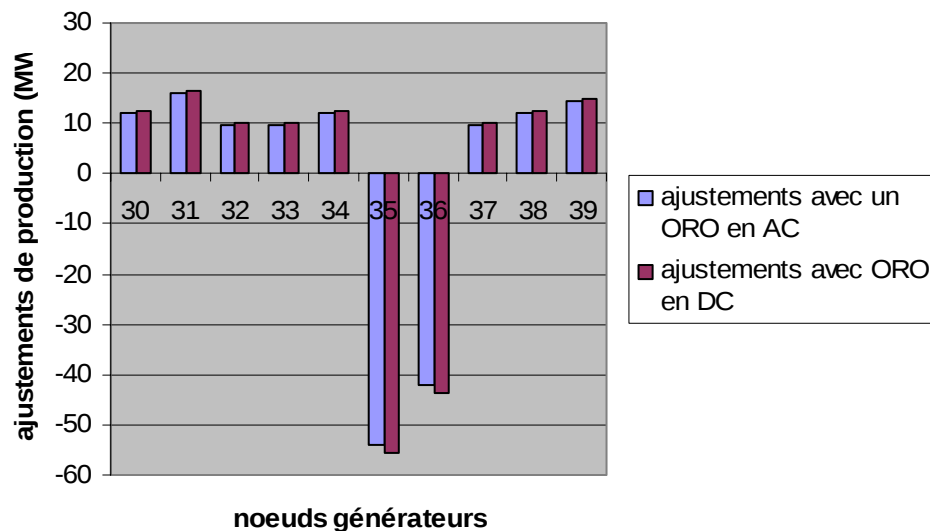


Figure III.10 : ajustements de production avec un ORO dans le modèle DC et dans le modèle AC

Nous avons toutefois noté que l'ORO écrit dans le modèle AC était beaucoup plus lent que celui écrit dans le modèle DC. Ce fait était d'autant plus notable sur le réseau 39 nœuds. En

effet, l'algorithme d'optimisation évalue à chaque itération de l'optimisation le gradient de la fonction objectif par rapport aux paramètres d'entrées (les productions ajustées) et en fonction des contraintes imposées. A chaque évaluation, il doit donc procéder à un nouveau calcul de répartition de charge, qui peut être d'autant plus contraignant que le modèle utilisé est compliqué et que le réseau est grand. La convergence de l'algorithme pour le réseau 39 nœuds était donc bien plus longue la modèle AC qu'avec le modèle DC. Ainsi, nous pouvons conclure sur le fait que l'ORO écrit dans le modèle DC nous assure une précision que nous jugeons suffisante pour notre étude, tout en étant plus rapide d'exécution. Cette rapidité d'exécution peut être utile lorsque l'on passe à un traitement des congestions coordonné qui demande plusieurs optimisations partielles successives (Chapitre V).

III.3.7) Synthèse des résultats obtenus par le buy back, coupures de transactions, et par l'usage de l'OPF

Au vu de l'étude que nous avons menée sur le réseau 9 nœuds où nous avons comparé successivement le traitement des congestions par buy back, coupures de transaction et OPF, il est clairement ressorti que le modèle du buy back est le mieux adapté. Le découplage marché de l'énergie/traitement des congestions introduit par le buy back présente trois principaux avantages :

- Le marché de l'énergie gagne en **transparence**. Avec des méthodes de type coupure de transactions, les souhaits des participants en matière de contrats peuvent être contrecarrés d'autant plus que les congestions sont fréquentes et sévères. Par ailleurs, les paiements dus au traitement des congestions sont *complètement séparés* des paiements du marché de l'énergie. Avec des méthodes de type OPF associé à une tarification marginale, les variations de prix nœaux peuvent être non seulement dues aux congestions sur le réseau, mais aussi au marché même. Il peut donc y avoir un manque de visibilité sur l'origine des variations de prix. De plus, il est difficile de calculer précisément sur une longue période de temps le surplus financier dégagé par la tarification marginale. Les participants sont en outre encouragés à se doter de droits de transport pour se prémunir des fluctuations de volume et de prix sur le marché. Toutefois, de tels outils financiers ne résolvent pas en eux-mêmes un problème technique (congestion).
- Cela est **cohérent** sur le plan technique : par delà les diverses formes contractuelles par lesquelles l'énergie peut se vendre et s'échanger, la réalité physique est unique et bien différente. Un consommateur ayant un contrat avec un fournisseur particulier soutire en réalité une énergie en provenance de divers centres de

production à des taux précisés [MAN1 03]. Couper des transactions pour les garder équilibrées n'a, donc, aucun fondement physique. Un traitement des congestions, si découplé du marché de l'énergie, ne devrait pas exiger que les productions et les consommations impliquées dans une transaction bilatérale ou multilatérale soient équilibrées. Seul l'état de charge du réseau sert de référence dans ce cas. De plus, nous avons mis en évidence les problèmes décisionnels pouvant surgir dans l'application de telles méthodes.

- Le **rôle** du traitement des congestions est clairement défini : le but de l'opérateur du système est d'accommoder les souhaits des participants tout en mettant en œuvre des moyens pour garantir la sécurité du réseau. Le traitement des congestions est un problème technique, tout comme l'équilibre production/consommation ou le réglage de la tension. Plutôt que de tenter de modifier ce qui a été conclu sur le marché de l'énergie, ou de permettre un profit au bénéfice des propriétaires du système, il doit, donc être inclut dans un *service système* complètement découplé de celui-ci. Les sorties du marché de l'énergie peuvent donc rester *confidentielles*. Ce service constitue aussi un *marché à part entière*, car il permet à des fournisseurs qualifiés pour ce service d'entrer en compétition en déposant leurs offres. Cette façon de faire est clairement plus transparente aux yeux des participants, car elle fait bien la différence entre l'achat/vente d'un produit (ici l'électricité) et les problèmes techniques liés au transport de ce produit.

En outre, la méthode du buy back nous permet d'espérer des coûts de congestion raisonnables (si les marges incluses dans les offres d'ajustement des producteurs sont modérées), contrairement à l'usage de l'OPF qui peut amener à des coûts de court terme relativement élevés.

Ceci dit, des problèmes peuvent surgir avec la méthode du buy back. Le premier est le manque d'offres d'ajustement. Les offres étant faites sur une base volontaire, un producteur peut choisir de ne pas déposer d'offres d'ajustement pour le traitement des congestions. Il peut alors y avoir des cas où l'opérateur du système se retrouve avec des ressources insuffisantes pour résoudre une situation de congestion. Ce point a aussi été soulevé par Tao et Gross [TAO 02]. Dans ces cas-là, l'opérateur du système devrait en tout état de cause bénéficier de l'autorité nécessaire pour solliciter les services d'un participant dont il connaît la disponibilité, et ce au nom de la sécurité du réseau.

Un autre défaut possible d'une telle approche est qu'elle peut permettre un comportement spéculatif venant des producteurs si on n'y prend pas garde. En effet, les producteurs peuvent faire deux offres successives, une pour le marché de l'énergie et une pour le traitement des congestions. Certains producteurs peuvent s'entendre de façon à engendrer artificiellement

des congestions pour faire du profit sur les contraintes du système [SEE 99]. Ce sont des comportements dits *spéculatifs*. Or, le traitement des congestions doit être avant tout un service rendu par des participants qualifiés, et ce pour le bien de tous. Il ne doit pas être un outil de spéculation.

En Californie, les offres d'ajustement sont limitées en prix à 750 \$ pour contrer le pouvoir de marché de certains producteurs bien placés sur le réseau. Toute offre qui dépasse le prix de 750 \$ est rejetée [CAI 03]. Cependant, imposer des limitations de prix pour les offres d'ajustement n'est pas une démarche suffisante en soi pour éviter les comportements spéculatifs. Il faut aussi une stratégie judicieuse d'allocation des coûts de congestion. La plupart du temps, les coûts des ajustements sont récupérés au prorata parmi les usagers comme cela se fait au Royaume-Uni sans véritable stratégie d'allocation. Malgré la simplicité de la démarche, celle-ci a le défaut de ne pas cibler correctement les responsabilités en cas de congestion. Cela a pour conséquence de ne pas décourager de façon efficace l'occurrence des scénarios congestionnels. D'autre part, l'absence de signaux économiques rend le traitement des congestions plus opaque et n'incite pas ceux qui ont particulièrement besoin de s'intéresser au développement du réseau. Or, il est largement reconnu qu'une méthode de traitement des congestions efficace doit envoyer des signaux économiques adéquats.

La suite de notre étude va se baser à présent exclusivement sur le traitement des congestions par la méthode du buy back. Dans le prochain chapitre, nous allons examiner les différentes stratégies d'allocation des coûts de congestion que l'on peut développer pour rendre ce modèle plus efficace en terme de transparence et de prévention des comportements spéculatifs.

CHAPITRE IV

Stratégies d'allocation des coûts de congestion basées sur la traçabilité de l'énergie

IV.1) Introduction

Dans ce chapitre, nous allons traiter de l'allocation des coûts de congestion sur une base physique résultant des ajustements de production (modèle du buy back). A partir d'un outil de traçabilité de l'énergie, on peut décliner plusieurs stratégies d'allocation innovantes dont on peut évaluer l'efficacité par rapport aux objectifs suivants : juste responsabilisation des usagers du réseau quant aux choix établis dans le marché de l'énergie, découragement des scénarios congestionnels et intérêt au développement du réseau. Il est nécessaire que l'allocation des coûts de congestion réponde à ces objectifs dans le nouveau contexte libéralisé du secteur de l'électricité. En effet, les consommateurs ont naturellement tendance à privilégier le prix de leur fourniture comme critère de choix de fournisseur plutôt que les conséquences techniques que peut avoir ce choix sur le réseau. De même, du fait de la séparation des activités de production et des activités de gestion de l'infrastructure de transport, l'opérateur du système et les producteurs n'ont plus les mêmes objectifs : l'opérateur du système a pour principal souci la conduite du réseau dans des conditions de sécurité satisfaisantes, tandis que les producteurs cherchent prioritairement à vendre la plus grande quantité d'énergie possible au prix le plus élevé. Enfin, il existe une tendance de plus en plus grande dans les marchés libéralisés à ce que les incitations au développement du réseau viennent des usagers même de ce réseau. Dans ce contexte, l'envoi de signaux économiques corrects peut amener de bonnes décisions d'investissements sur le réseau.

L'apparition récente de techniques de traçabilité de l'énergie [YAN 99] permet de déterminer aisément les taux d'utilisation des lignes de chaque usager du réseau de transport, et aussi l'origine de l'énergie consommée en chaque point du réseau. La méthode de traçabilité que nous avons utilisé, appelée Méthode des Images de Charges (MIC) [MAN1 03], [MAN2 03] se base exclusivement sur le modèle de calcul de répartition de charge.

Nous allons tout d'abord récapituler les principales méthodes d'allocation des coûts de congestion qui sont utilisées de nos jours dans les marchés de l'électricité. Ensuite, nous allons introduire la notion de *traçabilité de l'énergie* sur laquelle nous allons développer des stratégies d'allocation des coûts de congestion. Ces stratégies seront illustrées sur le cas que nous avons étudié au chapitre précédant et comparées à d'autres méthodes d'allocation ayant été récemment proposées et se basant sur l'usage des facteurs de distribution (modèle DC).

IV.2) Contexte et enjeux de l'allocation des coûts de congestion

Il existe depuis le début de la dérégulation trois types d'approches traditionnelles qui sont appliqués pour l'allocation des coûts de congestion :

- L'approche **nodale** par la méthode des prix nodaux dont nous avons parlé au chapitre précédent: le coût de congestion alloué au transfert d'1 MW d'un nœud à un autre du réseau est égal à la différence de prix entre ces deux nœuds. Cette propriété est destinée à envoyer des signaux économiques aux participants en présence et pour guider la localisation de nouveaux producteurs ou consommateurs. Comme nous l'avons déjà noté, elle souffre d'un manque de transparence quant à l'origine des fluctuations de prix (congestions, mécanismes de marché, etc...)
- L'approche **zonale** par régionalisation du marché (*Market Splitting*) : comme nous l'avons déjà dit, cette approche est difficilement applicable pour les réseaux fortement interconnectés. On peut aussi citer de nouveau l'approche zonale californienne, qui est très spécifique au fonctionnement du marché californien, ou l'approche scandinave
- le **timbre-poste** : les coûts de congestion sont alloués aux consommateurs sur la base du prorata. C'est une approche simple et très utilisée dans les marchés de l'électricité, mais elle ne reflète pas nécessairement l'usage de chacun du réseau et n'envoie donc pas les signaux économiques appropriés.

En Europe, un effort particulier est fait ces derniers temps pour encourager les différents opérateurs du système à fournir plus d'informations sur l'état des réseaux électriques. En effet, jusqu'à présent, ces informations sont jugées encore insuffisantes, et il y a un besoin de plus en plus grand pour les acteurs du marché de faire la lumière sur la réalité et les causes exactes des congestions [CRE 02]. De telles questions sont posées de la part d'acteurs du marché qui ont été particulièrement pénalisés par l'existence de congestions et qui désirent un traitement le moins arbitraire possible. Les techniques utilisées pour résoudre les congestions (ajustements de production) ne semblent pas être contestées ; par contre, de nombreuses discussions à l'heure actuelle portent sur la *transparence* de l'attribution de coûts de congestion à chaque acteur [CRE 03]. Il est donc nécessaire aujourd'hui de disposer d'un outil

d'allocation des coûts de congestion transparent, efficace, et qui envoie des signaux économiques appropriés à chaque usager du réseau en vue de réduire les congestions, voir d'inciter au développement du réseau par les usagers concernés. Le traitement des congestions étant surtout un problème de nature stratégique, la question de l'allocation des coûts de congestion est reconnu comme étant cruciale et de haute importance [RAU 00].

L'allocation du coût de congestion doit être beaucoup plus qu'une simple « participation aux frais » comme l'est le timbre-poste. Elle devrait répondre à plusieurs objectifs précis :

- Rendre compte de *l'impact physique* de chaque participant sur le réseau : cet impact n'est bien sûr pas déterminé par le moyen choisi pour vendre ou acheter de l'électricité (bourse ou contrat bilatéral) mais par des lois physiques.
- Distribuer des *payements équitables* qui rendent compte de ces impacts : beaucoup de contestations à l'heure actuelle viennent de règles d'allocation établies sans base solide. Cela dénote le besoin *d'un outil clair et impartial* qui dégage des responsabilités de façon indiscutable.

Sur la base des impacts dégagés, on va fournir aux participants des *signaux* pour endiguer le phénomène de congestion sur le long terme. Il s'agit surtout de décourager les comportements spéculatifs et les scénarios congestionnels. Les données techniques sur l'origine des congestions peuvent être notamment utiles à l'opérateur du système pour le placement de nouveaux producteurs ou pour le développement du réseau.

IV.3) Allocations des coûts de congestion basées sur la traçabilité de l'énergie

IV.3.1) la traçabilité de l'énergie

La traçabilité de l'énergie nous permet de déterminer :

- Le taux de participation de chaque production à chaque consommation (Fig. IV.1)
- Les liaisons transits-productions ou transits-consommations (Figure IV.2a et IV.2b)

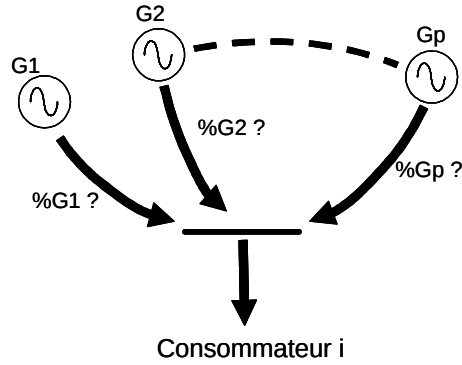


Figure IV.1 : consommation et taux de participation des productions associés

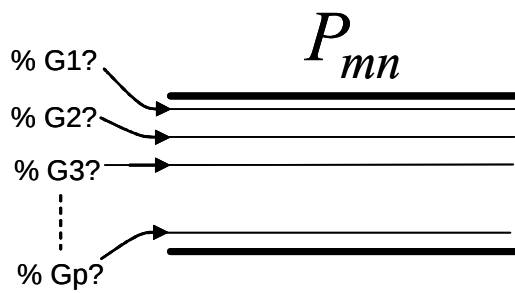


Figure IV.2a : Liaisons productions-transits

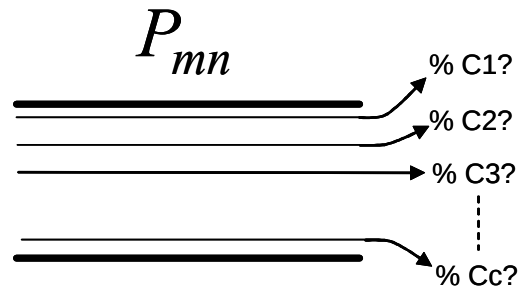


Figure IV.2b : Liaisons consommations-transits

Récemment, des méthodes de traçabilité basées sur une analyse topologique des transits de puissance ont été proposées [BIA 97], [KIR 98]. Elles se basent sur un principe de division proportionnelle qui est un postulat non démontrable physiquement. L'avantage de la MIC est qu'elle s'appuie exclusivement sur le modèle du calcul de répartition de charge et sur le principe de l'agrégation des productions et charges au même noeud. Originellement, la MIC a été développée à partir du modèle AC [MAN1 03], [MAN1 03]. Pour notre analyse, nous avons utilisé une version simplifiée de la MIC qui est basée sur le modèle DC. La description complète de la MIC dans le modèle DC est fournie en Annexe 3. A partir des données fournies par la MIC, nous pouvons donc déterminer :

- Les contributions des productions nettes aux consommations nettes

$$\mathbf{P}_G = \mathbf{C} * \mathbf{P}_C \quad (\text{IV.1})$$

où $\mathbf{C}$ est une matrice reliant les productions nettes aux consommations nettes

- Les liaisons productions-transits. On peut définir un Facteur d'Utilisation de Ligne (FUL) pour chaque production connectée à un nœud i et chaque ligne l

congestionnée. Il exprime le rapport entre l'apport P_l^i d'une production connectée à un nœud i à un transit d'une ligne l et la puissance totale P_l transitée par cette ligne:

$$F_l^i = \frac{P_l^i}{P_l} \quad (\text{IV.2})$$

où F_l^i est le FUL d'une production connectée au nœud i par rapport à la ligne l

On peut alors exprimer tout transit d'une ligne l comme étant la somme des apports individuels des productions

$$P_l = \sum_{i=1}^g F_l^i P_l \quad (\text{IV.3})$$

- Les liaisons consommations-transits. On peut comme pour les productions définir un FUL pour chaque charge connectée à un nœud j et chaque ligne l congestionnée. Il exprime le rapport entre l'apport d'une charge connectée à un nœud j à un transit d'une ligne l et le volume total transité par cette ligne :

$$E_l^j = \frac{P_l^j}{P_l} \quad (\text{IV.4})$$

où E_l^j est le FUL d'une production connectée au nœud j par rapport à la ligne l

On peut de même exprimer tout transit d'une ligne l comme étant la somme des apports individuels des charges :

$$P_l = \sum_{j=1}^c E_l^j P_l \quad (\text{IV.5})$$

IV.3.2) Formulations des allocations

Soit un état du réseau donné avec des productions et des charges nodales nettes et avec n congestions à traiter par des ajustements de production (méthode du buy back). Nous définissons le coût total de congestion Π de cette façon :

$$\Pi = \sum_i c_i^{+, -} \Delta G_i \quad (\text{IV.6})$$

On peut imaginer deux types d'approche basée sur les informations délivrées par la traçabilité de l'énergie pour l'allocation de coût de congestion:

- L'approche « *statique* » : on se réfère à *un état* du réseau, celui à la sortie du marché de l'énergie (avant traitement des congestions). On alloue le coût Π en fonction des taux d'utilisation des lignes congestionnées de chaque usager.
- L'approche « *différentielle* » : on se réfère à *deux états* du réseau (pré et post congestion), en comparant les taux délivrés par la traçabilité (taux des productions dans les consommations, taux dans les transits) avant et après traitement des congestions.

Nous allons dans la section suivante donner plus de détails sur ces approches et leur formulation.

IV.3.2.1) Approches statiques : allocations basées sur les contributions au transit

Dans un cas général où l'on est en présence de n congestions, le coût total Π est réparti entre les n lignes congestionnées de cette façon :

$$\Pi = \sum_{l=1}^n \Pi_l = \sum_{l=1}^n \psi_l \Pi \quad (\text{IV.7})$$

Les facteurs ψ_l peuvent être théoriquement choisis des façons suivantes :

- On peut choisir de répartir le coût uniformément sur toutes les lignes. Tous les ψ_l auront alors comme valeur $1/n$.
- On peut choisir de répartir le coût suivant la surcharge ΔP_l à éliminer. Alors, les ψ_l seront dans ce cas calculés de la manière suivante :

$$\psi_l = \frac{\Delta P_l}{\sum_l \Delta P_l} \quad (\text{IV.8})$$

- On peut aussi choisir de répartir le coût suivant le multiplicateur de Lagrange associé à la contrainte de transit sur la ligne l . C'est une méthodologie qui a été proposée par Singh. Les ψ_l seront dans ce cas calculés de la façon suivante :

$$\psi_l = \frac{\mu_l}{\sum_l \mu_l} \quad (\text{IV.9})$$

avec μ_l étant le multiplicateur de Lagrange associé à la contrainte de transit sur la ligne l tiré de la relation (II.19). Ce multiplicateur a une signification physique : c'est le coût épargné si la limite de puissance transitée était augmentée d'1 MW supplémentaire (Gedra, 1999). Cette solution a donc l'avantage de donner un signal économique.

On subdivise de nouveau le coût entre les catégories visées (producteurs et consommateurs) en introduisant un facteur de répartition α que l'on peut choisir entre 0 et 1. Ainsi, le coût associé à chaque ligne l congestionnée peut être exprimé comme la somme du coût alloué aux producteurs et aux consommateurs pour la congestion sur cette ligne :

$$\Pi_l = \Pi_l^P + \Pi_l^C = \alpha \Pi_l + (1 - \alpha) \Pi_l \quad (\text{IV.10})$$

avec Π_l^P étant le coût adressé aux producteurs pour la congestion sur la ligne l , et Π_l^C le coût adressé aux consommateurs pour la congestion sur la ligne l . Notons que le choix du facteur α est de nature purement stratégique. Nous discuterons plus loin de l'efficacité des signaux économiques selon la catégorie à laquelle ils sont envoyés.

Coût alloué aux usagers pour chaque ligne l congestionnée :

Le coût de congestion alloué aux producteurs pour une ligne l congestionnée est distribué au prorata des apports de chaque production sur le transit de la ligne l :

$$Pay_i^l = \frac{F_l^i}{\sum_{i \in \Omega_G} F_l^i} * \Pi_l^P = \frac{F_l^i}{\sum_{i \in \Omega_G} F_l^i} * \alpha \Pi_l \quad (\text{IV.11})$$

où Ω_G est l'ensemble des productions nodales dont la contribution au transit de la ligne l va dans le même sens que le flux dans la ligne l . Il est en effet possible dans un cas général qu'une production envoie un contre-transit dans une ligne l . Il peut en outre y avoir discussion sur la possibilité de rémunérer ces contre-transits qui tendent à soulager la contrainte sur une ligne congestionnée. Toutefois, ce principe peut dans ce cas mener à des difficultés si un usager paye pour une congestion, mais reçoit une rémunération pour une autre. Nous choisissons donc ici de n'allouer ni coût, ni rémunération aux contre-transits.

De même, le coût de congestion alloué aux consommateurs pour une ligne l congestionnée est distribué au prorata des apports de chaque charge sur le transit de la ligne l :

$$Pay_j^l = \frac{E_l^j}{\sum_{j \in \Omega_C} E_l^j} * \Pi_l^C = \frac{E_l^j}{\sum_{j \in \Omega_C} E_l^j} * (1 - \alpha) \Pi_l \quad (IV.12)$$

où Ω_C est l'ensemble des charges nodales dont la contribution au transit de la ligne l va dans le même sens que le flux dans la ligne l .

Ainsi, nous pouvons exprimer le coût total Π pour n congestions en fonction des paiements attribués aux consommateurs et aux producteurs :

$$\Pi = \sum_{l=1}^n \sum_{i=1}^g Pay_i^l + \sum_{l=1}^n \sum_{j=1}^c Pay_j^l = \sum_{l=1}^n \sum_{i=1}^g \frac{F_l^i}{\sum_{i \in \Omega_G} F_l^i} * \alpha \psi_l \Pi + \sum_{l=1}^n \sum_{j=1}^c \frac{E_l^j}{\sum_{j \in \Omega_C} E_l^j} * (1 - \alpha) \psi_l \Pi$$

(IV.13)

IV.3.2.2) Approche différentielle : allocations par destination des ajustements de production

Cette approche est complètement différente de la précédente. On n'alloue pas le coût en passant par les transits congestionnés, mais en adressant directement les ajustements de production aux consommations à l'aide des liaisons productions-consommations. On a donc besoin de l'état du réseau avant et après traitement des congestions.

Soit un vecteur d'ajustements $\Delta G^{+,-}$ décidés pour résoudre une situation de congestion. Ces ajustements peuvent s'exprimer en fonction des consommations grâce aux liaisons productions-consommations :

$$\Delta G^{+,-} = \mathbf{P}_G^{\text{post}} - \mathbf{P}_G^{\text{pré}} = (\mathbf{C}^{\text{post}} - \mathbf{C}^{\text{pré}}) * \mathbf{P}_C \quad (IV.14)$$

où $\mathbf{P}_G^{\text{pré}}$ et $\mathbf{P}_G^{\text{post}}$ sont les vecteurs de productions nodales avant et après traitement des congestions, et $\mathbf{C}^{\text{pré}}$ et $\mathbf{C}^{\text{post}}$ sont les matrices de liaison productions-consommations avant et après traitement des congestions.

La relation (IV.14) met donc directement en relation les ajustements de production avec les consommations et permet d'en connaître la destination. Ainsi, le coût alloué à l'ensemble des consommateurs connectés à un nœud j :

$$Pay_j = \sum_{i=1}^{ga} (C_{i,j}^{post} - C_{i,j}^{pré}) * P_{Gj} * c_i^{+,-} \quad (IV.15)$$

où ga est le nombre de producteurs ayant été ajustés (qui peut être inférieur ou égal au nombre total de producteurs connectés au réseau) . Notons qu'il est possible dans un cas général que le paiement d'un consommateur soit en fait une rémunération.

IV.4) Résultats numériques sur le réseau 9 nœuds

IV.4.1) Données fournies par la traçabilité de l'énergie

Nous reprenons le cas du réseau 9 nœuds étudié au chapitre précédent. La figure IV.3a et IV.3b nous rappellent les injections nodales et les transits avant et après traitement des congestions par la méthode du buy back:

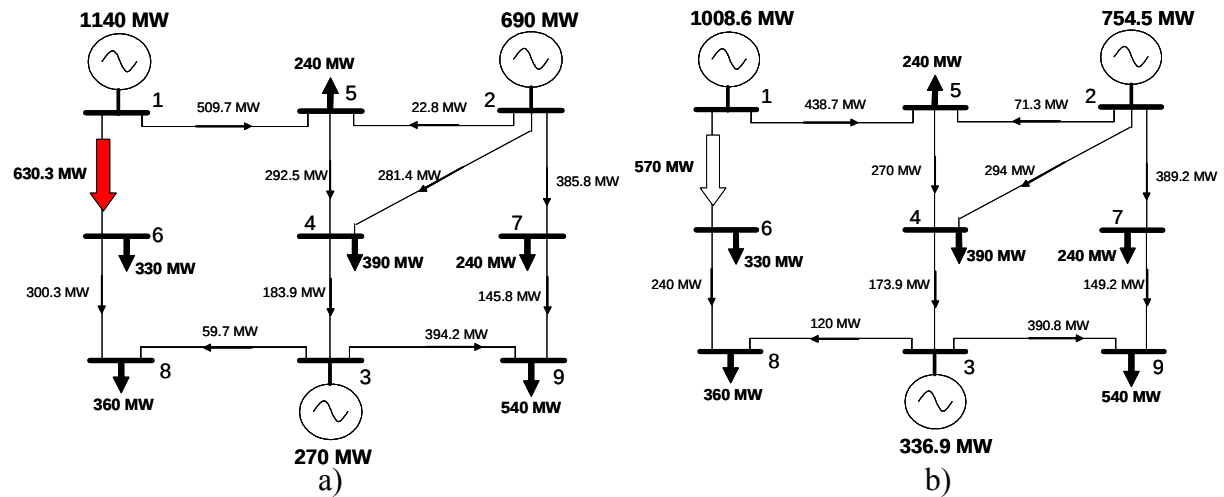


Figure IV.3a et IV.3b : Etat du réseau avant traitement des congestions a) et après traitement des congestions b)

A l'aide de la MIC, nous pouvons aisément déterminer l'origine de la puissance consommée dans chaque nœud de charge et les taux de présence de chaque usager sur chaque ligne du réseau. Les Figure IV.4 et IV.5 présentent les liaisons productions-consommations avant et après traitement des congestions. Nous rappelons que dans le réseau 9 noeuds que

nous avons choisi d'étudier, il y a deux producteurs à chaque nœud producteur et deux consommateurs à chaque nœud de charge. Pour simplifier la présentation des résultats, l'expression « G1 » désigne l'ensemble des deux producteurs connectés au nœud 1. De même, l'expression « C4 » désigne l'ensemble des deux consommateurs connectés au nœud 4 :

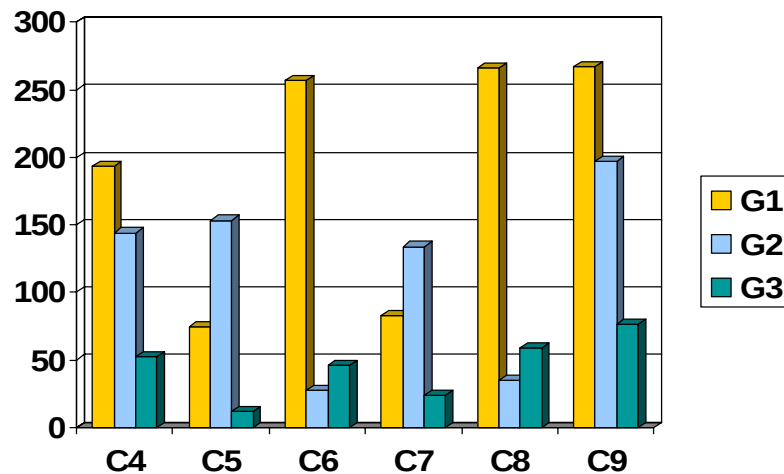


Figure IV.4 : apports des productions aux consommations (MW) dans l'état avant traitement des congestions

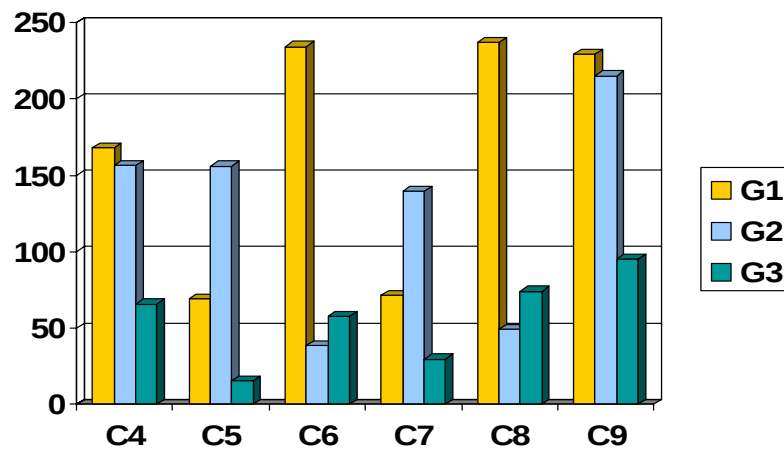


Figure IV.5 : apports des productions aux consommations (MW) dans l'état après traitement des congestions

Les données de la MIC nous ont permis aussi de déterminer le taux de présence de chaque usager sur le transit de la ligne 1-6 congestionnée. La Figure IV.6 nous donne le FUL (ramené en pourcent) de chaque charge nodale sur le transit de la ligne 1-6. De même, la Figure IV.7 nous donne les taux de participations des productions nodales au transit de cette ligne :

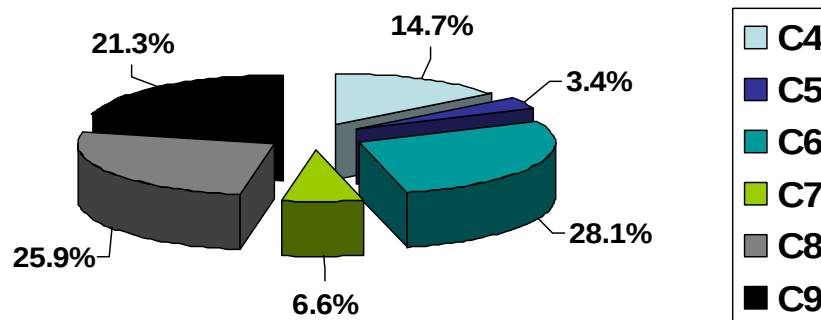


Figure IV.6 : apports des charges nodales au transit sur la ligne 1-6

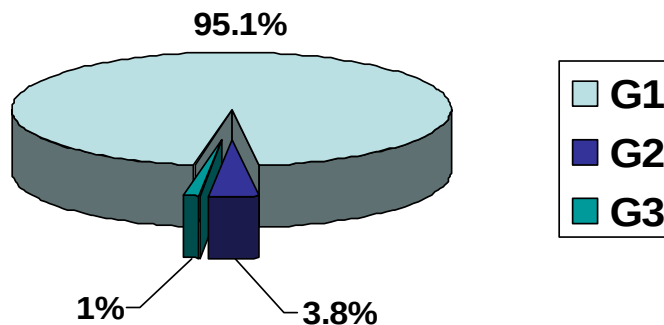


Figure IV.7 : apports des productions nodales au transit sur la ligne 1-6

IV.4.2) Résultats des allocations sur le réseau 9 nœuds basées sur la traçabilité

IV.4.2.1) Allocation du coût de la congestion sur la ligne 1-6

Dans le chapitre précédent, nous avons résolu la congestion sur la ligne 1-6 avec la méthode du buy back en ayant recours à des ajustements de production dont le coût total s'élevait à 2003.3 €/h. Ce coût est à la charge de l'opérateur du système, qui doit le récupérer sur les usagers du réseau. Nous avons alors laissé en suspens la question de savoir comment il pouvait facturer ce coût aux usagers.

A présent, nous pouvons appliquer les stratégies d'allocations basées sur la traçabilité de l'énergie que nous avons explicitées §IV.3.2. Nous allons illustrer l'allocation basée sur la contribution des usagers au transit (approche statique) pour deux cas précis :

- On alloue le coût entièrement aux consommateurs. On a alors $\alpha = 0$
- On alloue le coût entièrement aux producteurs. On a alors $\alpha = 1$

Les cas où l'on choisit une allocation mixte producteurs-consommateurs sont une simple modulation des deux cas précédents, nous n'allons pas les illustrer pour plus de concision. D'autre part, les allocations basées sur la traçabilité sont comparées à une allocation traditionnelle de type timbre-postale. Dans cette dernière, le coût est simplement alloué au prorata des puissances injectées ou soutirées sur le réseau. En outre, les paiements sont présentés sous deux formes différentes : d'une part, nous avons un paiement global en €/h qui est le paiement associé à l'ensemble de la production G_i ou de la consommation C_i en un nœud particulier. Ensuite, ce paiement peut être ramené à un paiement à l'unité de puissance injectée ou consommée en chaque nœud (paiement en €/MWh). Les Figures IV.8 et IV.9 présentent l'ensemble des résultats concernant l'allocation du coût aux consommateurs seulement ($\alpha = 0$) :

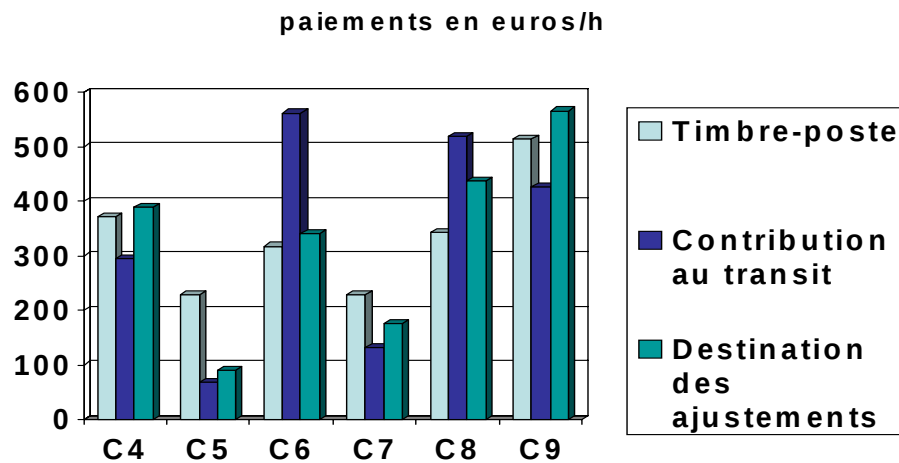


Figure IV.8 : allocation du coût de la congestion sur la ligne 1-6 aux consommateurs par timbre-poste, contribution au transit et destination des ajustements. Paiements à chaque nœud consommateur dans son ensemble (€/h)

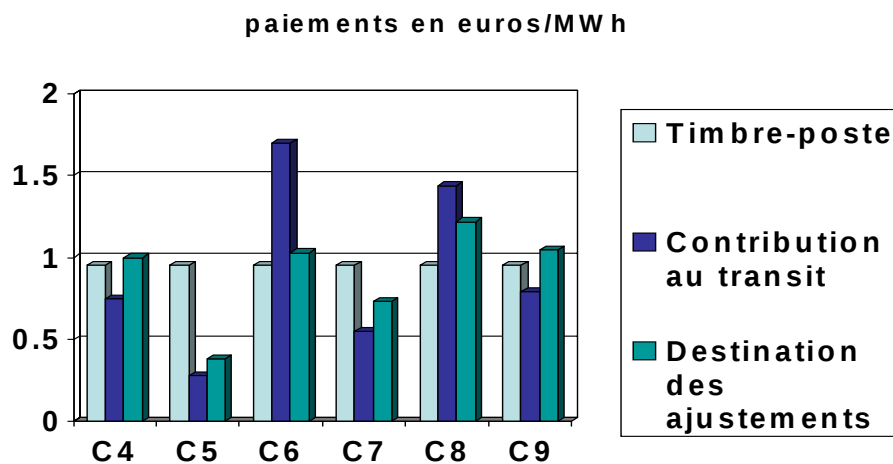


Figure IV.9 : allocation du coût de la congestion sur la ligne 1-6 aux consommateurs par timbre-poste, contribution au transit et destination des ajustements. Paiements ramenés à l'unité de puissance consommée en chaque nœud consommateur (€/h)

Lorsque nous examinons les Figures IV.8 et IV.9, nous constatons que l'allocation du coût de congestion suivant la contribution au transit offre le plus gros écart avec la méthode du timbre-poste. Parmi les méthodes basées sur la traçabilité, la méthode par destination des ajustements semble donner ici des signaux moins prononcés. Par exemple, les responsabilités au niveau purement physique de C6 et C8 sur le transit de la ligne 1-6 sont mieux mis en exergue dans l'allocation par contribution au transit que par l'allocation par destination des ajustements. Nous pouvons néanmoins observer la relative cohérence mutuelle des deux allocations basées sur la traçabilité de l'énergie. La cohérence des résultats obtenus s'observe aussi lorsque l'on compare la distribution des paiements alloués avec la localisation géographique des charges vis-à-vis de la congestion sur la ligne 1-6. Des charges situées loin ou en amont de la congestion (C5 et C7) payent relativement peu tandis que celles situées en aval du flux congestionné (C6, C8 et C9) payent plus. Les signaux économiques envoyés sont même très comparables aux signaux envoyés par la méthode des prix nodaux (voir Tableau III.10). Ainsi, on peut obtenir à l'aide de la traçabilité de l'énergie des signaux économiques cohérents avec les méthodes déjà appliquées, tout en ayant l'avantage de la séparation du marché de l'énergie et du traitement des congestions propre à la méthode du buy back.

Les Figures IV.10 et IV.11 présentent les résultats concernant l'allocation du coût de congestion pour le cas où on le destinerait uniquement aux producteurs ($\alpha = 1$). De la même façon que pour les consommateurs, l'allocation basée sur la contribution au transit est mise en parallèle avec une répartition timbre-postale.

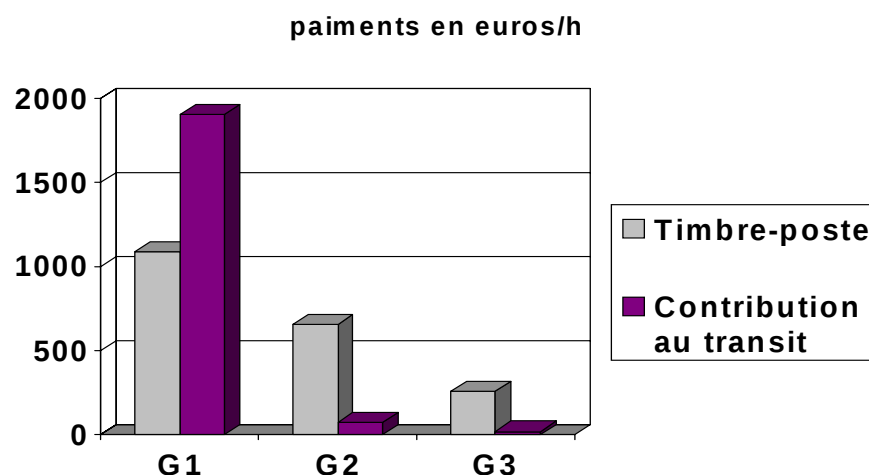


Figure IV.10 : allocation du coût de la congestion sur la ligne 1-6 aux producteurs par timbre-poste et contribution au transit. Paiements à chaque nœud consommateur dans son ensemble (€/h)

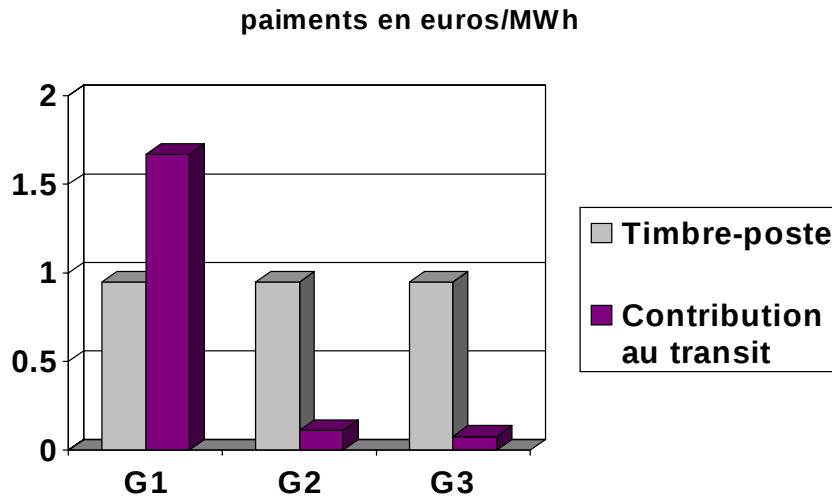


Figure IV.11 : allocation du coût de la congestion sur la ligne 1-6 aux producteurs par timbre-poste et contribution au transit. Paiements ramenés à l'unité de puissance consommée en chaque nœud consommateur (€/h)

Nous pouvons constater que pour le cas étudié les signaux économiques envoyés par la traçabilité de l'énergie sont encore plus visibles si on alloue le coût de congestion aux producteurs. Les producteurs situés au nœud 1 sont quasiment entièrement responsable de la congestion sur la ligne 1-6 du point de vue technique. Répartir les coûts de congestion au prorata des puissances injectées ne reflète pas correctement dans ce cas la réalité physique, et s'avère donc plus ou moins injuste.

D'autres stratégies d'allocations « physiques » des coûts de congestion ont été récemment proposées dans la littérature [PAN 00]. Ces méthodes se basent sur les facteurs de distribution tirés du modèle DC. Nous allons dans la section IV.4.3 appliquer au cas étudié deux de ces méthodes qui ont retenu notre attention.

IV.4.2.2) Cas présentant plusieurs lignes congestionnées

Dans cette partie, nous allons analyser la distribution du coût de congestion lorsque l'on est en présence de plusieurs lignes congestionnées. Nous allons nous concentrer sur l'analyse de la subdivision du coût total sur les différentes lignes congestionnées. Nous avons étudié le cas où l'on a la congestion sur la ligne 1-6 (la même que celle précédemment étudiée) et une congestion sur la ligne 5-4. Pour cette dernière, nous avons abaissé sa limite de transit à 280 MW, son transit prévisionnel à la clôture du marché de l'énergie étant de 292.5 MW (voir Fig.IV.3a). Nous avons ensuite traité cette situation de congestion en ayant recours aux offres d'ajustements des producteurs telles qu'elles sont définies dans le Chapitre III. Le transit de la ligne 1-6 après traitement des congestions a été ramené à sa limite de 570 MW et celui de la ligne 5-4 à 270 MW. Nous avons ensuite distribué le coût de congestion sur les deux lignes congestionnées suivant la relation (IV.7). Nous avons testé successivement une distribution

uniforme, au prorata de la surcharge, et suivant les multiplicateurs de Lagrange μ associée aux contraintes de transit des lignes congestionnées. Le résultat de cette distribution est donné par le Tableau IV.1 :

Tableau IV.1 : distribution du coût de congestion sur les deux lignes congestionnées

Coût réparti uniformément		Coût réparti au prorata de la surcharge ΔP_i		Coût réparti suivant les multiplicateurs μ	
Ligne 1-6	Ligne 5-4	Ligne 1-6	Ligne 5-4	Ligne 1-6	Ligne 5-4
1001.5	1001.5	1335	668	2003	0

Les ajustements trouvés ici sont exactement les mêmes que ceux trouvés pour le traitement de la ligne 1-6 seule. Le coût de congestion est donc égal à 2003.3 €/h, comme celui trouvé au Chapitre III. Cela s'explique par le fait que seule la ligne 1-6 est réellement contraignante dans ce cas. En effet, même si il n'y avait congestion que sur la ligne 1-6 seule, le transit de la ligne 5-4 aurait été de toute manière ramené à 270 MW (voir Fig. IV3b). Ceci explique la distribution du coût suivant les multiplicateurs de Lagrange que nous pouvons observer, où le multiplicateur μ_{1-6} de la ligne 1-6 est positif, tandis que celui de la ligne 5-4 est nul.

Nous faisons à présent un nouvel essai en ramenant transit maximum permis de la ligne 5-4 à 250 MW de manière à rendre la congestion sur cette ligne contraignante pour l'algorithme d'optimisation. Avec les mêmes offres d'ajustement que précédemment, nous obtenons cette fois un coût de congestion de 3130 €/h. Le résultat de la subdivision de ce coût est donné par le Tableau IV.2 :

Tableau IV.2 : distribution du coût de congestion sur les deux lignes congestionnées

Coût réparti uniformément		Coût réparti au prorata de la surcharge ΔP_i		Coût réparti suivant les multiplicateurs μ	
Ligne 1-6	Ligne 5-4	Ligne 1-6	Ligne 5-4	Ligne 1-6	Ligne 5-4
1565	1565	1832	1298	423	2706

La subdivision du coût suivant les multiplicateurs de Lagrange a encore particulièrement attiré notre attention. Dans ce cas, la congestion sur la ligne 1-6 est moins contraignante comparativement à la congestion sur la ligne 5-4. En effet, toute augmentation de la contrainte sur la ligne 5-4 a à présent un impact fort sur la fonction objectif du problème d'optimisation, ce qui est reflété par la valeur du multiplicateur μ_{5-4} associé. Les ajustements obtenus sont donc à présent surtout influencés par la contrainte sur la ligne 5-4. On a pu vérifier ce fait en déterminant les ajustements nécessaires à la résolution de la congestion sur la ligne 5-4 seule. Nous avons observé que le vecteur des ajustements que nous avons obtenu se rapprochait beaucoup de celui obtenu lorsque l'on a congestion sur les deux lignes.

Distribuer le coût sur les lignes congestionnées uniformément ou au prorata de la surcharge a l'avantage d'être simple, mais ne se justifie pas objectivement au niveau physique. La répartition du coût suivant les multiplicateurs μ peut refléter correctement l'influence de chaque contrainte sur le coût de congestion et donc donner des signaux plus appropriés, même si la méthodologie peut sembler plus complexe, et ses résultats moins prévisibles.

IV.4.3) Autres allocations physiques : résultats de méthodes basées sur les facteurs de distribution

IV.4.3.1) Allocation aux consommateurs par les facteurs de distribution

Considérons les facteurs de distribution donnés par la matrice **A** calculée pour tout réseau maillé (Annexe 1). Dans la méthode de Singh [SIN 98], les coûts de congestion sont alloués directement aux consommateurs sur la base de ces facteurs. Ainsi, les coûts de congestion sont répartis de cette façon pour chaque ligne congestionnée :

$$Pay_j^l = \frac{A_{l,j} * P_{Cj}}{\sum_{j \in \Omega_C} A_{l,j} * P_{Cj}} * \Pi_l \quad (IV.16)$$

où $A_{l,j}$ est le facteur de distribution reliant la charge P_{Cj} connectée au nœud j au transit de la ligne l , et où Ω_C est l'ensemble des charges dont la contribution calculée par le facteur de distribution associé va dans le même sens que le transit de la ligne l .

Une allocation individuelle avec les facteurs de distribution est problématique, car ils dépendent fortement du choix du nœud bilan. Nous avons appliqué cette méthode d'allocation au cas de la congestion sur la ligne 1-6 pour trois nœuds bilan différents : nœud 1, nœud 2 et nœud 3. Les résultats de cette allocation sont présentés dans le Tableau IV.1, où on peut noter

que le choix du nœud bilan influe de façon importante sur les résultats de l'allocation. Les résultats de cette allocation sont relativement proches de ceux obtenus par traçabilité si l'on choisit le nœud 1 comme nœud bilan. Toutefois, si l'on choisit d'autres nœuds bilan, tels que le nœud 2 ou le nœud 3, les résultats divergent fortement. Cela vient de l'interprétation physique de ces facteurs de distribution. En effet, le facteur de distribution $A_{l,i}$ représente le changement de transit de la ligne l si on modifie l'injection de +1MW au nœud i , avec une variation d'injection opposée située au nœud bilan. Ces facteurs varient donc fortement en fonction du choix du nœud bilan. A côté de cela, les FUL obtenus par la Méthode des Images de Charge ont l'avantage d'être uniques, et ils représentent la part réellement attribuable à toute injection ou charge nodale. Grâce à cette propriété, les allocations basées sur un outil de traçabilité de l'énergie sont plus « robustes ».

Tableau IV.1 : allocation aux consommateurs du coût de la congestion sur la ligne 1-6 suivant les facteurs de distribution. Impact du choix du nœud bilan sur l'allocation

	Nœud bilan N1		Nœud bilan N2		Nœud bilan N3	
	€/h	€/MWh	€/h	€/MWh	€/h	€/MWh
L4	340	0.87	159	0.41	0	0
L5	175	0.73	0	0	0	0
L6	400	1.21	819	2.5	1238	3.75
L7	208	0.87	94	0.39	0	0
L8	393	1.09	627	1.74	763	2.12
L9	485	0.90	303	0.56	0	0

IV.4.3.2) Allocation aux transactions par facteurs de distribution

Dans la méthode proposée par Chien [CHI 99], les coûts de congestion ne sont pas alloués individuellement aux consommateurs, mais sont alloués plutôt à des portefeuilles de participants au marché. Ces portefeuilles peuvent aller de la simple transaction bilatérale à la transaction multilatérale pouvant impliquer plusieurs producteurs et consommateurs. Le bilan électrique de ces portefeuilles est toutefois toujours équilibré. Chien introduit en outre des facteurs de distribution spécifiques, mais on peut obtenir les mêmes résultats avec les facteurs de distribution classiques tirés du calcul de la matrice **A**. La méthode d'allocation consiste d'abord à calculer à l'aide des facteurs de distribution la part du transit sur une ligne congestionnée attribuable à chaque portefeuille. Le coût de congestion est ensuite attribué à chaque portefeuille suivant sa contribution physique à la congestion. L'impact physique calculé ainsi pour une paire producteurs-consommateurs a l'avantage de ne pas varier en

fonction du nœud bilan, ce qui constitue un plus par rapport à la méthode préconisée par Singh. Une transaction bilatérale conclues entre un producteur i et un consommateur j peut se mettre sous la forme :

$$T_k = \{0 \quad 0 \quad P_{Gi} \quad 0 \quad \dots \quad 0 \quad 0 \quad P_{Cj} \quad 0 \quad \dots\} \quad (\text{IV.17})$$

$$\text{où } P_{Gi} + P_{Cj} = 0$$

De la même façon, on peut définir une transaction multilatérale comme un portefeuille équilibré de participants comprenant plus de un producteur ou consommateur :

$$T_k = \{0 \quad \dots \quad P_{Gi} \quad P_{Gi+1} \quad \dots \quad 0 \quad \dots \quad P_{Cj} \quad P_{Cj+1} \quad \dots\} \quad (\text{IV.18})$$

$$\text{où } \sum_i P_{Gi} + \sum_j P_{Cj} = 0$$

Une bourse de l'électricité équivaut dans ce cas à une transaction multilatérale du point de vue électrique. Soit une nouvelle transaction T_k conclue entre np producteurs et nc consommateurs. La part du transit d'une ligne l attribuable à la transaction T_k est :

$$P_l^{Tk} = \sum_{i=1}^{np} A_{l,i} * P_{Gi} - \sum_{j=1}^{nc} A_{l,i} * P_{Cj} \quad (\text{IV.19})$$

Le coût de congestion pour chaque ligne l congestionnée est donc alloué de cette façon :

$$Pay_Tk = \frac{P_l^{Tk}}{\sum_{Tk \in \Omega_T} P_l^{Tk}} * \Pi_l \quad (\text{IV.20})$$

Avec Ω_T étant l'ensemble des transactions dont la contribution va dans le même sens que le transit de la ligne congestionnée.

Nous avons appliqué cette méthode d'allocation au cas de la congestion sur la ligne 1-6. Les résultats sont donnés par le Tableau IV.2 :

Tableau IV.2 : résultats de l'allocation du coût de congestion sur la ligne 1-6 par portefeuille à l'aide des facteurs de distribution

Transaction	Impact sur ligne 1-6 (MW)	Payement €/h	Payement €/MWh
Contrat 1	138.3	439	1.83
Contrat 2	113.7	361	1.50
Contrat 3	23.1	74	1.23
Contrat 4	96	305	2.03
Contrat 5	8.4	27	0.11
Bourse	250.8	796	0.68

Nous pouvons remarquer en examinant les résultats du Tableau IV.2 que les contrats 1 à 4 signés avec le producteur G1T ont un impact relativement important sur la congestion sur la ligne 1-6. Cela concorde très bien avec les données de la traçabilité de l'énergie qui montrent que les producteurs connectés au nœud 1 ont une responsabilité physique quasi exclusive sur la congestion. Toutefois, les responsabilités sont plus difficilement visibles pour les portefeuilles larges avec ce genre d'allocation. Par exemple, la bourse de l'énergie reçoit dans son ensemble un paiement de 796 €/h, sans que la méthode ne suggère en elle-même comment répartir ce coût parmi les participants de la bourse. Or, si l'on examine les données fournies par la traçabilité de l'énergie, il ressort clairement que les participants à la bourse n'ont pas tous le même impact sur la contrainte. Une méthode d'allocation aux transactions peut donc donner de bons signaux économiques dans des marchés où les portefeuilles sont peu étendus, mais ne semble pas assez précise en présence de portefeuilles larges tels que les bourses d'énergie.

IV.4.3.3) Synthèse des résultats des allocations basées sur les facteurs de distribution

L'allocation individualisée aux consommateurs via les facteurs de distribution bien qu'utilisant des facteurs connus, peut différer de façon conséquente suivant le choix du nœud bilan, ce qui tend à rendre la méthode peu transparente. Toutefois, nous avons observé que pour un certain choix de nœud bilan (le nœud 1), les résultats se rapprochaient de ceux obtenus en allouant le coût aux consommateurs via les FUL.

L'allocation aux portefeuilles peut aussi donner des signaux économiques compatibles avec ceux tirés de la traçabilité de l'énergie, et quelque soit le choix du nœud bilan. Toutefois, ces signaux sont d'autant plus dilués que les portefeuilles sont larges, et la méthode ne précise pas comment répartir le coût de congestion au sein de ces larges portefeuilles. Ce fait est

particulièrement gênant en présence d'une bourse de l'énergie ayant un volume important. De plus, elle nécessite la connaissance exacte des contrats commerciaux existant entre les différents participants au marché. Or, dans un environnement déréglementé, l'opérateur du système ne connaît pas nécessairement tous les contrats commerciaux existant.

Nous pouvons ainsi souligner les avantages apportés par les nouvelles stratégies d'allocations basées sur la traçabilité :

- les solutions proposées sont « robustes » : elles ne dépendent pas du nœud bilan choisi et donnent des signaux économiques corrects et comparables.
- on peut mieux cibler l'allocation des coûts : Les stratégies d'allocation basées sur la traçabilité s'adaptent sur n'importe quel type de marché (bilatéral, multilatéral, marché spot) car elles sont individualisées. De plus, elles ne nécessitent pas la connaissance des liens commerciaux existants entre les participants. Elles ne requièrent que la connaissance des injections et soutirages en chaque point du réseau obtenus après clôture du marché de l'énergie.

IV.4.4) Evaluation approfondie des allocations basées sur la traçabilité de l'énergie

Les principaux objectifs que nous avons fixés pour l'allocation des coûts sont les suivants :

- juste responsabilisation des acteurs du marché par rapport aux conséquences que peuvent avoir sur le réseau les choix économiques décidés par le marché de l'énergie
- Découragement des scénarios congestionnels
- Intérêt au développement du réseau

Nous allons dans cette section discuter de l'efficacité des stratégies d'allocations des coûts de congestion que nous proposons par rapport à ces objectifs.

IV.4.4.1) Responsabilisation des acteurs du marché par rapport aux conséquences sur le réseau des choix économiques du marché

La liberté de choix de fournisseur ainsi que l'entrée de nouveaux arrivants sur le marché concurrençant le producteur historique se traduisent concrètement par de profonds changements dans la répartition de la production sur l'ensemble du réseau. Comme nous l'avons déjà dit au Chapitre I, ce sont ces profonds changements qui amènent de nouvelles

contraintes sur le réseau de transport et qui créent des situations de congestion là où elles n'existaient pas en situation de monopole. Une allocation efficace des coûts de congestion devrait cibler ces changements de production. Cela aurait alors pour effet de responsabiliser les usagers du réseau quant aux conséquences que peuvent avoir les préférences du marché sur la sécurité du système.

Dans ce cas, une allocation mettant plus l'accent sur les contributions des producteurs aux transits congestionnés semble plus appropriée. En effet, une allocation destinée exclusivement ou principalement aux consommateurs auraient tendance à plus cibler la *localisation géographique* des consommateurs que leur *choix économique*. Il est ainsi possible qu'un consommateur reçoive un paiement élevé pour une congestion relevant d'un changement de production, sans qu'il ait été forcé à l'origine de ce changement de production. Ceci peut avoir une certaine utilité, dans le sens où cela peut inciter les usagers de la ligne congestionnée à adapter leur consommation, en s'équipant par exemple de productions locales leur permettant de transiter moins par le réseau. Cependant, la réaction des consommateurs à un signal économique peut être relativement longue dans le temps, ce qui peut limiter l'efficacité d'une allocation trop ciblée aux consommateurs.

Or, si l'on cible plutôt les producteurs dans l'allocation des coûts de congestion, on tiendra mieux compte des changements de production affectant la sécurité du réseau. En particulier, dans le cas où ces changements de production sont la conséquence du choix de certains consommateurs seulement, on pourra responsabiliser indirectement ces derniers sur les conséquences de leur choix sur l'état du réseau. Ils pourront s'ils le désirent réviser une partie de leur fourniture en optant pour un producteur mieux placé sur le réseau, ce qui aura pour effet de soulager les contraintes du système. De même, un nouveau producteur gagnant des parts de marché importantes en bourse de l'énergie mais induisant de fortes contraintes sur le système pourra éventuellement réviser sa stratégie en réduisant le volume offert sur le marché.

Dans le cas étudié, on pouvait supposer que la congestion sur la ligne 1-6 était due à un transfert de puissance des nœuds générateurs 2 et 3 vers le nœud 1 pour des raisons économiques, vu que les producteurs en ce nœud sont les moins chers. L'allocation aux producteurs cible alors dans ce cas correctement ce transfert de puissance. Le paiement notamment alloué à $G1^T$ servira à responsabiliser indirectement les consommateurs qui par leur choix, ont contribué à générer la contrainte sur la ligne 1-6.

Ainsi, pour cibler efficacement les choix économiques du marché qui induisent des congestions, une allocation par contribution au transit destinée plutôt aux producteurs (facteur α proche de 1) semble plus judicieuse.

IV.4.4.2) Découragement des scénarios congestionnels

Dans certains cas, des usagers du réseau peuvent tirer bénéfice de l'existence de congestions sur le réseau. Des producteurs peuvent notamment s'allier pour provoquer volontairement des congestions dans le cadre du marché de l'énergie, et proposer des ajustements coûteux pour la résoudre [SEE 99], [CNN 02]. Ce comportement spéculatif est possible dans le cas où le traitement des congestions est successif à la clôture du marché de l'énergie, comme c'est le cas pour la méthode du buy back, ce qui représente son principal point faible.

Ainsi, un producteur ayant un fort impact sur une ligne pourrait proposer un grand volume d'énergie à petit prix en vue de congestionner le système. Il pourrait alors s'allier avec un deuxième producteur dont la localisation lui donne un pouvoir de marché par rapport à la congestion ; ce dernier pourra alors proposer un ajustement à la hausse très coûteux et partager les bénéfices avec le premier producteur. Dans ce genre de cas, une simple allocation timbre-postale, surtout si elle est destinée uniquement aux consommateurs, permet difficilement de se protéger contre de tels comportements. Toutefois, si on alloue les coûts en fonction de la contribution au transit congestionné, tout producteur voulant volontairement engendrer des congestions risque de s'en trouver pénalisé. Ainsi, une allocation prenant en compte en grande partie les contributions au transit des producteurs peut se révéler utile pour prévenir les comportements spéculatifs.

On doit tout de même noter que l'on risque aussi par ce procédé de pénaliser d'autres producteurs faisant un usage normal du réseau et n'ayant pas eu de comportement particulièrement spéculatif. Même si un paiement peut être d'une certaine façon utile pour les dissuader de vouloir à leur tour faire du profit sur les contraintes du système, il peut sembler suivant les cas inadéquat. Considérant ce fait, une allocation des coûts aux producteurs devrait aussi s'accompagner d'un plafonnement du prix des offres d'ajustement comme cela se fait en Californie pour rendre les tentatives de spéculation le moins rentables possibles, et donc espérer décourager leur occurrence.

IV.4.4.3) Intérêt au développement du réseau

Dans les anciens monopoles régulés, toute initiative d'investissement provenait de l'opérateur intégré, le réseau étant vu comme une « boîte noire » devant fournir à tout moment un service de transport efficace. Cependant, la restructuration du secteur de l'électricité a fait qu'il devient de plus en plus difficile pour une autorité centralisée de planifier les investissements réseau optimaux. Cette difficulté vient de l'incertitude liée à l'évolution du marché de l'énergie sur le long terme. En outre, la construction de nouvelles capacités de production est de plus en plus assurée par des producteurs indépendants, ce qui tend à

décorrélér les décisions liées au développement du réseau et celles liées au placement de nouvelles productions [TOR 99].

Pour toutes ces raisons, il existe une tendance de plus en plus forte à déléguer en partie la décision d'investir aux usagers du réseau qui en font la demande. Les sources de financement des investissements réseau proviendraient alors en grande partie des actions volontaires de la part des usagers, et non plus simplement de charges globales incluses dans les tarifs de transport. Les propositions récentes de la FERC⁹ (Federal Energy Regulation Commission) vont dans ce sens. Ces propositions consistent à encourager l'envoi de signaux économiques (par la méthode des prix nodaux) pour stimuler les demandes d'investissements réseau de la part des usagers [HOG2 99]. De même, la Commission de Régulation de l'Energie (CRE) a récemment proposé que le développement des capacités d'échanges internationales soit fait à la demande des utilisateurs de ces capacités [CRE 04].

Dans ce contexte, l'envoi de signaux économiques basés sur la traçabilité de l'énergie peut être grandement utile. Pour le traitement des congestions, ces signaux permettent de stimuler les éventuels besoins d'investissements de la part des usagers concernés. Ceux-ci peuvent par exemple choisir de financer un investissement permettant de résoudre une contrainte sur le réseau plutôt que de payer continuellement les coûts de congestion liés à cette contrainte. Ces coûts d'investissement seraient alors supportés uniquement par les usagers transitant par les interfaces saturées, plutôt que par l'ensemble des usagers du réseau.

Dans leur ensemble toutes les allocations basées sur la traçabilité, qu'elles soient destinées aux producteurs ou aux consommateurs peuvent encourager les initiatives d'investissements réseau.

IV.5) Conclusion sur l'allocation des coûts de congestion basée sur la traçabilité de l'énergie

Nos travaux sur le problème de l'allocation des coûts de congestion ont eu pour objectif de contribuer à un gain de transparence, de juste responsabilisation des usagers du réseau, et d'efficacité dans les signaux envoyés. L'usage de la traçabilité de l'énergie permet de répondre à des questions fondamentales concernant l'usage du réseau de chaque participant au marché de l'énergie. Il introduit plus d'objectivité, ce qui permet de résoudre d'éventuels problèmes de contestations, la réalité physique étant « indiscutable ». La Méthode des Images de Charges que nous avons développée à ces fins se base exclusivement sur le modèle bien connu de calcul de répartition de charge.

⁹ La FERC est l'autorité chargée de veiller au respect des règles qui dictent le fonctionnement des marchés aux Etats-Unis. Son équivalent française est la Commission de Régulation de l'Energie (CRE)

Les résultats que nous avons obtenus montrent que les signaux envoyés à l'aide de la traçabilité de l'énergie sont comparables à ceux fournis par la méthode reconnue des prix nodaux. Nous avons en outre comparé les résultats de nos allocations à d'autres propositions récentes d'allocation physique des coûts de congestion. Il en est également sorti une certaine cohérence sur les signaux économiques envoyés. Nous avons toutefois souligné les avantages des allocations que nous avons proposées, qui sont adaptables sur tous types de marchés (bourse, marché bilatéral, etc..) et qui sont basées sur des facteurs uniques pour un état donné du réseau.

Les répartitions des coûts de congestion de type timbre-postale, largement appliquées actuellement, ne sont pas réellement efficaces pour endiguer le phénomène de congestions sur le long terme. Les signaux économiques dégagés par les allocations proposées peuvent alors se révéler utiles pour guider les choix adoptés dans le marché de l'énergie, décourager l'occurrence volontaire des scénarios congestionnels, et inciter les usagers concernés à participer au développement du réseau.

CHAPITRE V

Vers une coordination supranationale du traitement des congestions

V.1) Introduction

Le plus souvent, lorsqu'un opérateur du système doit résoudre une contrainte de transit, il n'a en général accès qu'aux données topologiques de son réseau et aux ressources contenues dans sa zone de réglage, les réseaux voisins étant simplement modélisés par des équivalents plus ou moins appropriés [MEY 98]. Il n'y a pas encore actuellement de véritable stratégie de coordination du traitement des congestions, qui permettrait à un opérateur de faire appel non seulement aux ressources présentes sur sa zone, mais aussi à celles de réseaux voisins. Toutefois, en Europe, un effort de plus en plus grand est fait pour encourager la coordination du traitement des congestions au niveau supranational [ETSO 99], [ETSO 02].

Dans ce chapitre, nous allons présenter et analyser une méthodologie pour coordonner le traitement des congestions (ajustements de production et allocation des coûts) au niveau international. Le but recherché est de coordonner les actions de plusieurs opérateurs du système de telle sorte qu'ils traitent les congestions à la manière d'un « supra opérateur » virtuel qui aurait accès à toutes les informations concernant le grand réseau interconnecté [HOG 01]. Cette coordination pourrait être possible grâce à un algorithme de traitement des congestions décentralisé. En effet, un problème initial d'optimisation peut être décomposé en plusieurs sous-problèmes à traiter. La résolution coordonnée de chaque sous-problème permet d'arriver à la solution globale du problème initial. Ce principe sera appliqué à l'Outil de Redispatching Optimisé (ORO) tel qu'il a été défini au Chapitre II, en vue de coordonner la détermination des ajustements de productions optimaux sur plusieurs zones.

Nous allons appliquer l'algorithme coordonné sur le réseau IEEE RTS 96 composée de trois zones distinctes et contenant 72 nœuds. Nous allons montrer qu'un traitement des congestions coordonné, en faisant appel à toutes les ressources de l'ensemble du réseau, peut permet de réduire le coût de congestion dans certains cas. Nous allons en outre allouer les coûts de congestion en utilisant la traçabilité de l'énergie, ce qui nous permettra de mieux valider les conclusions tirées dans le Chapitre IV.

V.2) Principe de coordination du traitement des congestions entre plusieurs opérateurs du système

Les marchés actuels de l'énergie ont actuellement de plus en plus tendance à s'étendre au niveau international, gommant ainsi les frontières existantes entre les pays. Ceci est le cas de l'Europe, où il existe un marché de l'électricité à l'échelle européenne qui est en train de se développer. Cependant, comme nous l'avons souligné au Chapitre I de cette Thèse, l'internationalisation des échanges ne s'est pas véritablement accompagnée d'une adaptation des grands réseaux interconnectés, d'où un problème grandissant de congestions. Idéalement, le traitement des congestions à l'échelle supranationale pourrait être confiée à un Superviseur qui aurait accès à toutes les données topologiques des réseaux interconnectés, à tous les échanges d'énergie et à toutes les offres d'ajustements (sous réserve de confidentialité). Il détecterait toutes les contraintes apparaissant sur l'ensemble du système interconnecté et les résolverait de façon centralisée (ou décentralisée). Cette approche a toutefois plusieurs défauts majeurs. Elle nécessite en effet une masse énorme d'information à traiter. Il faudrait en outre que ce superviseur ait connaissance de toutes les données du marché de l'énergie et de tous les changements topologiques ayant pu avoir lieu sur un ensemble pouvant comporter plusieurs milliers de nœuds. En l'état actuel, une telle tâche semble impossible à surmonter par un seul Superviseur. D'autre part, certains participants au marché peuvent souhaiter que des données économiques telles que les paramètres des offres d'ajustement ne soient pas connus à l'extérieur. Ainsi, une approche consistant à coordonner les actions de plusieurs opérateurs régionaux semble plus appropriée à l'heure actuelle (Fig.V.1).

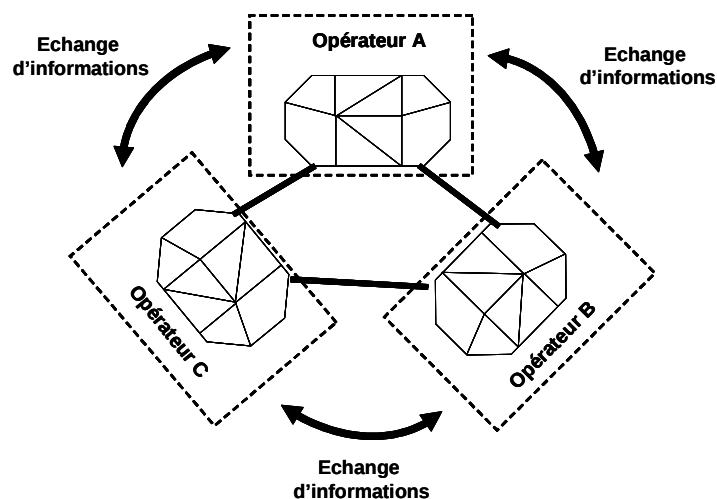


Figure V.1 : coordination de plusieurs opérateurs régionaux

L'échange d'informations nécessaire pour coordonner les actions des opérateurs devrait en outre être le plus minime possible et concerner des données qui sont non confidentielles.

V.3) Outil de Redispatching Optimisé (ORO) coordonné: formulation [BAK 03]

V.3.1) Problème initial (ORO global)

L'algorithme initial permettant de traiter les congestions sur un grand réseau interconnecté est le même que celui présenté au Chapitre II § I.2.5.2 :

$$\text{Min } \sum_i c_i^{+,-} \Delta G_i$$

Avec

$$\mathbf{g}(\Delta \mathbf{G}, \Delta \mathbf{P}, \Delta \mathbf{P}_G, \Delta \boldsymbol{\theta}) = 0$$

$$\mathbf{h}_{\min} \leq \mathbf{h}(\mathbf{P}_{ij}, \mathbf{P}_G) \leq \mathbf{h}_{\max}$$

où $\mathbf{g}$ représente l'ensemble des équations égalités correspondant au calcul de répartition de charge et à l'équilibre production-consommation (relations (II.22a) à (II.23)), et où $\mathbf{h}$ représente l'ensemble des équations inégalités concernant les limites sup/inf de transit et les limites sup/inf des productions ajustées (relations (II.24) et (II.25)).

V.3.2) Problème équivalent

A ce stade, nous allons procéder à un premier découplage. Nous pouvons reformuler le problème initial en rajoutant $2*N_{int}$ variables, N_{int} étant le nombre d'interconnexions présentes sur l'ensemble du système. Les nouvelles variables sont les suivantes : pour chaque interconnexion connectée entre un nœud i d'une zone A et le nœud j d'une zone B, on définit son transit T_{ij}^{AB} de la zone A vers une zone B ; de même la variable duale T_{ij}^{BA} définit le transit de cette interconnexion de la zone B vers la zone A (Fig V.2). Ainsi, le problème équivalent devient :

$$\text{Min } \sum_A \sum_{i \in A} c_i^{+,-} \Delta G_i^A \quad (\text{V.1})$$

Avec comme contraintes :

Equations du calcul de répartition de charge (pour chaque zone A)

$$\mathbf{P}_G^A - \mathbf{P}_C^A - \mathbf{R}^A * \mathbf{T}^{AB} = \mathbf{B}^A * \boldsymbol{\theta}^A \quad (\text{V.2})$$

$$\mathbf{P}_G^A = \mathbf{P}_G^{A^0} + \Delta \mathbf{G}^A \quad (\text{V.3})$$

Couplage avec les zones adjacentes

$$\frac{\theta_i^A - \theta_j^B}{x_{ij}} - T_{ij}^{AB} = 0 \quad (V.5)$$

Limites imposées aux lignes et interconnexions

$$-P_{ij}^{Amax} \leq P_{ij}^A \leq P_{ij}^{Amax} \quad (V.6)$$

$$-T_{ij}^{ABmax} \leq T_{ij}^{AB} \leq T_{ij}^{ABmax} \quad (V.7)$$

avec $\mathbf{R}^A$ étant le vecteur incident défini ainsi pour l'ensemble des interconnexions de la zone A :

$$R_{k,1}^A = 1 \text{ si il y a une interconnexion reliée au nœud } k$$

$$R_{k,1}^A = 0 \text{ s'il n'y a pas d'interconnexion reliée au nœud } k$$

$\mathbf{T}^{AB}$ est le vecteur des transits sur les interconnexions de la zone A vers les zones B adjacentes

La solution optimale du problème équivalent est la même que celle du problème initial. En outre, la matrice des admittances nodales $\mathbf{B}^A$ est une sous matrice de la matrice $\mathbf{B}$ correspondant au système interconnecté en entier. Le bilan de puissance est effectué ici par le troisième terme de (V.2) $\mathbf{R}^A * \mathbf{T}^{AB}$ qui représente les imports/exports de la zone A. Pour la référence des phases nodales, il faut fixer un $\theta_{ref} = 0$ dans une seule zone seulement.

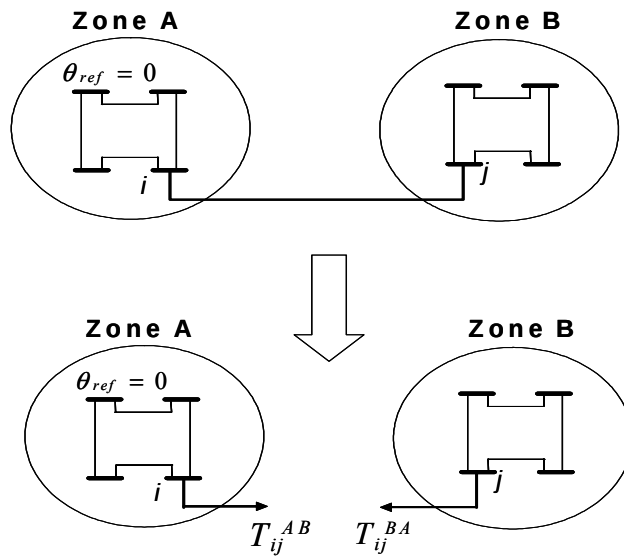


Figure V.2 : principe de découplage de deux zones reliées par une interconnexion

V.3.3) Problème coordonné

Nous cherchons à présent à découpler le problème équivalent en ses sous-problèmes associés à chaque zone. Pour chaque zone, les contraintes associées sont les mêmes que celles listées de (V.2) à (V.7). Le Lagrangien d'un sous-problème comporte bien sûr ces contraintes, mais on doit lui ajouter les contraintes de couplage des zones adjacentes pour reconstituer les conditions de Karush-Kuhn-Tucker du problème équivalent (voir Annexe 6). Ainsi, si $\hat{\theta}_j^B$ et $\hat{\alpha}_{ij}^B$ sont les valeurs optimales obtenue dans une zone B adjacente, le sous-problème à résoudre pour chaque zone A devient :

$$\text{Min } \sum_{i \in A} C_i(\Delta G_i) + \sum_{ij \in \Gamma_A} \hat{\alpha}_{ij}^B \left(\frac{\hat{\theta}_j^B - \theta_i^A}{x_{ij}} - \hat{T}_{ij}^{BA} \right) \quad (\text{V.8})$$

Avec comme contraintes :

Equations du load-flow (pour chaque zone A)

$$\mathbf{P}_G^A - \mathbf{P}_C^A - \mathbf{R}^A * \mathbf{T}^{AB} = \mathbf{B}^A * \boldsymbol{\theta}^A \quad (\text{V.9})$$

$$\mathbf{P}_G^A = \mathbf{P}_G^{A^0} + \Delta \mathbf{G}^A \quad (\text{V.10})$$

Couplage avec les zones adjacentes

$$\frac{\theta_i^A - \hat{\theta}_j^B}{x_{ij}} - T_{ij}^{AB} = 0 \quad (\text{V.11})$$

Limites imposées aux lignes et interconnexions

$$-\mathbf{P}_{ij}^{Amax} \leq \mathbf{P}_{ij}^A \leq \mathbf{P}_{ij}^{Amax} \quad (\text{V.12})$$

$$-\mathbf{T}_{ij}^{ABmax} \leq \mathbf{T}_{ij}^{AB} \leq \mathbf{T}_{ij}^{ABmax} \quad (\text{V.13})$$

Γ_A est l'ensemble des interconnexions de la zone A. $\hat{\theta}_j^B$ est la valeur optimale de la phase nodale au nœud j d'une zone adjacente obtenue après résolution du sous-problème associée à cette zone adjacente. De même, $\hat{\alpha}_{ij}^B$ est la valeur optimale du multiplicateur de Lagrange associée à la contrainte de couplage de la zone adjacente en question.

Lorsque tous les sous-problèmes associés à chaque zone sont formulés suivant la procédure décrite ici, nous pouvons observer que les conditions de Karush-Kuhn-Tucker (KKT) du

problème équivalent sont identiques. Ainsi, nous avons en principe la garantie de converger vers la solution optimale du problème global, tout en résolvant chaque sous-problème indépendamment les uns des autres.

Ainsi, la solution globale du problème décrit au §V.3.1 peut être trouvée via un algorithme découplé tel que celui-ci :

INITIALISATION : $T_{ij}^{AB0} = 0$, $\alpha_{ij}^A = 0$ pour toutes les interconnexions ij d'une zone A

TANT QUE $\varepsilon < 0$

Résoudre les sous-problèmes (V.8)→(V.13) pour chaque zone.

SI $|T_{ij}^{AB} + T_{ij}^{BA}| < \varepsilon$ pour toutes les interconnexions pour toutes les zones → FIN

SINON

Echanger les $\hat{T}_{ij}^{BA}$, les $\hat{\alpha}_{ij}^{BA}$ et les $\hat{\theta}_{ij}^B$ entre toutes les zones adjacentes

Nitérations = Nitérations+1

FIN TANT QUE

Lors de la résolution de chaque sous-problème, les valeurs $\hat{T}_{ij}^{BA}$, $\hat{\alpha}_{ij}^{BA}$ et $\hat{\theta}_{ij}^B$ sont traitées comme des constantes. Ainsi, à l'itération k , la zone A a besoin des valeurs $\hat{T}_{ij}^{BA}$, $\hat{\alpha}_{ij}^{BA}$ et $\hat{\theta}_{ij}^B$ obtenues dans toutes zone B adjacente à l'itération $k-1$. Il est possible qu'un sous-problème puisse ne pas converger, surtout au lancement de l'algorithme. Dans ce cas, il peut tout de même fournir les valeurs liées à ses interconnexions avec ses voisins qu'il a pu trouver. Ces valeurs pourront normalement être utilisées lors de l'itération suivante.

Notons qu'un opérateur du système a un minimum d'informations à s'échanger avec les opérateurs voisins. Il n'a besoin d'aucune donnée topologique concernant les réseaux voisins et n'a pas besoin d'avoir accès aux offres d'ajustement déposées ailleurs que dans sa zone de contrôle.

D'autres approches coordonnées ont été proposées et appliquées au cas de l'OPF [CON 97], [KIM 97]. Cependant, l'algorithme présenté a l'avantage d'être relativement pratique d'utilisation pour notre cas d'étude, et s'est révélé particulièrement efficace.

V.4) Etude du réseau RTS96

Le réseau test RTS 96 a été développé en 1979 dans sa première version. Le but était de disposer d'un réseau test sur lequel on pouvait mener des études dédiées à la sûreté de fonctionnement des réseaux de transport. Le réseau est composé de trois modules identiques comportant chacun 24 nœuds, et reliés entre eux par des interconnexions. Toutes les données topologiques du réseau, les informations sur les unités de production et les charges, ainsi que les limites de transit imposées aux lignes sont données dans la référence [IEEE 99]. Son schéma complet est donné dans l'Annexe 4. Les trois zones composant le réseau seront appelées ici zone A, zone B et zone C, comme cela est indiqué dans l'Annexe 4.

V.4.1) Etat initial du réseau

Chacune des trois zones doit fournir une charge de 2850 MW. La répartition initiale de la production sur l'ensemble du réseau (que l'on peut trouver en Annexe 4), n'occasionne aucun dépassement de transit sur aucune ligne du réseau. Le bilan de puissance de chaque zone est équilibré, de sorte que le flux entrant dans une zone est égal au flux sortant. La Figure V.3 nous montre les transits sur les interconnexions à l'état initial :

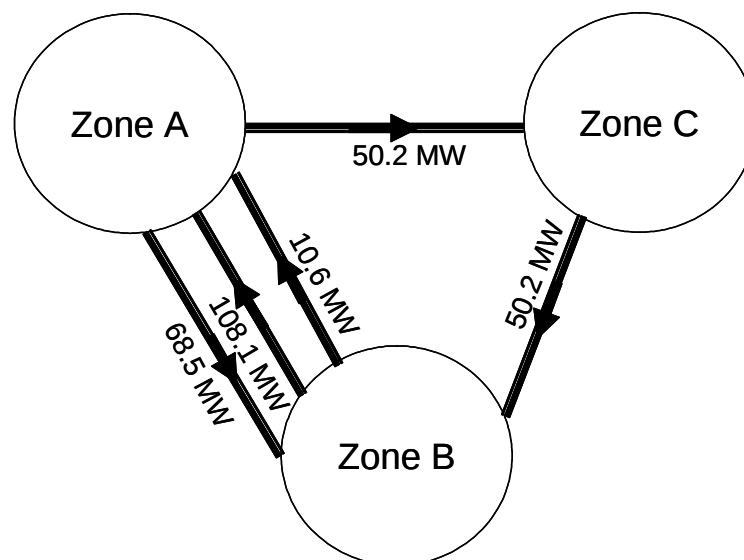


Figure V.3 : état initial du réseau

Pour les besoins de notre étude, nous supposons que cet état du réseau est celui pour lequel le réseau a été conçu dans une logique « monopoliste ».

V.4.2) Cas contraints : présentation des cas de congestion étudiés

Pour pouvoir créer des contraintes sur le système, nous allons effectuer des changements de production et faire quelques modifications dans les données du réseau. Nous allons limiter l'analyse pour trois cas précis de congestion : un cas de congestion intra zonale provoqué par des changements de production extérieur à la zone (cas 1), un cas de congestion double intra zonale (cas2) et cas de congestion sur une interconnexion (cas3).

V.4.2.1) Cas 1 : congestion sur la ligne connectant le nœud 220 au nœud 223

Nous avons simulé un changement de production en déplaçant 800 MW de la zone A vers la zone C. Ce changement peut être assimilé à un changement de fournisseurs effectué par certains consommateurs de la zone A voulant quitter leur ancien fournisseur local pour importer de l'énergie d'une zone adjacente. De plus, on suppose que le chemin contractuel choisi passe uniquement par l'interconnexion reliant la zone A à la zone C. Ce changement de production occasionne alors un nouveau flux parallèle partant de la zone C et allant à la zone A via la zone B. En même temps, l'impédance de l'interconnexion reliant la zone A à la zone C a été légèrement augmenté de manière à renforcer l'effet de ce nouveau flux parallèle. Ce déséquilibre crée une congestion sur la ligne 220-223 en zone B dont le transit s'élève à 533 MW, dépassant les 500 MW de puissance maximale permise. Ce cas est schématisé par la Figure V.4 :

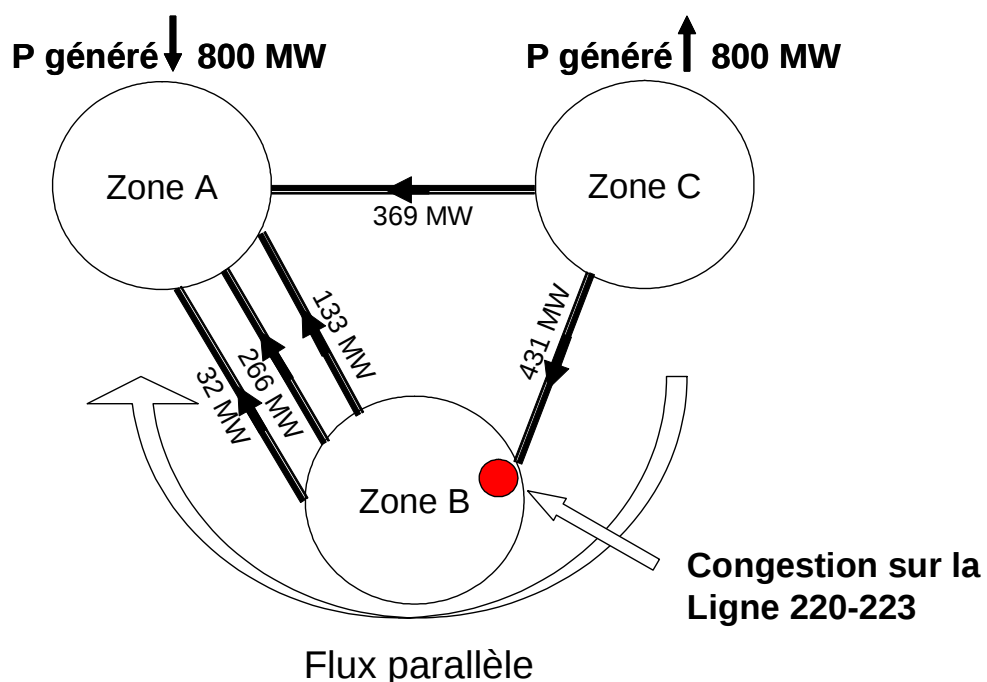


Figure V.4 : congestion sur la ligne 220-223 (cas 1)

Ci-dessous les quantités déplacées et les nœuds où elles l'ont été :

<u>Zone A</u>	<u>Zone B</u>
Nœud 101 : -100 MW	Nœud 315 : +200 MW
Nœud 102 : -100 MW	Nœud 316 : +100 MW
Nœud 107 : -100 MW	Nœud 318 : +300 MW
Nœud 113 : -100 MW	Nœud 321 : +100 MW
Nœud 118 : -100 MW	Nœud 322 : +100 MW
Nœud 121 : -200 MW	
Nœud 122 : -100 MW	

V.4.2.2) Cas 2 : congestion sur la ligne 220-223 et sur la ligne 221-215

Ce cas est inspiré du cas précédent. Tout en gardant la nouvelle configuration de la production adoptée pour le cas 1, nous avons ramené la limite de la ligne reliant le nœud 221 au nœud 215 en zone B à 450 MW de manière à créer aussi une congestion sur cette ligne.

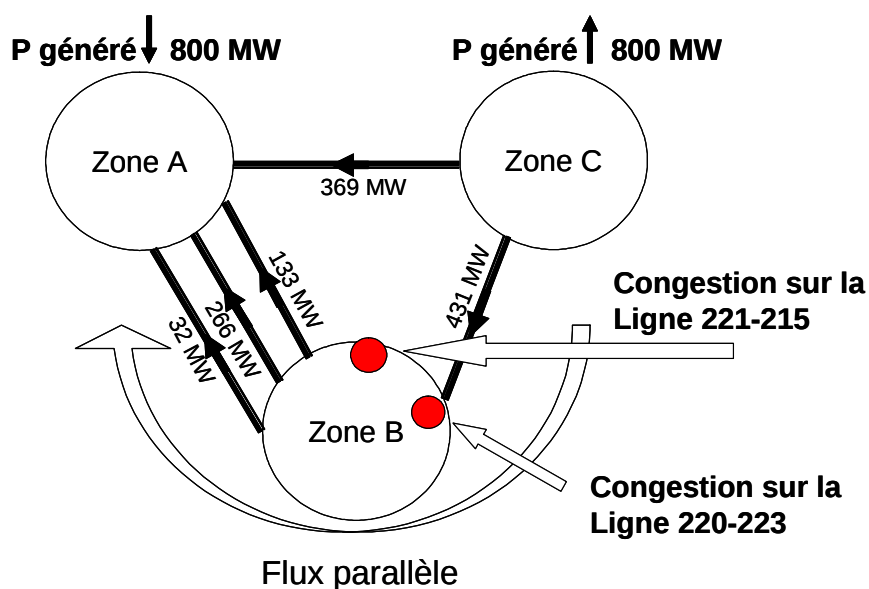


Figure V.5 : double congestion en zone B (cas 2)

V.4.2.3) Cas 3 : congestion sur l'interconnexion 107-203 reliant la zone A à la zone B

Pour ce dernier cas étudié, nous avons créé un changement de production entre la zone A et la zone B de façon à congestionner l'une des trois interconnexions reliant ces deux zones (Fig. V.6). L'interconnexion contrainte est celle qui relie le nœud 107 en zone A au nœud 203 en zone B.

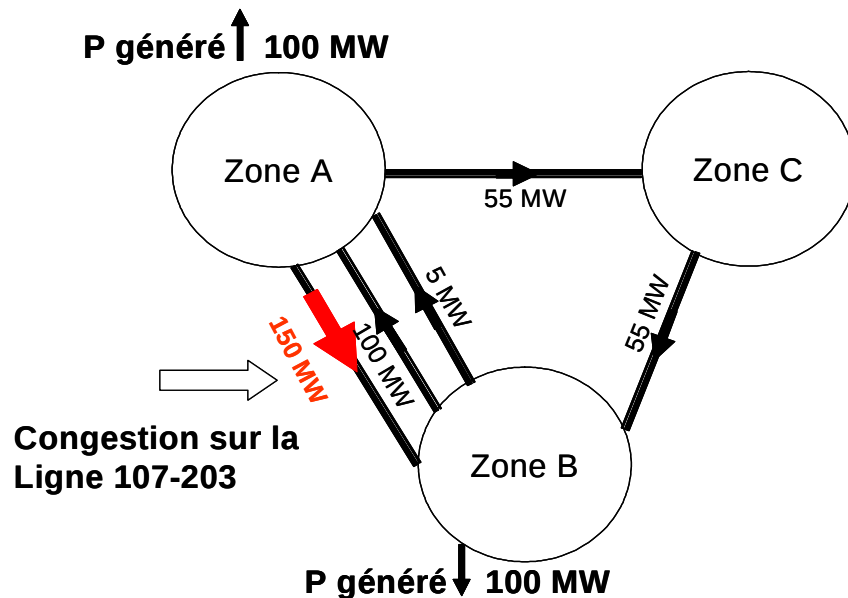


Figure V.6 : congestion sur une interconnexion reliant la zone A à la zone B (cas 3)

La limite de transit de l'interconnexion a été en outre ramenée à 110 MW. Ci-dessous le changement de production qui a été effectué :

<u>Zone A</u>	<u>Zone B</u>
Nœud 107 : +100 MW	Nœud 201 : +50 MW
	Nœud 202 : +50 MW

Comme pour le cas 1, on peut supposer que ce changement de production est du à la volonté de certains consommateurs de la zone B de quitter leur ancien fournisseur en zone B pour un fournisseur en zone A.

V.4.3) Paramètres des offres d'ajustements

Comme le réseau est constitué de trois zones identiques, nous avons constitué un jeu d'offres d'ajustements identiques aussi pour les trois zones considérées. Pour notre scénario, nous avons six producteurs ayant fourni des offres d'ajustement (nous avons supposé ici que certains n'avaient pas fait d'offres dans le cadre du traitement des congestions, leur puissance de sortie restant alors constante). De plus, ce jeu d'offres d'ajustement est identique pour toutes les simulations conduites.

Les paramètres a_i et b_i de chaque offre d'ajustement sont donnés dans le Tableau V.1. Ils sont donnés pour une seule zone. Les limites de variation des ajustements sont fixés à +/- 50 MW maximum.

Tableau V.1 : paramètres des offres d'ajustement

Nœud de connection	a_i (€/h)	b_i (€/MWh)
1	25	0.4
2	26	0.42
7	25	0.4
16	26	0.38
22	25	0.35
23	20	0.3

V.4.4) Traitement des congestions dans les trois cas présentés : comparaison de différents ORO (local, global, coordonné)

Nous avons traité les trois cas présentés avec trois algorithmes différents :

- Un ORO *local* : pour les cas de congestions intrazonales (cas 1 et cas 2), nous avons supposé que l'opérateur de la zone considéré (la zone B) devait résoudre le problème seul. Il n'a donc accès qu'aux offres d'ajustement utilisables sur sa zone. Pour tenir compte des flux de bouclage aux interconnexions, nous avons laissé les zones A et C connectées à la zone B, tout en gardant dans l'esprit que les opérateurs en pratique modélisent les zones adjacentes par des équivalents réseaux. Pour le cas particulier d'une congestion sur une interconnexion (cas3), ce sont

uniquement les opérateurs de chaque côté de la congestion qui participent à son traitement, le troisième n'étant pas impliqué.

- L'ORO *global* : toute congestion est traitée au niveau global sur les trois zones considérées. Cet algorithme suppose l'existence d'un Superviseur ayant connaissance de toutes les données topologique et de toutes les offres d'ajustement de l'ensemble interconnecté.
- L'ORO *coordonné* : pour toute congestion, l'ORO coordonné décrit au §V.3.3 est lancé.

Cette comparaison entre ces trois algorithmes a deux principaux objectifs :

- Premièrement, il est intéressant de voir dans quelle mesure le fait d'étendre le traitement des congestions à l'ensemble interconnecté permet de réduire le coût de congestion
- Ensuite, elle permet de tester si l'ORO coordonné converge dans chaque cas correctement vers la solution globale du problème

Pour les trois cas étudiés, nous avons relevé le coût de congestion délivré par les trois algorithmes. Pour le cas 1 et 2, le coût de congestion de l'ORO local correspond au coût total des ajustements effectués dans la zone B seule :

$$\Pi_{local} = \sum_{i \in B} C_i(\Delta G_i^B) \quad (V.14)$$

De même, le coût de congestion de l'ORO global correspond au coût total des ajustements effectués dans les trois zones :

$$\Pi_{global} = \sum_{i \in A, B, C} C_i(\Delta G_i) \quad (V.15)$$

Enfin, le coût de congestion de l'ORO coordonné est calculé en faisant la somme algébrique des coûts de congestions obtenus pour chaque zone :

$$\Pi_{coord} = \sum_{i \in A} C_i(\Delta G_i^A) + \sum_{i \in B} C_i(\Delta G_i^B) + \sum_{i \in C} C_i(\Delta G_i^C) \quad (V.16)$$

Pour le cas 3, le coût de congestion de l'ORO local correspond au coût total des ajustements effectués dans les zones A et B seulement :

$$\Pi_{local} = \sum_{i \in A, B} C_i(\Delta G_i^{A, B}) \quad (V.17)$$

Les résultats sont présentés par le Tableau V.2 :

Tableau V.2 : impact du choix de l'ORO (local, global, coordonné) sur le coût de congestion

	Coût de congestion (€/h)		
	ORO local	ORO global	ORO coordonné
Cas 1	1553	800	800
Cas 2	Pas de solution	2650	2650
Cas 3	1053	1012	1012

Convergence de l'ORO coordonné

Nous observons que dans les trois cas, l'ORO coordonné a convergé vers la même solution que l'ORO global. Ainsi, l'ORO coordonné se comporte comme un Superviseur virtuel qui aurait accès à toutes les ressources et toutes les données topologiques de l'ensemble interconnecté. Les Figures V.7, V.8 et V.9 nous fournissent une illustration de la convergence de l'ORO coordonné pour le cas 1. La Figure V.7 donne l'évolution du coût de congestion durant le processus d'optimisation. La Figure V.8 donne l'évolution du transit de l'interconnexion 223-318 reliant la zone B à la zone C. Sur cette figure, on a d'une part le transit sortant de la zone B (correspondant à la variable $T_{223-318}^{BC}$ utilisée dans le sous-problème de la zone B) et le transit sortant de la zone C (correspondant à la variable $T_{318-223}^{CB}$ utilisée dans le sous-problème de la zone C). Enfin, la Figure V.9 nous donne l'évolution des multiplicateurs de Lagrange $\alpha_{223-318}^{BC}$ et $\alpha_{318-223}^{CB}$ correspondant aux contraintes de couplage des zone B et C.

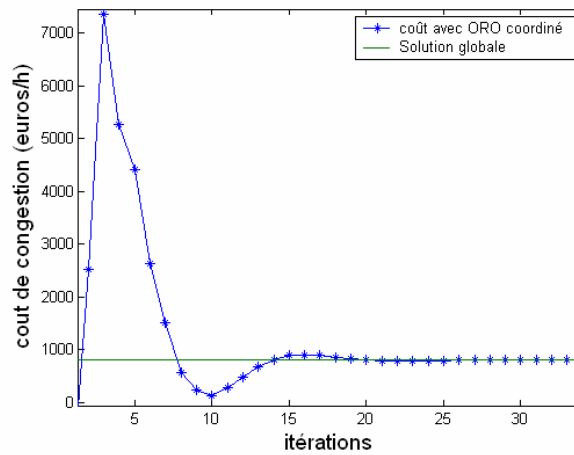


Figure V.7 : évolution du coût de congestion durant le processus d'optimisation de l'ORO coordonné dans le cas 1

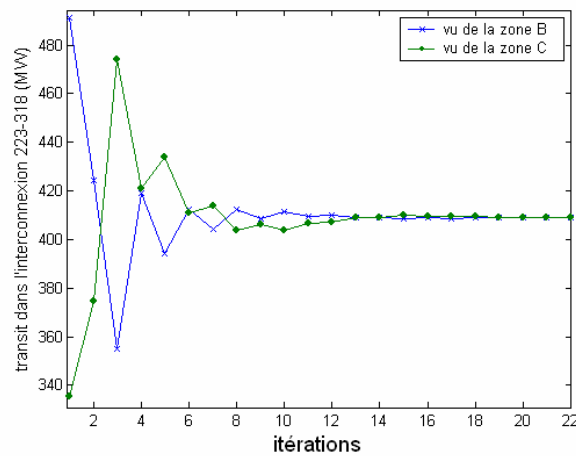


Figure V.8 : évolution du transit de l'interconnexion 223-318 (zone B-zone C) dans le cas 1

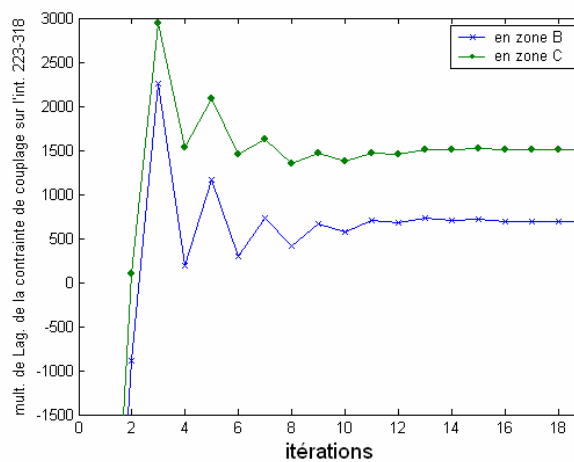


Figure V.9 : évolution des multiplicateurs de Lagrange associés aux contraintes de couplage en zone B et en zone C (cas1)

Les Figures V.10 et V.11 présentent l'évolution du coût de congestion durant le processus d'optimisation de l'ORO coordonné pour les cas 2 et 3 :

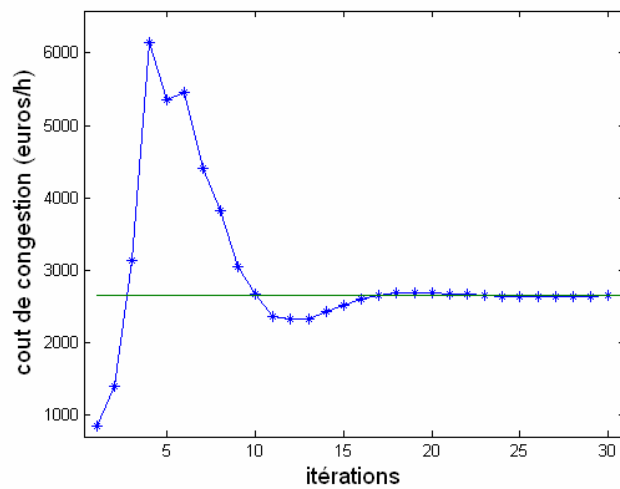


Figure V.10 : évolution du coût de congestion durant le processus d'optimisation de l'ORO coordonné dans le cas 2

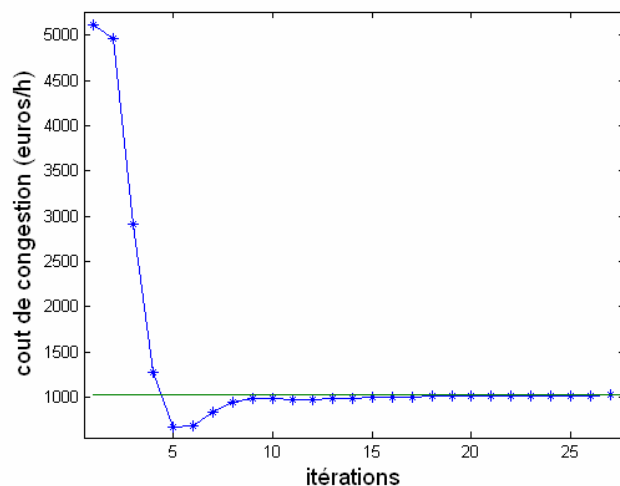


Figure V.11 : évolution du coût de congestion durant le processus d'optimisation de l'ORO coordonné dans le cas 3

Dans les trois cas, la solution globale est pratiquement trouvée au bout de 15 itérations environ. Au lancement de l'algorithme coordonné, le couplage entre les zones ne se fait pas encore correctement (Fig. V.8). Cela se traduit par un bilan global des ajustements de production non nul. Cependant, une fois que le couplage entre zones est suffisamment correct, le bilan global des ajustements s'équilibre et l'on satisfait les conditions de KKT du problème

global, ce qui nous amène à la détermination de sa solution. Nous pouvons aussi remarquer que les deux multiplicateurs de Lagrange associés à la contrainte de couplage au niveau d'une interconnexion ne convergent pas nécessairement vers la même valeur. Dans [BAK 03], ces multiplicateurs sont interprétés comme des prix d'imports/exports. Toutefois, dans notre cas, comme nous n'avons pas choisi de fonder l'allocation des coûts de congestion sur les coûts marginaux, nous ne leur donnons, pour notre étude, aucune signification économique.

Analyse en terme de réduction de coût de congestion et de faisabilité

Nous pouvons noter que dans le cas 1, le fait d'étendre le traitement des congestions aux 3 zones a permis de réduire de façon conséquente le coût de congestion sur la ligne 220-223. Il y a plusieurs raisons qui permettent d'expliquer ce gain en coût très appréciable. Tout d'abord, la congestion sur la ligne 220-223 est principalement due à un grand flux parallèle allant de la zone C à la zone A via la zone B (Fig.V.4). L'opérateur de la zone B peut alors se trouver en difficulté pour résoudre cette congestion dont la cause est en fait extérieure à sa zone. Le fait de faire appel aux ressources extérieures à la zone B permet alors plus de marge de manœuvre à l'algorithme de traitement des congestions qui pourra minimiser le coût de façon plus efficace. D'autre part, la difficulté de résoudre cette congestion uniquement avec les ressources de la zone B se traduit aussi par des ajustements ayant un volume important. Les producteurs ajustés dans le cas de l'ORO global (ou coordonné) sont certes plus nombreux, mais le volume ajusté est en moyenne plus petit. Le coût de congestion étant une fonction quadratique du volume des ajustements, le coût de congestion est donc moins élevé en traitant le cas 1 avec un algorithme global plutôt qu'avec un algorithme local.

Le cas 2 est un cas de figure intéressant car dans une approche seulement locale, l'algorithme n'a pu trouver de solution faisable. Dans un cas pratique, cela aurait pu obliger l'opérateur de la zone B à solliciter davantage les producteurs, qui auraient offert leur service mais probablement à un coût très élevé. L'opérateur aurait aussi pu choisir de délester certaines consommations, ce qui ne représente pas non plus une solution satisfaisante en terme de bien-être social. Toutefois, si l'on traite les deux contraintes du cas 2 en faisant appel aux ressources présentes sur les trois zones, on arrive à trouver une solution faisable. Ainsi, étendre le traitement de congestions intrazonales à plusieurs zones permet dans certains cas de gagner en faisabilité.

Les cas 1 et 2 montrent aussi qu'un traitement des congestions basé uniquement sur le caractère interzonal/intrazonal d'une congestion ne se justifie ni en terme de réduction de coût, ni en terme de faisabilité. La complexité physique des réseaux fait que dans certains cas, le traitement d'une congestion interzonale demande le déploiement de ressources externes pour être plus efficace.

Toutefois, nous ne pouvons généraliser nos observations sur la diminution du coût de congestion. Il peut en effet y avoir un nombre non négligeable de cas où faire intervenir l'ensemble des ressources du système permet de gagner peu en terme de coût de congestion. Dans le cas 3, nous passons d'un coût de congestion de 1053 €/h en faisant intervenir les ressources des zones A et B à un coût de 1012 €/h faisant appel à toutes les ressources. Nous constatons donc un gain de coût d'à peine 4%, ce qui ne justifie pas dans ce cas-là l'emploi d'un algorithme coordonné sur les trois zones. Il est très difficile de déterminer statistiquement, pour un réseau donné, le pourcentage de cas où l'emploi d'un algorithme coordonné peut être justifié en terme de gain de coût. En effet, cela dépend fortement de la répartition et des paramètres choisis pour les offres d'ajustement, ce qui ferait un nombre considérable de combinaisons à analyser.

On pourrait préconiser d'utiliser de façon systématique pour résoudre les problèmes de transit à grande échelle. Toutefois, toute coordination à l'échelle supranationale constitue une approche plus lourde, plus difficile à mettre en œuvre et plus demandeuse en temps de calcul qu'une simple approche locale. En outre, elle demande une synchronisation du traitement des congestions au niveau international et probablement aussi une uniformisation du format des offres d'ajustements. L'emploi d'une méthode coordonnée doit donc être justifié par le caractère récurrent, difficile de résolution et coûteux d'une congestion. Ainsi, nous suggérons de traiter toute nouvelle congestion d'abord à l'aide de ressources locales. Ensuite, en cas de persistance du problème, on peut appliquer un outil coordonné de traitement des congestions pour résoudre la contrainte. En complément, une allocation des coûts de congestion basée sur la traçabilité de l'énergie permettra d'envoyer les bons signaux économiques en vue d'endiguer le plus rapidement possible le problème.

V.4.5) Allocation des coûts de congestion sur le réseau RTS 96

V.4.5.1) Une allocation décentralisée des coûts de congestion

Comme nous l'avons déjà signalé, chaque opérateur ne possède qu'un modèle simplifié des zones adjacentes pour pouvoir modéliser les flux de bouclages à ses frontières. Or, si une congestion survenant sur sa zone provient en partie de l'extérieur (exemple du cas 1), il ne peut rétribuer le coût aux usagers externes par manque d'information. Toutefois, en utilisant les propriétés de la traçabilité de l'énergie, il est possible de contourner cette difficulté.

Par exemple, soit une zone A et une zone B adjacente l'une de l'autre, avec B qui importe de l'énergie de la zone A. Si l'opérateur en A constate qu'une congestion dans sa zone provient en grande partie du flux importé de B, il peut considérer ce flux comme une production équivalente. En utilisant les liaisons productions-transits délivrées par la MIC, il

peut assigner un coût à ce flux en provenance de la zone B et transmettre le montant à l'opérateur de la zone B. Ce dernier pourra alors répercuter ce coût sur les producteurs concernés en modélisant le flux sortant de sa zone comme une charge équivalente et en utilisant les liaisons productions-consommations de la MIC. Ce principe est illustré par la Figure V.12 :

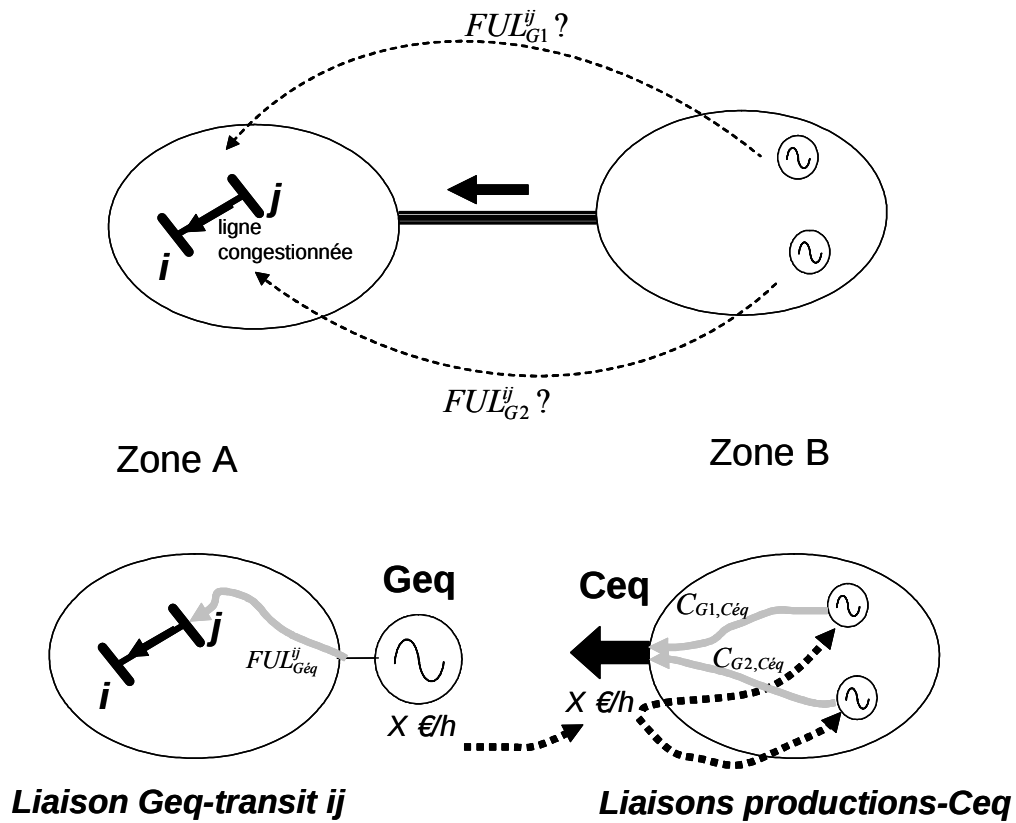


Figure V.12 : principe d'une allocation décentralisée

Ce principe peut être de la même manière appliqué pour tous les types d'utilisateurs (producteurs, consommateurs) et pour toute congestion. Le même principe peut être appliqué pour l'allocation par destination des ajustements présentée au Chapitre IV. Ainsi, l'exploitation des propriétés de la traçabilité de l'énergie permet de nous passer d'un Superviseur qui aurait à manier une grande masse de données en vue d'allouer les coûts de congestion à chaque utilisateur sur l'ensemble du réseau interconnecté.

Dans les paragraphes suivants, nous allons présenter les résultats des allocations de coût de congestion pour les trois cas étudiés. Vu la grande taille du système étudié, nous avons choisi de ne présenter que des données générales ou vraiment significatives. Pour avoir une

meilleure vision de la répartition géographique du coût de congestion pour chaque cas ou obtenir certaines données numériques bien précises, nous invitons le lecteur à se référer à l'Annexe 4. Nous avons choisi de commenter ici plus précisément les allocations pour le cas 1 et cas 3, le cas 2 étant juste une extension du cas 1.

V.4.5.2) Allocation du coût de congestion dans le cas 1

Nous avons le coût de congestion obtenu par ORO coordonné (801 €/h) suivant les méthodes présentées au Chapitre IV. Les Figures V.13 et V.14 donnent la répartition de ce coût par zone si on alloue le coût entièrement aux consommateurs (qui sont assimilés aux charges nodales du réseau) par contribution au transit et par destination des ajustements.

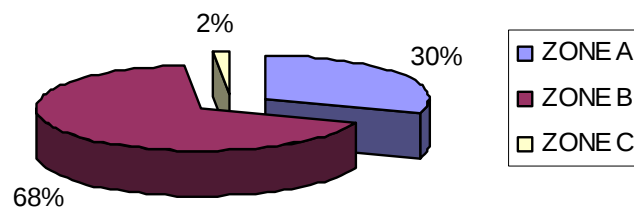


Figure V.13 : répartition du coût de congestion du cas 1 par zone avec une allocation aux consommateurs par contribution au transit

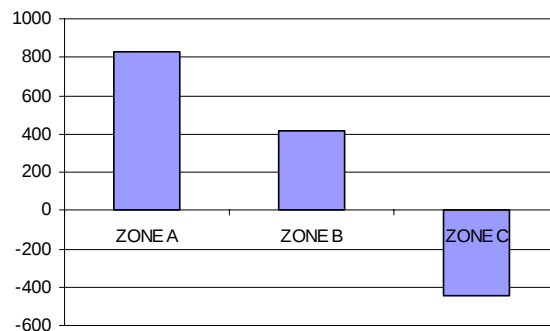


Figure V.14 : répartition du coût de congestion du cas 1 par destination des ajustements

Nous pouvons aisément expliquer ces résultats obtenus d'un point de vue purement physique. Avec l'allocation aux consommateurs par contribution au transit (Fig.V.13), la zone B et la zone A payent la quasi-totalité du coût de congestion. Ces paiements s'expliquent par le flux parallèle du au changement de production de la zone A vers la zone B. Ce flux parallèle (passant par la ligne 220-223) vient d'abord alimenter les consommateurs en zone B, ce qui

explique le paiement attribué à la zone B (68%). Il vient ensuite alimenter les consommateurs de la zone A, la production locale en zone A étant devenue insuffisante.

Avec l'allocation par destination des ajustements, nous voyons apparaître un paiement « négatif » pour la zone C, ce qui signifie que l'ensemble des consommateurs reçoit en réalité un crédit ! Ceci vient de la propriété de l'allocation par destination des ajustements telle que nous l'avons défini au Chapitre IV. En effet, pour un consommateur donné, le crédit dû aux ajustements à la baisse peut très bien être supérieur au paiement dû aux ajustements à la hausse. Ceci peut être aussi aisément interprété en terme de signal économique. En effet, rémunérer les consommateurs en zone C tendrait à les encourager à augmenter leur consommation, ce qui aurait pour effet de détourner une partie du flux parallèle et de soulager la contrainte sur la ligne 220-223. Toutefois, nous devons aussi noter que le fait de rémunérer certains consommateurs peut être moins bien accepté. C'est pour cette raison d'ailleurs que nous avons choisi de ne pas rémunérer les contre-transits pour l'allocation par contribution au transit (voir Ch. IV).

Le Tableau V.1 donne la répartition de ce coût pour les principaux usagers concernés si on alloue cette fois le coût entièrement aux producteurs par contribution au transit :

Tableau V.1 : coût alloué aux producteurs principalement concernés

Producteurs (nœuds de connect.)	Coût alloué (€/h)
N°223	433
N°318	87
N°321	90
N°322	136
Reste	54

En examinant l'Annexe 4 qui nous donne une meilleure vision de la répartition géographique du coût aux producteurs, nous pouvons noter que l'allocation aux producteurs est beaucoup plus ciblée que les allocations aux consommateurs. Le changement de production est globalement assez bien ciblé. En effet, les producteurs connectés aux nœuds 318, 321, 322 totalisent 313 €/h. Ce paiement permet aussi de cibler indirectement les consommateurs en zone A liés à ces producteurs par contrat et qui ont été à l'origine du

changement de production provoquant la congestion sur la ligne 220-223. Il est important de noter que l'allocation aux consommateurs illustrée plus haut permet uniquement de faire porter la responsabilité de ce choix aux consommateurs de la zone A dans leur ensemble.

On peut cependant discuter de façon plus approfondie du paiement important attribué au producteur du nœud 223 qui, du fait de sa position, est le principal utilisateur de la ligne congestionnée. Ce paiement peut sembler d'une certaine manière inapproprié car ce producteur est pénalisé en fonction de sa localisation, alors qu'il n'a pas modifié son injection avant l'apparition de la congestion. Nous n'avons toutefois pas jugé judicieux d'inclure dans nos stratégies d'allocation proposées une règle de type « premier arrivé, premier servi », car elle ne constitue pas dans un cas général une règle objective en soi. Par ailleurs, on devrait toujours garder à l'esprit qu'une allocation physique des coûts de congestion a aussi un rôle préventif en plus d'un rôle purement de responsabilisation. Dans ce cas, le paiement attribué au producteur au nœud 223 peut le dissuader d'augmenter sa production, chose qui serait très préjudiciable pour la contrainte sur la ligne congestionnée. On pourrait aussi considérer deux cas de figure types:

- Ce producteur est un nouvel arrivant sur le marché qui n'a pas encore contribué au développement local du réseau : étant à égalité avec les autres producteurs, il n'y a donc aucune raison objective pour ne pas lui attribuer le paiement qui lui est dû étant l'utilisateur principal de la ligne congestionnée.
- Ce producteur est un producteur historique qui a déjà beaucoup contribué au développement local du réseau : on peut alors envisager exceptionnellement de réduire son paiement, voire de le reporter entièrement sur les autres producteurs concernés. Cette entorse à la règle peut se justifier par le fait que le producteur aurait une préséance sur la partie du réseau qu'il aurait déjà financé. Toutefois, ce genre de traitement devrait se faire uniquement s'il peut être trouvé des raisons objectives suffisantes, qui doivent être clairement indiquées aux usagers du réseau. De plus, ce traitement peut favoriser un producteur historique face à un nouvel arrivant, ce qui peut être contesté.

En guise de synthèse pour l'allocation du coût de congestion dans le cas 1, nous pouvons dire qu'une allocation plutôt aux producteurs est plus ciblée qu'une allocation plutôt aux consommateurs. Quelques points particuliers ont surgi. Avec l'allocation par destination des ajustements, nous avons observé l'attribution d'un crédit plutôt que d'un coût aux usagers de la zone C. Avec l'allocation aux producteurs, un paiement important attribué au producteur connecté au nœud 223 peut amener une discussion à laquelle nous avons tenté de donner quelques suggestions de réponse.

V.4.5.3) Allocation dans le cas 3

Les résultats pour ce cas sont plus équivoques que pour le cas 1. Avec l'allocation aux consommateurs par contribution au transit, ce sont logiquement l'ensemble des consommateurs de la zone B qui se partagent le coût de la congestion, qui était de 1012 €/h. Avec l'allocation aux consommateurs par destination des ajustements, on voit apparaître des paiements en zone A. Cela peut être gênant vu que ces consommateurs ne transitent pas par l'interconnexion, ni ne sont responsable de la congestion en terme de choix (ils sont supposé ne pas avoir changé de fournisseur). Toutefois, ce paiement attribué à l'ensemble de la zone A est relativement peu élevé (168 €/h) comparé à celui attribué à la zone B (844 €/h). On pourra se reporter à l'Annexe 4 pour avoir une meilleure vision géographique de ces allocations.

L'allocation aux producteurs donne un résultat très clair : elle cible clairement le changement de production en attribuant le coût presque entièrement au producteur connecté au nœud 107. Le montant de son paiement s'élève à 912 €/h, le reste du paiement étant éparpillé parmi le reste des producteurs présents sur le réseau.

V.4.5.4) Synthèse des résultats des allocations obtenues sur le réseau RTS 96

Au vu, des résultats que nous avons observé sur les cas 1 et cas 3, nous avons pu constater que les allocations aux producteurs étaient mieux ciblées que celles aux consommateurs. En effet, dans les allocations aux consommateurs, tout changement de fournisseur provoquant une congestion pénalise l'ensemble des consommateurs d'une zone donnée. Avec l'allocation aux producteurs, non seulement les paiements sont moins dispersés, mais on arrive globalement à mieux cibler les changements de fournisseurs étant à l'origine de congestions. Ce faisant, on peut cibler indirectement les consommateurs qui, par leur choix, sont à l'origine de ces changements, plutôt que de pénaliser l'ensemble de la zone à laquelle appartiennent ces consommateurs.

Ainsi, les essais effectués sur le réseau RTS 96 nous donnent de bonnes raisons de penser une nouvelle fois qu'une allocation plutôt destinée aux producteurs semble plus judicieuse, ce qui rejoint nos conclusions du Chapitre IV. Toutefois, étant donné que les signaux envoyés par l'allocation aux consommateurs par contribution au transit sont corrects du point de vue physique, on peut choisir aussi de moduler l'allocation des coûts afin d'en tenir compte. L'allocation par destination des ajustements ne permet malheureusement pas ce genre de modulation, vu qu'elle s'adresse qu'aux consommateurs exclusivement. Bien que donnant des signaux globalement corrects, son usage semble donc moins intéressant qu'une allocation par contribution au transit qui a l'avantage d'être plus flexible.

V.5) Conclusion sur l'approche supranationale du traitement des congestions

Dans ce chapitre, nous avons présenté une nouvelle méthodologie pour coordonner le traitement des congestions au niveau supranational.

Cette coordination permet à différents opérateurs adjacents de trouver la solution optimisée du problème en cas de congestion, tout en se passant d'un Superviseur qui aurait du avoir la connaissance de toutes les données du réseau interconnecté. La coordination est rendue possible grâce à un algorithme de traitement des congestions décentralisé qui intègre dans chaque fonction objectif d'un sous-problème les contraintes de couplage entre zone. Les conditions de KKT du problème global étant réunies, un tel algorithme doit donc en principe converger vers la solution optimale du problème. Les opérateurs ont un minimum d'informations à s'échanger, et surtout, chose très importante, ils n'ont pas à s'échanger des données économiques qui peuvent ainsi rester *confidentielles*. Nous avons constaté que dans tous les cas étudiés sur le réseau RTS 96, la solution de l'algorithme coordonné convergeait bien vers la solution globale. Cela vient du fait que, mathématiquement, les conditions de Karush-Kuhn-Tucker (KKT) du problème coordonné coïncident avec celles du problème initial lorsque la solution optimisée est trouvée (Annexe 6). De plus, nous avons montré que dans certains cas, cette coordination peut être préférable à une résolution locale en réduisant de façon notable le coût de congestion, et en gagnant en faisabilité. Toutefois, il peut y avoir des cas où cette coordination peut ne pas être justifiée si le gain en réduction de coût est jugé trop faible. Vu la difficulté de prévoir statistiquement, pour un réseau donné, le pourcentage de cas où la réduction de coût est réellement significative. Un retour d'expérience est donc nécessaire à chaque nouveau cas de congestion pour décider de l'usage d'une approche coordonnée. Néanmoins, nous pouvons dire que la méthode coordonnée présentée ici est un outil efficace qui ouvre des perspectives intéressantes dans le contexte actuel. En effet, bien qu'il soit reconnu qu'une coordination peut être bénéfique, il n'y a pas encore à l'heure actuelle de stratégie claire de coordination implantée dans nos réseaux interconnectés.

Pour les cas étudiés sur le réseau RTS 96, nous avons aussi alloué le coût de congestion suivant les méthodologies proposées au Chapitre IV, après avoir suggéré comment on pouvait utiliser la traçabilité de l'énergie en présence de plusieurs opérateurs adjacents. Nous avons constaté une fois de plus que les allocations proposées donnaient des signaux économiques cohérents avec la réalité physique. Au vu des résultats, nous pouvons néanmoins ajouter qu'une allocation plutôt aux producteurs semble mieux cibler les changements de production dus au choix des participants au marché. Les allocations aux consommateurs semblent en effet plutôt rétribuer la conséquence d'un changement de fournisseur au niveau collectif, sur l'ensemble d'une zone donnée. Avec une allocation plutôt aux producteurs, on peut arriver à

cibler à la fois le changement de production en lui-même, et indirectement les consommateurs étant précisément à l'origine de ce changement.

CONCLUSION GENERALE

L'objectif de cette Thèse était de définir une méthodologie de traitement des congestions transparente, techniquement faisable, efficace économiquement et coordonnable sur plusieurs pays. Nous avons donc considéré le traitement des congestions comme un service séparé du marché de l'énergie, basé sur l'usage d'offres d'ajustements volontaires et dont le coût devait être alloué sur une base physique en vue d'envoyer des signaux économiques.

Notre objectif a d'abord été de déterminer quel modèle de base nous allions définir pour le traitement des congestions. Notre réflexion, appuyée par des résultats de simulation obtenus sur le cas du réseau 9 noeuds, nous a conduit à montrer que la méthode du buy back est techniquement faisable, qu'elle est adaptable sur les différentes structures de marché, et qu'elle a l'avantage d'être transparente. Nous avons montré que la méthode basée sur les coupures de transactions pouvait, dans certains cas, poser des problèmes techniques en engendrant de nouvelles contraintes sur le système lors de son usage. La méthode du buy back, du fait qu'elle tient compte des interactions au sein de l'ensemble du réseau, parvient beaucoup mieux à trouver une solution faisable en cas de contrainte. La méthode basée sur la tarification marginale crée rapidement des surplus financiers importants dépendant de la répartition des prix nodaux, qui ne donnent pas les bonnes incitations aux propriétaires du réseau pour son développement. Elle souffre en outre d'un manque de transparence sur la détermination de ces prix. Cependant, le caractère volontaire du dépôt d'une offre d'ajustement peut dans, certains cas, être un inconvénient, car l'opérateur du système peut se trouver avec un nombre d'offres insuffisant. Si ces cas se produisent, il convient de mettre la sécurité du réseau comme étant la priorité essentielle à assurer. Des accords préalables entre les fournisseurs potentiels du service et les opérateurs du système peuvent en outre éviter ces situations. D'autre part, la méthode du buy back peut faciliter les comportements spéculatifs du fait de la séparation du traitement des congestions et du marché de l'énergie. Cependant, une allocation des coûts de congestion sur une base physique peut contribuer à dissuader l'occurrence de ces comportements.

Dans un second temps, nous avons élaboré de nouvelles stratégies d'allocation des coûts de congestion basées sur l'usage de la traçabilité de l'énergie. Nous avons noté que les signaux économiques envoyées par ces nouvelles allocations étaient comparables à ceux délivrés par la méthode des prix nodaux. Nous avons en outre montré que les allocations proposées

délivrent des résultats cohérents avec ceux d'autres allocations physiques, tout en étant plus robustes et flexibles d'utilisation. Nos résultats (notamment ceux obtenus sur le réseau RTS 96) ont aussi montré qu'une allocation plutôt destinée aux producteurs parvenait à mieux cibler les changements de production résultants des choix faits sur le marché de l'énergie.

Enfin, le traitement des congestions est un service rendu aux usagers du réseau qui doit être le moins cher possible. Il convient alors de mettre en œuvre le maximum de ressources, sous réserve de faisabilité, en vue d'atteindre cet objectif. Nos travaux sur la coordination du traitement des congestions au niveau supranational ont montré que cette coordination peut réduire de façon substantielle le coût de congestion, voire même améliorer la faisabilité du traitement. La méthode de coordination présentée est basée sur un échange itératif d'informations non confidentielles et en nombre réduit entre opérateurs adjacents. Pour résoudre les sous-problèmes liés à leur zone, chaque opérateur n'a besoin que des données internes de son réseau ainsi que d'informations non confidentielles provenant des zones voisines. L'avantage de cette méthode se base donc sur le nombre réduit de données à traiter par chaque opérateur et sur le maintien de la confidentialité des données économiques.

Les perspectives de ce travail de Thèse s'inscrivent sur une gestion à long terme des congestions. En effet, la gestion des congestions par ajustements de production avec son allocation ciblée représente une solution de court terme, destinée à inciter les usagers à adapter rapidement leur usage du réseau et à les intéresser au développement du réseau. Des développements futurs du réseau pourront ainsi progressivement remplacer les ajustements de production au fur et à mesure qu'ils se feront, en réponse aux bonnes incitations données par les signaux économiques. Au vu de la difficulté actuelle de construire de nouvelles lignes, la possibilité d'utiliser des transformateurs déphaseurs apparaît comme une solution prometteuse. De nos jours, ces transformateurs déphaseurs sont uniquement utilisés pour un réglage local des transits, mais on envisage de plus en plus de faire appel à eux pour contrôler les transits sur une échelle internationale. Toutefois, la multiplication de ces déphaseurs ou d'autres actionneurs réseaux risquera de poser des problèmes d'interaction qu'il conviendra de résoudre soigneusement. En outre, avec l'accroissement des échanges internationaux et les réseaux portés toujours au plus près de leurs limites, gérer les flux uniquement avec des déphaseurs peut dans certains cas présenter des problèmes de faisabilité. Aussi, la gestion des congestions par ajustements de production et par allocation ciblée de leur coût pourra toujours se révéler être solution pratique pouvant être remise à l'ordre du jour et opérationnelle relativement rapidement.

RÉFÉRENCES BIBLIOGRAPHIQUES

- [ABB 01] ABB, New Concepts for Transmission Grids, DOE Workshop on Analysis and Concepts to address Electric Infrastructure Needs, Washington DC, Août 2001
- [ALO 99] Alomoush M.I., Shahidehpour, S.M., "Fixed Transmission Rights for Zonal Congestion Management", IEE Proceedings-Generation, Transmission, Distribution, vol.146, n°5, p.471-476, Septembre 1999
- [BAK 03] Bakirtzis A.G., Biskas P.N., "A Decentralized Solution to the DC-OPF of Interconnected Power Systems", IEEE Transactions on Power Systems, vol.18, n°3, p. 1007-1013, Août 2003
- [BIA 97] Bialek J., "Topological generation and Load Distribution Factors for Supplement Charge Allocation in Transmission Open Access", IEEE Transactions on Power Systems, vol.12, n°3, p. 1185-1193, Août 1997
- [CAD 01] Cadwalader M.D., Harvey S.M., Pope S.L., Hogan W.W., "Market Coordination of Transmission Loading Relief across Multiple Regions", Havard University, Décembre 1998
- [CAI 99] California Independent System Operator, "Facilitating Congestion Management in California", Avril 1999
<http://www.aiso.com>
- [CAI 03] California ISO, Congestion Settlement Guide, Février 2003
<http://www.aiso.com>
- [CUL 94] Culioli J.C., Introduction à l'optimisation, ellipses, 1994
- [CHE 02] Chen L., Suzuki H., Wachi T., Shimura Y., "Components of Nodal Prices for Electric Power Systems", IEEE Transactions on Power Systems, vol. 17, n°1, p. 41-49, Février 2002
- [CHI 99] Chien-Ning YU, Marija D.Ilić, "Congestion Clusters-Based Markets for Transmission Management", Proceedings of the IEEE PES Winter Meeting, p. 821-832, New York 1999
- [CHR 00] Christie R.D., Wollenberg B.F., Wangenstein I., "Transmission Management in the Deregulated Environment", Proceedings of the IEE, vol.88, n°2, Février 2000, p.170-194

- [CON 97] Conejo A.J., Aguado J.A., “Multi-Area Coordinated Decentralized DC Optimal Power Flow”, IEEE Transactions on Power Systems, vol.13, n°4, p. 1272-1278, Novembre 1998

- [CNN 02] CNN, “Feinstein wants DOJ probe of ENRON’s CA dealings”, CNN politics, 7 mai 2002
<http://www.cnn.com>

- [CRE 02] Commission de Régulation de l’Energie (CRE), Rapport d’activité 2002
<http://www.cre.fr>

- [CRE 03] Commission de Régulation de l’Energie, Synthèse de la consultation organisée par la CRE et la CREG relative à la gestion de l’interconnexion France-Belgique, Janvier 2003

- [CRE 04] Commission de Régulation de l’Energie, Consultation publique sur les principes et la structure du tarif d’utilisation des réseaux publics de transport et de distribution d’électricité, Février 2004

- [EEI 04] Edison Electric Institute, “New York ISO-Transmission Tariffs, Pricing&Agreements”, Wholesale Markets Guide, 2004

- [ETSO 99] ETSO, “Evaluation of congestion management methods for cross-border transmission, Florence Regulators Meeting”, Novembre 1999
<http://www.etsso-net.org>

- [ETSO 01] ETSO, “Co-ordinated use of power exchanges for congestion management”, Avril 2001

- [ETSO 00] ETSO, “Net Transfer Capacities (NTC) and Available Transfer Capacities (ATC) in the Internal Market of Electricity in Europe (IEM)”, Mars 2000.

- [ETSO 01] ETSO, “Definitions of Transfer Capacities in liberalised electricity markets”, Avril 2001

- [ETSO 02] ETSO, “Coordinated Congestion Management-An ETSO vision”, Février 2002

- [FAN1 99] Fang R.S., David A.K., “Optimal Dispatch under Transmission Contracts”, IEEE Transactions on Power Systems, vol.14, n°2, p. 877-883, Mai 1999

- [FAN2 99] Fang R.S., David A.K., “Transmission Congestion Management in an Electricity Market”, IEEE Transactions on Power Systems, vol.14, n°3, p. 732-737, Août 1999

- [GAL1 02] Galiana F.D., Kockar I., Cuervo Franco P., "Combined Pool/Bilateral Dispatch-Part1: Performance of Trading Strategies", IEEE Transactions on Power Systems, vol.17, n°1, Février 2002, p.92-99
- [GAL2 02] Galiana F.D., Kockar I., Cuervo Franco P., "Combined Pool/Bilateral Dispatch-Part2: Curtailment of Firm and non Firm Contracts", IEEE Transactions on Power Systems, vol.17, n°4, Novembre 2002, p.1184-1190
- [GE 99] Ge S.Y., Chung T.S., "Optimal Active Power Flow Incorporating Power Flow Control Needs in Flexible AC Transmission Systems", IEEE Transactions on Power Systems, vol.14, n°2, Mai 1999, p.738-744
- [GED 99] Gedra T.W., "On Transmission Congestion and Pricing", IEEE Transactions on Power Systems, vol.14, n°1, Février 1999, p.241-248
- [GLA 02] Glachant J.M., Pignon V., "Nordic Congestion's Arrangement as a model for Europe? Physical constraints and Economic Incentives", ISNIE'02 conference
- [GRI 98] Gribik P.R., George A., Angelidis, Kovacs R. R., "Transmission Access with Multiple Separate Energy Forward Markets", IEEE PES Conference Proceeding, Winter 1998
- [HAD 92] Nouredine Hadj-Saïd, Contribution à l'automatisation de l'analyse de sécurité des grands réseaux de transport et d'interconnexion par une approche locale-frontière, Thèse de doctorat, INPG, Avril 1992
- [HOG 97] Hogan W.W., "Nodes and zones in electricity markets: seeking simplified congestion pricing", IAEE Conference Session « Creating and designing electricity markets », 18th annual North American Conference of the USAEE/IAEE, San Francisco, Septembre 1997
- [HOG1 99] Hogan W.W., "Electric Transmission Adequacy and Market Institutions", Harvard University, Juillet 1999
<http://ksghome.harvard.edu/~whogan.cbg.Ksg/>
- [HOG2 99] Hogan W.W., "Market Based Transmission Investment and Competitive Electricity Markets", Harvard University, Août 1999
<http://ksghome.harvard.edu/~whogan.cbg.Ksg/>
- [HOG 01] Hogan W.W., "Interregional Coordination of Electricity Markets", Harvard University, Juin 2001

- [HUA 03] Huazhong Y.H., Changsha D.X., "A New clustering Method for Network Partitionning for Zonal Pricing", 38th International Universities Power Engineering Conference, 2003 Proceedings vol.2
- [IEEE 99] IEEE RTS Task force of APM Subcommittee, "The IEEE Reliability Test System 1996", IEEE Transactions on Power Systems, vol. 14, n°3, p. 1010-1020, Août 1999
- [KIM 97] Kim B. H., Baldick R., "Coarse-Grained distributed optimal power flow", IEEE Transactions on Power Systems, vol. 12, n°3, p. 932-939, Mai 1997
- [KIR 97] Kirschen D., Allan R., Strbac G., "Contributions of Individual Generators to Loads and Flows", IEEE Transactions on Power Systems, vol. 12, n°1, p. 52-57, Février 1997
- [KIR 98] Kirsch L.D., "ISO economics: How California flubbed it on Transmission Pricing", Public Utilities Fortnightly, 15 Octobre 1998
- [LAF 03] Laffaye H., Tesseron J.M., Coulon M., " Gestion des interconnexions électriques en Europe", Techniques de l'Ingénieur, art. D4085, Février 2003
- [MAN1 03] Manescu L.G., Ciontu M., Hadjsaïd N., Sabonnadière J.C., "La Traçabilité de l'énergie dans les réseaux électriques : Partie I: Méthodes des Images de Charges", Revue Internationale de Génie Electrique, vol. 6, n°3-4, 2003
- [MAN2 03] Manescu L.G., Ciontu M., Hadjsaïd N., Sabonnadière J.C., "La Traçabilité de l'énergie dans les réseaux électriques : Partie II: Utilisation dans les réseaux à accès ouverts ", Revue Internationale de Génie Electrique, vol. 6, n°3-4, 2003
- [MAT 01] MATLAB, Optimisation Toolbox, User'guide version 2, 2001
- [MEY 98] Meyer B., Jerosolimski M., Stubbe M., "Outils de simulation dynamique des réseaux électriques", Techniques de l'Ingénieur, art. D4120, Novembre 1998
- [MOM 01] Momoh J.A., Zhu J.Z., Boswell G.D., Hoffman S., "Power System Security Enhancement by OPF with Phase Shifter", IEEE Transactions on Power Systems, vol. 16, n°2, p. 287-293, Mai 2001
- [MUR 98] Murray B., Electricity Markets: Investment, Performance and Analysis, Wiley, 1998
- [NAS 99] Nasar S.A., Trutt F.C., Electric Power Systems, CRC Press, 1999
- [NER 03] North America Reliability Council, Operating Manual, Juillet 2003
<http://www.nerc.com/~oc/pds.html>

- [NG 81] Ng Wai Y., "Generalised Generation Distribution Factors for Power System Security Evaluations", Transactions on Power Apparatus and Systems, vol. PAS-100, n03, p. 1001-1005, Mars 1981
- [PAN 00] Pan J., Teklu Y., Rahman S., Jun K., "Review of Usage-Based Transmission Cost Allocation Methods under Open Access", IEEE Transactions on Power Systems, vol. 15, n°4, p. 1218-1224, Novembre 2000
- [PAT 99] Paterni P., Vitet S., Bena M., Yokoyama A., "Optimal location of phase shifters in the French network by genetic algorithm", IEEE Transactions on Power Systems, vol. 14, n°1, p. 37-42, Février 1999
- [PER 95] Perez-Arriaga I.J., Rubio F.J., Puerta J.F., "Marginal Pricing of Transmission Services : an Analysis os Cost Recovery", IEEE Transactions on Power Systems, vol. 10, n°1, p. 543-553, Février 1995
- [PEE 01] Power Engineering Education Committee
http://www.ece.mtu.edu/faculty/ljbohman/peec/Dig_Rsor.htm
- [PJM 00] PJM Interconnexion, PJM Open Access-Transmission Tariff, Mai 2000
<http://www.pjm.com>
- [QUI 99] Quintana V.H., Torres G.L., "Introduction to Interior Point Methods", IEEE-PES, PICA '99
<http://thunderbox.uwaterloo.ca/~ieee-ipm/IEEE-TF.ps>
- [QUI 00] Quintana V.H., "Interior Point Methods and Their Applications to Power System: a classification of Publications and Software codes", IEEE Transactions on Power Systems, vol. 15, n°1, p. 170-176, Février 2000
- [RAI1 01] Raikar S., Illié M., Interruptible Physical Transmission Contracts for Congestion Management, Energy Laboratory publication, Massachusetts Institute of Technology, MIT EL 01-010WP, Février 2001
- [RAI2 01] Raikar S., Illié M., Assessment of Transmission Congestion for Major Electricity markets in the US, Energy Laboratory publication, Massachusetts Institute of Technology, MIT EL 01-009WP, Février 2001
- [RAU 00] Rau N., S., "Transmission Loss and Congestion Cost Allocation-An approach Based on Responsibility", IEEE Transactions on Power Systems, vol. 15, n°4, p. 1401-1409, November 2000
- [RUD 95] Rudnick H., Palma R., Fernandez J.E., "Marginal Pricing and Supplement Cost Allocation in Transmission Open Access", IEEE Transactions on Power Systems, vol. 10, n°2, p. 1125-1142, Mai 1995

- [RTE1 02] Réseau de Transport d'Electricité (RTE), Etude comparative de différentes méthodes d'allocation des capacités de transport d'électricité à l'interconnexion France-Belgique, Juillet 2002
<http://www.rte-france.com>
- [RTE2 02] RTE, Mémento de la sûreté du système électrique, édition 2002
- [SEE 99] Seeley K., Lawarree J., Liu C. C., "Coupling Market Clearing with Congestion Management to Prevent Strategic Bidding" Proceedings of IEEE PES Summer Meeting, Juillet 1999.
- [SEE 02] SEE, Les transformateurs déphaseurs: comment doper le réseau électrique THT face aux attentes nouvelles du marché?, ELEC 2002, Paris, décembre 2002
- [SIN 98] Singh H., Hao S., Papalexopoulos A., "Transmission Congestion Management in Competitive Electricity Markets", IEEE Transactions on Power Systems, vol. 13, n°2, p. 672-680, Mai 1998
- [SCH 88] Schweeppe F.C., Caramanis M.C., Tabors R.D., Bohn R.E., Spot Pricing of Electricity, Kuwer Academic, 1988
- [SVE 03] Svenska Kraftnät, The Swedish Electricity Market and the Role of Svenska Kraftnät, Décembre 2003
<http://www.svk.se>
- [TAO 02] Tao S., Gross G., "A Congestion Management Allocation Mechanism for Multiple Transaction Networks", IEEE Transactions on Power System, vol.17, n°3, p. 826-833, Août 2002
- [TOR 99] Torre D.L.T., Feltes J.W., Gomez T.S.R, Merrill H.M., "Deregulation, Privatization And Competition: Transmission Planning under Uncertainty", IEEE Transactions on Power Systems, vol.14, n°2, p. 460-465, Mai 1999
- [TOR 01] Torres G.L. Quintana V.H., "On a Nonlinear Multiple-Centraly-Corrections Interior Point Method for Optimal Power Flow", IEEE Transactions on Power Systems, vol.16, n°2, p. 222-228, Mai 2001
- [WAN 00] Wang X., Song Y.H., "Advanced Real-time Congestion Management through both Pool Balancing Market and Bilateral Market", IEEE Power Engineering Review, p.47-49, Février 2000
- [YAN 99] Yang J., Anderson M.D., "Tracing the flows of power in transmission networks for use-of-transmission –system charges and congestion management", PES Winter Meeting IEEE vol.1, 1999

ANNEXE 1

Calcul de répartition de charge dit « à courant continu » (Modèle DC)

A1.1) les modèles de calcul de répartition de charge

Il existe deux modèles de calcul de répartition de charge :

- Le modèle dit « à courant alternatif » ou modèle « AC » : il calcule les transits actifs et réactifs, détermine le profil de tension, les phases des tensions nodales et l'énergie réactive fournie par chaque nœud à tension fixée. Le modèle AC est un modèle quadratique non linéaire.
- Le modèle dit à « courant continu » ou modèle « DC » : c'est une linéarisation du modèle procédant, qui ne modélise pas les échanges de puissance réactive, et qui ne tient pas en compte des différences de tension entre chaque nœud, ni des pertes actives (pour le modèle DC sans pertes). De par sa simplicité et sa linéarité, c'est celui auquel nous allons nous intéresser plus particulièrement.

A1.1.1) Expression de la puissance transitée dans le modèle AC [NAS 99]

Soit une ligne d'un réseau de transport telle que celle de la figure A1.1. Nous pouvons définir à ses deux extrémités le nœud « émetteur » i et le nœud « récepteur » j , les tensions V_i et V_j aux nœuds i et j , ainsi que l'impédance de la ligne Z_{ij} . Cette impédance est composée de sa partie résistive r_{ij} et de sa partie inductive jx_{ij} :

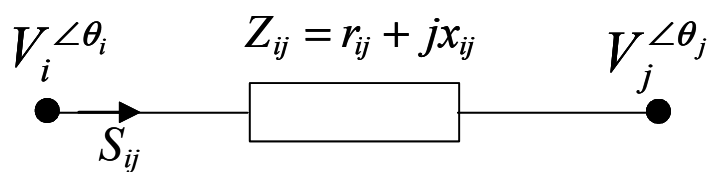


Figure A1.1 : modèle d'une ligne

L'impédance de la ligne peut aussi se mettre sous cette forme complexe:

$$Z_{ij} = |Z_{ij}| * e^{j\varphi}$$

$$\text{avec } \varphi = \arctan\left(\frac{x_{ij}}{r_{ij}}\right)$$

La puissance apparente sortant du nœud i et transitant dans la ligne ij peut s'écrire comme la somme de la puissance active P_{ij} et la puissance réactive complexe jQ_{ij} :

$$S_{ij} = P_{ij} + jQ_{ij}$$

La puissance active dans le modèle AC est donc égale à :

$$P_{ij} = \frac{|V_i|^2 * \cos \varphi}{|Z_{ij}|} - \frac{|V_i||V_j| * \cos(\theta_i - \theta_j + \varphi)}{|Z_{ij}|} \quad (\text{A1.1})$$

avec θ_i et θ_j étant les phases de la tension au nœud i et au nœud j .

Pour calculer les transits actifs, il faut donc connaître les tensions nodales complexes, ce qui nécessite de recourir à un processus itératif. Les méthodes de résolution les plus connues sont les méthodes de Gauss-Seidel et la méthode de Newton-Raphson.

A1.1.2) Modèle DC (sans pertes)

Le modèle DC sans pertes est basé sur les hypothèses suivantes :

- Pour tout nœud i et j , le module de tension au nœud i est très proche du module de tension au nœud j . Dans la pratique, les plages de variation de tension permises autour de la tension de consigne sont en général resserrées (+/- 5% en HTA), ce qui rend acceptable l'hypothèse
- Les différences de phases nodales ne sont pas très élevées
- Les pertes par effet joules dues à la résistance r_{ij} étant dans les réseaux de transport faibles par rapport au reste de la puissance active transitée, la résistance des lignes n'est pas représentée¹⁰

¹⁰ Bien qu'il soit possible d'introduire les pertes dans le modèle DC, nous allons exclusivement nous baser sur le modèle DC sans pertes par la suite

Ainsi, dans le cadre de ces hypothèses, on fait les approximations suivantes :

$$V_i = V_j$$

$$\sin(\theta_i - \theta_j) = \theta_i - \theta_j$$

Dérivée de la relation (A1.2), l'expression la puissance active transitée dans la ligne reliant le nœud i au nœud j dans le modèle DC sans pertes est donc :

$$P_{ij} = \frac{\theta_i - \theta_j}{x_{ij}} = b_{ij}(\theta_i - \theta_j) \quad (\text{A1.2})$$

avec b_{ij} définie comme l'admittance de la ligne connectant le nœud i au nœud j

Ainsi, dans le modèle DC, la puissance active transitée entre deux nœuds du réseau est directement proportionnelle à la différence de phases nodales entre ces nœuds. Si on étend la relation (A1.5) pour tous les transits d'un réseau à N nœuds et comportant Nb branches, on peut l'exprimer sous cette forme :

$$\mathbf{P}_{ij} = \mathbf{H} * \boldsymbol{\theta} \quad (\text{A1.3})$$

avec

$\mathbf{P}_{ij}$ vecteur des transits de dimension Nb

$\mathbf{H}$ matrice reliant les phases des tensions nodales aux transits de dimension $Nb * N$

$\boldsymbol{\theta}$ vecteur des phases des tensions nodales de dimension N

Ainsi, pour le transit du nœud i au nœud j , les coefficients de la matrice $\mathbf{H}$ prennent les valeurs suivantes :

$$H(ij, n) = \frac{1}{x_{ij}} \text{ pour } n = i$$

$$H(ij, n) = -\frac{1}{x_{ij}} \text{ pour } n = j$$

$$H(ij, n) = 0 \text{ pour les autres nœuds du réseau}$$

D'autre part l'injection nette au nœud i est définie comme la différence algébrique entre les puissances générées en ce nœud et les puissances consommées en ce nœud. Cette injection

nette au nœud i représente aussi le bilan des puissances transitées par ce nœud (puissances entrantes et puissances sortantes) :

$$P_i = P_{Gi} - P_{Ci} = \sum_{j \in C} P_{ij} = \sum_{j \in C} \frac{\theta_i - \theta_j}{x_{ij}} \quad (\text{A1.4})$$

avec

P_i injection nette au nœud i

P_{Gi} production au nœud i

P_{Ci} charge au nœud i

C le domaine des nœuds connectés au nœud i

La relation (A1.4) eut aussi se mettre sous forme matricielle :

$$\mathbf{P} = \mathbf{B} * \boldsymbol{\theta} \quad (\text{A1.5})$$

avec

$\mathbf{P}$ vecteur des injections nettes nodales de dimension N

$\mathbf{B}$ matrice des admittances nodales de dimension $N*N$

Les coefficients de la matrice $\mathbf{B}$ prennent ces valeurs-ci :

$$B(i, j) = -b_{ij} = -\frac{1}{x_{ij}} \text{ pour } j \in C$$

$$B(i, i) = \sum_{j \in C} b_{ij}$$

$$B(i, j) = 0 \text{ pour } j \notin C$$

La matrice $\mathbf{B}$ est fonction des admittances des lignes composant le réseau, ainsi que de la topologie du réseau. Toutefois, telle quelle, elle n'est pas inversible. Pour rendre la matrice inversible, on doit choisir un nœud de référence, appelé nœud bilan. Sa phase est fixée à 0 et

sert de référence pour le calcul des autres phases nodales. En outre, il effectue le bilan des puissances actives générées et consommées sur le réseau. Ainsi, si nous choisissons le nœud 1 comme nœud de référence dans un réseau à N nœuds, nous obtenons :

$$\begin{bmatrix} \theta_2 \\ \vdots \\ \theta_N \end{bmatrix} = \begin{bmatrix} & & \\ & \mathbf{X} & \\ & & \end{bmatrix} * \begin{bmatrix} P_2 \\ \vdots \\ P_N \end{bmatrix} \quad (\text{A1.8})$$

Phase de référence $\theta_1 = 0$

Bilan des injections nettes $P_1 = \sum_{i=2}^N P_i$

Ici, la matrice $\mathbf{X}$ est l'inverse de la matrice $\mathbf{B}$, dont on a retiré la ligne et la colonne correspondant au nœud bilan. Nous avons donc une relation linéaire entre les transits et les phases nodales (A1.3) et une relation linéaire entre les injections nodales nettes et les phases nodales (A1.5). Nous pouvons donc relier linéairement les injections nodales nettes aux transits :

$$\mathbf{P}_{ij} = \mathbf{A} * \mathbf{P} \quad (\text{A1.9})$$

avec

$$\mathbf{A} = \mathbf{H} * \mathbf{X} \quad (\text{A1.10})$$

La matrice $\mathbf{A}$ est la matrice des *coefficients de sensibilité*, coefficients appelés aussi *facteurs de distribution*. Ces facteurs dépendent de la topologie du réseau et du choix du nœud bilan. La matrice $\mathbf{A}$ est de dimension $Nb*(N-1)$ comme on a exclu le nœud bilan. Toutefois, on peut tenir compte du nœud bilan en rajoutant une colonne de zéros à la matrice $\mathbf{A}$ à l'emplacement du nœud bilan.

Ainsi, deux propriétés fondamentales se dégagent du modèle DC :

- La *linéarité* : si les transferts de puissance au sein du réseau sont augmenté linéairement d'un facteur x , les transits augmenteront tous aussi d'un facteur x .
- La *superposition* : soit les transits issus respectivement d'un état 1 et d'un état 2 du réseau. Si on superpose ces deux états, les nouveaux transits obtenus sont la somme algébrique des transits de l'état 1 et de l'état 2.

Pour illustrer le calcul de répartition de charge dit à courant continu et ses propriétés, considérons un petit réseau de 3 nœuds tel que celui de la figure A1.3. Les admittances des lignes sont données en unités relatives (u.r). Nous voulons conduire sur ce réseau deux transactions T1 et T2 de 100 MW.

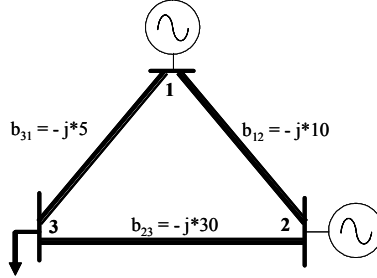


Figure A1.3 : réseau à 3 nœuds et 3 branches

Nous avons défini dans l'exemple les transits de la façon suivante :

$$P_{12} = b_{12} * (\theta_1 - \theta_2)$$

$$P_{23} = b_{23} * (\theta_2 - \theta_3) \quad (A1.11)$$

$$P_{31} = b_{31} * (\theta_3 - \theta_1)$$

On choisit le nœud 1 comme nœud de référence. D'où les matrices calculables à l'aide de la topologie de ce réseau qui sont:

$$\mathbf{H} = \begin{bmatrix} 10 & -10 & 0 \\ 0 & 30 & -30 \\ -5 & 0 & 5 \end{bmatrix}$$

$$\mathbf{B} = \begin{bmatrix} 15 & -10 & -5 \\ -10 & 40 & -30 \\ -5 & -30 & 35 \end{bmatrix}$$

$$\mathbf{A} = \begin{bmatrix} 0 & -0.7 & -0.6 \\ 0 & 0.3 & -0.6 \\ 0 & 0.3 & 0.4 \end{bmatrix}$$

On peut alors déterminer les phases nodales ainsi que les transits en fonction des injections nodales. La figure A1.4 nous montre les flux induits par le transfert de puissance T1 seul, les flux induits par le transfert de puissance T2 seul, et les flux de puissance obtenus si on voulait conduire ces deux transactions en même temps.

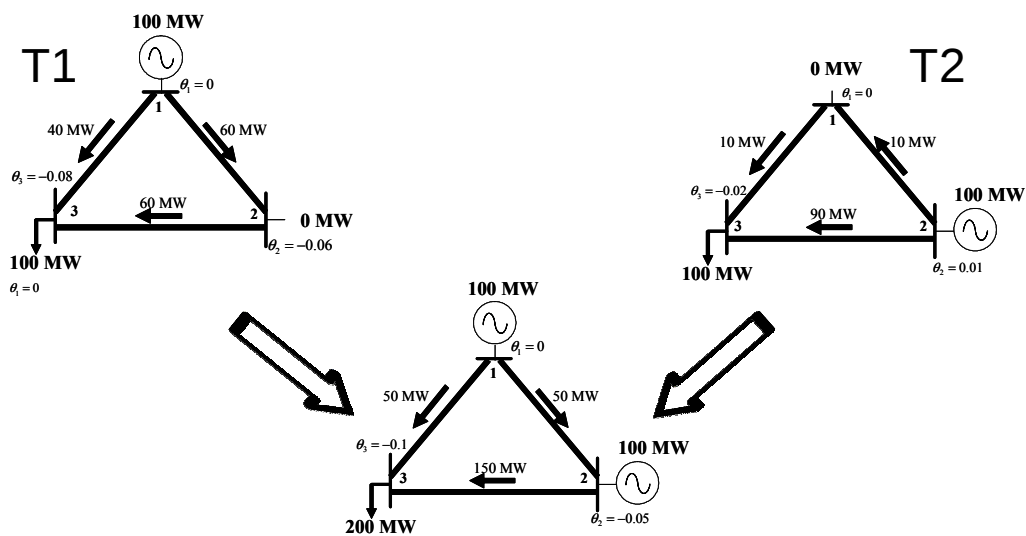


Figure A1.4 : transaction T1 et T2 et état du réseau et des transits

Les propriétés du modèle DC sont utilisées dans beaucoup de méthodes de traitement des congestions. Par exemple, on peut analyser l'impact d'une transaction sur une congestion à l'aide des facteurs de distribution, comme cela se fait dans les méthodes par coupure de transaction (nous reviendrons plus précisément sur cette méthode plus loin).

On peut même aisément déterminer une reconfiguration de la production (ou *redispatching*) pour modifier la répartition des transits, tout en maintenant le bilan électrique global constant. Dans l'exemple précédant, si on voulait limiter le transit du nœud 2 au nœud 3 à 140 MW, on doit donc réduire son transit de 10 MW. A partir de la relation A1.9, on trouve :

$$\Delta P_{12} = \begin{bmatrix} 0 & 0.3 & -0.6 \end{bmatrix} * \begin{bmatrix} -33.33 \\ +33.33 \\ 0 \end{bmatrix} = -10 \text{ MW}$$

Il faut donc dans cet exemple transférer 33.33 MW du nœud 2 au nœud 1 pour pouvoir réduire le flux transité du nœud 2 au nœud 3 de 10 MW. Cette technique de contrôle des transits par reconfiguration de la production est fondamentale pour le traitement des congestions et sert de base aux principales méthodes appliquées actuellement.

Les facteurs de distribution contenus dans la matrice A ont une signification physique particulière. Le facteur de distribution $A_{l,i}$ représente le changement de transit de la ligne l si on modifie l'injection de +1 MW au nœud i , avec une variation d'injection opposée située au nœud bilan. Si on reprend l'exemple du réseau à 3 nœuds, on peut voir que le facteur $A(1,2)$ est égal à -0.7. Cela signifie qu'une augmentation de +1 MW au nœud 2 associée à une baisse de -1 MW au nœud bilan (nœud 1) provoque une baisse de transit de -0.7 MW dans le sens conventionnel de transit de la ligne 1 (sens qui va du nœud 1 au nœud 2, voir relation (A1.11)). On peut aisément vérifier ce résultat en effectuant un calcul de répartition de charge sur le réseau à 3 nœuds avec 99 MW au nœud 1, 101 MW au nœud 2, et -100 MW de consommation au nœud 3. Le transit de la ligne 1 du nœud 1 au nœud 2 est alors de 49.3 MW au lieu de 50 MW (Fig. A1.4).

Le modèle AC est certes plus complet que le modèle DC, et aussi celui qui offre la plus grande précision. Toutefois, il est plus opaque, plus difficile à manier et donne des temps de calcul beaucoup plus élevés que le modèle DC, qui lui est plus « direct » d'application grâce à sa linéarité. En outre, une autre propriété exploitable dans notre cas est le découplage physique qu'il est possible de faire entre l'énergie active et l'énergie réactive. En effet, les puissances actives transitées sont surtout sensibles aux différences de phases nodales θ , et peu sensibles aux modules des tensions $|V|$, ce qui est l'inverse pour les transits d'énergie réactive. Ainsi, il est donc possible de traiter à part les problèmes liés aux transits actifs (dont le traitement des congestions) des problèmes où les transits d'énergie réactive jouent un rôle majeur (dont le réglage de la tension, qui peut être traité localement). Cette propriété est notamment utilisée dans le *Fast Decoupled Load-Flow* ou méthode « découplée rapide ». En effet, suivant la méthode Newton-Raphson [NAS 99], on peut exprimer les variations de puissance active et réactive en fonction des variations des phases nodales et des tensions par l'intermédiaire d'un Jacobien de cette façon

$$\begin{bmatrix} \Delta \mathbf{P} \\ \Delta \mathbf{Q} \end{bmatrix} = \begin{bmatrix} \mathbf{J} \end{bmatrix} \begin{bmatrix} \Delta \boldsymbol{\theta} \\ \Delta \mathbf{V} \end{bmatrix} = \begin{bmatrix} \frac{\partial \mathbf{P}}{\partial \boldsymbol{\theta}} & \frac{\partial \mathbf{P}}{\partial \mathbf{V}} \\ \frac{\partial \mathbf{Q}}{\partial \boldsymbol{\theta}} & \frac{\partial \mathbf{Q}}{\partial \mathbf{V}} \end{bmatrix} \begin{bmatrix} \Delta \boldsymbol{\theta} \\ \Delta \mathbf{V} \end{bmatrix} \quad (\text{A1.11})$$

La méthode découplée rapide, notamment présentée en [HAD 92] néglige les grandeurs $\partial \mathbf{P} / \partial \mathbf{V}$ et $\partial \mathbf{Q} / \partial \boldsymbol{\theta}$ ce qui lui permet de gagner en temps de calcul, tout en restant suffisamment précise. La faible corrélation physique entre les transits actifs et réactifs fait que les transits actifs calculés à l'aide du modèle DC sont en pratique très proches de ceux calculés à l'aide du modèle AC.

ANNEXE 2

Théorie de l'optimisation

A2.1) Optimisation sans contraintes [CUL 94]

Un problème général d'optimisation sans contraintes peut se mettre sous cette forme :

$$\min_{\mathbf{x} \in U^{ad}} f(\mathbf{x}) \quad (\text{A2.1})$$

avec :

$\mathbf{x}$: vecteur de n paramètres

$f(\mathbf{x})$: fonction objectif

U^{ad} : domaine d'admissibilité du vecteur $\mathbf{x}$

On définit un minimum local de f sur U^{ad} par un vecteur $\mathbf{x}^*$ tel qu'il existe une boule $B(\mathbf{x}^*, \rho)$ de centre $\mathbf{x}^*$ et de rayon $\rho > 0$ telle que :

$$\forall \mathbf{x} \in B(\mathbf{x}^*, \rho) \cap U^{ad}, \quad f(\mathbf{x}) \geq f(\mathbf{x}^*)$$

Un minimum global de f sur U^{ad} est un vecteur $\mathbf{x}^{**}$ tel que

$$\forall \mathbf{x} \in U^{ad}, \quad f(\mathbf{x}) \geq f(\mathbf{x}^{**})$$

Dans un cas général, avec les méthodes de type descente telle que la méthode de Newton, on doit se contenter seulement d'un minimum local. Toutefois, si des hypothèses peuvent être faites sur les propriétés de f (optimisation linéaire, convexité, fonction quadratique), l'optimum trouvé peut être global. Pour trouver un minimum, toutes les méthodes de descente partent du point de vue suivant : on connaît un point $\mathbf{x}$ ainsi que la valeur $f(\mathbf{x})$; en se déplaçant localement à partir de $\mathbf{x}$, on peut déterminer si f augmente ou diminue. Le déplacement le plus simple étant la ligne droite, on peut définir un certain nombre de déplacements autour de $\mathbf{x}$ dans des directions d en vue d'améliorer la valeur du critère f .

A2.1.1) Conditions d'optimalité

On appelle direction admissible d en $\mathbf{x}$ un vecteur d le long duquel on pourra se déplacer en partant de $\mathbf{x}$ tout en restant dans le domaine U^{ad} .

Supposons que f est différentiable sur U^{ad} , son gradient étant noté $\nabla f(\mathbf{x})$. Alors si $\mathbf{x}^*$ est un minimum local de f sur U^{ad} , pour toute direction d admissible en $\mathbf{x}^*$ on a :

$$\nabla f(\mathbf{x}^*)d \geq 0 \quad (\text{A2.2})$$

Les conditions d'optimalité sont illustrées par la Figure A2.1 :

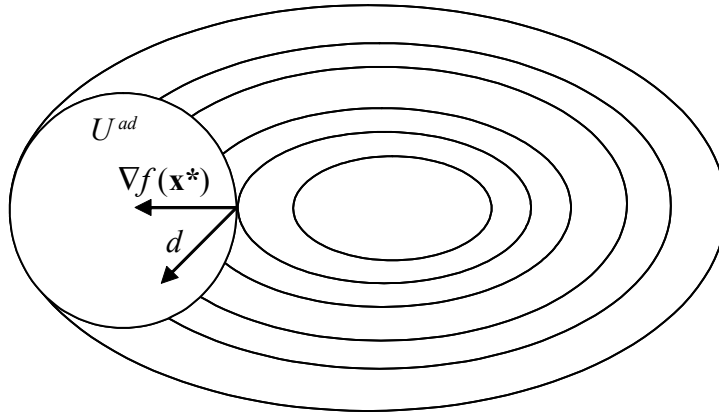


Figure A2.1 : illustration des conditions d'optimalité

A2.1.2) Cas des fonctions convexes

Si l'on veut se placer dans des cas où l'optimum est global et non simplement local, il faut que la fonction f soit convexe. Une fonction f est définie convexe sur un ensemble Y si :

$$\forall x_1, x_2 \in Y, \forall t \in [0,1], f(tx_1 + (1-t)x_2) \leq tf(x_1) + (1-t)f(x_2) \quad (\text{A2.3})$$

Les hypothèses de forte convexité sont notamment vérifiées pour les fonctions quadratiques de la forme :

$$f(\mathbf{x}) = \frac{1}{2} \mathbf{x}^T \mathbf{A} \mathbf{x} - \mathbf{b}^T \mathbf{x} \quad (\text{A2.4})$$

La fonction objectif choisie pour le traitement des congestions par ajustements de production (Chapitre III) se ramène à cette forme. Dans le cas des fonctions convexes, les conditions d'optimalité font que tout minimum local est aussi un minimum global, d'où on peut en déduire que l'optimum trouvé avec un algorithme de descente est unique dans le cas des fonctions convexes.

A2.1.3) Algorithme de Newton

Le principe de la méthode de Newton pour l'optimisation est de minimiser successivement les approximations au second ordre de la fonction f . Soit $\mathbf{x} \in U$, par un développement de Taylor au second ordre au voisinage de $\mathbf{x}$, on obtient :

$$f^0(\mathbf{x} + \boldsymbol{\varepsilon}) = f(\mathbf{x}) + \nabla f(\mathbf{x})^T \boldsymbol{\varepsilon} + \frac{1}{2} \boldsymbol{\varepsilon}^T \nabla^2 f(\mathbf{x}) \boldsymbol{\varepsilon} \quad (\text{A2.5})$$

On minimise la fonction quadratique f^0 , ce qui fournit un vecteur $\mathbf{x}_1$:

$$\nabla^2 f(\mathbf{x}) \mathbf{x}_1 = \nabla^2 f(\mathbf{x}) \mathbf{x} - \nabla f(\mathbf{x}) \quad (\text{A2.6})$$

A l'itération k , on construit f^k approximation quadratique de f au voisinage de $\mathbf{x}_k$, que l'on minimise pour obtenir $\mathbf{x}_{k+1}$ défini par :

$$\nabla^2 f(\mathbf{x}_k) \boldsymbol{\delta}_k = -\nabla f(\mathbf{x}_k), \quad \mathbf{x}_{k+1} = \mathbf{x}_k + \boldsymbol{\delta}_k \quad (\text{A2.7})$$

La méthode de Newton est la généralisation multi-dimensionnelle de la méthode de Newton-Raphson appliquée à la recherche des racines de $G(\mathbf{x}) = \nabla f(\mathbf{x})$. Cette méthode fonctionne très bien pour des problèmes de petites dimensions (quelques dizaines de variables) lorsque le calcul du Hessien $\nabla^2 f(\mathbf{x})$ est facile. Dans les autres cas, on préfère une méthode du Gradient conjugué ou de Quasi-Newton. La prochaine section sera consacrée à cette dernière.

A2.1.4) Algorithme de Quasi-Newton

La méthode de Newton consiste à imiter l'algorithme de Newton, mais sans calculer le Hessien de f ni son inverse. Au lieu de procéder à l'itération théorique $\mathbf{x}_{k+1} = \mathbf{x}_k - [\nabla^2 f(\mathbf{x}_k)]^{-1} \nabla f(\mathbf{x}_k)$, on calcule :

$$\mathbf{x}_{k+1} = \mathbf{x}_k - \alpha_k \mathbf{S}^k \nabla f(\mathbf{x}_k)$$

avec $\mathbf{S}^k$ une approximation symétrique définie positive de $[\nabla^2 f(\mathbf{x}_k)]^{-1}$ et α_k un paramètre positif fourni par une recherche linéaire le long de la direction $\mathbf{d}^k = -\mathbf{S}^k \nabla f(\mathbf{x}_k)$. Plus $\mathbf{S}^k$ sera proche au sens d'une norme matricielle de $[\nabla^2 f(\mathbf{x}_k)]^{-1}$, plus l'algorithme convergera rapidement. A l'étape k , la remise à jour de la matrice $\mathbf{S}^k$ se fait avec une formule simple additive :

$$\mathbf{S}^{k+1} = \mathbf{S}^k + \mathbf{C}^k \quad (\text{A2.8})$$

$\mathbf{C}^k$ étant une matrice de correction qui intègre mieux la nouvelle information fournie par $\mathbf{x}_{k+1}$ et $\nabla f(\mathbf{x}_{k+1})$. $\mathbf{C}^k$ sera donc choisie de telle manière que $\mathbf{S}^{k+1}$ satisfasse l'équation suivante, dite « équation de la sécante » ou « condition de Quasi-Newton » :

$$\mathbf{S}^{k+1} \boldsymbol{\gamma}^k = \boldsymbol{\delta}^k \text{ avec } \boldsymbol{\gamma}^k = (\nabla f(\mathbf{x}_{k+1})) - \nabla f(\mathbf{x}_k) \quad (\text{A2.9})$$

Correction de rang 1 :

Supposons que la mise à jour de la matrice $\mathbf{S}^k$ se fasse de la façon suivante :

$$\mathbf{S}^{k+1} = \mathbf{S}^k + \alpha_k \mathbf{v}_k (\mathbf{v}_k)^T \quad (\text{A2.10})$$

avec α_k une constante et $\mathbf{v}_k$ un vecteur. Nous pouvons montrer que pour un certain choix de α_k et de $\mathbf{v}_k$, la condition exprimée en (A2.9) est satisfaite. Dans ce cas-là, on doit alors avoir :

$$\mathbf{S}^{k+1} \boldsymbol{\gamma}^k = \boldsymbol{\delta}^k \Rightarrow \mathbf{S}^k \boldsymbol{\gamma}^k + \alpha_k \mathbf{v}_k [(\mathbf{v}_k)^T \boldsymbol{\gamma}^k] = \boldsymbol{\delta}^k \quad (\text{A2.11})$$

De là, on en déduit que pour que la condition en (A2.9) soit satisfaite, on doit prendre $\mathbf{v}_k$ proportionnel à $\boldsymbol{\delta}^k - \mathbf{S}^k \boldsymbol{\gamma}^k$ et α_k tel que $\alpha_k [(\mathbf{v}_k)^T \boldsymbol{\gamma}^k] = 1$. Ainsi, si l'on prend $\mathbf{v}_k = \boldsymbol{\delta}^k - \mathbf{S}^k \boldsymbol{\gamma}^k$, on a :

$$\mathbf{S}^{k+1} = \mathbf{S}^k + \frac{(\boldsymbol{\delta}^k - \mathbf{S}^k \boldsymbol{\gamma}^k)(\boldsymbol{\delta}^k - \mathbf{S}^k \boldsymbol{\gamma}^k)^T}{(\boldsymbol{\delta}^k - \mathbf{S}^k \boldsymbol{\gamma}^k)^T \boldsymbol{\gamma}^k} \quad (\text{A2.12})$$

D'autres méthodes de remise à jour de la matrice $\mathbf{S}^k$ existent, telle que la méthode DFP (David-Fletcher-Powell) et BFGS (Broyden-Fletcher-Goldfarb-Shanno). Celles-ci sont décrites plus en détail dans [CUL 94].

A2.2) Optimisation sous contraintes

Un problème général d'optimisation sous contraintes peut se mettre sous cette forme :

$$\min_{\mathbf{x} \in U^{ad}} f(\mathbf{x}) \quad (\text{A2.13})$$

avec comme contraintes :

$$\underline{\mathbf{h}} \leq h(\mathbf{x}) \leq \bar{\mathbf{h}} \quad (\text{A2.14})$$

et

$$g(\mathbf{x}) = 0 \quad (\text{A2.15})$$

où $h(\mathbf{x})$ est un vecteur représentant m contraintes inégalités, $\underline{\mathbf{h}}$ le vecteur des bornes inférieures des contraintes inégalité de longueur m , $\bar{\mathbf{h}}$ le vecteur des bornes supérieures des contraintes inégalités de longueur m , et $g(\mathbf{x})$ un vecteur représentant n contraintes égalité.

Pour résoudre le problème formulé ci-dessus, on introduit une nouvelle fonctionnelle, le Lagrangien, défini par :

$$\Gamma(\mathbf{x}, \mathbf{p}) = f(\mathbf{x}) + \mathbf{p}^T \theta(\mathbf{x}) \quad (\text{A2.16})$$

avec

- $\mathbf{p}$ multiplicateur de Lagrange (contraintes égalités) ou multiplicateur de Karush-Kuhn-Tucker (contraintes inégalités) de dimension $m + n$
- $\theta(\mathbf{x})$ vecteur des contraintes égalité et inégalité

Lorsque f et θ sont convexes, les conditions obtenues sont automatiquement nécessaires et suffisantes. Les conditions d'optimalité du premier ordre sont les suivantes : si $\mathbf{x}^*$ est un minimum local, alors il existe un vecteur $\mathbf{p}^*$ tel que :

$$\nabla f(\mathbf{x}^*) + \sum_{i=1}^{m+n} p_i^* \nabla \theta_i(\mathbf{x}^*) = 0 \quad (\text{A2.17}) \quad \text{et} \quad (\mathbf{p}^*)^T \theta(\mathbf{x}^*) = 0 \quad (\text{A2.18})$$

L'expression ci-dessus nous illustre l'interprétation économique des multiplicateurs de Lagrange qui est faite notamment dans la tarification marginale : si f est une fonction de coût à minimiser et que θ contienne les équations égalité du calcul de répartition de charge, on voit que le vecteur $\mathbf{p}^*$ exprime la variation de ce coût par rapport à la variation du niveau de contrainte. En particulier, si on suppose une variation unitaire de consommation en un nœud donné du réseau (modifiant ainsi le niveau de contrainte du problème), on peut facilement tirer à partir des multiplicateurs de Lagrange la variation de coût correspondante affectant la fonction objectif.

Une autre propriété du Lagrangien vient d'une condition dite « des écarts complémentaires » qui se décompose en $m + n$ conditions $p_j^* \theta_j(\mathbf{x}^*) = 0$. En effet, chaque p_j^* et $\theta_j(\mathbf{x}^*)$ associés sont toujours de signes contraires, de sorte que tous les termes de $\mathbf{p}^T \boldsymbol{\theta}(\mathbf{x})$ sont de signe positif. Chacun doit être nul pour que leur somme le soit (en effet, minimiser une fonction objectif sous contraintes revient en fait à minimiser cette fonction additionnée de son Lagrangien). De là viennent les propriétés suivantes des multiplicateurs considérés :

- Les multiplicateurs de Lagrange (contraintes égalité) sont toujours non nuls
- Un multiplicateur de Karush-Kuhn-Tucker (contraintes inégalité) est nul si la contrainte n'est pas activée, c'est-à-dire si $\theta_j(\mathbf{x}^*) < 0$
- Si une contrainte inégalité est activée, c'est-à-dire si $\theta_j(\mathbf{x}^*) = 0$, alors p_j^* est non nul

Les conditions de Karush-Kuhn-Tucker (KKT) se prêtent à une interprétation géométrique très simple : le vecteur $-\nabla f(\mathbf{x}^*)$ (direction de descente) est combinaison linéaire des gradients des contraintes $\nabla \theta_i(\mathbf{x}^*)$ (Fig. A2.2).

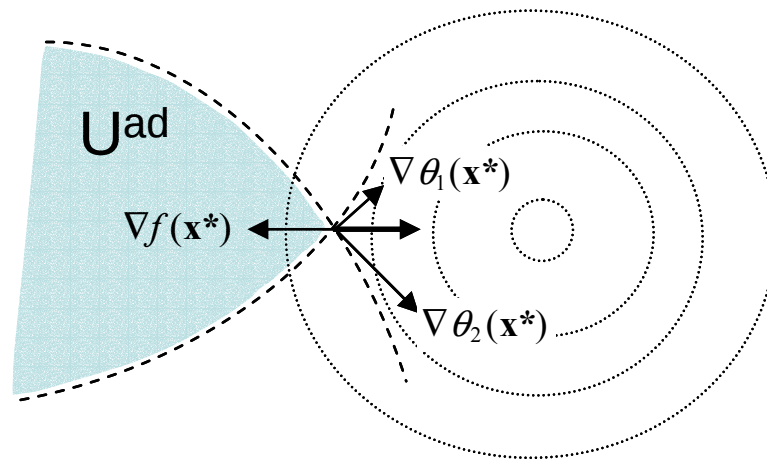


Figure A2.2 : illustration des conditions de Karush-Kuhn-Tucker :
 $-\nabla f(\mathbf{x}^*) = p_1 \nabla \theta_1(\mathbf{x}^*) + p_2 \nabla \theta_2(\mathbf{x}^*)$

A2.3) Méthode de type Point Intérieur : algorithme primal-dual [QUI 99], [TOR 01]

A l'origine, les méthodes de type « Point Intérieur » ont été conçues pour résoudre les problèmes de programmation non linéaire. Des recherches plus approfondies sur ces méthodes ont montré qu'elles donnaient de très bonnes performances en terme de vitesse de convergence pour les problèmes de grande échelle. L'algorithme présenté dans cette section,

connu sous le nom d' « algorithme primal-dual » est l'un des plus utilisés. Le principe de cette méthode est de rajouter à la fonction objectif une fonction logarithmique « barrière » incluant des contraintes et qui décroît progressivement au fil de l'optimisation pour tendre vers 0. Typiquement, considérons un problème de la forme :

$$\min_{\mathbf{x} \in U^{ad}} f(\mathbf{x}) \quad \text{avec } h(\mathbf{x}) \geq 0$$

On peut théoriquement ce problème contraint en incorporant les contraintes inégalités dans la fonction objectif, transformant le problème ci-dessus en un problème non contraint :

$$\min_{\mathbf{x} \in U^{ad}} f_{\mu}(\mathbf{x}, \mu^k) \quad \text{avec} \quad f_{\mu}(\mathbf{x}, \mu^k) = f(\mathbf{x}) - \mu^k \sum_i \ln h_i(\mathbf{x}) \quad (\text{A2.19})$$

où $\mu^k > 0$ est un paramètre de pénalisation qui tend vers 0 au fil des itérations par remise à jour appropriée. Le choix de la valeur initiale de μ^0 ainsi que sa procédure de remise à jour doivent être choisies de manière judicieuses pour éviter les problèmes de divergence.

A2.3.1) Formulation du problème et conditions d'optimalité

La méthode de Point Intérieur décrite ici transforme d'abord les contraintes inégalités contenues en (A2.14) en contraintes égalités en rajoutant des vecteurs tampons, comme suit :

$$\min_{\mathbf{x} \in U^{ad}} f(\mathbf{x})$$

avec

$$-\mathbf{s} - \mathbf{z} - \underline{\mathbf{h}} + \bar{\mathbf{h}} = \mathbf{0}$$

$$-\mathbf{z} - h(\mathbf{x}) + \bar{\mathbf{h}} = \mathbf{0}, \quad \mathbf{s} \geq 0 \quad \text{et} \quad \mathbf{z} \geq 0$$

$$g(\mathbf{x}) = 0$$

Ensuite, les conditions de non-négativité ($\mathbf{s} \geq 0$ et $\mathbf{z} \geq 0$) sont incorporées dans une fonction logarithmiques de cette façon :

$$\min_{\mathbf{x} \in U^{ad}} f(\mathbf{x}) - \mu^k \sum_i (\ln s_i + \ln z_i) \quad (\text{A2.20})$$

avec

$$-\mathbf{s} - \mathbf{z} - \underline{\mathbf{h}} + \bar{\mathbf{h}} = \mathbf{0} \quad (\text{A2.21})$$

$$-\mathbf{z} - h(\mathbf{x}) + \bar{\mathbf{h}} = \mathbf{0}, \quad \mathbf{s} > 0 \quad \text{et} \quad \mathbf{z} > 0 \quad (\text{A2.22})$$

$$g(\mathbf{x}) = 0 \quad (\text{A2.23})$$

Les termes logarithmiques imposent une condition de stricte positivité des vecteurs $\mathbf{s}$ et $\mathbf{z}$.

Pour résoudre le problème décrit par les relations (A2.20) à (A2.23), on peut utiliser la méthode de Newton. On associe donc au problème formulé un Lagrangien qui est donné par :

$$L(\mathbf{y}) = f(\mathbf{x}) - \mu_k \sum_i (\ln(s_i) + \ln(z_i)) - \lambda^T \mathbf{g}(\mathbf{x}) - \pi^T (-\mathbf{s} - \mathbf{z} - \underline{\mathbf{h}} + \bar{\mathbf{h}}) - \mathbf{v}^T (-\mathbf{h}(\mathbf{x}) - \mathbf{z} + \bar{\mathbf{h}})$$

$$\mathbf{y} = [\mathbf{s}, \mathbf{z}, \boldsymbol{\pi}, \mathbf{v}, \mathbf{x}, \boldsymbol{\lambda}]$$

(A2.24)

où $\boldsymbol{\lambda}, \boldsymbol{\pi}$ et $\mathbf{v}$ sont les vecteurs des multiplicateurs de Lagrange (aussi appelés variables duales) et $\mathbf{y}$ le vecteur définissant l'état du point courant ainsi que des différentes variables ou multiplicateurs.

Ce Lagrangien doit bien sûr satisfaire les conditions de KKT qui s'écrivent:

$$\nabla L(\mathbf{y}) = \begin{pmatrix} \boldsymbol{\pi} - \mu_k \mathbf{S}^{-1} \mathbf{e} \\ \hat{\mathbf{v}} - \mu_k \mathbf{Z}^{-1} \mathbf{e} \\ -\mathbf{s} - \mathbf{z} - \underline{\mathbf{h}} + \bar{\mathbf{h}} \\ -\mathbf{h}(\mathbf{x}) - \mathbf{z} + \bar{\mathbf{h}} \\ \nabla f(\mathbf{x}) - \mathbf{J}_g(\mathbf{x})^T \boldsymbol{\lambda} + \mathbf{J}_h(\mathbf{x})^T \mathbf{v} \\ -\mathbf{g}(\mathbf{x}) \end{pmatrix} = 0 \quad (\text{A2.25})$$

où $\mathbf{S} = \text{diag}(s_1, \dots, s_m)$, $\mathbf{Z} = \text{diag}(z_1, \dots, z_m)$, $\mathbf{e} = (1, \dots, 1)^T$ et $\hat{\mathbf{v}} = \mathbf{v} + \boldsymbol{\pi}$.

Les conditions de KKT exprimées en (A2.25) peuvent être interprétées comme suit : le troisième, quatrième et sixième terme de (A2.25) avec la condition ($\mathbf{s} \geq 0$ et $\mathbf{z} \geq 0$) assure la faisabilité dite « primale ». Le cinquième terme avec la condition ($\boldsymbol{\pi} \geq 0, \hat{\mathbf{v}} \geq 0$) assure la faisabilité dite « duale » ; enfin, le premier et le second terme représentent les conditions complémentaires avec $\mu^k \neq 0$. L'algorithme primal-dual ne réclame pas nécessairement un point initial strictement faisable, mais des conditions de stricte positivité pour les vecteurs $\mathbf{s}, \mathbf{z}, \boldsymbol{\pi}$ et $\hat{\mathbf{v}}$ doivent être respectées pour tous les points. Pour préserver cette condition les itérations successives de l'algorithme suivent une trajectoire placée dans l'aire positive de l'espace défini par le produit $\mathbf{s}_i \mathbf{z}_i$.

Les méthodes de Point Intérieur primales-duales appliquent la méthode de Newton pour calculer le point suivant au cours de l'optimisation et résoudre le système de KKT exprimé en (A2.25), remettent à jour les variables et réduisent μ^k .

A2.3.2) Résolution suivant les directions de Newton

Bien que le système de KKT exprimé en (A2.25) soit non linéaire, on calcule en général une valeur approchée de sa solution en effectuant une itération de la méthode de Newton. On obtient alors le système suivant :

$$\begin{bmatrix} \mathbf{\Pi} & \mathbf{0} & \mathbf{S} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \hat{\mathbf{Y}} & \mathbf{Z} & \mathbf{Z} & \mathbf{0} & \mathbf{0} \\ \mathbf{I} & \mathbf{I} & \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{I} & \mathbf{0} & \mathbf{0} & \mathbf{J}_h & \mathbf{0} \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{J}_h^T & \nabla_x^2 L_\mu & -\mathbf{J}_g^T \\ \mathbf{0} & \mathbf{0} & \mathbf{0} & \mathbf{0} & -\mathbf{J}_g & \mathbf{0} \end{bmatrix} * \begin{bmatrix} \Delta \mathbf{s} \\ \Delta \mathbf{z} \\ \Delta \boldsymbol{\pi} \\ \Delta \mathbf{v} \\ \Delta \mathbf{x} \\ \Delta \boldsymbol{\lambda} \end{bmatrix} = \begin{bmatrix} \xi_s \\ \xi_z \\ \xi_\pi \\ \xi_v \\ \xi_x \\ \xi_\lambda \end{bmatrix} \quad (\text{A2.26})$$

où $\mathbf{\Pi} = \text{diag}(\pi_1, \dots, \pi_m)$, $\hat{\mathbf{Y}} = \text{diag}(\hat{v}_1, \dots, \hat{v}_p)$ et :

$$\begin{aligned} \xi_s &= -\mathbf{S}\boldsymbol{\pi} + \mu^k \mathbf{e} \\ \xi_z &= -\mathbf{Z}\boldsymbol{\pi} + \mu^k \mathbf{e} \\ \xi_\pi &= -\mathbf{s} - \mathbf{z} - \underline{\mathbf{h}} + \bar{\mathbf{h}} \\ \xi_v &= -h(\mathbf{x}) - \mathbf{z} + \bar{\mathbf{h}} \\ \xi_x &= -\nabla_x f(\mathbf{x}) + \mathbf{J}_g(\mathbf{x})^T \boldsymbol{\lambda} - \mathbf{J}_h(\mathbf{x})^T \mathbf{v} \\ \xi_\lambda &= -g(\mathbf{x}) \end{aligned} \quad (\text{A2.27})$$

Le calcul de $\nabla_x^2 L_\mu$ implique de connaître le Hessien de la fonction objectif f , ainsi que celui des contraintes égalités et inégalités :

$$\nabla_x^2 L_\mu(\mathbf{y}) = \nabla_x^2 f(\mathbf{x}) - \sum_{i=1}^n \lambda_i \nabla_x^2 g_i(\mathbf{x}) + \sum_{j=1}^m \nu_j \nabla_x^2 h_j(\mathbf{x}) \quad (\text{A2.28})$$

A2.3.3) Remise à jour des variables

Les nouvelles variables sont calculées comme suit :

$$\begin{aligned}
\mathbf{x}^{k+1} &= \mathbf{x}^k + \alpha_P^k \Delta \mathbf{x} & \boldsymbol{\lambda}^{k+1} &= \boldsymbol{\lambda}^k + \alpha_D^k \Delta \boldsymbol{\lambda} \\
\mathbf{s}^{k+1} &= \mathbf{s}^k + \alpha_P^k \Delta \mathbf{s} & \boldsymbol{\pi}^{k+1} &= \boldsymbol{\pi}^k + \alpha_D^k \Delta \boldsymbol{\pi} \\
\mathbf{z}^{k+1} &= \mathbf{z}^k + \alpha_P^k \Delta \mathbf{z} & \mathbf{v}^{k+1} &= \mathbf{v}^k + \alpha_D^k \Delta \mathbf{v}
\end{aligned} \tag{A2.29}$$

où les scalaires $\alpha_P^k \in (0,1]$ et $\alpha_D^k \in (0,1]$ sont déterminés par ces relations :

$$\alpha_P^k = \min \left\{ 1, \gamma^* \min_i \left\{ \frac{-s_i^k}{\Delta s_i}, \Delta s_i < 0, \quad \frac{-z_i^k}{\Delta z_i}, \Delta z_i < 0 \right\} \right\} \tag{A2.30}$$

$$\alpha_D^k = \min \left\{ 1, \gamma^* \min_i \left\{ \frac{-\pi_i^k}{\Delta \pi_i}, \Delta \pi_i < 0, \quad \frac{-\hat{v}_i^k}{\Delta \hat{v}_i}, \Delta \hat{v}_i < 0 \right\} \right\}$$

Le scalaire $\gamma \in (0,1)$ est le facteur de « sûreté » destiné à assurer que le point suivant respecte les conditions de strictes positivité ; une valeur typique est $\gamma = 0.99995$.

Pour la remise à jour de μ^k , on peut utiliser une procédure qui a prouvé son efficacité pour les problèmes de programmation non linéaire en général. Tout d'abord, on calcule un « écart complémentaire » ρ qui s'exprime ainsi :

$$\rho^k = (\mathbf{s}^k)^T \boldsymbol{\pi}^k + (\mathbf{z}^k)^T \hat{\mathbf{v}}^k \tag{A2.31}$$

Le paramètre ρ^k tend vers 0 au fur et à mesure que l'algorithme converge vers $\mathbf{x}^*$. La réduction de μ^k peut alors se faire sur la base d'une réduction prédite de l'écart complémentaire comme ceci :

$$\mu^{k+1} = \sigma^k \frac{\rho^k}{2m} \tag{A2.32}$$

où σ^k est le paramètre de « centrage » compris entre 0 et 1. Choisir $\sigma^k = 1$ définit une direction « centrale », tandis que prendre $\sigma^k = 0$ revient à définir une évolution basée

uniquement sur le principe de la méthode de Newton. Un bon compromis entre ces deux extrêmes serait de choisir σ^k tel que $\sigma^k = \max\{0.99\sigma^{k-1}, 0.1\}$ avec $\sigma^0 = 0.2$.

A2.3.4) Test de convergence

L'algorithme a convergé si on remplit les conditions suivantes :

$$\nu_1^k \leq \varepsilon_1$$

$$\mu^k \leq \varepsilon_\mu$$

$$\nu_2^k \leq \varepsilon_2$$

Ou bien

$$\|\Delta \mathbf{x}\|_\infty \leq \varepsilon_2$$

$$\nu_3^k \leq \varepsilon_3$$

$$\|g(x^k)\|_\infty \leq \varepsilon_1$$

$$\nu_4^k \leq \varepsilon_4$$

$$\nu_4^k \leq \varepsilon_2$$

Où

$$\nu_1 = \max\left\{\max\{\underline{\mathbf{h}} - h(\mathbf{x})\}, \max\{h(\mathbf{x}) - \bar{\mathbf{h}}\}, \|g(\mathbf{x})\|_\infty\right\}$$

$$\nu_2 = \frac{\|\nabla_x f(\mathbf{x}) - \mathbf{J}_g(\mathbf{x})^T \boldsymbol{\lambda} + \mathbf{J}_h(\mathbf{x})^T \mathbf{v}\|_\infty}{1 + \|\mathbf{x}\|^2}$$

(A2.33)

$$\nu_3 = \frac{\rho}{1 + \|\mathbf{x}\|^2}$$

$$\nu_4 = \frac{|f(\mathbf{x}^k) - f(\mathbf{x}^{k-1})|}{1 + |f(\mathbf{x}^k)|}$$

Les tolérances typiques que l'on peut appliquer sont $\varepsilon_1 = 10^{-4}$, $\varepsilon_2 = 10^{-2}$, $\varepsilon_3 = 10^{-4}$ et $\varepsilon_\mu = 10^{-10}$.

A2.4) Optimisation sous MATLAB [MAT 01]

Sous l'environnement informatique Matlab, il est possible de programmer un problème donné d'optimisation à l'aide de fonctions d'optimisations déjà intégrées à l'environnement de travail. La fonction d'optimisation que nous avons utilisé pour programmer des OPF et des ORO est la fonction *fmincon*. Elle est basée sur l'algorithme de Quasi-Newton, et peut

basculer pour des problèmes de grande taille sur des méthodes de type Point Interieur. Fmincon peut trouver un optimum aux problèmes contraints de la forme suivante :

Min $F(X)$ avec $A * X \leq B$ $Aeq * X \leq Beq$ (contraintes linéaires)

$C(X) \leq 0$ $Ceq(X) = 0$ (contraintes non linéaires)

L'appel de la fonction d'optimisation peut se faire de la façon suivante :

$[X, FVAL, EXITFLAG, OUTPUT, LAMBDA] = \text{fmincon} (@Fonc, X0, A, B, Aeq, Beq, LB, UB, @Nonlcon)$

avec en arguments d'entrée:

- Fonc le nom de la fonction objectif. Elle est stockée dans un fichier du même nom Fonc.m
- X0 est le vecteur d'initialisation
- LB est la limite inférieures des valeurs permises pour les éléments de X
- UB est la limite supérieures des valeurs permises pour les éléments de X
- Nonlcon qui est une fonction qui retourne les vecteurs C et Ceq. Elle est stockée dans le fichier Nonlcon.m

et en arguments de sortie :

- X qui est la valeur finale du vecteur X optimisé
- FVAL qui est la valeur de la fonction objectif après convergence
- EXITFLAG qui retourne un chiffre indiquant l'état de convergence de l'algorithme (>0 si une solution a été trouvée, <0 si l'algorithme n'a pas trouvé de solution, et 0 si le nombre max d'itérations est dépassé)
- OUTPUT qui est une structure qui renferme des arguments concernant le nombre d'itérations, d'évaluations de la fonction objectif, etc ...
- LAMBDA qui renvoie les multiplicateurs de Lagrange finaux associés aux contraintes

Nous avons adapté l'ORO que nous avons défini de façon théorique au Chapitre II au format de la fonction *fmincon*, avec les considérations suivantes :

- Dans le vecteur de paramètres X , nous avons choisi de ranger les p ajustements de production ΔG_i et les $N-1$ phases nodales θ_i . Le vecteur d'optimisation X est donc de la forme suivante :

$$[X]^T = ([\theta_2 \quad \dots \quad \theta_N \quad \Delta G_1 \quad \dots \quad \Delta G_p])$$

- L'équation décrivant l'équilibre des ajustements peut être écrite en tant que contrainte égalité linéaire
- Les inégalités quant aux transits des lignes sont non linéaires : en effet, comme on a considéré dans notre cas que $P_{ij}^{\min} = -P_{ij}^{\max}$, les contraintes de transit sur les lignes peuvent se ramener à $|P_{ij}| \leq P_{ij}^{\max}$ pour toute ligne ij .
- Si l'on fait intervenir les phases nodales, il est plus simple de rentrer l'équations du load-flow $\mathbf{P}_G - \mathbf{P}_C = \mathbf{B} * \boldsymbol{\theta}$ en tant que contrainte égalité non linéaire

Dans un fichier principal, après avoir chargé toutes les données du réseau à étudier et les données sur les offres d'ajustement, nous devons attribuer des valeurs aux vecteurs d'entrées que l'on va utiliser avant l'appel de la fonction. Dans notre cas, nous devons définir les vecteurs $X0$, LB et UB :

- Nous avons en général fixé les éléments de $X0$ à 0.
- $LB(i) = -\Delta G_i^{\max}$ pour tout ajustement i . En ce qui concerne les phases nodales, les limites sont fixées à $\pm \pi$. On a donc :

$$LB = ([-\pi \quad \dots \quad -\pi \quad -\Delta G_1^{\max} \quad \dots \quad -\Delta G_1^{\max}])$$

- De même, UB est défini ainsi :

$$UB = ([+\pi \quad \dots \quad +\pi \quad +\Delta G_1^{\max} \quad \dots \quad +\Delta G_1^{\max}])$$

- Pour l'équation régissant le bilan des ajustements $\sum_i \Delta G_i = 0$, nous avons défini la matrice Aeq et Beq comme suit :

$$Aeq = ([0 \quad \dots \quad 0 \quad 1 \quad \dots \quad 1]) \quad Beq = 0$$

Il nous reste à définir la fonction objectif, les contraintes inégalités portant sur le transit des lignes et les contraintes égalités liées aux équations du load-flow. Ceux-ci devront être défini dans deux fichiers de fonction : `Fonc.m` et `Nonlcon.m`

Fonction Fonc.m

En première ligne, on doit définir un fichier comme étant une fonction en spécifiant son(ses) argument(s) d'entrée et son(ses) argument(s) de sortie. Dans notre cas, nous avons écrit :

$$[f_DC] = \text{Fonc}(X) ;$$

X étant le vecteur à optimiser et f_DC une variable locale stockant la valeur de la fonction objectif à chaque itération. Le coût total des ajustements est donc défini ainsi :

$$f_DC = \sum_{i=1}^p (a_i + b_i X(i)) * X(i)$$

a_i et b_i étant les paramètres des fonctions d'ajustement définis dans le fichier principal.

Fonction Nonlcon.m

En ce qui concerne la contrainte inégalité appliquée sur les lignes, nous avons défini pour chaque transit :

$$C(i) = |P_{ij}| - P_{ij}^{\max}$$

Les transits P_{ij} sont calculés à chaque itération à l'aide de la matrice **H** reliant les transits aux phases nodales :

$$P_{ij} = H_{ij} * X(1 \dots N-1)$$

Nous avons rentré l'équation du load-flow reliant les injections nodales aux phases nodales par l'intermédiaire de la matrice B de cette façon :

$$Ceq(i) = P_{Gi} - P_{Ci} - B_i * X(1 \dots N-1) \text{ pour } i \neq \text{nœud bilan}$$

A chaque itération les productions nodales sont calculés par :

$P_{Gi} = P_{Gi}^0 + X(N+i-1)$ avec P_{Gi}^0 étant la production initiale au nœud à la sortie du marché.

ANNEXE 3

Méthodes des Images de Charge basée sur le modèle DC

A3.1) Description théorique

Soit un réseau maillé à N nœuds. Ce réseau possède g nœuds générateurs et c nœuds consommateurs. L'expression qui relie les injections nodales aux phases nodales peut être mise sous cette forme :

$$\begin{bmatrix} P_{G1} \\ \vdots \\ P_{Gg} \\ P_{C1} \\ \vdots \\ P_{Cc} \end{bmatrix} = \begin{bmatrix} b_{1,1} & \cdots & b_{1,g} \\ \vdots & \ddots & \vdots \\ b_{g,1} & \cdots & b_{g,g} \\ b_{g+1,1} & \cdots & b_{g+1,g} \\ \vdots & \ddots & \vdots \\ b_{g+c,1} & \cdots & b_{g+c,g} \end{bmatrix} \begin{bmatrix} b_{1,g+1} & \cdots & b_{1,g+c} \\ \vdots & \ddots & \vdots \\ b_{g,g+1} & \cdots & b_{g,g+c} \\ b_{g+1,g+1} & \cdots & b_{g+1,g+c} \\ \vdots & \ddots & \vdots \\ b_{g+c,g+1} & \cdots & b_{g+c,g+c} \end{bmatrix} * \begin{bmatrix} \theta_{G1} \\ \vdots \\ \theta_{Gg} \\ \theta_{C1} \\ \vdots \\ \theta_{Cc} \end{bmatrix} \quad (\text{A3.1})$$

$$\text{ou bien plus simplement } \begin{bmatrix} \mathbf{P}_G \\ \mathbf{P}_C \end{bmatrix} = \begin{bmatrix} \mathbf{B}_{GG} & \mathbf{B}_{GC} \\ \mathbf{B}_{CG} & \mathbf{B}_{CC} \end{bmatrix} * \begin{bmatrix} \boldsymbol{\theta}_G \\ \boldsymbol{\theta}_C \end{bmatrix} \quad (\text{A3.2})$$

avec :

$\mathbf{P}_G$ vecteur des productions aux nœuds générateurs

$\mathbf{P}_C$ vecteur des charges aux nœuds consommateurs

$\boldsymbol{\theta}_G$ vecteur des phases nodales des nœuds générateurs

$\boldsymbol{\theta}_C$ vecteur des phases nodales des nœuds consommateurs

$\mathbf{B}_{GG}$, $\mathbf{B}_{GC}$, $\mathbf{B}_{CG}$, $\mathbf{B}_{CC}$ sous matrices de la matrice des admittances nodales

A partir de la relation (IV.2), grâce à une réduction de Kron, on peut arriver à exprimer les productions en fonction des charges :

$$\mathbf{P}_G = \mathbf{B}_g^{\text{réd}} * \boldsymbol{\theta}_G + \mathbf{D} * \mathbf{P}_C \quad (\text{A3.3})$$

avec

$$\mathbf{B}_g^{\text{réd}} = \mathbf{B}_{GG} - \mathbf{B}_{GC} * \mathbf{B}_{CC}^{-1} * \mathbf{B}_{CG} \quad (\text{A3.4})$$

et

$$\mathbf{D} = \mathbf{B}_{GC} * \mathbf{B}_{CC}^{-1} \quad (\text{A3.5})$$

La matrice $\mathbf{B}_g^{\text{réd}}$ de dimension $g * g$ est la matrice des admittances nodales du réseau réduit aux g nœuds générateurs et la matrice $\mathbf{D}$ de dimension $g * c$ est la matrice qui multipliée au vecteur des charges donne les images de charges (Fig. IV.3).

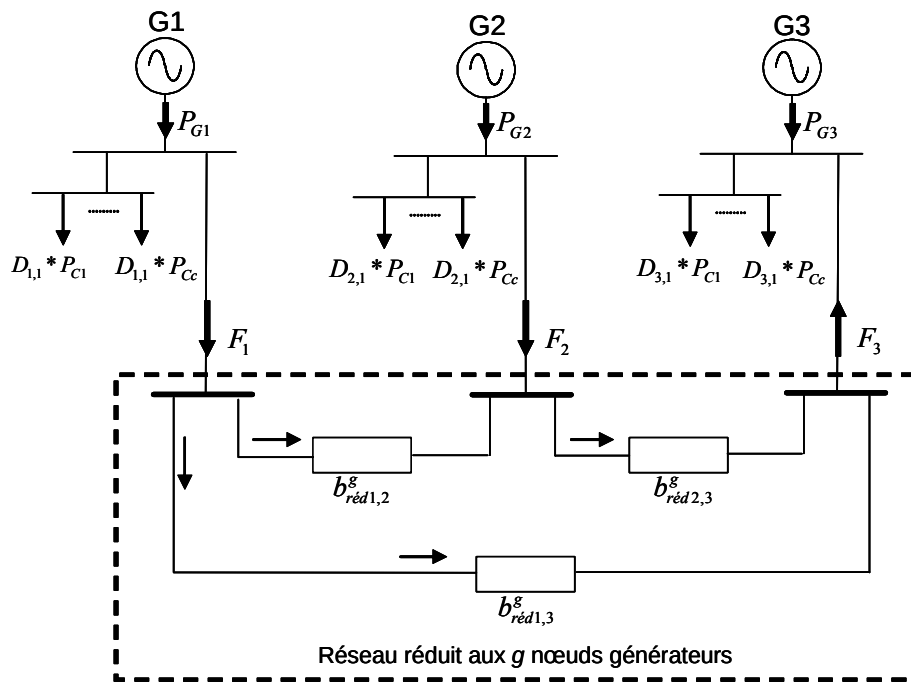


Figure A3.1 : exemple de réduction d'un réseau comportant trois nœuds générateurs avec images de charge associées

On peut définir des injections nodales à travers le réseau réduit et qui s'exprime de la façon suivante :

$$\mathbf{F} = \mathbf{B}_g^{\text{réd}} * \boldsymbol{\theta}_G \quad (\text{A3.6})$$

La relation (A3.3) peut donc se réécrire :

$$\mathbf{P}_G = \mathbf{F} + \mathbf{D} * \mathbf{P}_C \quad (\text{A3.7})$$

Ces injections nodales à travers le réseau réduit peuvent aussi s'exprimer en chaque nœud comme la somme des injections vers les autres nœuds sources ou provenant des autres nœuds sources :

$$F_i = \sum_{k=1}^g F_i^k \quad \text{pour } k \notin i$$

Les images de charge doivent être agrégées avec la production en chaque nœud générateur. Si en un nœud générateur i la somme des images de charge est inférieure à la production en ce nœud, on obtient une production équivalente ($F_i > 0$). Si la somme des images de charge est supérieure à la production du nœud i , on obtient une charge équivalente ($F_i < 0$). Dans ce cas, l'énergie consommée en ce nœud provient à la fois de la source locale au point i , et des injections provenant d'autres nœuds sources. Ce procédé d'agrégation est illustré par les Figures A3.2 et A3.3 :

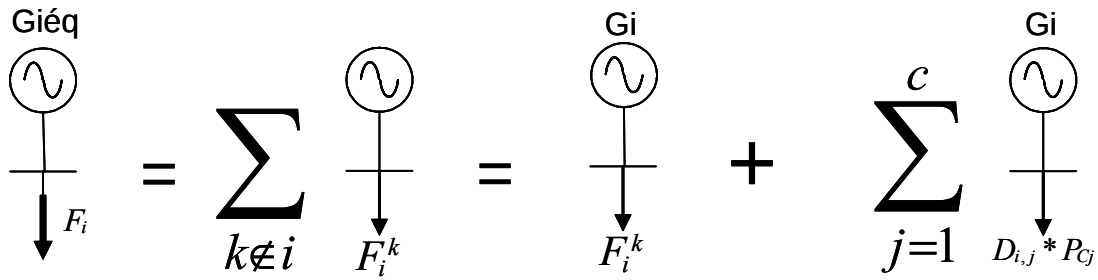


Figure A3.2 : décomposition du groupe agrégé $Giéq$ (production équivalente)

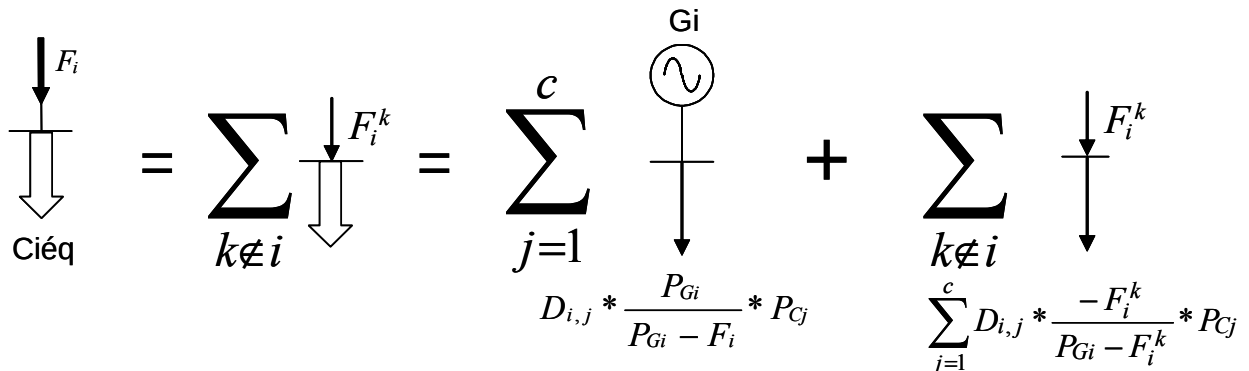


Figure A3.3 : décomposition du groupe agrégé $Ciéq$ (charge équivalente)

On peut donc représenter le réseau réduit de la Figure A.I.4 avec ses productions et charges équivalentes :

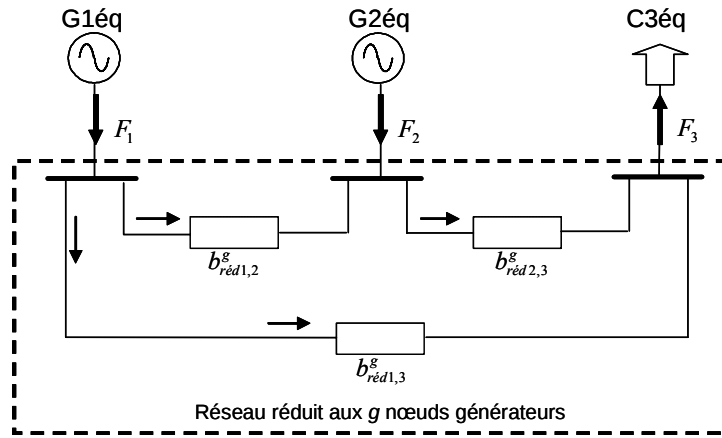


Figure A3.4 : réseau réduit avec productions et charges équivalentes agrégées

Après avoir effectué une première réduction du réseau d'origine, on peut refaire une nouvelle réduction en partant en reprenant la méthode décrite par les relations (A3.1) à (A3.3). Une fois que le réseau d'origine a été entièrement réduit, toutes les charges ont été agrégées aux productions, et l'on peut alors définir une matrice **C** qui lie les charges aux productions par des coefficients compris entre 0 et 1 :

$$\mathbf{P}_G = \mathbf{C} * \mathbf{P}_C \quad (\text{A3.7})$$

Nous pouvons alors aisément déterminer les liaisons productions-transits en faisant un calcul de répartition de charge avec uniquement la production en question (les autres étant mises à 0) et ses charges associées (déterminées à partir des coefficients de la matrice **C**). En utilisant les propriétés du modèle DC, on peut calculer les transits attribuables à la production située à un nœud *i* de la façon suivante :

$$\mathbf{P}_{ij}^{P_{Gi}} = [\mathbf{A}] * \left[\begin{bmatrix} \dots & 0 & P_{Gi} & 0 & \dots \end{bmatrix} \begin{bmatrix} C_{i,1} * P_{C1} & \dots & C_{i,c} * P_{Cc} \end{bmatrix} \right]^T \quad (\text{A3.8})$$

De la même façon, on peut déterminer les contributions de chaque charge aux transits en simulant uniquement cette charge avec ses productions associées.

A3.2) Application sur le réseau 9 nœuds

Nous pouvons illustrer la MIC sur le réseau 9 nœuds en détaillant les différentes étapes de réduction du réseau permettant d'établir les productions en fonction des charges selon la relation (A3.4). Nous partons donc du cas suivant :

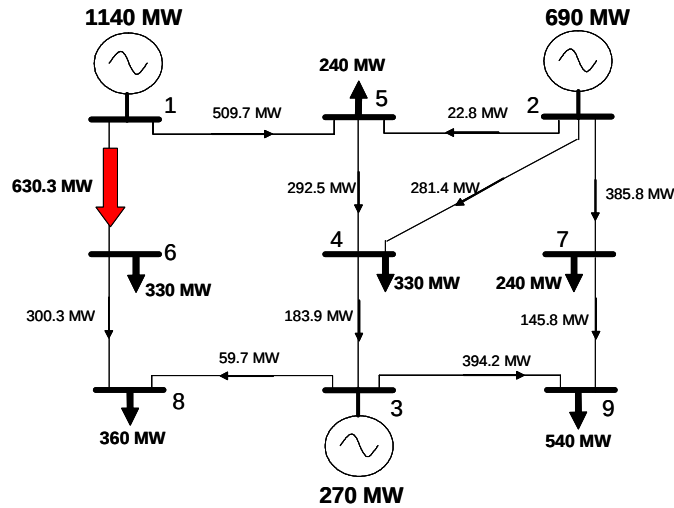


Figure A3.5 : cas du réseau 9 nœuds

La relation (A3.1) appliquée au cas du réseau 9 nœuds devient :

$$\begin{bmatrix} 1.14 \\ 6.9 \\ 2.7 \\ -3.3 \\ -2.4 \\ -3.3 \\ -2.4 \\ -3.6 \\ -5.4 \end{bmatrix} = \begin{bmatrix} 34.5 & 0 & 0 & 0 & -16.7 & -17.9 & 0 & 0 & 0 & 0 \\ 0 & 143.4 & 0 & -45.5 & -66.7 & 0 & -31.3 & 0 & 0 & 0 \\ 0 & 0 & 299.2 & -166.7 & 0 & 0 & 0 & -76.9 & -55.6 & 0 \\ 0 & -45.5 & -166.7 & 262.12 & -50 & 0 & 0 & 0 & 0 & 0 \\ -16.7 & -66.7 & 0 & -50 & 133.3 & 0 & 0 & 0 & 0 & 0 \\ -17.9 & 0 & 0 & 0 & 0 & 117.9 & 0 & -100 & 0 & 0 \\ 0 & -31.3 & 0 & 0 & 0 & 0 & 102.7 & 0 & -71.5 & 0 \\ 0 & 0 & -76.9 & 0 & 0 & -100 & 0 & 176.9 & 0 & 0 \\ 0 & 0 & -55.6 & 0 & 0 & 0 & -71.5 & 0 & 127 & 0 \end{bmatrix} \begin{bmatrix} 0 \\ -0.302 \\ -0.375 \\ -0.364 \\ -0.306 \\ -0.353 \\ -0.426 \\ -0.383 \\ -0.446 \end{bmatrix} \quad (A3.6)$$

Les puissances nodales sont données en unité réduite (1 u.r. = 100 MW). Nous avons isolé les sous-matrices $\mathbf{B}_{gg}$, $\mathbf{B}_{gc}$, $\mathbf{B}_{cg}$ et $\mathbf{B}_{cc}$ comme cela est suggéré par les pointillés

1^{ère} réduction

A partir de ce modèle, nous pouvons réduire le réseau aux nœuds générateurs en calculant les matrices $\mathbf{B}_{\text{red}}^g$ et $\mathbf{D}$ à l'aide des relations (A3.4) et (A3.5). On obtient donc :

$$\mathbf{B}_{\text{red}}^g = \begin{bmatrix} 27.1 & -10.1 & -16.9 \\ -10.1 & 74.0 & -63.9 \\ -16.9 & -63.9 & 80.8 \end{bmatrix} \quad (\text{A3.7})$$

$$\mathbf{D} = \begin{bmatrix} -0.026 & -0.135 & -0.291 & 0 & -0.164 & 0 \\ -0.289 & -0.609 & 0 & -0.500 & 0 & -0.281 \\ -0.685 & -0.257 & -0.709 & -0.500 & -0.835 & -0.719 \end{bmatrix} \quad (\text{A3.8})$$

On peut donc calculer les injections à travers le réseau réduit aux 3 nœuds générateurs de cette façon :

$$\mathbf{F} = \begin{bmatrix} 27.1 & -10.1 & -16.9 \\ -10.1 & 74.0 & -63.9 \\ -16.9 & -63.9 & 80.8 \end{bmatrix} * \begin{bmatrix} 0 \\ -0.302 \\ -0.375 \end{bmatrix} = \begin{bmatrix} 9.42 \\ 1.59 \\ -11.01 \end{bmatrix} \quad (\text{A3.9})$$

La relation (A3.7) en application numérique s'écrit donc :

$$\begin{bmatrix} 11.40 \\ 6.9 \\ 2.7 \end{bmatrix} = \begin{bmatrix} 9.42 \\ 1.59 \\ -11.01 \end{bmatrix} + \begin{bmatrix} -0.026 & -0.135 & -0.291 & 0 & -0.164 & 0 \\ -0.289 & -0.609 & 0 & -0.500 & 0 & -0.281 \\ -0.685 & -0.257 & -0.709 & -0.500 & -0.835 & -0.719 \end{bmatrix} * \begin{bmatrix} -3.3 \\ -2.4 \\ -3.3 \\ -2.4 \\ -3.6 \\ -5.4 \end{bmatrix} \quad (\text{A3.10})$$

La matrice des images de charge $\mathbf{D}$ représente les charges à agréger directement en chaque nœud générateur, tandis que $\mathbf{F}$ représente les injections passant par d'autres nœuds générateurs du réseau réduit vers d'autres images de charges. Ainsi, si un coefficient $D(i,j)$ de la matrice $\mathbf{D}$ est nul, cela signifie que le générateur i ne peut atteindre la charge j sans passer par le réseau réduit. Par exemple, si l'on considère le générateur connecté au nœud 1, celui-ci ne peut venir alimenter la charge située au nœud 7 sans passer par les nœuds générateurs 2 ou 3 quelque soit le chemin supposé. Cela explique pourquoi le coefficient $D(1,7)$ est nul. La Figure (A3 .6) représente le réseau réduit aux 3 nœuds générateurs à la première réduction :

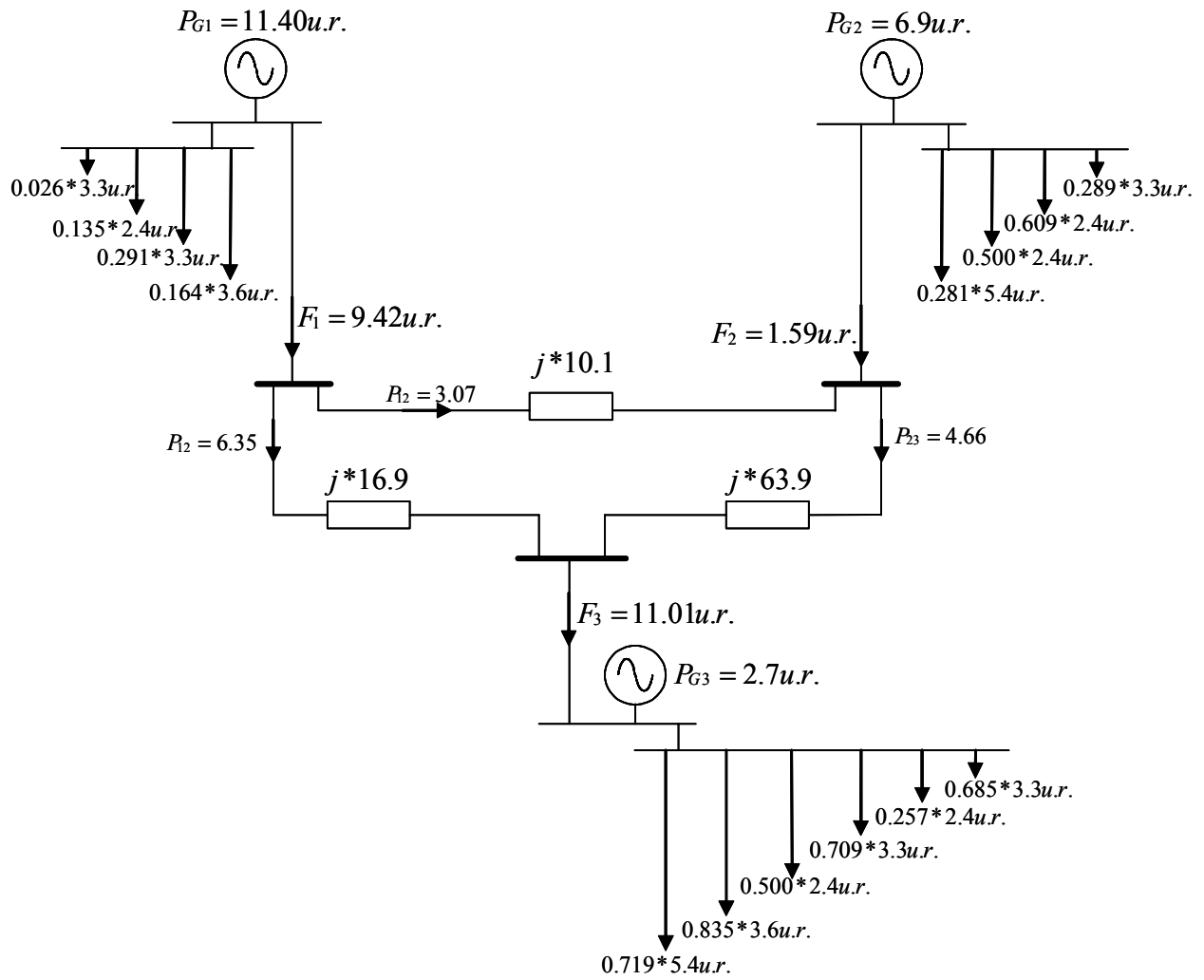


Figure A3.6 : réseau réduit aux trois nœuds générateurs à la première réduction

Les images de charge doivent être agrégées à la source pour chaque nœud générateur. On va alors stocker les résultats des agrégations dans une matrice Ag à 3 lignes (pour les générateurs) et 6 colonnes (pour les charges). Cette matrice sera incrémentée à chaque nouvelle réduction et agrégation. Ainsi, pour la première réduction et agrégation, cette matrice nous donne :

$$Ag = \begin{bmatrix} 0.100 & 0.323 & 0.961 & 0 & 0.592 & 0 \\ 1.129 & 1.461 & 0 & 1.2 & 0 & 1.519 \\ 0.526 & 0.121 & 0.461 & 0.236 & 0.592 & 0.764 \end{bmatrix} \begin{matrix} \leftarrow G1 \\ \leftarrow G2 \\ \leftarrow G3 \end{matrix}$$

$$\begin{matrix} \uparrow & \uparrow & \uparrow & \uparrow & \uparrow & \uparrow \\ C4 & C5 & C6 & C7 & C8 & C9 \end{matrix}$$

Chaque coefficient $Ag(i,j)$ est calculé comme suit :

$$Ag(i, j) = D(i, j) * Cj \quad (A3.11)$$

Nous pouvons remarquer qu'à ce stade, le générateur G3 a déjà été entièrement agrégé à ses charges associées. Le réseau réduit avec ses productions et charges équivalentes nodales est donc :

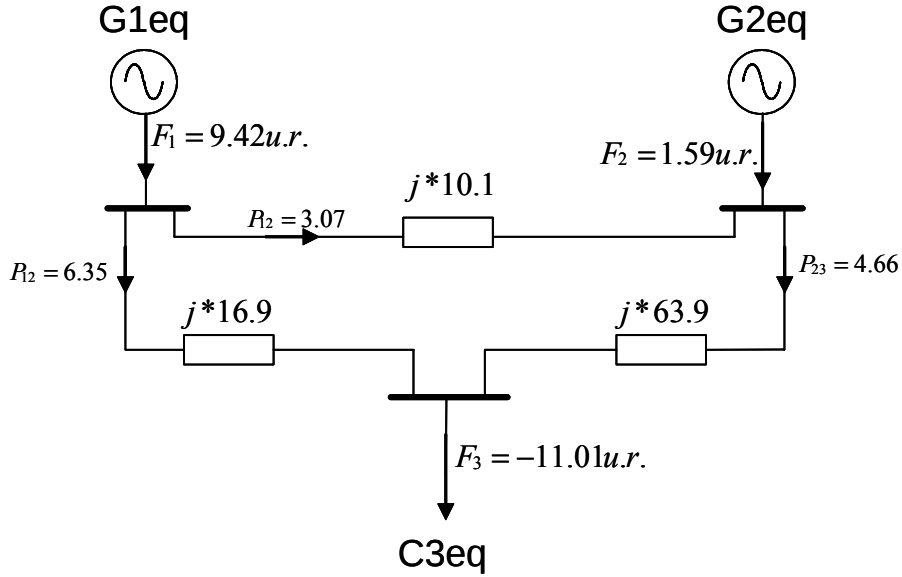


Figure A3.7 : réseau réduit avec injections équivalentes

La consommation équivalente au nœud représente la somme des images de charges qui restent encore à agréger aux productions restantes aux nœuds 1 et 2. On peut exprimer C3eq en fonction des charges initiales Cj de cette façon :

$$C3eq = \mathbf{K} * \mathbf{P}_C = [K_1 \quad \dots \quad K_c] * \begin{bmatrix} C_4 \\ \vdots \\ C_9 \end{bmatrix} \quad (A3.12)$$

$$\text{avec } K_j = \frac{F_3}{P_{G3} - F_3} * D(3, j) \quad (A3.13)$$

2^{ème} réduction

Le réseau de la Figure A3.7 peut être réduit de la même façon que le réseau originel. On repart donc de l'équation reliant les injections nodales aux phases nodales :

$$\begin{bmatrix} 9.42 \\ 1.59 \\ -11.01 \end{bmatrix} = \begin{bmatrix} 27.1 & -10.1 & -16.9 \\ -10.1 & 74.0 & -63.9 \\ -16.9 & -63.9 & 80.8 \end{bmatrix} * \begin{bmatrix} 0 \\ -0.302 \\ -0.375 \end{bmatrix} \quad (\text{A3.14})$$

La matrice du réseau réduit est la matrice des images de charges sont:

$$\mathbf{B}_{\text{red}}^g = \begin{bmatrix} 23.53 & -23.53 \\ -23.53 & 23.53 \end{bmatrix} \quad (\text{A3.15}) \quad \mathbf{D} = \begin{bmatrix} -0.210 \\ -0.790 \end{bmatrix} \quad (\text{A3.16})$$

Le vecteur des injections nodales est :

$$\mathbf{F} = \begin{bmatrix} 23.53 & -23.53 \\ -23.53 & 23.53 \end{bmatrix} * \begin{bmatrix} 0 \\ -0.302 \end{bmatrix} = \begin{bmatrix} 7.11 \\ -7.11 \end{bmatrix} \quad (\text{A3.17})$$

On a donc (relation (A3.7)):

$$\begin{bmatrix} 9.42 \\ 1.59 \end{bmatrix} = \begin{bmatrix} 7.11 \\ -7.11 \end{bmatrix} + \begin{bmatrix} -0.210 \\ -0.790 \end{bmatrix} * [-11.01] \quad (\text{A3.18})$$

A l'aide des relations (A3.12) et (A3.13), nous pouvons exprimer directement les productions restantes en fonction des injections dans le réseau réduit et du reste des charges originelles à agréger :

$$\begin{bmatrix} 9.42 \\ 1.59 \end{bmatrix} = \begin{bmatrix} 7.11 \\ -7.11 \end{bmatrix} + \begin{bmatrix} -0.115 & -0.043 & -0.119 & -0.084 & -0.140 & -0.121 \\ -0.435 & -0.163 & -0.450 & -0.317 & -0.530 & -0.456 \end{bmatrix} * \begin{bmatrix} -3.3 \\ -2.4 \\ -3.3 \\ -2.4 \\ -3.6 \\ -5.4 \end{bmatrix} \quad (\text{A3.19})$$

Le réseau réduit à cette deuxième réduction est donc :

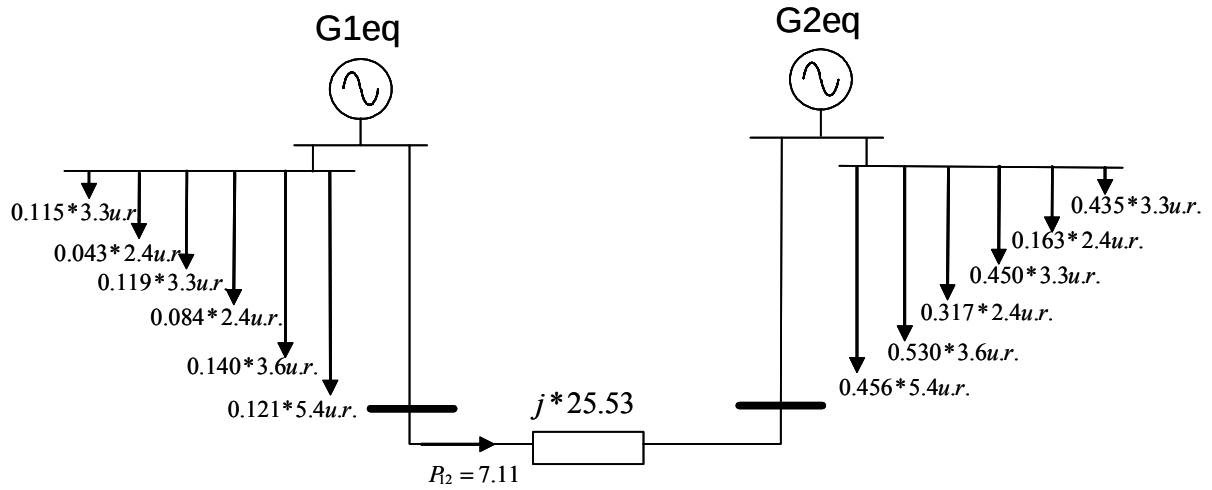


Figure A3.8 : réseau réduit (deuxième réduction)

Les images de charge aux nœuds 1 et 2 sont agrégées aux sources, ce qui nous permet d'incrémenter les deux premières lignes de la matrice Ag qui devient :

$$Ag = \begin{bmatrix} 0.550 & 0.427 & 1.355 & 0.202 & 1.099 & 0.653 \\ 1.439 & 1.532 & 0.271 & 1.340 & 0.349 & 1.969 \\ 0.526 & 0.121 & 0.461 & 0.236 & 0.592 & 0.764 \end{bmatrix} \quad (A3.20)$$

Le réseau réduit avec ses injections équivalentes devient :

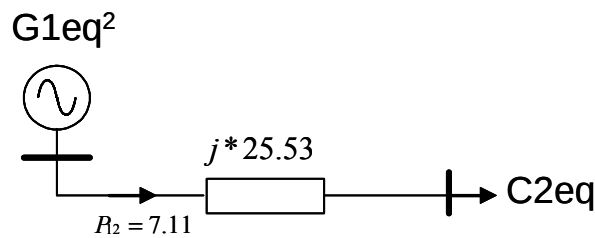


Figure A3.9 : réseau réduit (deuxième réduction) avec injections équivalentes

En calculant la matrice K , on détermine du coup les dernières images de charges à agréger à la production restante au nœud 1. On obtient donc :

$$[7.11] = [0.355 \quad 0.133 \quad 0.368 \quad 0.259 \quad 0.433 \quad 0.373] * \begin{bmatrix} -3.3 \\ -2.4 \\ -3.3 \\ -2.4 \\ -3.6 \\ -5.4 \end{bmatrix} \quad (A3.21)$$

Cela correspond au schéma suivant :

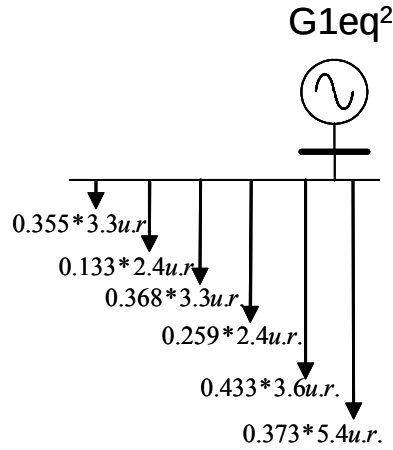


Figure A3.10 : dernières images de charges restantes

La matrice Ag complète est donc :

$$Ag = \begin{bmatrix} 1.935 & 0.747 & 2.568 & 0.825 & 2.659 & 2.667 \\ 1.439 & 1.532 & 0.271 & 1.340 & 0.349 & 1.969 \\ 0.526 & 0.121 & 0.461 & 0.236 & 0.592 & 0.764 \end{bmatrix} \quad (A3.22)$$

A ce stade, toutes les images de charges ont été agrégées aux productions. Nous pouvons noter qu'il nous a fallu très peu de réductions (deux au total) pour arriver à ce résultat. La méthode possède une rapidité d'exécution comparable sur des réseaux plus étendus tels que le réseau New England ou le RTS 96.

On peut donc exprimer le vecteur des productions en fonction du vecteur des consommations par une relation linéaire du type :

$$P_G = C * P_C \quad (A3.23)$$

avec $C(i, j) = Ag(i, j) / P_{C(j)}$

Nous pouvons aussi déterminer les Facteurs d'Utilisation de Ligne associés à chaque production ou consommation. Pour la production G1 par exemple, il suffit juste de faire un calcul de répartition de charge sur l'ensemble du réseau, en plaçant la production G1 et ses charges associées définies par la matrice Ag . Toutes les autres productions et consommations sont mises à zéro. Dans ce cas, les transits trouvés indiquent le cheminement empruntée par la puissance produite au nœud 1 à travers le réseau jusqu'aux charges associées. En utilisant la relation reliant les transits aux injections nodales via la matrice A des facteurs de distribution, on peut écrire que :

$$[P_{ij}^{G_1}] = [A]^* \begin{bmatrix} P_{G_1} \\ 0 \\ \vdots \\ -Ag(1,1) \\ \vdots \\ -Ag(1,c) \end{bmatrix} \quad (A3.24)$$

Le cheminement de la puissance produite au nœud 1 est représenté par la Figure A3.11 :

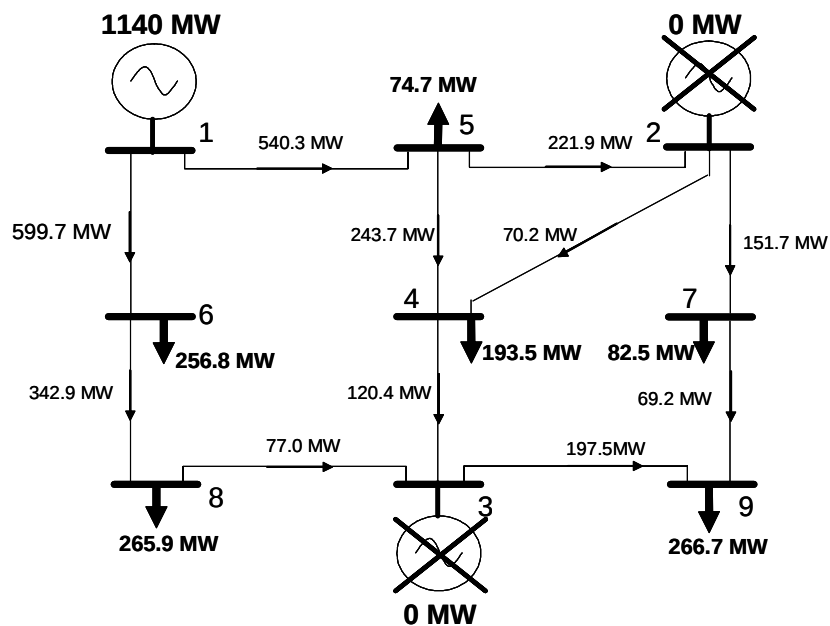


Figure A3.11 : transits attribuables à la production générée au nœud 1 uniquement

De la même façon, on peut déterminer les transits dus à la charge C4 en faisant un calcul de répartition de charge avec C4 et ses productions associées uniquement. Les transits obtenus sont représentés par la Figure (A3.12) :

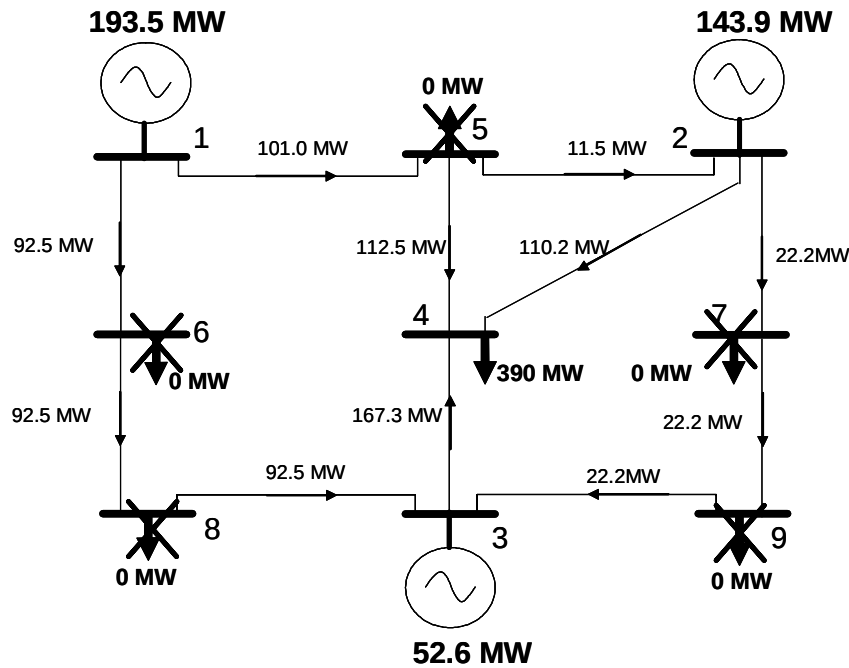


Figure A3.12 : transits attribuables à la charge C4 uniquement

L'état initial du réseau peut être aisément reconstruit en superposant les différentes paires production/consommations associées.

Les résultats de la MIC peuvent être comparés à d'autres méthodes ayant pour but de déterminer les transits attribuables à chaque injection nodale. En 1981, Ng introduit les facteurs de distribution généralisés appelés aussi *Generalized Generation Distribution Factors* (GGDF) pour les productions, et *Generalized Load Distribution Factors* (GLDF) pour les charges [NG 81]. Ces facteurs sont obtenus à partir des facteurs de distribution de la matrice A et sont indépendants du choix du nœud bilan. Bialek a de son côté proposé une méthode de traçabilité de l'énergie basée sur le principe de proportionnalité [BIA 97]. Ce principe fait d'abord la distinction pour chaque nœud du réseau entre les flux « entrants » (productions et transit entrant du nœud) et les flux « sortants » (consommations et transits sortant du nœud). Il stipule ensuite que si la contribution des productions amont aux transits entrants est connue, alors la contribution relative de ces productions aux transits sortant est conservée (Fig.A3.13). Les facteurs calculés par cette méthode sont les facteurs de distribution topologiques appelés aussi *Topological Generation Distribution Factors* (TGDF) pour les productions.

Le Tableau A3.1 nous donne la contribution des productions G1, G2 et G3 au transit de la ligne 1-6, déterminée par la MIC, par les GGDF et par les TGDF.

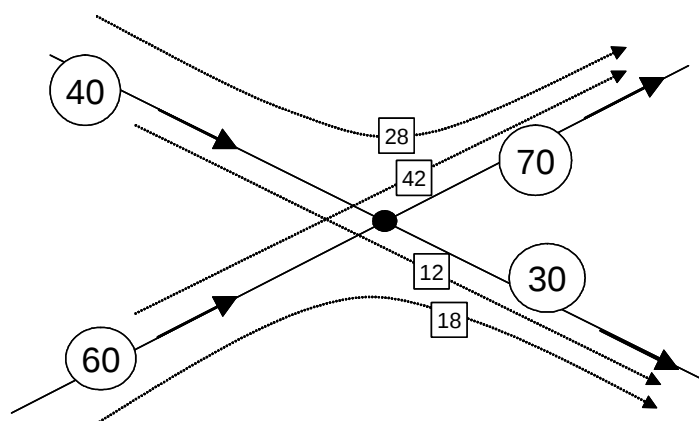


Figure A3.13 : principe de division proportionnelle utilisée dans la méthode de Bialek

Tableau A3.1 : contributions des productions au transit de la ligne 1-6 suivant la MIC, les GGDF et les TGDF (en MW)

	MIC	GGDF	TGDF
G1	599.7	573.7	630.3
G2	24.2	53.8	0
G3	6.4	2.8	0

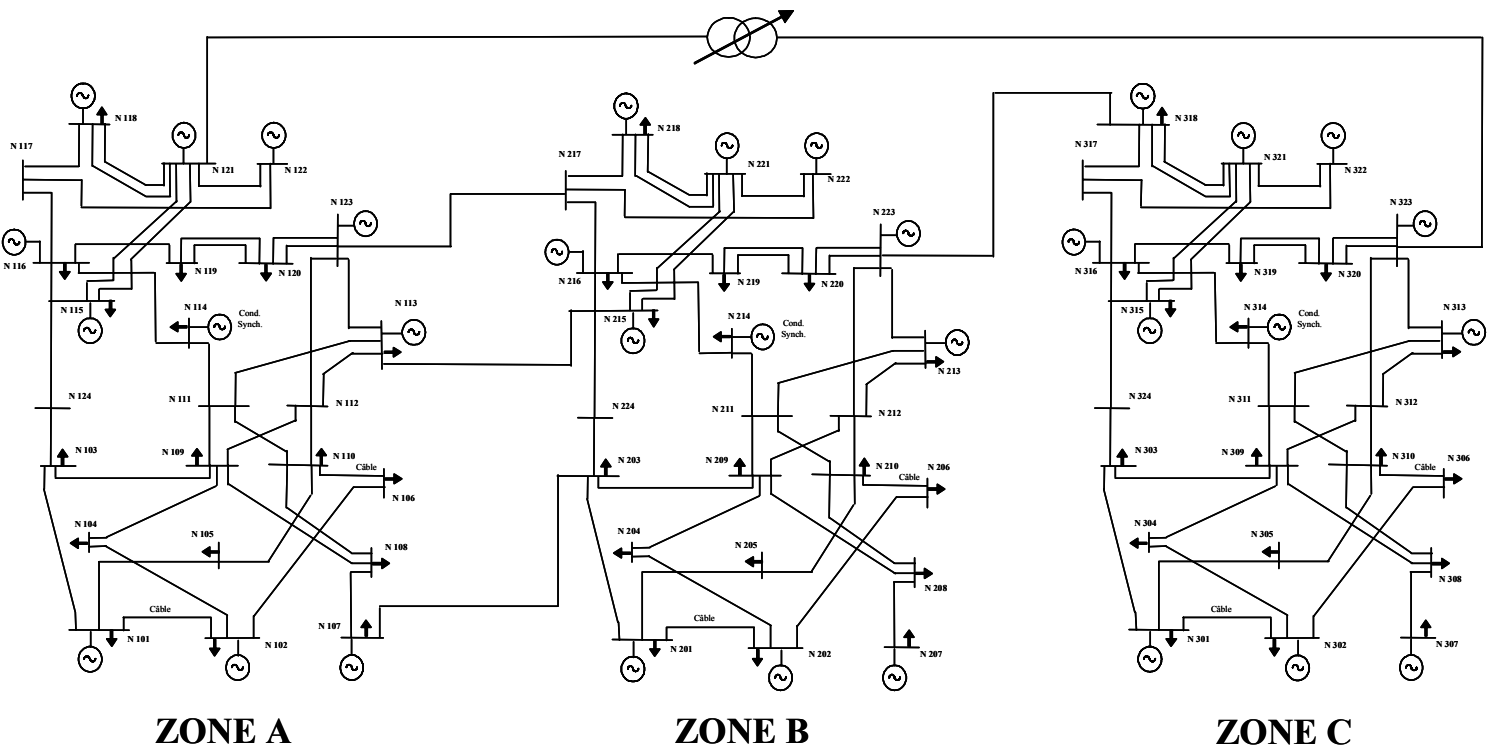
Nous remarquons que les résultats obtenus pour ce cas précis sont très proches, ce qui valide la MIC par rapport aux méthodes ayant été proposées avant elle. Toutefois, nous devrions souligner que la méthode des GGDF peut se révéler très approximative, car elle sous-entend en réalité que chaque MW produit en un nœud se distribue au prorata des puissances consommées en chaque nœud du réseau, ce qui n'est jamais vraiment le cas. La méthode proposée par Bialek a l'inconvénient de contrevenir au principe d'agrégation des charges et des productions au même nœud, ce qui fait qu'en présence de charges et de productions au même nœud, la MIC et la méthode de Bialek peuvent différer énormément. En outre, le problème est que le principe de proportionnalité a été introduit pour résoudre un problème à la base indéterminé, et ne peut être démontré, tandis que le principe d'agrégation constitue une des bases du « load-flow ».

La MIC offre ainsi des résultats compatibles avec les outils d'analyse existants (calcul de répartition de charge) et cohérents avec la réalité physique.

Annexe 4

Réseau RTS 96 et résultats complémentaires sur les allocations de coût

A4.1) Schéma du réseau :



A4.2) Données du réseau

Tableau A4.1 : Productions, charges à l'état initial dans une zone

Nœud n°	Production (MW)	Consommation (MW)
1	172	108
2	172	97
3	0	180
4	0	74
5	0	71
6	0	136
7	240	125
8	0	171
9	0	175
10	0	195
11	0	0
12	0	0
13	291	265
14	0	194
15	60	317
16	155	100
17	0	0
18	400	333
19	0	181
20	0	128
21	400	0
22	300	0
23	660	0
24	0	0

A4.3) Allocation du coût de congestion pour le cas 1 (congestion sur la ligne 220-223)

A4.3.1) Allocation aux consommateurs

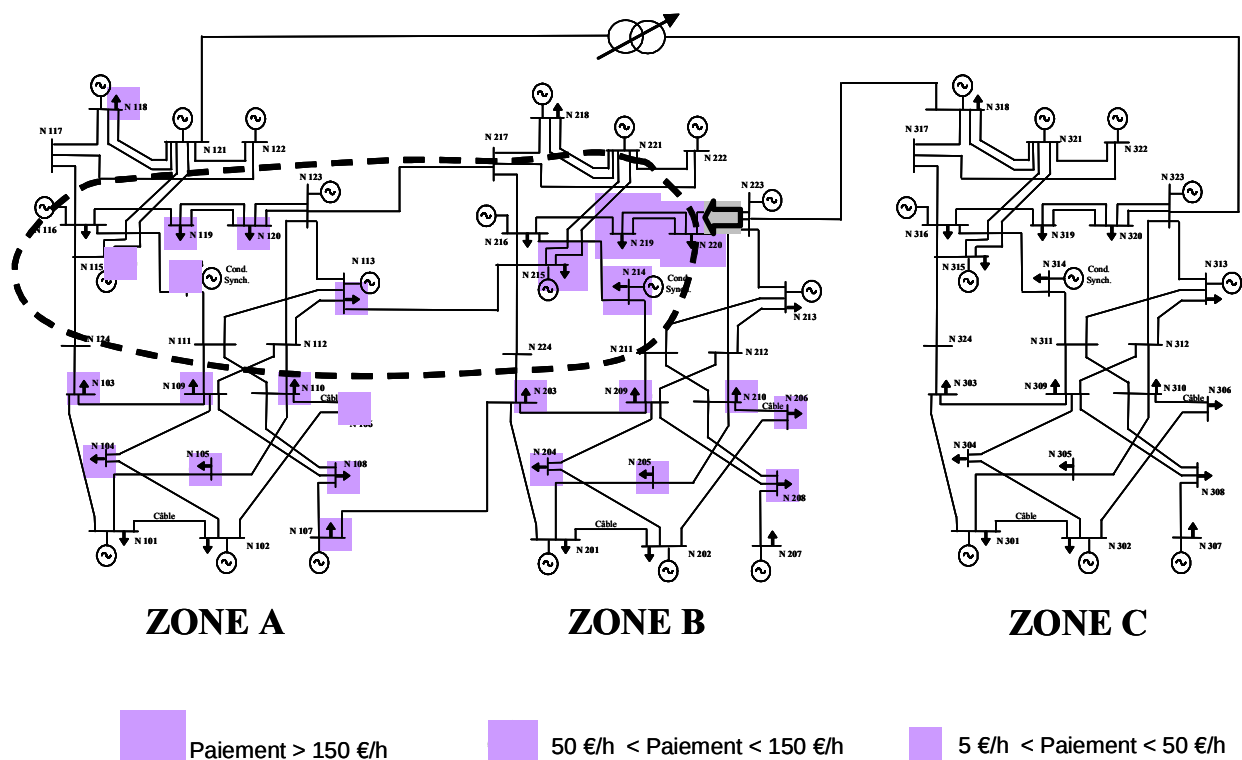


Figure A4.1 : allocation aux consommateurs par contribution au transit (cas1)

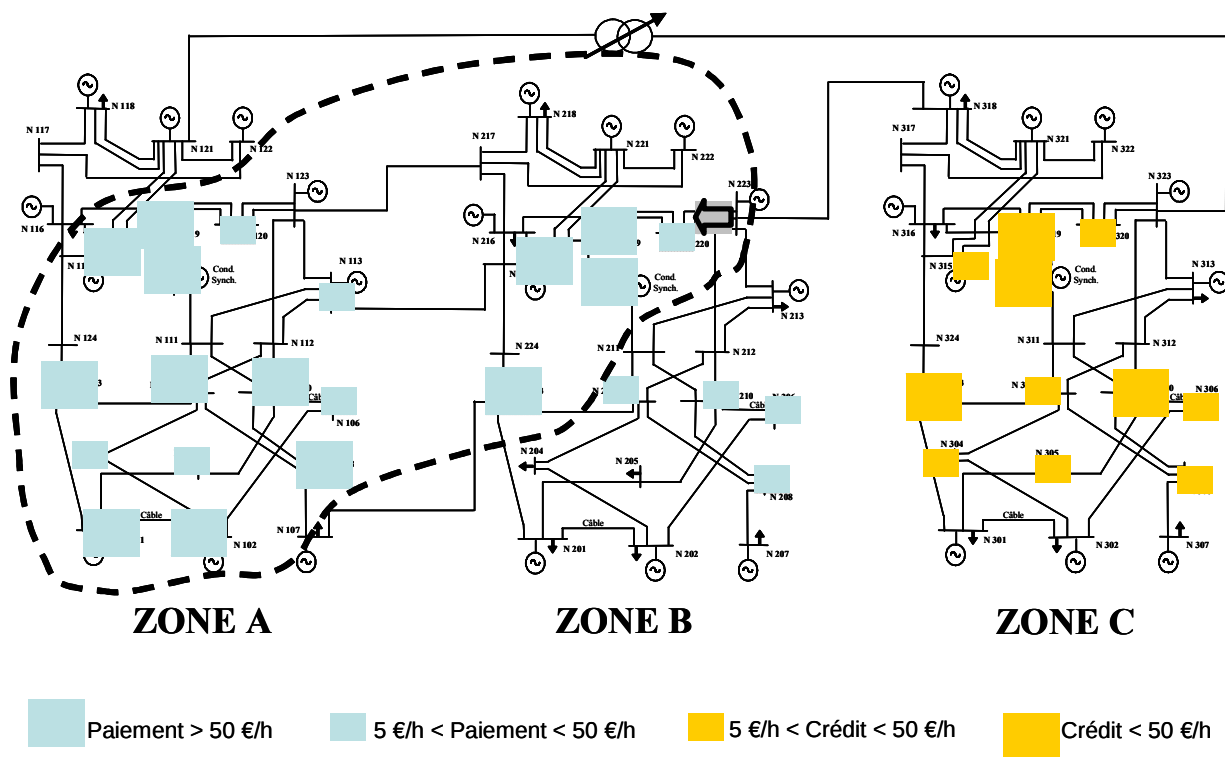


Figure A4.2 : allocation aux consommateurs par destination des ajustements (cas1)

A4.3.2) Allocation aux producteurs

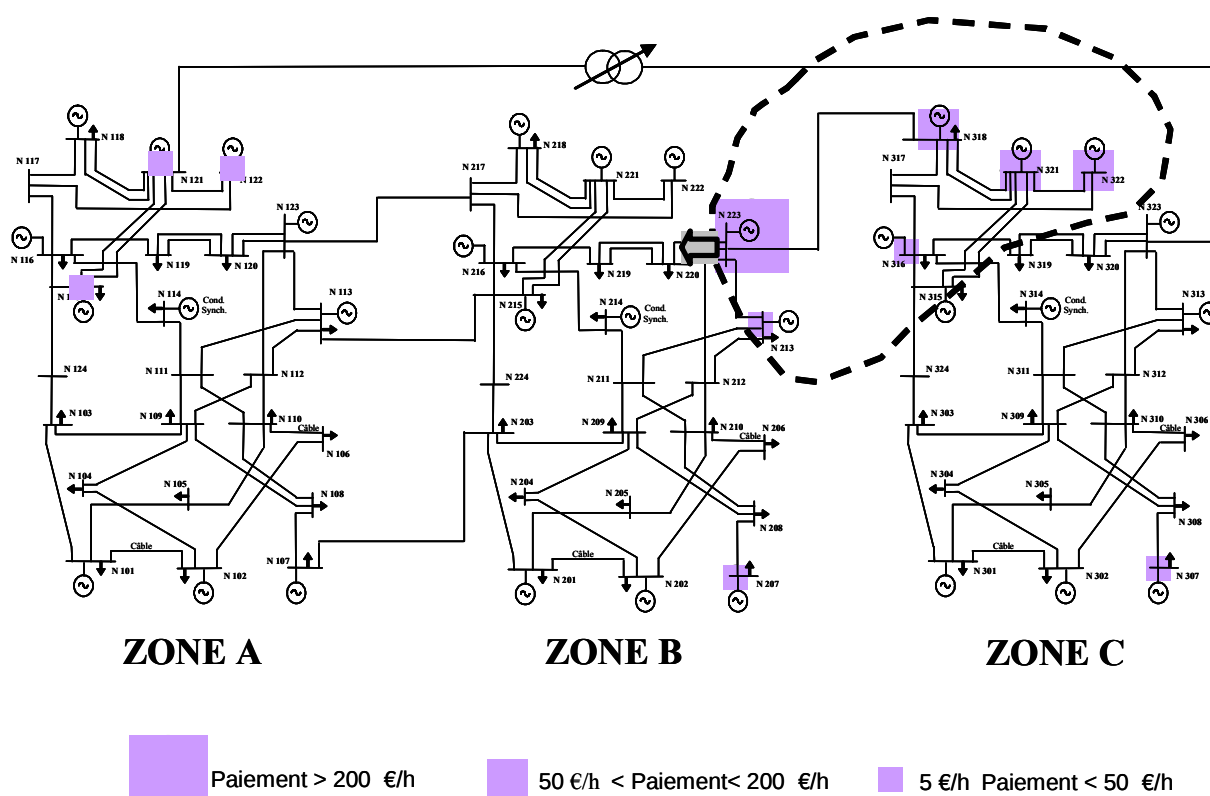


Figure A5.3 : allocation du coût de congestion aux producteurs suivant la contribution au transit (cas1)

A4.4) Allocation du coût de congestion pour le cas 2 (interconnexion 107-203)

A4.4.1) Allocation aux consommateurs

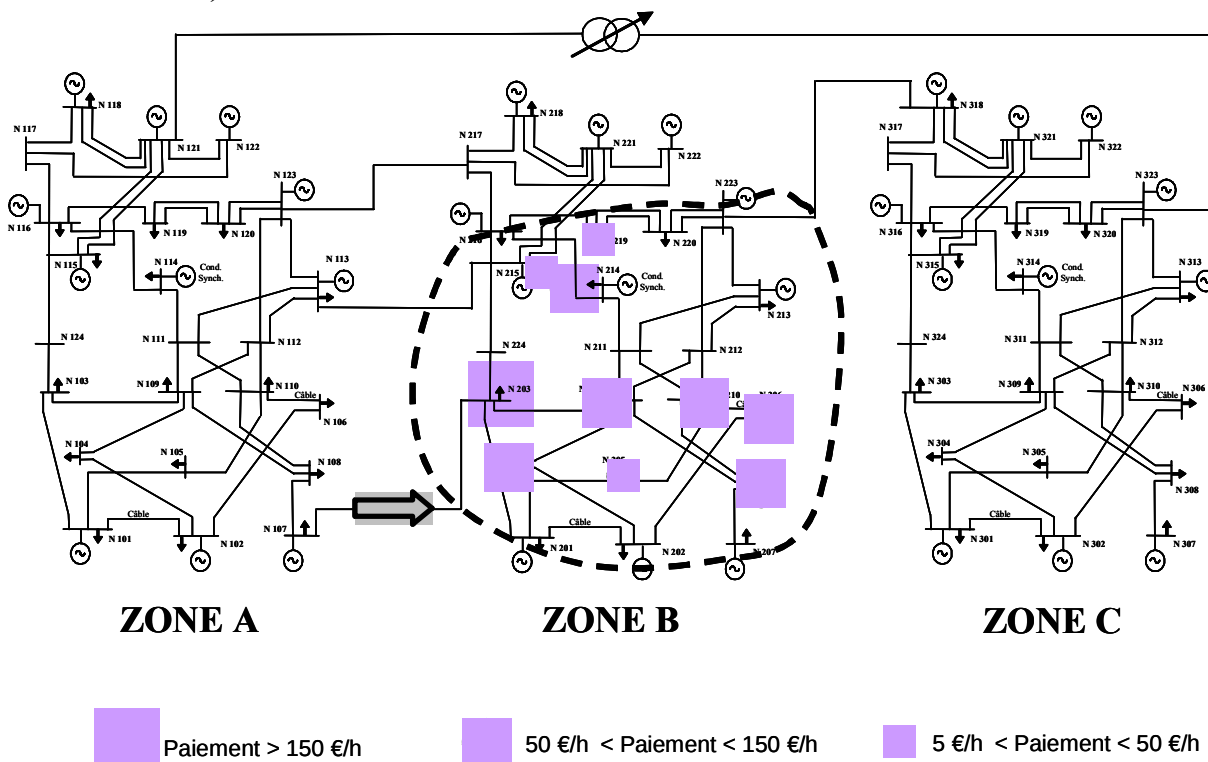


Figure A4.4 : allocation aux consommateurs par contribution au transit

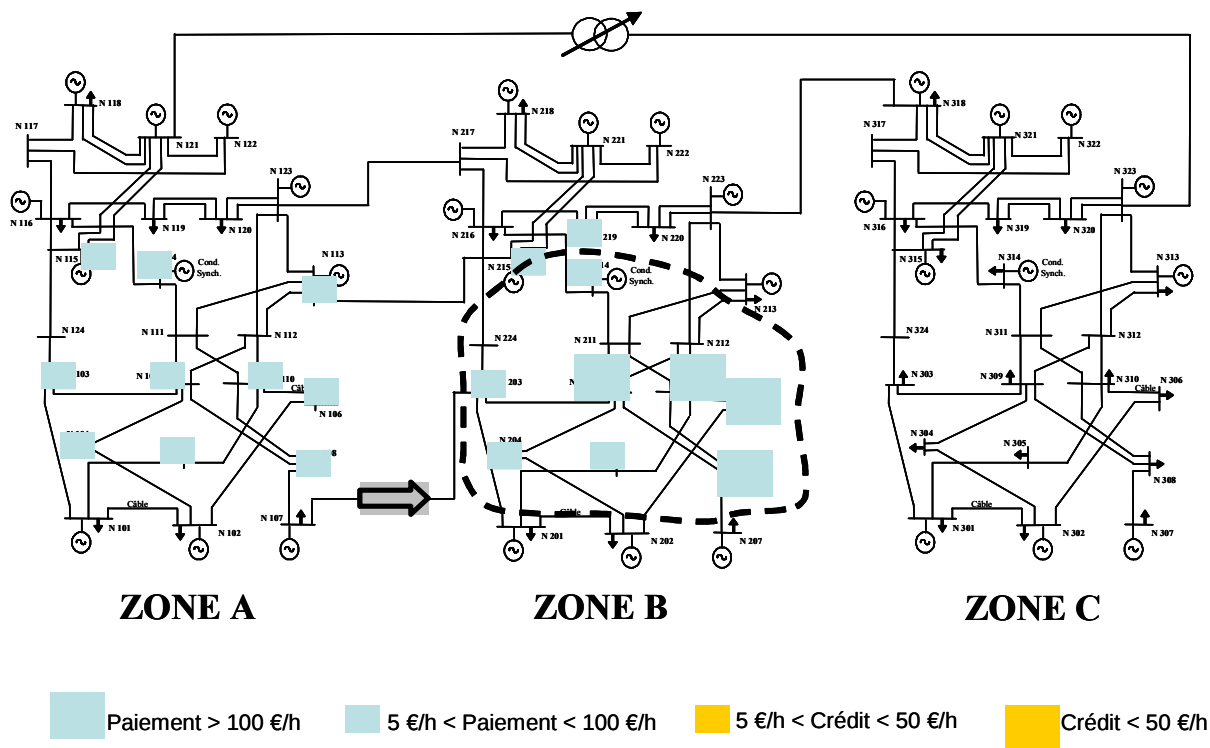


Figure A4.5 : allocation aux consommateurs suivant la destination des ajustements

ANNEXE 5

Les transformateurs déphaseurs

A5.1) Introduction

Une solution de long terme pour résoudre les problèmes de congestion serait d'accroître les capacités de transport par des renforcements au niveau de l'infrastructure du réseau de transport. Cependant, les contraintes environnementales font qu'il est de plus en plus difficile de construire de nouvelles liaisons ; aussi une solution alternative consiste à placer des transformateurs déphaseurs sur les lignes de transport susceptibles d'être congestionnées en vue de réguler leur transit et à optimiser l'usage des capacités de transport.

Le transformateur déphaseur permet, en injectant une tension série sur la liaison où il est connecté, de modifier le transit de puissance active traversant cette liaison. Il est alors possible de décharger cette liaison au profit d'autres liaisons moins chargées. Il constitue à l'image du réajustement des programmes de production un moyen de maîtriser la répartition des transits dans un réseau maillé.

Dans le cas français, plusieurs installations ont déjà vu le jour : à Pragnères, un transformateur déphaseur a été installé en 1998 sur la liaison 225 kV Pragnères-Biescas qui risquait de sévères surcharges en cas d'ouverture de l'une des interconnexions voisines. D'autres appareils du même type ont aussi été installés, comme à la Rance en Ille-et-Villaine, et enfin à la Praz en 2002 pour sécuriser la liaison France-Italie [SEE 02].

A5.2) Modélisation d'un transformateur déphaseur [GE 99]

Le modèle d'un transformateur déphaseur est montré par la Figure A5.1. Il consiste en une tension série insérée sur la ligne ij :

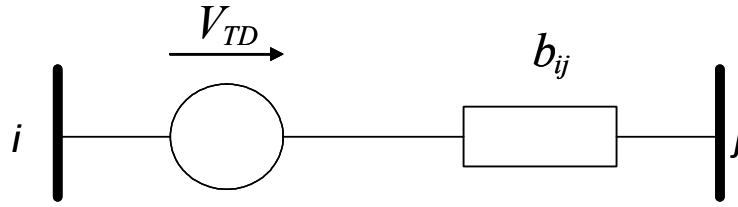


Figure A5.1 : modèle équivalent d'un transformateur déphaseur

Le paramètre de contrôle est ici le déphasage angulaire des tensions aux nœuds i et j .

Si nous considérons un déphasage variable et le modèle DC, nous pouvons écrire le transit de la ligne ij en fonction de l'écart angulaire de tension entre le nœud i et le nœud j , de l'admittance de ligne et du déphasage ψ introduit :

$$P_{ij} = b_{ij} (\theta_i - \theta_j + \psi) \quad (\text{A5.1})$$

On peut modéliser l'effet du déphasage par une injection équivalente à travers la ligne ij :

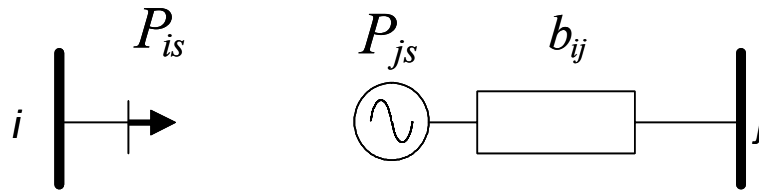


Figure A5.2 : injection équivalente du déphasage Ψ

Cette injection équivalente de puissance active s'exprime aisément en fonction du déphasage Ψ et de l'admittance de ligne b_{ij} :

$$P_{js} = b_{ij} \psi \quad (\text{A5.2})$$

$$P_{is} = -b_{ij} \psi$$

Prenons un exemple simple pour illustrer l'action d'un transformateur déphaseur : supposons une interface composée de deux lignes en parallèle traversant par une puissance totale de 120 MW (Fig. A5.3). La première ligne a une impédance de X u.r. et la deuxième ligne une impédance de $2X$. Il en résulte que la première ligne supporte un transit de 80 MW, tandis que la deuxième est traversée par un transit de seulement 40 MW.

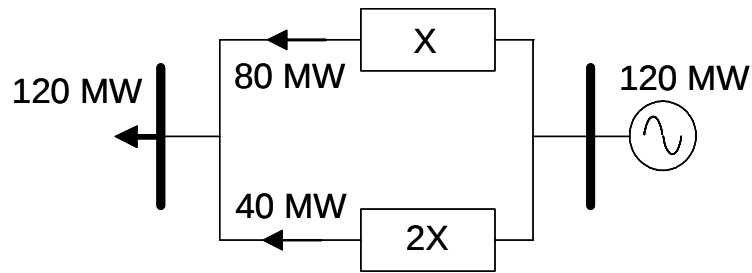


Figure A5.3 : interface sans transformateur déphaseur

En introduisant un transformateur déphaseur en série sur la première ligne, il est possible de rééquilibrer les flux de manière à ce que les deux lignes supportent le même transit de puissance :

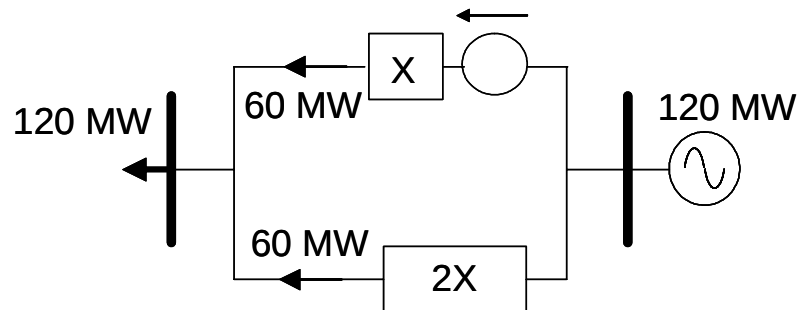


Figure A5.4 : redistribution des flux

Dans cet exemple, on a réduit le transit de 20 MW sur la première ligne, quantité qui s'est reportée sur la deuxième ligne moins chargée.

A5.3) Application sur le cas 1 du réseau RTS 96

La version la plus récente du réseau RTS 96 comprend un transformateur déphaseur inséré sur l'interconnexion 323-121 reliant la zone C à la zone A. Ce dernier peut constituer un bon moyen pour traiter la congestion du cas 1 située sur la ligne 220-223. Il permettrait en effet de détourner en partie le flux parallèle traversant le réseau de la zone C vers la zone A via la zone B en sollicitant plus l'interconnexion 323-121 dont la réserve en capacité disponible est encore grande. Ainsi, l'introduction d'un déphasage supplémentaire de 5° sur l'interconnexion 323-121 a pour effet de ramener le transit de la ligne 220-223 de 533 MW à 481 MW, ce qui est en dessous de sa limite de transit (500 MW). Les figures A5.5 et A5.6 nous montre les changements qu'entraîne l'introduction de ce déphasage sur les principaux transits du réseau :

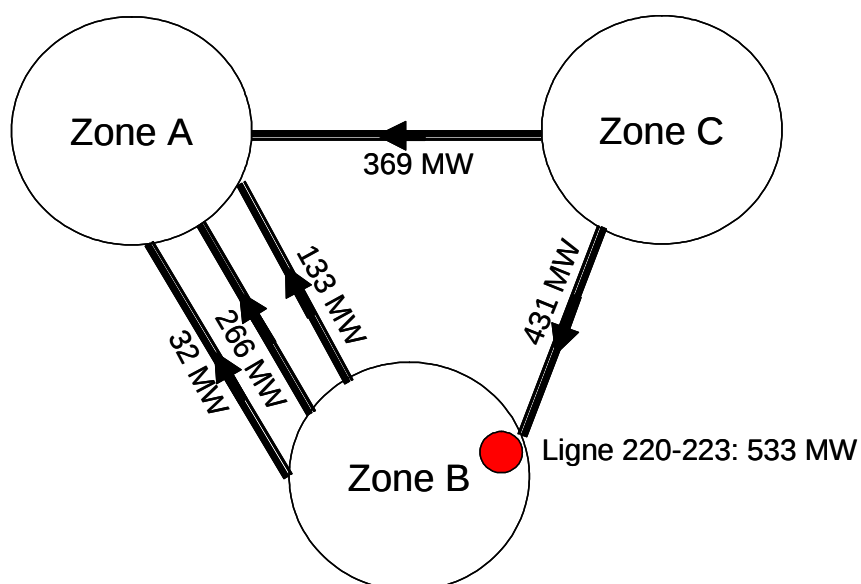


Figure A5.5 : état des transits avant actionnement du déphaseur

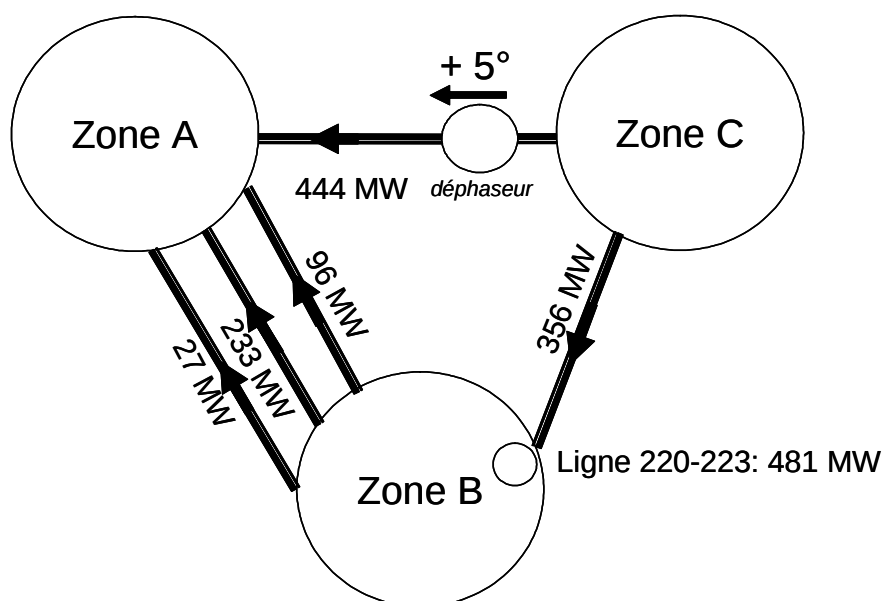


Figure A5.6 : état des transits après actionnement du déphaseur

ANNEXE 6

Preuve de la convergence de l'algorithme d'optimisation découplé [BAK 03]

Reprenons le cas où nous avons deux zones séparées par une interconnexion :

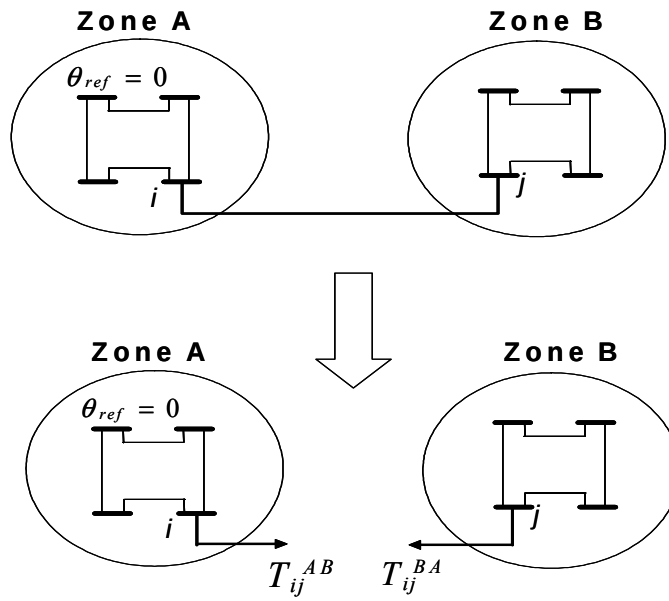


Figure A6.1 : principe de découplage de deux zones reliées par une interconnexion

Le problème d'optimisation reformulé pour ces deux zones peut s'écrire de cette façon :

$$\text{Min } f(\mathbf{x}_A) + f(x_B)$$

Avec

$$g_A(\mathbf{x}_A) \leq 0$$

$$g_B(\mathbf{x}_B) \leq 0$$

$$h_A(\mathbf{x}_A, \mathbf{x}_B) = 0$$

$$h_B(\mathbf{x}_A, \mathbf{x}_B) = 0$$

où g_A et g_B sont les contraintes égalités et inégalités propres à chaque zones, et h_A et h_B sont les contraintes de couplage qui sont associées aux multiplicateurs de Lagrange α_{ij}^{AB} et α_{ij}^{BA} de cette façon :

$$h_A(\mathbf{x}_A, \mathbf{x}_B) = \frac{1}{x_{ij}}(\theta_i - \theta_j) - T_{ij}^{AB} = 0$$

$$h_B(\mathbf{x}_A, \mathbf{x}_B) = \frac{1}{x_{ij}}(\theta_j - \theta_i) - T_{ij}^{BA} = 0$$

Ainsi, si on considère les conditions de KKT des deux sous-problèmes, on a :

$$\text{Min } f(\mathbf{x}_A) + \hat{\alpha}_{ij}^{BA} \left[\frac{1}{x_{ij}}(\hat{\theta}_j - \theta_i) - \hat{T}_{ij}^{BA} \right]$$

$$g_A(\mathbf{x}_A) \leq 0$$

$$h_A(\mathbf{x}_A, \hat{\mathbf{x}}_B) = 0$$

et

$$\text{Min } f(\mathbf{x}_B) + \hat{\alpha}_{ij}^{AB} \left[\frac{1}{x_{ij}}(\hat{\theta}_i - \theta_j) - \hat{T}_{ij}^{AB} \right]$$

$$g_B(\mathbf{x}_B) \leq 0$$

$$h_B(\hat{\mathbf{x}}_A, \mathbf{x}_B) = 0$$

Ces conditions de KKT combinées coïncident avec les conditions de KKT du problème global, d'où la convergence de l'algorithme découplé vers la solution globale du problème.